



(51) International Patent Classification:
G06F 3/06 (2006.01)

(21) International Application Number:
PCT/US2023/077225

(22) International Filing Date:
18 October 2023 (18.10.2023)

(25) Filing Language: English

(26) Publication Language: English

(30) Priority Data:
63/386,938 12 December 2022 (12.12.2022) US
18/449,428 14 August 2023 (14.08.2023) US

(71) Applicant: **WESTERN DIGITAL TECHNOLOGIES, INC.** [US/US]; 5601 Great Oaks Parkway, San Jose, California 95119 (US).

(72) Inventors: **YANG, Niles**; c/o Western Digital Technologies, Inc., 5601 Great Oaks Parkway, San Jose, California 95119 (US). **LINNEN, Daniel J.**; c/o Western Digital Technologies, Inc., 5601 Great Oaks Parkway, San Jose, California 95119 (US). **HAHN, Judah Gamliel**; c/o Western Digital Technologies, Inc., 5601 Great Oaks Parkway, San Jose, California 95119 (US).

(74) Agent: **ALTER, Scott M.** et al.; Michael Best & Friedrich LLP, (Western Digital), 790 N Water St, Suite 2500, Milwaukee, WI 53202 (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, CV, CZ, DE, DJ, DK, DM, DO, DZ, EC, EE, EG, ES, FI, GB, GD, GE, GH, GM, GT, HN, HR, HU, ID, IL, IN, IQ, IR, IS, IT, JM, JO, JP, KE, KG, KH, KN, KP, KR, KW, KZ, LA, LC, LK, LR, LS, LU, LY,

(54) Title: SEGREGATING LARGE DATA BLOCKS FOR DATA STORAGE SYSTEM

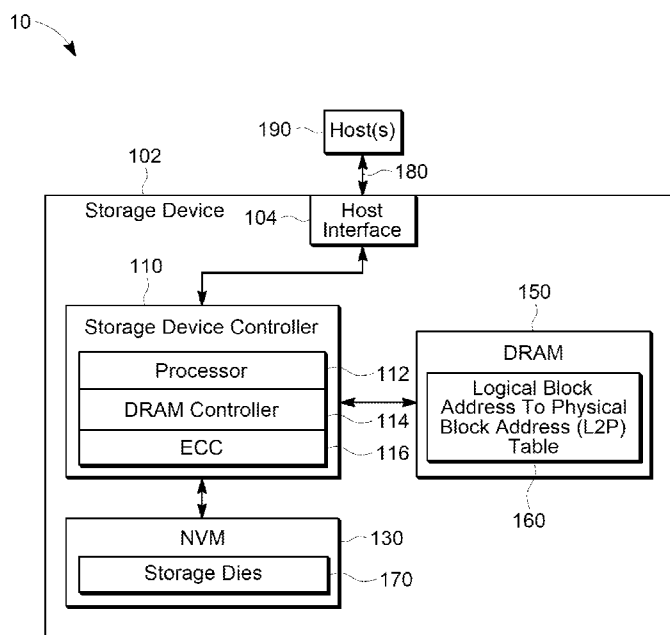


FIG. 1

(57) **Abstract:** Methods and apparatus for efficiently handling large data files and their updates in NAND memory. In one example, provided is a data-storage system configured to reduce the frequency of data relocations by segregating a large data file into a plurality of subfiles. The size of such subfiles is appropriately selected to reduce the probability of occurrence for host-relocation conflicts and the magnitude of write amplification, thereby enabling the data-storage system to provide better quality of service while substantially maintaining acceptable levels of other pertinent performance characteristics. In some examples, a sequence of host read-modify-write commands is handled by generating a copy of implicated subfiles in a data buffer, applying subfile updates to the copy in the data buffer in accordance with the sequence, and relocating the implicated subfiles in the NAND memory using the updated versions thereof from the data buffer.



MA, MD, MG, MK, MN, MU, 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, ST, SV, SY, TH,
TJ, TM, TN, TR, TT, TZ, UA, UG, US, UZ, VC, VN, WS,
ZA, ZM, ZW.

(84) Designated States (*unless otherwise indicated, for every kind of regional protection available*): ARIPO (BW, CV, GH, GM, KE, LR, LS, MW, MZ, NA, RW, SC, 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, ME, 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))*

SEGREGATING LARGE DATA BLOCKS FOR DATA STORAGE SYSTEM**CROSS-REFERENCE TO RELATED APPLICATION**

[0001] This application claims the benefit of and hereby incorporates by reference, for all purposes, the entirety of the contents of U.S. Nonprovisional Application No 18/449,428, entitled “SEGREGATING LARGE DATA BLOCKS FOR DATA STORAGE SYSTEM” and filed in the United States Patent & Trademark Office on August 14, 2023, which is claims priority of U.S. Provisional application number 63/386,938, filed December 12, 2022.

FIELD

[0002] This application relates generally to data storage devices, and more particularly but not exclusively, to handling updates of blocks of data having relatively large sizes.

BACKGROUND

[0003] This section introduces aspects that may help facilitate a better understanding of the disclosure. Accordingly, the statements of this section are to be read in this light and are not to be understood as admissions about what is in the prior art or what is not in the prior art.

[0004] Write amplification (WA) is an undesirable phenomenon associated with flash memory and solid-state drives (SSDs) due to which the actual amount of information physically written to the storage media can be significantly larger than the logical amount of data intended to be written. For example, because the erase operation has a coarser granularity than the write operation, a process of performing erase and write operations on user data typically results in moving (or rewriting) user data and metadata more than once. As a result, in some examples, rewriting a block of data causes an already-used-portion of flash memory to be read, updated, and written to a new location, which is preceded by erasing the new location if that location was previously used at some point in time. Due to different respective granularities of the erase and write operations, larger portions of flash are typically erased and rewritten compared to the amount of new data. This granularity mismatch effectively increases the number of writes, which can shorten the lifespan of the corresponding data storage device. In addition, the increased volume of writes can consume bandwidth, thereby detrimentally affecting the throughput performance of the data storage device.

SUMMARY

[0005] Disclosed herein are various embodiments of methods and apparatus for efficiently handling large data files and their updates in NAND memory. In one example, provided is a data-storage system configured to reduce the frequency of data relocations by segregating a large data file into a plurality of subfiles. The size of such subfiles is appropriately selected to reduce the probability of occurrence for host-relocation conflicts and the magnitude of write amplification, thereby enabling the data-storage system to provide better quality of service (QoS)

while substantially maintaining acceptable levels of other pertinent performance characteristics. In some examples, a sequence of host read-modify-write commands is handled by generating a copy of implicated subfiles in a data buffer, applying subfile updates to the copy in the data buffer in accordance with the sequence, and relocating the implicated subfiles in the NAND memory using the updated versions thereof from the data buffer.

[0006] According to an example embodiment, provided is a data storage device, comprising: a nonvolatile memory to store data; a second memory to store a logical-to-physical (L2P) table and a subfile mapping table; and a controller coupled to the nonvolatile memory and the second memory and configured to: in response to a host command to write a large data file to the nonvolatile memory, segregate the large data file into a plurality of subfiles, the large data file having a file size larger than a first fixed size, each subfile of the plurality of subfiles having a respective size smaller than or equal to the first fixed size; store a first mapping of the plurality of subfiles in the subfile mapping table, the first mapping including, for each subfile of the plurality of subfiles, a respective word-line physical address; cause the plurality of subfiles to be written to the nonvolatile memory based on the first mapping; and access the L2P table and the subfile mapping table to control a memory operation on the large data file in the nonvolatile memory.

[0007] According to another example embodiment, provided is a method performed by a data storage device, the method comprising: in response to a host command to write a large data file to a nonvolatile memory, segregating, via a controller, the large data file into a plurality of subfiles, the large data file having a file size larger than a first fixed size, each subfile of the plurality of subfiles having a respective size smaller than or equal to the first fixed size; storing, via the controller, a first mapping of the plurality of subfiles in a subfile mapping table, the first mapping including, for each subfile of the plurality of subfiles, a respective word-line physical address; causing, via the controller, the plurality of subfiles to be written to the nonvolatile memory based on the first mapping; and accessing, via the controller, an L2P table and the subfile mapping table to control a memory operation on the large data file in the nonvolatile memory.

[0008] According to yet another example embodiment, provided is a data storage device, comprising: means for segregating a large data file into a plurality of subfiles in response to a host command to write the large data file to a nonvolatile memory, the large data file having a file size larger than a first fixed size, each subfile of the plurality of subfiles having a respective size smaller than or equal to the first fixed size; means for storing a first mapping of the plurality of subfiles in a subfile mapping table, the first mapping including, for each subfile of the plurality of subfiles, a respective word-line physical address; means for causing the plurality of

subfiles to be written to the nonvolatile memory based on the first mapping; and means for accessing an L2P table and the subfile mapping table to control a memory operation on the large data file in the nonvolatile memory.

