

TMS320C64x DSP Two-Level Internal Memory Reference Guide

Literature Number: SPRU610C
February 2006



IMPORTANT NOTICE

Texas Instruments Incorporated and its subsidiaries (TI) reserve the right to make corrections, modifications, enhancements, improvements, and other changes to its products and services at any time and to discontinue any product or service without notice. Customers should obtain the latest relevant information before placing orders and should verify that such information is current and complete. All products are sold subject to TI's terms and conditions of sale supplied at the time of order acknowledgment.

TI warrants performance of its hardware products to the specifications applicable at the time of sale in accordance with TI's standard warranty. Testing and other quality control techniques are used to the extent TI deems necessary to support this warranty. Except where mandated by government requirements, testing of all parameters of each product is not necessarily performed.

TI assumes no liability for applications assistance or customer product design. Customers are responsible for their products and applications using TI components. To minimize the risks associated with customer products and applications, customers should provide adequate design and operating safeguards.

TI does not warrant or represent that any license, either express or implied, is granted under any TI patent right, copyright, mask work right, or other TI intellectual property right relating to any combination, machine, or process in which TI products or services are used. Information published by TI regarding third-party products or services does not constitute a license from TI to use such products or services or a warranty or endorsement thereof. Use of such information may require a license from a third party under the patents or other intellectual property of the third party, or a license from TI under the patents or other intellectual property of TI.

Reproduction of information in TI data books or data sheets is permissible only if reproduction is without alteration and is accompanied by all associated warranties, conditions, limitations, and notices. Reproduction of this information with alteration is an unfair and deceptive business practice. TI is not responsible or liable for such altered documentation.

Resale of TI products or services with statements different from or beyond the parameters stated by TI for that product or service voids all express and any implied warranties for the associated TI product or service and is an unfair and deceptive business practice. TI is not responsible or liable for any such statements.

Following are URLs where you can obtain information on other Texas Instruments products and application solutions:

Products		Applications	
Amplifiers	amplifier.ti.com	Audio	www.ti.com/audio
Data Converters	dataconverter.ti.com	Automotive	www.ti.com/automotive
DSP	dsp.ti.com	Broadband	www.ti.com/broadband
Interface	interface.ti.com	Digital Control	www.ti.com/digitalcontrol
Logic	logic.ti.com	Military	www.ti.com/military
Power Mgmt	power.ti.com	Optical Networking	www.ti.com/opticalnetwork
Microcontrollers	microcontroller.ti.com	Security	www.ti.com/security
		Telephony	www.ti.com/telephony
		Video & Imaging	www.ti.com/video
		Wireless	www.ti.com/wireless

Mailing Address: Texas Instruments
Post Office Box 655303 Dallas, Texas 75265

Copyright © 2006, Texas Instruments Incorporated

Read This First

About This Manual

The TMS320C64x™ digital signal processors (DSPs) of the TMS320C6000™ DSP family have a two-level memory architecture for program and data. The first-level program cache is designated L1P, and the first-level data cache is designated L1D. Both the program and data memory share the second-level memory, designated L2. L2 is configurable, allowing for various amounts of cache and SRAM. This document discusses the C64x™ two-level internal memory.

Notational Conventions

This document uses the following conventions.

- Hexadecimal numbers are shown with the suffix h. For example, the following number is 40 hexadecimal (decimal 64): 40h.
- Registers in this document are shown in figures and described in tables.
 - Each register figure shows a rectangle divided into fields that represent the fields of the register. Each field is labeled with its bit name, its beginning and ending bit numbers above, and its read/write properties below. A legend explains the notation used for the properties.
 - Reserved bits in a register figure designate a bit that is used for future device expansion.

Related Documentation From Texas Instruments

The following documents describe the C6000™ devices and related support tools. Copies of these documents are available on the Internet at www.ti.com.
Tip: Enter the literature number in the search box provided at www.ti.com.

TMS320C6000 CPU and Instruction Set Reference Guide (literature number SPRU189) describes the TMS320C6000™ CPU architecture, instruction set, pipeline, and interrupts for these digital signal processors.