[0009] Various aspects of the present disclosure provide for improvements in data storage devices, for example, optimizing the processes in which host device inputs/outputs are handled by data storage devices. The present disclosure can be embodied in various forms, including hardware or circuits controlled by software, firmware, or a combination thereof. The foregoing summary is intended solely to give a general idea of various aspects of the present disclosure and does not limit the scope of the present disclosure in any way.

BRIEF DESCRIPTION OF THE DRAWINGS

[0010] FIG. 1 is a block diagram illustrating a data-storage system in which various embodiments may be practiced.

[0011] FIG. 2 is a block diagram illustrating a memory layout that can be used in the data-storage system of FIG. 1 according to an embodiment.

[0012] FIG. 3 is a block diagram illustrating an example subfile mapping table that can be used in the data-storage system of FIG. 1 according to an embodiment.

[0013] FIG. 4 is a flowchart illustrating a host write operation performed in the data-storage system of FIG. 1 according to an embodiment.

[0014] FIG. 5 is a flowchart illustrating read-modify-write operations performed in the data-storage system of FIG. 1 according to an embodiment.

[0015] FIGS. 6-9 are block diagrams illustrating data layouts in the memory layout of FIG. 2 according to various examples.

DETAILED DESCRIPTION

[0016] In the following description, numerous details are set forth, such as data storage device configurations, controller operations, and the like, in order to provide an understanding of one or more aspects of the present disclosure. It will be readily apparent to one skilled in the art that these specific details are merely exemplary and not intended to limit the scope of this application. In particular, the functions associated with the controller can be performed by hardware (for example, analog or digital circuits), a combination of hardware and software (for example, program code or firmware stored in a non-transitory computer-readable medium that is executed by a processor or control circuitry), or any other suitable means. The following description is intended solely to give a general idea of various aspects of the present disclosure and does not limit the scope of the disclosure in any way. Furthermore, it will be apparent to those of skill in the art that, although the present disclosure refers to NAND flash, the concepts

discussed herein may be applicable to other types of solid-state memory, such as NOR, PCM (“Phase Change Memory”), ReRAM, etc.

[0017] FIG. 1 is a block diagram illustrating a data-storage system 10 in which example embodiments can be practiced. System 10 comprises a data storage device 102 connected to a host device 190 by way of a communication path 180. In an example embodiment, communication path 180 can be implemented using an electrical bus, a wireless connection, or any other suitable data link. Data storage device 102 can be a flash memory device, e.g., an SSD, a flash memory, or other suitable nonvolatile memory (NVM).

[0018] In some embodiments, data storage device 102 may be embedded within host device 190. In some other embodiments, data storage device 102 may be removable from host device 190, e.g., may be removably coupled to host device 190 in accordance with a removable universal serial bus (USB) configuration. In some embodiments, data storage device 102 may be used as an embedded storage drive, e.g., a mobile embedded storage drive, an enterprise storage drive (ESD), a client storage device, a cloud storage drive, or other suitable data storage device.

[0019] As shown in FIG. 1, data storage device 102 comprises a host-device interface 104, an electronic controller 110, an NVM 130, and a volatile memory (e.g., a DRAM) 150. In operation, host-device interface 104 enables communications between data storage device 102 and host device(s) 190. Such communications may include, *inter alia*, transmission of data between NVM 130 and host device(s) 190. NVM 130 comprises storage dies 170, which may include any one type or any suitable combination of NAND flash devices, NOR flash devices, and other suitable nonvolatile memory devices. Storage dies 170 may be organized into channels, each of the channels being based on a corresponding bus, e.g., an 8-bit bus, connecting the corresponding subset of storage dies 170 to controller 110. Individual ones of storage dies 170 may further be organized into a hierarchy of planes, blocks, and pages. NVM 130 and/or individual ones of the storage dies 170 thereof may also include support circuitry (not explicitly shown in FIG. 1), such as read and write circuitry. Such read/write circuitry may be implemented in a single component or may be divided into separate components, such as a read-circuitry component and a separate write-circuitry component. In an example embodiment, DRAM 150 is used, *inter alia*, to store an L2P table 160. In operation, data storage device 102 uses L2P table 160 to translate logical addresses of input/output (I/O) requests to corresponding flash-memory addresses. The layer that performs the translation is referred to as the flash translation layer (FTL).

[0020] Controller 110 includes components, such as circuits, firmware, and software, that bridge NVM 130 to host-device interface 104, with only some of such components being explicitly shown in FIG. 1 for better clarity. For example, controller 110 may include: (i) an

embedded processor 112; (ii) an electrically erasable firmware read-only memory (ROM, not explicitly shown in FIG. 1); (iii) a DRAM controller 114; (iv) an error-correction code (ECC) circuit or module 116; and (v) a flash component interface (not explicitly shown in FIG. 1). In some embodiments, controller 110 may also incorporate DRAM 150 or other functionally similar volatile-memory having stored therein the L2P table 160.

[0021] Processor 112 is configured to support, e.g., some or all of the following operations: wear leveling, bad-block management, data scrambling, error-correction coding, garbage collection, trim, address mapping, and other pertinent operations. DRAM controller 114 operates as an electronic controller of DRAM 150. ECC entity 116 may typically include two submodules, with a first of the submodules running a first ECC applied to data stored in NVM 130, and with a second of the submodules running a different second ECC applied to data stored in DRAM 150, including the data of L2P table 160 and other tables and/or control-data structures that may be stored therein.

[0022] Host device 190 may address data stored in NVM 130 using a logical address. However, the data are stored in NVM 130 at physical addresses. The L2P table 160 is used to map from the logical address used by host device 190 to the physical address where the data are actually stored in NVM 130. When host device 190 requests to read data, controller 110 may obtain the physical address from L2P table 160 to access the requested data. When host device 190 requests to write or update data in NVM 130, controller 110 may update L2P table 160 accordingly.

[0023] A master L2P table may be stored in NVM 130 but, to be efficiently used or updated by controller 110, at least a pertinent portion of the master L2P table is stored in DRAM 150, e.g., in L2P table 160. When L2P table 160 is updated, the updates need to be written back into the master L2P table stored in NVM 130. During the design process of data storage device 102, it may be stipulated whether DRAM 150 will be specified to hold a copy of the whole master L2P table for the entire logical address space. A decision to not hold the whole master L2P table in DRAM 150 may be due to, e.g., a size constraint of DRAM 150. Similar design decisions may also be made with respect to other tables that may be used instead of, in addition to, or in conjunction with L2P table 160 according to various embodiments.

[0024] In some examples of page-level L2P mapping, a physical address may include a channel number, a die number, a block number, and a page number. For a 4-KB page size, a 32-TB NVM 130 may require approximately 32 GB of such L2P entries. With the added ECC overhead, this volume of L2P entries can be accommodated, e.g., by a 40-GB DRAM 150. This estimate of the DRAM capacity scales approximately linearly with the data capacity of NVM 130.

[0025] Although FIG. 1 shows controller 110 as being connected to and employing DRAM 150, alternative embodiments are also possible. For example, in one such embodiment, controller 110 may employ, for similar functions, a built-in cache memory, e.g., a static random-access memory (static RAM or SRAM), which is a part of the controller circuitry. In another embodiment, instead of DRAM 150, controller 110 may use a dedicated portion of NVM 130 as storage for L2P table 160 and any associated mapping tables.

[0026] In a data center employing data-storage system 10, data files with the size in the range between 1 GB and 10 TB are often encountered. In operation, multiple terabytes of data can typically be accessed and rewritten by multiple users at a given pace. With the 4 KB read/write granularity, frequent changes from various users on the same large files may lead to frequent read-modify-write (RMW) operations, which are sometimes characterized by significant and undesirable effects of WA and the corresponding deterioration of the QoS metrics. In some examples, such frequent changes on the same files are due to the host data model training and/or data iteration at the same memory address. Another undesirable impact of WA is on the operation of the Data Processing Unit (DPU) wherein data access slowdowns may result in further performance and/or QoS constraints on the corresponding data-storage system 10.

[0027] At least some of the above-indicated and possibly some other related problems in the state of the art can beneficially be addressed using various examples, aspects, features, and/or embodiments disclosed herein. In some examples, data-storage system 10 is configured to reduce the frequency of data relocations by segregating large data files into respective pluralities of subfiles and then using the subfiles for updates associated with the host RMW operations. The size of subfiles is appropriately selected to beneficially reduce the probability of occurrence for host-relocation conflicts and the magnitude of WA, and thereby, data-storage system 10 provides better QoS while substantially maintaining acceptable levels of other pertinent performance characteristics.