TMS320C6000 DSP Peripherals Overview Reference Guide (literature number SPRU190) describes the peripherals available on the TMS320C6000™ DSPs.

TMS320C64x Technical Overview (SPRU395) The TMS320C64x technical overview gives an introduction to the TMS320C64x™ digital signal processor, and discusses the application areas that are enhanced by the TMS320C64x VelociTI™.

TMS320C6000 DSP Cache User's Guide (literature number SPRU656) explains the fundamentals of memory caches and describes how to efficiently utilize the two-level cache-based memory architecture in the digital signal processors (DSPs) of the TMS320C6000™ DSP family. It shows how to maintain coherence with external memory, how to use DMA to reduce memory latencies, and how to optimize your code to improve cache efficiency.

TMS320C6000 Programmer's Guide (literature number SPRU198) describes ways to optimize C and assembly code for the TMS320C6000™ DSPs and includes application program examples.

TMS320C6000 Code Composer Studio Tutorial (literature number SPRU301) introduces the Code Composer Studio™ integrated development environment and software tools.

Code Composer Studio Application Programming Interface Reference Guide (literature number SPRU321) describes the Code Composer Studio™ application programming interface (API), which allows you to program custom plug-ins for Code Composer.

TMS320C6x Peripheral Support Library Programmer's Reference (literature number SPRU273) describes the contents of the TMS320C6000™ peripheral support library of functions and macros. It lists functions and macros both by header file and alphabetically, provides a complete description of each, and gives code examples to show how they are used.

TMS320C6000 Chip Support Library API Reference Guide (literature number SPRU401) describes a set of application programming interfaces (APIs) used to configure and control the on-chip peripherals.

Trademarks

Code Composer Studio, C6000, C62x, C64x, C67x, TMS320C6000, TMS320C62x, TMS320C64x, TMS320C67x, and VelociTI are trademarks of Texas Instruments.

Contents

1	Memory Hierarchy Overview	9
2	Cache Terms and Definitions	13
3	Level 1 Data Cache (L1D)	19
3.1	L1D Parameters	19
3.2	L1D Performance	20
3.2.1	L1D Memory Banking	20
3.2.2	L1D Miss Penalty	22
3.2.3	L1D Write Buffer	23
3.2.4	L1D Miss Pipelining	24
4	Level 1 Program Cache (L1P)	27
4.1	L1P Parameters	27
4.2	L1P Performance	27
4.2.1	L1P Miss Penalty	27
4.2.2	L1P Miss Pipelining	28
5	Level 2 Unified Memory (L2)	30
5.1	L2 Cache and L2 SRAM	30
5.2	L2 Operation	32
5.3	L2 Bank Structure	34
5.4	L2 Interfaces	35
5.4.1	L1D/L1P-to-L2 Request Servicing	35
5.4.2	EDMA-to-L2 Request Servicing	36
5.4.3	L2 Request Servicing Using EDMA	36
5.4.4	EDMA Access to Cache Controls	37
5.4.5	HPI and PCI Access to Memory Subsystem	37
6	Registers	38
6.1	Cache Configuration Register (CCFG)	39
6.2	L2 EDMA Access Control Register (EDMAWEIGHT)	41
6.3	L2 Allocation Registers (L2ALLOC0–L2ALLOC03)	42
6.4	L2 Writeback Base Address Register (L2WBAR)	43
6.5	L2 Writeback Word Count Register (L2WWC)	43
6.6	L2 Writeback–Invalidate Base Address Register (L2WIBAR)	44
6.7	L2 Writeback–Invalidate Word Count Register (L2WIWC)	44