[0028] Herein, RMW is a class of atomic operations that both read a memory location and write a new value into the memory location, either with a completely new value or some function of the previous value. RMW operations can be used to, e.g., read a specific page or modify a few words, bytes, or bits in a page. A read command of an RMW operation causes the desired data to be copied from the flash memory to a data buffer. The modify command of the RMW operation then causes the data in the data buffer to be modified as needed. A write command of the RMW operation causes the modified data to be written back to the flash memory. In some examples, the RMW operates on a minimum write unit, which can be smaller than a page. In some additional examples, multiple pages are programmed at once.

[0029] Some embodiments disclosed herein implement a tradeoff between write amplification and performance. For example, smaller map unit sizes may result in reduced write amplification, as less data is being read, modified, and then rewritten on the NAND. In contrast, increasing the size of the map unit in some embodiments may increase write amplification somewhat, e.g., by rewriting more information on every modification/write by the host to a special purpose data area, at the cost of endurance, but with the benefit of better performance at some specific times. It should also be noted that write amplification is the quotient resulting from the total writes to the NAND divided by the host writes to the device. As such, write amplification may be affected by additional factors beyond the above-mentioned granularity mismatches, with the example additional factors including media management policies related to read disturb, data retention, error mitigation, redundancy, metadata, etc.

[0030] At least some approaches disclosed herein are transferrable to hard disk drives (HDDs).

[0031] FIG. 2 is a block diagram illustrating an example memory layout 200 that can be used in data storage device 102 according to an embodiment. In other embodiments, other layouts can also be used. At least some example values mentioned herein below in reference to FIG. 2 may be different in other examples or embodiments of memory layout 200.

[0032] According to memory layout 200, a first meta block, labeled 202₁, has memory blocks in a first physical die, labeled Die 0. A second meta block, labeled 202₂, has memory blocks in a second physical die, labeled Die 1. Taken together, meta blocks 202₁ and 202₂ form a jumbo block 204. Each of the physical dies Die 0 and Die 1 is partitioned into four respective logical dies, labeled Logical Die 0 through Logical Die 3, respectively. Logical Dies 0, 1, 2, 3 can be accessed using channels 0, 1, 2, 3, respectively. Each of the logical dies has four respective planes, labeled Plane 0 through Plane 3, respectively. Each of the planes is connected to word lines WL0, WL1, ..., WL19. Each word line (WL) typically includes a plurality of (for example, four) respective string lines (not explicitly shown in FIG. 2). Each of the planes includes twenty respective WL blocks, one of which is labeled 210 in FIG. 2. Memory layout 200 has a total of 640 WL blocks 210. Each of meta blocks 202₁ and 202₂ has 320 respective WL blocks 210.

[0033] In one illustrative example, WL block 210 comprises triple-level cells (TLCs), each capable of storing three bits. In this example, WL block 210 has a size of 48 KB and can store twelve 4-KB pages. Each row of WL blocks connected to the same word line WL_n ($n=0, 1, \dots, 19$) in the meta block 202_{*i*} ($i=1, 2$) has a size of 768 KB. Each row of WL blocks connected to the same word line WL_n in the jumbo block 204 has a size of 1536 KB. The read/write granularity in memory layout 200 is one page. The erase granularity in memory layout 200 is one plane. In some other examples, more than two meta blocks 202 can be logically combined to form a corresponding jumbo block. In various examples, a WL block may comprise various

types of NAND memory cells, such as single-level cells (SLCs), multi-level cells (MLC), TLCs, or quad-level cells (QLCs). In such various examples, a corresponding WL block has a corresponding size that may be different from 48 KB.

[0034] In an example embodiment, data storage device 102 operates to segregate a large data file into a plurality of WL-level subfiles. The host may indicate the size of the input file to be written to NVM 130 using the Non-Volatile Memory Express (NVMe) protocol or an equivalent logical storage interface protocol. The data storage device may indicate a “large” file using Optimal Read/Write/Deallocate Size parameters as described in storage interface protocols. It should be noted that the host may use multiple write commands or operations to transfer a “large” file, each of which is individually smaller than the “large” file size. Herein, a “large” file is a file having a size that is larger than the combined size of the WL blocks (such as 210, FIG. 2) connected to the same word line WL_n in meta block 202_i or in jumbo block 204. In the above illustrative example, a file that is “large” in this sense has a size larger than the meta-block word-line (MBW) capacity of 768 KB or the jumbo-block word-line (JBW) capacity of 1536 KB.

[0035] In some examples, the MBW capacity is not limited to a single physical die and can similarly be distributed over multiple physical dies, such as two physical dies or four physical dies. A corresponding JBW capacity may thus comprise, e.g., four MBW capacities in four times more physical dies. In some examples, a single physical die may have multiple planes, such as two, four, or more planes, e.g., as indicated in FIG. 2. When a physical die has only one plane or two planes, the MBW capacity may span more physical dies. In contrast, when a physical die has four or more planes, then the MBW capacity may span fewer physical dies, and with the JBW capacity comprising more units of the MBW capacity. In various examples, MBW sizing may also depend on the specific NAND die architecture used in the corresponding NVM.

[0036] For a given large data file from host 190, a mapping table is defined for every such MBW/JBW at the FTL level by specifying the MBW/JBW-sized subfiles of that data file. In other words, a large data file is logically partitioned into MBW (or JBW) subfiles, with the corresponding L2P mappings stored in the controller memory. For any subfile that is smaller than the MBW or JBW capacity, a plurality of such smaller subfiles from different original data files can be combined to form a “virtual” subfile that fills up the MBW or JBW capacity. In some examples, to maintain an acceptable level of performance, such virtual subfiles are written to a memory partition having SLCs therein. When concurrent operations (e.g., including RMW operations) are needed for two or more host files associated with a virtual subfile, the SLC partition is typically more conducive to being able to efficiently handle concurrent operations. Data storage device 102 further operates to generate L2P mappings for various MBW (or JBW)

subfiles of large data files and store the generated mappings in L2P table 160 and/or in a supplemental subfile mapping table (e.g., 300, FIG. 3) linked with the L2P table 160.

[0037] FIG. 3 is a block diagram illustrating an example subfile mapping table 300 according to an embodiment. For illustration purposes and without any implied limitations, only a portion of the table 300 corresponding to two host files is explicitly shown in FIG. 3. The host file identifiers (IDs) of those two files are listed in column 302 of table 300. The first host file has the host file ID=0 and the size of 1936 KB. The second host file has the host file ID=1 and the size of 1136 KB. The first host file (host file ID=0) is segregated into three subfiles having the subfile IDs 0, 1, and 2, respectively, which are listed in column 304 of table 300. The sizes of those three subfiles are 768, 768, and 400 KB, respectively, which are listed in column 306 of the table 300. The second host file (host file ID=1) is segregated into two subfiles having the subfile IDs 0 and 1, respectively, which are listed in column 304 of the table 300. The sizes of those two subfiles are 368 and 768 KB, respectively, which are listed in column 306 of the table 300.

[0038] The example table 300 shown in FIG. 3 is an MBW mapping table, with the corresponding attributes of the MBW mapping being specified in columns 308-316 of the table. Based on the provided description, a person of ordinary skill in the pertinent art will be able to make and use a similar JBW mapping table without any undue experimentation. In the example shown, the columns 308, 310, 312, 314, and 316 list MBW IDs, MBW offsets, MBW physical addresses, MBW cell types, and MBW logical die IDs, respectively. Note that the subfiles having the (Host File ID, Subfile ID) values of (0, 2) and (1, 0) form a virtual MBW subfile. The size of the virtual subfile is 768 KB (= 400 KB + 368 KB). The two portions of that virtual MBW subfile are located at different respective MBW offsets of 0 and 100, respectively, as specified in MBW offset column 310. The corresponding MBW physical address of the two portions is the same, i.e., 100h, as specified in the MBW physical address column 312. Also note that the virtual MBW subfile is located in an SLC partition for the above-indicated reasons, whereas the other listed MBW subfiles are located in TLC partitions. The logical dies with SLCs therein have the MBW logical die IDs = 8, 9, 10, and 11, as specified in column 316 of table 300. The logical dies with TLCs therein have the MBW logical die IDs = 0, 1, 2, 3, 4, 5, 6, 7, 12, 13, 14, and 15, as further specified in column 316 of table 300.