6.8	L2 Invalidate Base Address Register (L2IBAR)	45
6.9	L2 Invalidate Word Count Register (L2IWC)	45
6.10	L1P Invalidate Base Address Register (L1PIBAR)	46
6.11	L1P Invalidate Word Count Register (L1PIWC)	46
6.12	L1D Writeback-Invalidate Base Address Register (L1DWIBAR)	47
6.13	L1D Writeback-Invalidate Word Count Register (L1DWIWC)	47
6.14	L1D Invalidate Base Address Register (L1DIBAR)	48
6.15	L1D Invalidate Word Count Register (L1DIWC)	48
6.16	L2 Writeback All Register (L2WB)	49
6.17	L2 Writeback-Invalidate All Register (L2WBINV)	50
6.18	L2 Memory Attribute Registers (MAR0–MAR255)	51
7	Memory System Control	52
7.1	Cache Mode Selection	52
7.1.1	L1D Mode Selection Using DCC Field in CSR	52
7.1.2	L1P Mode Selection Using PCC Field in CSR	53
7.1.3	L2 Mode Selection Using L2MODE Field in CCFG	53
7.2	Cacheability Controls	56
7.3	Program-Initiated Cache Operations	60
7.3.1	Global Cache Operations	62
7.3.2	Block Cache Operations	63
7.3.3	Effect of L2 Commands on L1 Caches	65
7.4	L2-to-EDMA Request Control	67
7.5	EDMA Access Into L2 Control	68
8	Memory System Policies	69
8.1	Memory System Coherence	69
8.2	EDMA Coherence in L2 SRAM Example	71
8.3	Memory Access Ordering	76
8.3.1	Program Order of Memory Accesses	76
8.3.2	Strong and Relaxed Memory Ordering	77
	Revision History	81

Figures

1	TMS320C64x DSP Block Diagram	9
2	TMS320C64x Two-Level Internal Memory Block Diagram	12
3	L1D Address Allocation	19
4	Address to Bank Number Mapping	20
5	Potentially Conflicting Memory Accesses	21
6	L1P Address Allocation	27
7	L2 Address Allocation, 256K Cache (L2MODE = 111b)	31
8	L2 Address Allocation, 128K Cache (L2MODE = 011b)	31
9	L2 Address Allocation, 64K Cache (L2MODE = 010b)	31
10	L2 Address Allocation, 32K Cache (L2MODE = 001b)	31
11	Cache Configuration Register (CCFG)	39
12	L2 EDMA Access Control Register (EDMAWEIGHT)	41
13	L2 Allocation Registers (L2ALLOC0–L2ALLOC3)	42
14	L2 Writeback Base Address Register (L2WBAR)	43
15	L2 Writeback Word Count Register (L2WWC)	43
16	L2 Writeback-Invalidate Base Address Register (L2WIBAR)	44
17	L2 Writeback-Invalidate Word Count Register (L2WIWC)	44
18	L2 Invalidate Base Address Register (L2IBAR)	45
19	L2 Invalidate Word Count Register (L2IWC)	45
20	L1P Invalidate Base Address Register (L1PIBAR)	46
21	L1P Invalidate Word Count Register (L1PIWC)	46
22	L1D Writeback-Invalidate Base Address Register (L1DWIBAR)	47
23	L1D Writeback-Invalidate Word Count Register (L1DWIWC)	47
24	L1D Invalidate Base Address Register (L1DIBAR)	48
25	L1D Invalidate Word Count Register (L1DIWC)	48
26	L2 Writeback All Register (L2WB)	49
27	L2 Writeback-Invalidate All Register (L2WBINV)	50
28	L2 Memory Attribute Register (MAR)	51
29	CPU Control and Status Register (CSR)	52
30	Block Cache Operation Base Address Register (BAR)	63
31	Block Cache Operation Word Count Register (WC)	63
32	Streaming Data Pseudo-Code	72
33	Double Buffering Pseudo-Code	72
34	Double-Buffering Time Sequence	73
35	Double Buffering as a Pipelined Process	74

Tables

1	TMS320C621x/C671x/C64x Internal Memory Comparison	10
2	Terms and Definitions	13
3	Cycles Per Miss for Different Numbers of L1D Misses That Hit L2 Cache	26
4	Cycles Per Miss for Different Numbers of L1D Misses that Hit L2 SRAM	26
5	Average Miss Penalties for Large Numbers of Sequential Execute Packets	29
6	Internal Memory Control Registers	38
7	Cache Configuration Register (CCFG) Field Descriptions	39
8	L2 EDMA Access Control Register (EDMAWEIGHT) Field Descriptions	41
9	L2 Allocation Registers (L2ALLOC0–L2ALLOC3) Field Descriptions	42
10	L2 Writeback Base Address Register (L2WBAR) Field Descriptions	43
11	L2 Writeback Word Count Register (L2WWC) Field Descriptions	43
12	L2 Writeback–Invalidate Base Address Register (L2WIBAR) Field Descriptions	44
13	L2 Writeback–Invalidate Word Count Register (L2WIWC) Field Descriptions	44
14	L2 Invalidate Base Address Register (L2IBAR) Field Descriptions	45
15	L2 Invalidate Word Count Register (L2IWC) Field Descriptions	45
16	L1P Invalidate Base Address Register (L1PIBAR) Field Descriptions	46
17	L1P Invalidate Word Count Register (L1PIWC) Field Descriptions	46
18	L1D Writeback–Invalidate Base Address Register (L1DWIBAR) Field Descriptions	47
19	L1D Writeback–Invalidate Word Count Register (L1DWIWC) Field Descriptions	47
20	L1D Invalidate Base Address Register (L1DIBAR) Field Descriptions	48
21	L1D Invalidate Word Count Register (L1DIWC) Field Descriptions	48
22	L2 Writeback All Register (L2WB) Field Descriptions	49
23	L2 Writeback–Invalidate All Register (L2WBINV) Field Descriptions	50
24	Memory Attribute Register (MAR) Field Descriptions	51
25	L1D Mode Setting Using DCC Field	52
26	L1P Mode Setting Using PCC Field	53
27	L2 Mode Switch Procedure	55
28	Memory Attribute Registers	57
29	Summary of Program-Initiated Cache Operations	60
30	L2ALLOC Default Queue Allocations	68
31	Coherence Assurances in the Two-Level Memory System	70
32	Program Order for Memory Operations Issued From a Single Execute Packet	76
33	Document Revision History	81