[0039] During an RMW operation, the data corresponding to an MBW can be changed in a portion thereof having any size between the read/write granularity (which is 4 KB in the above example) and the full MBW capacity (which is 768 KB in the above example). In one example, data storage device 102 is configured to perform data modifications using a data buffer physically located, e.g., in DRAM 150 or in an internal RAM of controller 110. Thus, the

corresponding portion of the data buffer is at least 768 KB in size. In another example, data storage device 102 is configured to perform data modifications using a host memory buffer (HMB) located in host 190. Once the corresponding portion of the RAM is full or another request for using the same RAM resource has come in, MBW relocation can take place as part of an RMW operation. The latter relocation can be implemented, e.g., using a regular relocation thread of data storage device 102. It should also be noted that, when multiple MBWs are undergoing RMW operations at substantially the same time, the corresponding relocations can be such that a new (i.e., the relocation target) MBW may be configured to contain multiple 4 KB units from several original (i.e., the relocation source) MBWs. In such cases, the corresponding subfile mapping table, such as table 300, is updated accordingly, e.g., similar to the table updates performed for the write operations initiated by host 190.

[0040] FIG. 4 is a flowchart illustrating a method 400 of performing a host write operation in data storage device 102 according to an embodiment. Method 400 includes the controller 110 receiving a write command from host 190 (in block 402). The received write command specifies, *inter alia*, the size of the data file to be written to NVM 130.

[0041] Method 400 also includes the controller 110 determining whether the data file to be written to NVM 130 is a “large” data file (in decision block 404). In one embodiment, such determination in decision block 404 is made by comparing the file size with the MBW capacity, which is 768 KB in the above example. In another embodiment, such determination in decision block 404 is made by comparing the file size with the JBW capacity, which is 1536 KB in the above example.

[0042] When the file size is larger than the MBW (or JBW) capacity (“Yes” at decision block 404), method 400 includes the controller 110 segregating the data file into a plurality of subfiles (in block 406). The plurality of subfiles includes one or more subfiles of a size equal to the MBW (or JBW) capacity. For some data files, the plurality of subfiles also includes at least one subfile of a size smaller than the MBW (or JBW) capacity. As already indicated above, the smaller subfiles from different host data files can be used to form a virtual subfile that fills the corresponding MBW or JBW to capacity.

[0043] Method 400 also includes the controller 110 generating an MBW (or JBW) mapping of the subfiles defined in block 406 and further includes the controller 110 generating an MBW (or JBW) mapping table detailing the generated MBW (or JBW) mapping (in block 408). In some examples, the generated MBW mapping table is similar to table 300 (FIG. 3). Method 400 also includes the controller 110 storing the generated MBW (or JBW) mapping table in DRAM 150 (in block 408).

[0044] When the file size is smaller than or equal to the MBW (or JBW) capacity (“No” at decision block 404), method 400 includes the controller 110 mapping the whole host data file to a physical address in NVM 130 (in block 410). Method 400 also includes the controller 110 updating L2P table 160 in DRAM 150 to store therein the generated mapping (in block 410).

[0045] Method 400 also includes the controller 110 configuring the corresponding write circuitry of data storage device 102 to write the host data file to NVM 130 based on the applicable mapping (in block 412). When the operations preceding block 412 are the operations of block 408, the applicable mapping is the corresponding MBW (or JBW) mapping. When the operations preceding block 412 are the operations of block 410, the applicable mapping is the corresponding conventional L2P mapping.

[0046] FIG. 5 is a flowchart illustrating a method 500 of performing RMW operations in data storage device 102 according to an embodiment. In some examples, method 500 can be used, e.g., after method 400 (FIG. 4), wherein a large host data file is segregated into subfiles, with the corresponding MBW (or JBW) mapping specified in table 300 or a functionally similar table. For illustration purposes and without any implied limitations, method 500 is described below in reference to MBW subfiles.

[0047] Method 500 includes the controller 110 receiving a sequence of RMW commands from host 190 for a data file stored in NVM 130 and copying one or more implicated MBW subfiles from NVM 130 to a data buffer (in block 502). For the above-described specific example, the data-buffer space allocated to each of the copied MBW subfiles is at least 768 KB. Method 500 also includes applying updates tracing the RMW operations of the RMW-command sequence to the copy of the MBW subfiles in the data buffer (in block 504). The RMW operations of the sequence are not applied in block 504 to the source MBW subfiles stored in NVM 130.

[0048] Method 500 also includes determining (in decision block 506) whether or not the data buffer space allocated to the copy of the MBW subfiles is to be re-allocated for a different purpose. In a representative example, such reallocations are governed by the applicable buffer-management policies. When it is determined that the re-allocation is not yet coming (“No” in decision block 506), the processing of method 500 is looped back to continue performing operations of block 504. Otherwise (“Yes” in decision block 506), the processing of method 500 is advanced forward to perform operations of block 508.

[0049] Method 500 includes operations directed at relocating the implicated MBW subfiles to new respective MBW-sized locations in NVM 130 (in block 508). In one example, controller 110 configures the corresponding write circuitry of data storage device 102 to write a current copy (i.e., a copy updated via operations of block 504) of an MBW subfile stored in the data buffer to a new MBW-sized location in NVM 130. The controller 110 may also mark the data of

that MBW subfile stored in the previous location in NVM 130 as invalid. In some examples, the controller 110 may also generate a new virtual MBW subfile, wherein the constituent smaller subfiles are from a set of host data files that is different from the set of host data files in the corresponding previous virtual MBW subfile.

[0050] Method 500 also includes the controller 110 updating the MBW mapping table (in block 510). More specifically, the controller 110 operates (in block 510) to perform updates of the MBW mapping table to log therein the MBW subfile mapping changes corresponding to the operations performed in block 508 of method 500. For example, in table 300 shown in FIG. 3, the table entries for the MBW subfiles having the (Host File ID, Subfile ID) values of (0, 0) and (0, 1) will be updated in block 510 to change the respective old MBW physical addresses listed in column 312 of table 300 to the new respective MBW physical addresses when those MBW subfiles are relocated in NVM 130 due to operations performed in block 508 of method 500. Other entries in table 300 may also be appropriately updated to reflect the corresponding operations performed in block 508 of method 500.

[0051] FIGS. 6-9 are block diagrams illustrating data layouts in jumbo block 204 according to various examples. Different data layouts can be produced, e.g., with different respective scheduling algorithms used to implement operations in blocks 408 and 410 of method 400. In the notation used in FIGS. 6-9, “JPC N_1 - N_2 ” denotes WL-block-sized portions of JBW subfiles, wherein the number N_1 is the JBW-subfile identifier, and the number N_2 is the subfile-portion identifier. An example WL block is block 210 shown in FIG. 2. JBW subfiles are mapped to respective portions of jumbo block 204 using operations of block 408 of method 400. “Normal” denotes a WL block containing data of “small” files mapped thereto using operations of block 410 of method 400.

[0052] FIG. 6 illustrates an example data layout produced with forced ordering. In a representative example, forced ordering is used to ensure that the data are laid down evenly and orderly across jumbo block 204 such that host or garbage collection writes are not interleaved with JBW-subfile writes. In the example of FIG. 6, forced ordering results in: (i) three JBW subfiles having been written to word lines WL0, WL1, and WL2, respectively, and (ii) conventional writes having been written to word lines WL3, WL4, and WL5. In one example, forced ordering is achieved by placing a higher priority on writing JBW subfiles and selecting a priority threshold such as to prevent conventional writes from being interleaved with JBW-subfile writes. Additionally, dies can be prevented from getting ahead of one another in terms of the programmed volume by setting the value of allowances for how far a die can be ahead of another die to exactly zero.

[0053] FIG. 7 illustrates another example data layout resulting from conventional host or garbage collection writes being allowed to interleave with JBW-subfile writes. In the example of FIG. 7, the interleaving of such writes results in: (i) three JBW subfiles having been written to word lines WL0, WL2, and WL5, respectively, and (ii) conventional writes having been written to word lines WL1, WL3, WL4, and WL6. In some examples, correlated parallelism and ordering among different dies can still be enforced in a manner similar to that used in the example of FIG. 6 so that there is little or no detrimental impact of the interleaving on the performance. Breaks between JBW-subfile writes are provided to accommodate the queued conventional writes and/or other activity.

[0054] FIG. 8 illustrates yet another example data layout resulting from writes being scheduled in the first in/first out (FIFO) order of WL blocks. In the example of FIG. 8, the FIFO scheduling results in each of the three JBW subfiles having respective portions thereof on two different word lines. More specifically, the WL blocks JPC0-N₂ (where N₂ = 0, 1, ..., 31) of the first (N₁ = 0) JBW subfile are spread over word lines WL0 and WL1. The WL blocks JPC1-N₂ (where N₂ = 0, 1, ..., 31) of the second (N₁ = 1) JBW subfile are spread over word lines WL1 and WL2. The WL blocks JPC2-N₂ (where N₂ = 0, 1, ..., 31) of the third (N₁ = 2) JBW subfile are spread over word lines WL2 and WL3. Even with such spread, the level of parallelism can typically be relatively high. Relatively more-favorable data layouts can be achieved by doing this type of FIFO scheduling during periods of relatively low activity. With not allowing any die to program further than other dies and/or prioritization of JPC writes, the occurrence of the data layout pattern illustrated in FIG. 8 is made more likely.