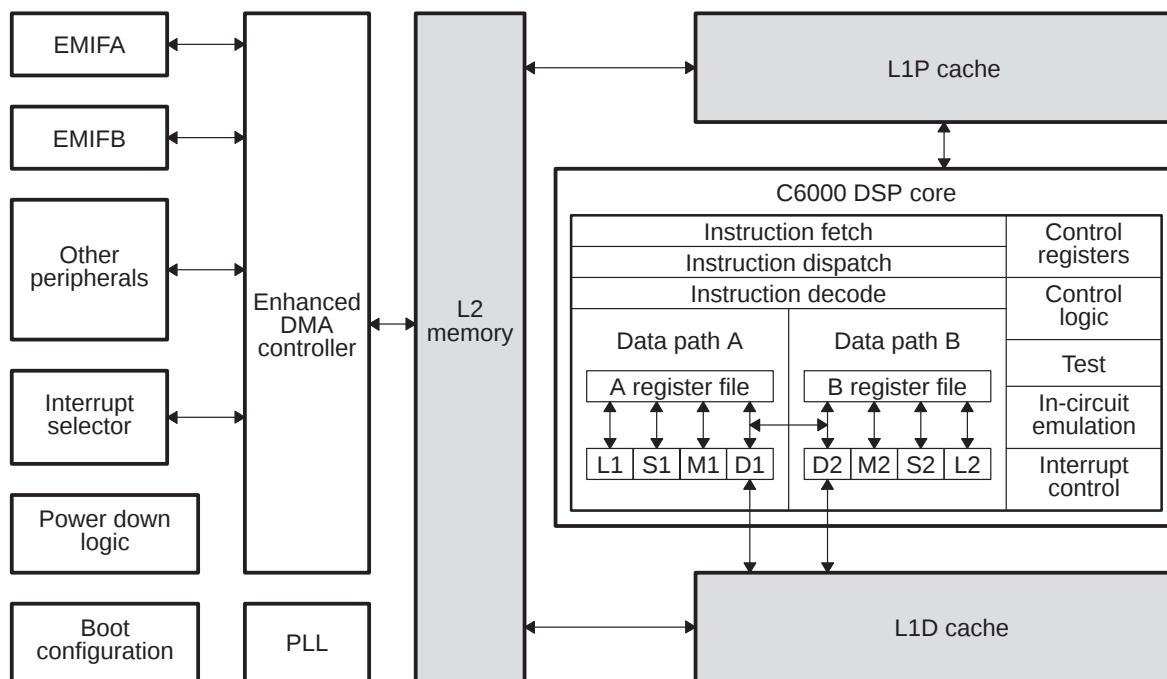
TMS320C64x Two-Level Internal Memory

The TMS320C621x, TMS320C671x, and TMS320C64x™ digital signal processors (DSPs) of the TMS320C6000™ DSP family have a two-level memory architecture for program and data. The first-level program cache is designated L1P, and the first-level data cache is designated L1D. Both the program and data memory share the second-level memory, designated L2. L2 is configurable, allowing for various amounts of cache and SRAM. This document discusses the C64x™ two-level internal memory. For a discussion of the C621x/C671x two-level internal memory, see *TMS320C621x/C671x DSP Two-Level Internal Memory Reference Guide* (SPRU609).

1 Memory Hierarchy Overview

Figure 1 shows the block diagram of the C64x DSP. Table 1 summarizes the differences between the C621x/C671x and C64x internal memory. Figure 2 illustrates the bus connections between the CPU, internal memories, and the enhanced DMA (EDMA) of the C6000™ DSP.

Figure 1. TMS320C64x DSP Block Diagram



Note: EMIFB is available only on certain C64x devices. Refer to the device-specific data sheet for the available peripheral set.

Table 1. TMS320C621x/C671x/C64x Internal Memory Comparison

	TMS320C621x/C671x DSP	TMS320C64x DSP
Internal memory structure	Two Level	
L1P size	4 Kbytes	16 Kbytes
L1P organization	Direct mapped	
L1P CPU access time	1 cycle	
L1P line size	64 bytes	32 bytes
L1P read miss action	1 line allocated in L1P	
L1P read hit action	Data read from L1P	
L1P write miss action	L1P writes not supported	
L1P write hit action	L1P writes not supported	
L1P → L2 request size	2 fetches/L1P line	1 fetch/L1P line
L1P protocol	Read Allocate	Read Allocate; Pipelined Misses
L1P memory	Single-cycle RAM	
L1P → L2 single request stall	5 cycles for L2 hit	8 cycles for L2 hit
L1P → L2 minimum cycles between pipelined misses	Pipelined misses not supported	1 cycle
L1D size	4 Kbytes	16 Kbytes
L1D organization	2-way set associative	
L1D CPU access time	1 cycle	
L1D line size	32 bytes	64 bytes
L1D replacement strategy	2-way Least Recently Used	
L1D banking	64-bit-wide dual-ported RAM	8 × 32 bit banks
L1D read miss action	1 line allocated in L1D	
L1D read hit action	Data read from L1D	
L1D write miss action	No allocation in L1D, data sent to L2	
L1D write hit action	Data updated in L1D; line marked dirty	
L1D protocol	Read Allocate	Read allocate; Pipelined Misses
L1D → L2 request size	2 fetches/L1D line	

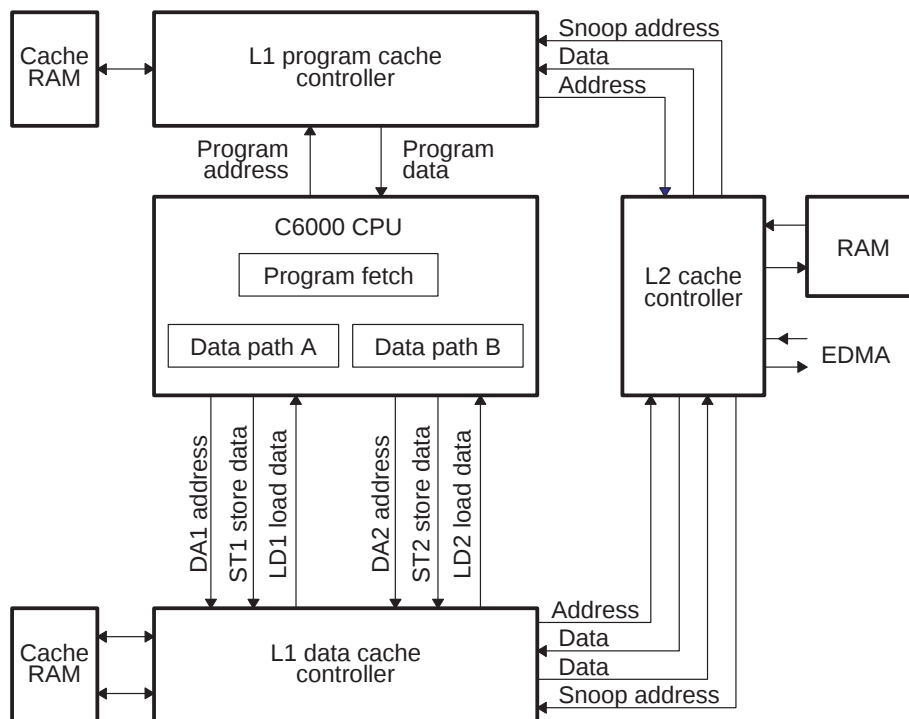
† Some C64x devices may not support the 256K cache mode. Refer to the device-specific datasheet.

Table 1. TMS320C621x/C671x/C64x Internal Memory Comparison (Continued)

	TMS320C621x/C671x DSP	TMS320C64x DSP
L1D → L2 single request stall	4 cycles for L2 hit	6 cycles/L2 SRAM hit
L1D → L2 minimum cycles between pipelined misses	Pipelined misses not supported	2 cycles
L2 total size	Varies by part number. Refer to the datasheet for the specific device.	
L2 SRAM size	Varies by part number. Refer to the datasheet for the specific device.	
L2 cache size [†]	0/16/32/48/64 Kbytes	0/32/64/128/256 Kbytes
L2 organization	1/2/3/4-way set associative	4-way set associative cache
L2 line size	128 bytes	
L2 replacement strategy	1/2/3/4-way Least Recently Used	4-way Least Recently Used
L2 banking	4 × 64 bit banks	8 × 64 bit banks
L2-L1P protocol	Coherency invalidates	
L2-L1D protocol	Coherency snoop-invalidates	Coherency snoops and snoop-invalidates
L2 protocol	Read and Write Allocate	
L2 read miss action	Data is read via EDMA into newly allocated line in L2; requested data is passed to the requesting L1	
L2 read hit action	Data read from L2	
L2 write miss action	Data is read via EDMA into newly allocated line in L2; write data is then written to the newly allocated line.	
L2 write hit action	Data is written into hit L2 location	
L2 → L1P read path width	256 bit	
L2 → L1D read path width	128 bit	256 bit
L1D → L2 write path width	32 bit	64 bit
L1D → L2 victim path width	128 bit	256 bit
L2 → EDMA read path width	64 bit	
L2 → EDMA write path width	64 bit	

[†] Some C64x devices may not support the 256K cache mode. Refer to the device-specific datasheet.

Figure 2. TMS320C64x Two-Level Internal Memory Block Diagram



2 Cache Terms and Definitions

Table 2 lists the terms used throughout this document that relate to the operation of the C64x two-level memory hierarchy.