[0055] FIG. 9 illustrates yet another example data layout resulting from FIFO scheduling wherein some dies are allowed to program further than other dies, and the system is relatively busy with mixed workloads. In the example of FIG. 9, such scheduling results in: (i) the logical dies 1, 3, 5, and 7 having been programmed down to WL3; (ii) the logical dies 0 and 6 having been programmed down to WL4; and (iii) the logical dies 2 and 4 having been programmed down to WL5.

[0056] In some examples, for power-loss safety and/or performance enhancement, an update of a small portion of a JBW subfile is scheduled to be done first, which then triggers a corresponding rewrite of a larger data segment at a later time to improve the level of parallelism. In such examples, the following example sequence of events may occur:

- (a) Host 190 issues a write command for a 4-KB page;
- (b) Controller 110 looks up the 4-KB page in L2P table 160 and/or subfile mapping table 300 and determines that the 4-KB page is in a JBL subfile;
- (c) Controller 110 configures the write circuitry to update the 4-KB page in NVM 130;

- (d) Controller 110 reports the completion of the write command to host 190;
- (e) Controller 110 schedules the corresponding update of the JBL subfile having the updated 4-KB page; and
- (f) Controller 110 configures the write circuitry to perform the JBL subfile update in the background, e.g., when workload permits, and the write circuitry relocates the JBL subfile to achieve a more favorable data layout from the parallelism viewpoint, e.g., similar to the data layout illustrated in FIG. 6.

[0057] In some examples, an update can happen in two stages, wherein the first stage is a write of the piece of data being modified by the host to a nonvolatile memory (which is a safe option against power loss and relatively easy to implement for enhanced performance), and wherein the second stage is a write to a more permanent location with the entire chunk of the data being handled. The two-stage update is also expected to work well with coalescing, e.g., with multiple modifications being grouped together and updated in a single action, which can be beneficial in terms of performance and write amplification.

[0058] With regard to the processes, systems, methods, heuristics, etc. described herein, it should be understood that, although the steps of such processes, etc. have been described as occurring according to a certain ordered sequence, such processes could be practiced with the described steps performed in an order other than the order described herein. It further should be understood that certain steps could be performed simultaneously, that other steps could be added, or that certain steps described herein could be omitted. In other words, the descriptions of processes herein are provided for the purpose of illustrating certain implementations and should in no way be construed to limit the claims.

[0059] Accordingly, it is to be understood that the above description is intended to be illustrative and not restrictive. Many embodiments and applications other than the examples provided would be apparent upon reading the above description. The scope should be determined, not with reference to the above description, but should instead be determined with reference to the appended claims, along with the full scope of equivalents to which such claims are entitled. It is anticipated and intended that future developments will occur in the technologies discussed herein, and that the disclosed systems and methods will be incorporated into such future embodiments. In sum, it should be understood that the application is capable of modification and variation.

[0060] All terms used in the claims are intended to be given their broadest reasonable constructions and their ordinary meanings as understood by those knowledgeable in the technologies described herein unless an explicit indication to the contrary is made herein. In

particular, use of the singular articles such as “a,” “the,” “said,” etc. should be read to recite one or more of the indicated elements unless a claim recites an explicit limitation to the contrary.

[0061] Unless explicitly stated otherwise, each numerical value and range should be interpreted as being approximate as if the word “about” or “approximately” preceded the value or range.

[0062] Reference herein to “one embodiment” or “an embodiment” means that a particular feature, structure, or characteristic described in connection with the embodiment can be included in at least one embodiment of the disclosure. The appearances of the phrase “in one embodiment” in various places in the specification are not necessarily all referring to the same embodiment, nor are separate or alternative embodiments necessarily mutually exclusive of other embodiments. The same applies to the term “implementation.”

[0063] Unless otherwise specified herein, the use of the ordinal adjectives “first,” “second,” “third,” etc., to refer to an object of a plurality of like objects merely indicates that different instances of such like objects are being referred to and is not intended to imply that the like objects so referred-to have to be in a corresponding order or sequence, either temporally, spatially, in ranking, or in any other manner.

[0064] Unless otherwise specified herein, in addition to its plain meaning, the conjunction “if” may also or alternatively be construed to mean “when” or “upon” or “in response to determining” or “in response to detecting,” which construal may depend on the corresponding specific context. For example, the phrase “if it is determined” or “if [a stated condition] is detected” may be construed to mean “upon determining” or “in response to determining” or “upon detecting [the stated condition or event]” or “in response to detecting [the stated condition or event].”

[0065] The functions of the various elements shown in the figures, including any functional blocks labeled as “processors” and/or “controllers,” may be provided through the use of dedicated hardware as well as hardware capable of executing software in association with appropriate software. When provided by a processor, the functions may be provided by a single dedicated processor, by a single shared processor, or by a plurality of individual processors, some of which may be shared. Moreover, explicit use of the term “processor” or “controller” should not be construed to refer exclusively to hardware capable of executing software, and may implicitly include, without limitation, digital signal processor (DSP) hardware, network processor, application specific integrated circuit (ASIC), field programmable gate array (FPGA), read only memory (ROM) for storing software, random access memory (RAM), and nonvolatile storage. Other hardware, conventional and/or custom, may also be included. Similarly, any switches shown in the figures are conceptual only. Their function may be carried out through the operation of program logic, through dedicated logic, through the interaction of program control

and dedicated logic, or even manually, the particular technique being selectable by the implementer as more specifically understood from the context.

[0066] As used in this application, the term “circuitry” may refer to one or more or all of the following: (a) hardware-only circuit implementations (such as implementations in only analog and/or digital circuitry); (b) combinations of hardware circuits and software, such as (as applicable): (i) a combination of analog and/or digital hardware circuit(s) with software/firmware and (ii) any portions of hardware processor(s) with software (including digital signal processor(s)), software, and memory(ies) that work together to cause an apparatus, such as a mobile phone or server, to perform various functions); and (c) hardware circuit(s) and/or processor(s), such as a microprocessor(s) or a portion of a microprocessor(s), that requires software (e.g., firmware) for operation, but the software may not be present when it is not needed for operation.” This definition of circuitry applies to all uses of this term in this application, including in any claims. As a further example, as used in this application, the term circuitry also covers an implementation of merely a hardware circuit or processor (or multiple processors) or portion of a hardware circuit or processor and its (or their) accompanying software and/or firmware. The term circuitry also covers, for example and if applicable to the particular claim element, a baseband integrated circuit or processor integrated circuit for a mobile device or a similar integrated circuit in server, a cellular network device, or other computing or network device.

[0067] “SUMMARY” in this specification is intended to introduce some example embodiments, with additional embodiments being described in “DETAILED DESCRIPTION” and/or in reference to one or more drawings. “SUMMARY” is not intended to identify essential elements or features of the claimed subject matter, nor is it intended to limit the scope of the claimed subject matter.

[0068] “ABSTRACT” is provided to allow the reader to quickly ascertain the nature of the technical disclosure. It is submitted with the understanding that it will not be used to interpret or limit the scope or meaning of the claims. In addition, in the foregoing “DETAILED DESCRIPTION,” it can be seen that various features are grouped together in various embodiments for the purpose of streamlining the disclosure. This method of disclosure is not to be interpreted as reflecting an intention that the claimed embodiments require more features than are expressly recited in each claim. Rather, as the following claims reflect, inventive subject matter lies in less than all features of a single disclosed embodiment. Thus, the following claims are hereby incorporated into “DETAILED DESCRIPTION,” with each claim standing on its own as a separately claimed subject matter.

CLAIMS

What is claimed is:

1. A data storage device, comprising:
a nonvolatile memory to store data;
a second memory to store a logical-to-physical (L2P) table and a subfile mapping table;
and
a controller coupled to the nonvolatile memory and the second memory and configured to:
 - in response to a host command to write a large data file to the nonvolatile memory, segregate the large data file into a plurality of subfiles, the large data file having a file size larger than a first fixed size, each subfile of the plurality of subfiles having a respective size smaller than or equal to the first fixed size;
 - store a first mapping of the plurality of subfiles in the subfile mapping table, the first mapping including, for each subfile of the plurality of subfiles, a respective word-line physical address;
 - cause the plurality of subfiles to be written to the nonvolatile memory based on the first mapping; and
 - access the L2P table and the subfile mapping table to control a memory operation on the large data file in the nonvolatile memory.
2. The data storage device of claim 1, wherein the first fixed size is a meta-block word-line capacity or a jumbo-block word-line capacity, the jumbo-block word-line capacity being an integer multiple of the meta-block word-line capacity, the meta-block word-line capacity being a combined data capacity of NAND memory blocks connected to a same word line in a single physical die of the nonvolatile memory or an integer multiple of said combined data capacity in one or more physical dies of the nonvolatile memory.
3. The data storage device of claim 1, wherein the controller is further configured to:
in response to a host command to write a small data file to the nonvolatile memory, store a mapping of the small data file in the L2P table, the small data file having a file size smaller than the first fixed size; and
cause the small data file to be written to the nonvolatile memory based on said mapping of the small data file.

4. The data storage device of claim 3, wherein the first mapping causes the plurality of subfiles to be interleaved with a plurality of small data files in the nonvolatile memory.
5. The data storage device of claim 1, wherein, in response to a sequence of host read-modify-write (RMW) commands, the controller is further configured to:
 - generate a copy of one or more implicated subfiles of the plurality of subfiles in a data buffer;
 - apply subfile updates corresponding to the sequence of host RMW commands to the copy; and
 - in response to re-allocation of the data buffer, cause relocation of the one or more implicated subfiles in the nonvolatile memory.
6. The data storage device of claim 5, wherein the controller is further configured to update the first mapping in the subfile mapping table in accordance with the relocation.
7. The data storage device of claim 5, wherein the controller is configured not to apply the subfile updates corresponding to the sequence of host RMW commands to the one or more implicated subfiles in the nonvolatile memory until the relocation.
8. The data storage device of claim 1, wherein the first mapping causes an individual subfile of the plurality of subfile to occupy, in the nonvolatile memory, an array of NAND memory blocks connected to a single word line in a single physical die of the nonvolatile memory or two or more instances of said array in one or more physical dies of the nonvolatile memory.
9. The data storage device of claim 8,
 - wherein at least one NAND memory block in the array of NAND memory blocks has stored therein both a smaller subfile of a first large data file and a smaller subfile of a second large data file; and
 - wherein the smaller subfile of the first large data file and the smaller subfile of a second large data file have a combine size equal to the first fixed size.
10. The data storage device of claim 1, wherein the first mapping causes an individual subfile of the plurality of subfiles to occupy, in the nonvolatile memory, a first array of NAND memory blocks connected to a first word line in a physical die of the nonvolatile memory and a second

array of NAND memory blocks connected to a second word line in the physical die of the nonvolatile memory.

11. The data storage device of claim 1, wherein the first mapping causes the plurality of subfiles to be interleaved with a plurality of other data files or other subfiles in the nonvolatile memory.

12. A method performed by a data storage device, the method comprising:
in response to a host command to write a large data file to a nonvolatile memory, segregating, via a controller, the large data file into a plurality of subfiles, the large data file having a file size larger than a first fixed size, each subfile of the plurality of subfiles having a respective size smaller than or equal to the first fixed size;

storing, via the controller, a first mapping of the plurality of subfiles in a subfile mapping table, the first mapping including, for each subfile of the plurality of subfiles, a respective word-line physical address;

causing, via the controller, the plurality of subfiles to be written to the nonvolatile memory based on the first mapping; and

accessing, via the controller, a logical-to-physical (L2P) table and the subfile mapping table to control a memory operation on the large data file in the nonvolatile memory.

13. The method of claim 12, wherein the first fixed size is a meta-block word-line capacity or a jumbo-block word-line capacity, the jumbo-block word-line capacity being an integer multiple of the meta-block word-line capacity, the meta-block word-line capacity being a combined data capacity of NAND memory blocks connected to a same word line in a single physical die of the nonvolatile memory.

14. The method of claim 12, further comprising:

in response to a host command to write a small data file to the nonvolatile memory, storing, via the controller, a mapping of the small data file in the L2P table, the small data file having a file size smaller than the first fixed size; and

causing, via the controller, the small data file to be written to the nonvolatile memory based on said mapping of the small data file.

15. The method of claim 12, further comprising:

receiving, via the controller, a sequence of host read-modify-write (RMW) commands;

generating, via the controller, a copy of one or more implicated subfiles of the plurality of subfiles in a data buffer, the one or more implicated subfiles being implicated by the sequence of host RMW commands;

applying, via the controller, subfile updates corresponding to the sequence of host RMW commands to the copy; and

in response to re-allocation of the data buffer, causing, via the controller, relocation of the one or more implicated subfiles in the nonvolatile memory.

16. The method of claim 15, wherein the controller is configured to update the first mapping in the subfile mapping table in accordance with the relocation.

17. The method of claim 15, wherein the controller is configured not to apply the subfile updates corresponding to the sequence of host RMW commands to the one or more implicated subfiles in the nonvolatile memory until the relocation.

18. The method of claim 12, wherein the first mapping causes an individual subfile of the plurality of subfiles to occupy, in the nonvolatile memory, a first array of NAND memory blocks connected to a first word line in a physical die of the nonvolatile memory and a second array of NAND memory blocks connected to a second word line in the physical die of the nonvolatile memory.

19. A data storage device, comprising:

means for segregating a large data file into a plurality of subfiles in response to a host command to write the large data file to a nonvolatile memory, the large data file having a file size larger than a first fixed size, each subfile of the plurality of subfiles having a respective size smaller than or equal to the first fixed size;

means for storing a first mapping of the plurality of subfiles in a subfile mapping table, the first mapping including, for each subfile of the plurality of subfiles, a respective word-line physical address;

means for causing the plurality of subfiles to be written to the nonvolatile memory based on the first mapping; and

means for accessing a logical-to-physical (L2P) table and the subfile mapping table to control a memory operation on the large data file in the nonvolatile memory.

20. The data storage device of claim 19, further comprising:

means for receiving a sequence of host read-modify-write (RMW) commands;

means for generating a copy of one or more implicated subfiles of the plurality of subfiles in a data buffer, the one or more implicated subfiles being implicated by the sequence of host RMW commands;

means for applying subfile updates corresponding to the sequence of host RMW commands to the copy; and

means for causing relocation of the one or more implicated subfiles in the nonvolatile memory in response to re-allocation of the data buffer.

1/19

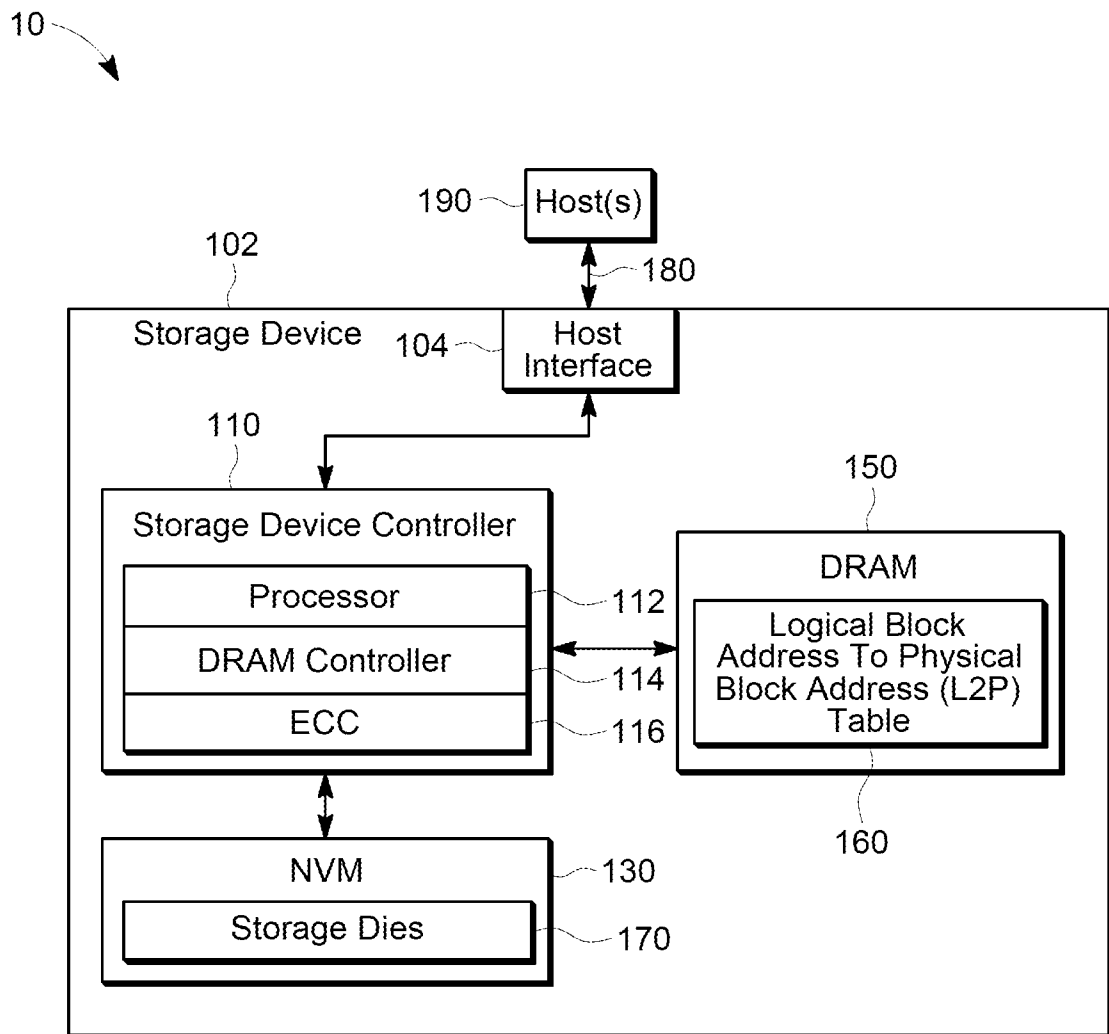


FIG. 1

200

Meta Block	Channel 0				Channel 1				Channel 2				Channel 3			
	Die 0				Die 0				Die 0				Die 0			
	Logical Die 0				Logical Die 1				Logical Die 2				Logical Die 3			
Wordline	Plane 0	Plane 1	Plane 2	Plane 3	Plane 0	Plane 1	Plane 2	Plane 3	Plane 0	Plane 1	Plane 2	Plane 3	Plane 0	Plane 1	Plane 2	Plane 3
0																
1																
2																
3																
4																
5																
6																
7																
8																
9																
10																
11																
12																
13																
14																
15																
16																
17																
18																
19																

210

204

FIG. 2

2022

[illegible]

FIG. 2(Continued)

300

302	304	306	308	310	312	314	316
Host File ID	Subfile ID	Subfile Size (KB)	MBW ID	MBW Offset	MBW Address	MBW Cell Type	MBW Logical Die ID
0	0	768	0	0	0Ah	TLC	0,1,2,3
0	1	768	1	0	01h	TLC	4,5,6,7
0	2	400	2	0	100h	SLC	8,9,10,11
1	0	368	0	100	100h	SLC	8,9,10,11
1	1	768	1	0	200h	TLC	12,13,14,15
...	...	...	...	...	...	...	...

FIG. 3

6/19

400

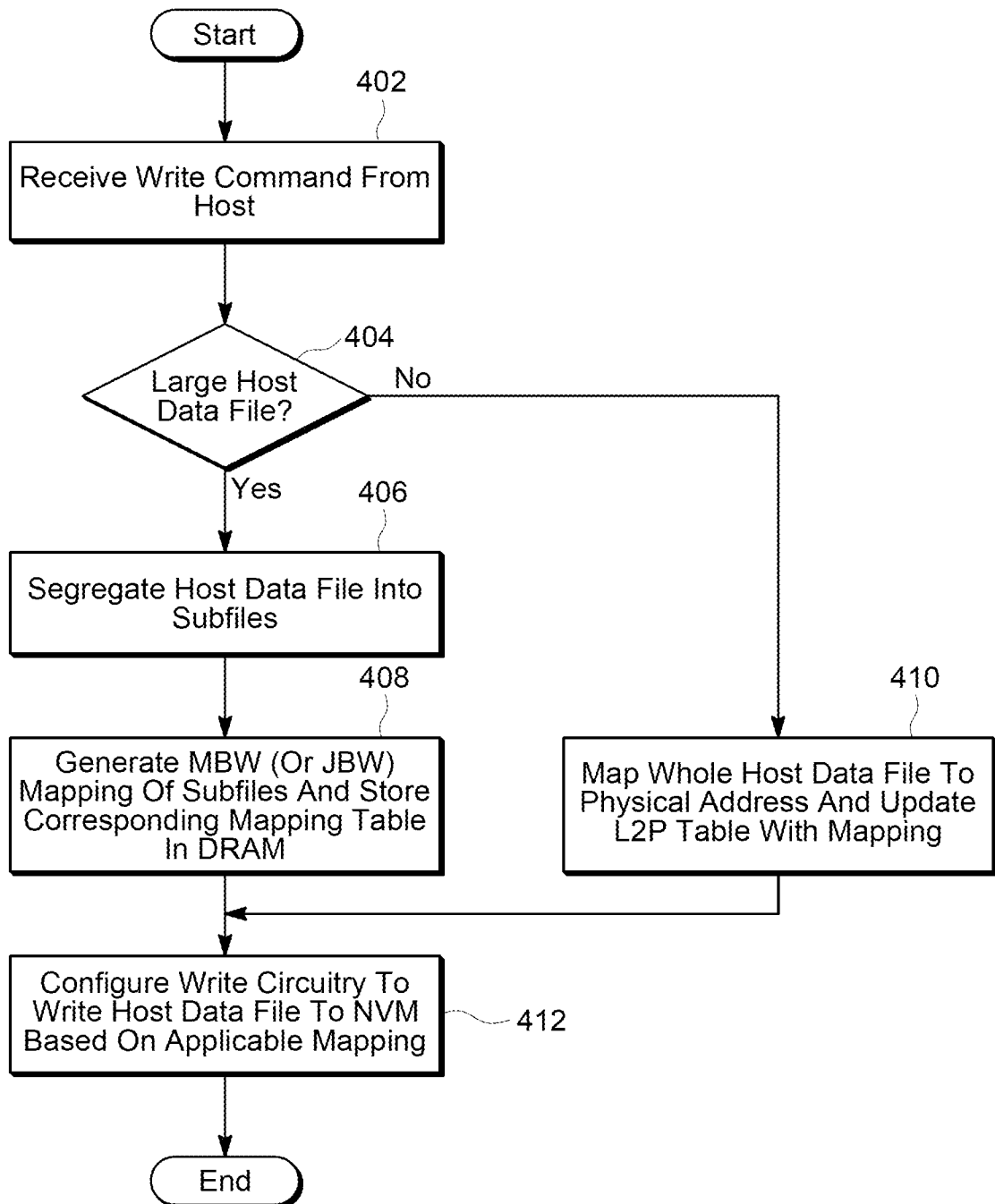


FIG. 4

7/19

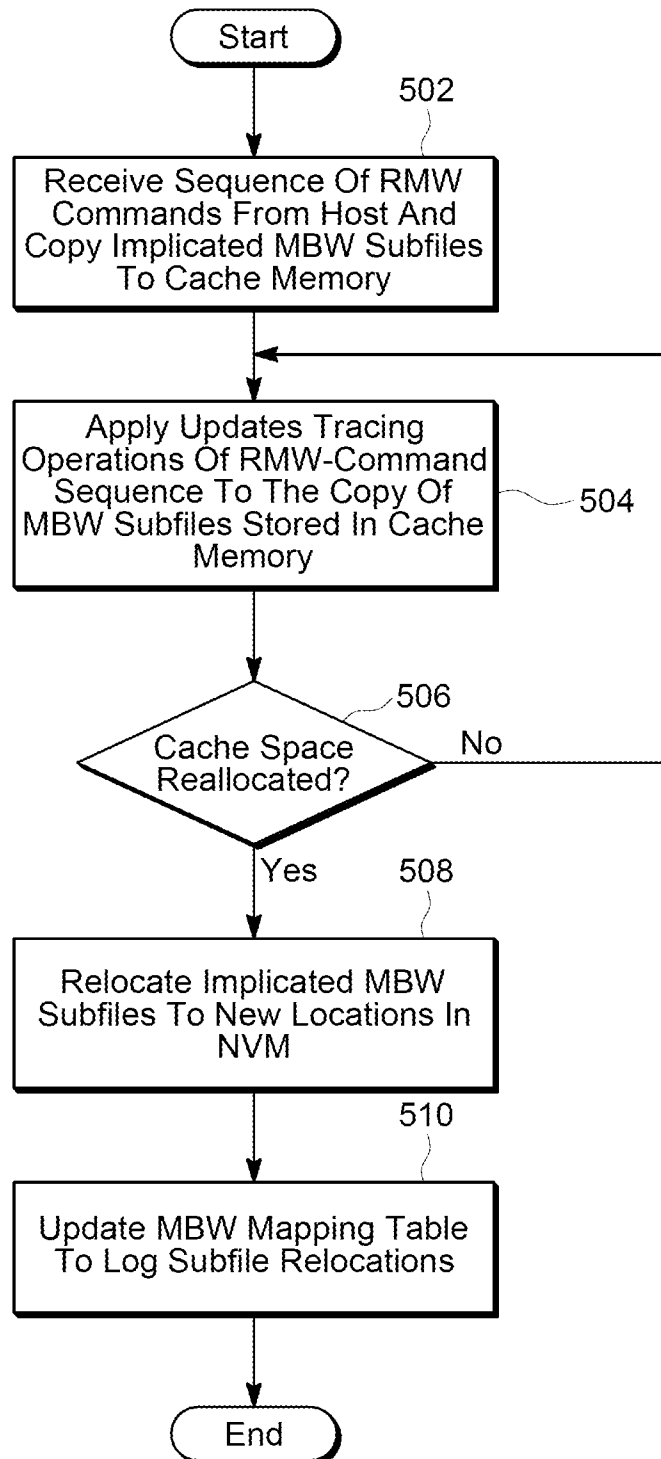
500

FIG. 5

Meta Block	Channel 0				Channel 1				Channel 2				Channel 3	
	Die 0				Die 0				Die 0				Die 0	
	Logical Die 0				Logical Die 1				Logical Die 2				Logical Die 3	
Wordline	Plane 0	Plane 1	Plane 2	Plane 3	Plane 0	Plane 1	Plane 2	Plane 3	Plane 0	Plane 1	Plane 2	Plane 3	Plane 0	Plane 1
0	JPC0-0	JPC0-1	JPC0-2	JPC0-3	JPC0-4	JPC0-5	JPC0-6	JPC0-7	JPC0-8	JPC0-9	JPC0-10	JPC0-11	JPC0-12	JPC0-13
1	JPC1-0	JPC1-1	JPC1-2	JPC1-3	JPC1-4	JPC1-5	JPC1-6	JPC1-7	JPC1-8	JPC1-9	JPC1-10	JPC1-11	JPC1-12	JPC1-13
2	JPC2-0	JPC2-1	JPC2-2	JPC2-3	JPC2-4	JPC2-5	JPC2-6	JPC2-7	JPC2-8	JPC2-9	JPC2-10	JPC2-11	JPC2-12	JPC2-13
3	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal
4	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal
5	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal
6														
7														
8														
9														
10														
11														
12														
13														
14														
15														
16														
17														
18														
19														

204

FIG. 6

[illegible]

FIG. 6 (Continued)

Meta Block	Channel 0				Channel 1				Channel 2				Channel 3	
	Die 0				Die 0				Die 0				Die 0	
	Logical Die 0				Logical Die 1				Logical Die 2				Logical Die 3	
Wordline	Plane 0	Plane 1	Plane 2	Plane 3	Plane 0	Plane 1	Plane 2	Plane 3	Plane 0	Plane 1	Plane 2	Plane 3	Plane 0	Plane 1
0	JPC0-0	JPC0-1	JPC0-2	JPC0-3	JPC0-4	JPC0-5	JPC0-6	JPC0-7	JPC0-8	JPC0-9	JPC0-10	JPC0-11	JPC0-12	JPC0-13
1	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal
2	JPC1-0	JPC1-1	JPC1-2	JPC1-3	JPC1-4	JPC1-5	JPC1-6	JPC1-7	JPC1-8	JPC1-9	JPC1-10	JPC1-11	JPC1-12	JPC1-13
3	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal
4	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal
5	JPC2-0	JPC2-1	JPC2-2	JPC2-3	JPC2-4	JPC2-5	JPC2-6	JPC2-7	JPC2-8	JPC2-9	JPC2-10	JPC2-11	JPC2-12	JPC2-13
6	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal
7														
8														
9														
10														
11														
12														
13														
14														
15														
16														
17														
18														
19														

204  FIG. 7

[illegible]

FIG. 7 (Continued)

[illegible]

FIG. 7 (Continued)

Meta Block	Channel 0				Channel 1				Channel 2			
	Die 0				Die 0				Die 0			
	Logical Die 0				Logical Die 1				Logical Die 2			
Wordline	Plane 0	Plane 1	Plane 2	Plane 3	Plane 0	Plane 1	Plane 2	Plane 3	Plane 0	Plane 1	Plane 2	Plane 3
0	Normal	JPC0-0	JPC0-1	JPC0-2	JPC0-3	JPC0-4	JPC0-5	JPC0-6	JPC0-7	JPC0-8	JPC0-9	JPC0-10
1	JPC0-31	Normal	Normal	Normal	JPC1-0	JPC1-1	JPC1-2	JPC1-3	JPC1-4	JPC1-5	JPC1-6	JPC1-7
2	JPC1-28	JPC1-29	JPC1-30	JPC1-31	Normal	Normal	Normal	Normal	Normal	JPC2-0	JPC2-1	JPC2-2
3	JPC2-23	JPC2-24	JPC2-25	JPC2-26	JPC2-27	JPC2-28	JPC2-29	JPC2-30	JPC2-31	Normal	Normal	Normal
4	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal	Normal
5												
6												
7												
8												
9												
10												
11												
12												
13												
14												
15												
16												
17												
18												
19												

204

FIG. 8

[illegible]

FIG. 8 (Continued)

[illegible]

FIG. 8 (Continued)

204

Meta Block	Channel 0				Channel 1				Channel 2			
	Die 0				Die 0				Die 0			
	Logical Die 0				Logical Die 1				Logical Die 2			
Wordline	Plane 0	Plane 1	Plane 2	Plane 3	Plane 0	Plane 1	Plane 2	Plane 3	Plane 0	Plane 1	Plane 2	Plane 3
0	Normal	JPC0-0	JPC0-1	JPC0-2	JPC0-3	JPC0-4	JPC0-5	JPC0-6	JPC0-7	JPC0-8	JPC0-9	JPC0-10
1	JPC0-31	Normal	Normal	Normal	JPC1-0	JPC1-1	JPC1-2	JPC1-3	JPC1-4	JPC1-5	JPC1-6	JPC1-7
2	JPC1-28	JPC1-29	JPC1-30	JPC1-31	Normal	Normal	Normal	Normal	Normal	JPC2-0	JPC2-1	JPC2-2
3	JPC2-23	JPC2-24	JPC2-25	JPC2-26	Normal	Normal	Normal	Normal	JPC2-31	Normal	Normal	Normal
4	JPC2-27	JPC2-28	JPC2-29	JPC2-30					Normal	Normal	Normal	Normal
5									Normal	Normal	Normal	Normal
6												
7												
8												
9												
10												
11												
12												
13												
14												
15												
16												
17												
18												
19												

FIG. 9

204

[illegible]

[illegible]

FIG. 9 (Continued)

204

INTERNATIONAL SEARCH REPORT

International application No.

PCT/US2023/077225

A. CLASSIFICATION OF SUBJECT MATTER G06F 3/06(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 3/06(2006.01); G06F 11/10(2006.01); G06F 12/02(2006.01); G06F 12/10(2006.01) 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: nonvolatile memory, logical-to-physical (L2P) table, mapping, segregate, write, read-modify-write (RMW), re-allocation, word line		
C. DOCUMENTS CONSIDERED TO BE RELEVANT		
Category*	Citation of document, with indication, where appropriate, of the relevant passages	Relevant to claim No.
Y	US 2019-0272118 A1 (SK HYNIX INC.) 05 September 2019 (2019-09-05) paragraphs [0031], [0048], [0060], [0063], [0070], [0076], [0080]; and claims 1-4	1-20
Y	US 2018-0267706 A1 (INTEL CORPORATION) 20 September 2018 (2018-09-20) paragraphs [0015], [0031]; and claim 1	1-20
Y	US 2018-0321874 A1 (ALIBABA GROUP HOLDING LIMITED) 08 November 2018 (2018-11-08) claim 1	3-4,9,11,14
A	US 2015-0347026 A1 (SANDISK TECHNOLOGIES INC.) 03 December 2015 (2015-12-03) paragraphs [0037]-[0040]; and figure 7	1-20
A	US 2020-0233610 A1 (SILICON MOTION, INC.) 23 July 2020 (2020-07-23) paragraphs [0051]-[0062]; and figures 5A-5B	1-20
☐ 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 “D” document cited by the applicant in the international application “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 14 February 2024		Date of mailing of the international search report 14 February 2024
Name and mailing address of the ISA/KR Korean Intellectual Property Office 189 Cheongsa-ro, Seo-gu, Daejeon 35208, Republic of Korea Facsimile No. +82-42-481-8578		Authorized officer YANG, JEONG ROK Telephone No. +82-42-481-5709

INTERNATIONAL SEARCH REPORT
Information on patent family members

International application No.

PCT/US2023/077225

Patent document cited in search report			Publication date (day/month/year)	Patent family member(s)			Publication date (day/month/year)
US	2019-0272118	A1	05 September 2019	CN	110232035	A	13 September 2019
				CN	110232035	B	27 June 2023
				KR	10-2503177	B1	24 February 2023
				US	10698635	B2	30 June 2020
US	2018-0267706	A1	20 September 2018	US	10078453	B1	18 September 2018
US	2018-0321874	A1	08 November 2018	None			
US	2015-0347026	A1	03 December 2015	US	9383927	B2	05 July 2016
US	2020-0233610	A1	23 July 2020	CN	111459844	A	28 July 2020
				CN	111459844	B	11 November 2022
				TW	202028982	A	01 August 2020