Table 2. Terms and Definitions

Term	Definition
Allocation	The process of finding a location in the cache to store newly cached data. This process can include <i>evicting</i> data that is presently in the cache to make room for the new data.
Associativity	The number of <i>line frames</i> in each <i>set</i> . This is specified as the number of <i>ways</i> in the cache.
Capacity miss	A cache <i>miss</i> that occurs because the cache does not have sufficient room to hold the entire <i>working set</i> for a program. Compare with <i>compulsory miss</i> and <i>conflict miss</i> .
Clean	A cache <i>line</i> that is <i>valid</i> and that has not been written to by upper levels of memory or the CPU. The opposite state for a valid cache line is <i>dirty</i> .
Coherence	Informally, a memory system is coherent if any read of a data item returns the most recently written value of that data item. This includes accesses by the CPU and the EDMA. Cache coherence is covered in more detail in section 8.1.
Compulsory miss	Sometimes referred to as a <i>first-reference miss</i> . A compulsory miss is a cache <i>miss</i> that must occur because the data has had no prior opportunity to be allocated in the cache. Typically, compulsory misses for particular pieces of data occur on the first access of that data. However, some cases can be considered compulsory even if they are not the first reference to the data. Such cases include repeated write misses on the same location in a cache that does not <i>write allocate</i> , and cache misses to noncacheable locations. Compare with <i>capacity miss</i> and <i>conflict miss</i> .
Conflict miss	A cache <i>miss</i> that occurs due to the limited <i>associativity</i> of a cache, rather than due to capacity constraints. A <i>fully-associative cache</i> is able to allocate a newly cached <i>line</i> of data anywhere in the cache. Most caches have much more limited associativity (see <i>set-associative cache</i>), and so are restricted in where they may place data. This results in additional cache misses that a more flexible cache would not experience.
Direct-mapped cache	A direct-mapped cache maps each address in the <i>lower-level memory</i> to a single location in the cache. Multiple locations may map to the same location in the cache. This is in contrast to a multi-way <i>set-associative cache</i> , which selects a place for the data from a <i>set</i> of locations in the cache. A direct-mapped cache can be considered a single-way set-associative cache.
Dirty	In a <i>writeback cache</i> , writes that reach a given level in the memory hierarchy may update that level, but not the levels below it. Thus, when a cache <i>line</i> is <i>valid</i> and contains updates that have not been sent to the next lower level, that line is said to be dirty. The opposite state for a valid cache line is <i>clean</i> .

Table 2. Terms and Definitions (Continued)

Term	Definition
DMA	Direct Memory Access. Typically, a DMA operation copies a block of memory from one range of addresses to another, or transfers data between a peripheral and memory. On the C64x DSP, DMA transfers are performed by the enhanced DMA (EDMA) engine. These DMA transfers occur in parallel to program execution. From a cache <i>coherence</i> standpoint, EDMA accesses can be considered accesses by a parallel processor.
Eviction	The process of removing a <i>line</i> from the cache to make room for newly cached data. Eviction can also occur under user control by requesting a <i>writeback-invalidate</i> for an address or range of addresses from the cache. The evicted line is referred to as the <i>victim</i> . When a victim line is <i>dirty</i> (that is, it contains updated data), the data must be written out to the next level memory to maintain <i>coherency</i> .
Execute packet	A block of instructions that begin execution in parallel in a single cycle. An execute packet may contain between 1 and 8 instructions.
Fetch packet	A block of 8 instructions that are fetched in a single cycle. One fetch packet may contain multiple <i>execute packets</i> , and thus may be consumed over multiple cycles.
First-reference miss	A cache <i>miss</i> that occurs on the first reference to a piece of data. First-reference misses are a form of <i>compulsory miss</i> .
Fully-associative cache	A cache that allows any memory address to be stored at any location within the cache. Such caches are very flexible, but usually not practical to build in hardware. They contrast sharply with <i>direct-mapped caches</i> and <i>set-associative caches</i> , both of which have much more restrictive <i>allocation</i> policies. Conceptually, fully-associative caches are useful for distinguishing between <i>conflict misses</i> and <i>capacity misses</i> when analyzing the performance of a direct-mapped or set-associative cache. In terms of set-associative caches, a fully-associative cache is equivalent to a set-associative cache that has as many <i>ways</i> as it does <i>line frames</i> , and that has only one <i>set</i> .
Higher-level memory	In a hierarchical memory system, higher-level memories are memories that are closer to the CPU. The highest level in the memory hierarchy is usually the Level 1 caches. The memories at this level exist directly next to the CPU. Higher-level memories typically act as caches for data from <i>lower-level memory</i> .
Hit	A cache hit occurs when the data for a requested memory location is present in the cache. The opposite of a hit is a <i>miss</i> . A cache hit minimizes stalling, since the data can be fetched from the cache much faster than from the source memory. The determination of hit versus miss is made on each level of the memory hierarchy separately—a miss in one level may hit in a lower level.

Table 2. Terms and Definitions (Continued)

Term	Definition
Invalidate	The process of marking <i>valid</i> cache <i>lines</i> as invalid in a particular cache. Alone, this action discards the contents of the affected cache lines, and does not write back any updated data. When combined with a <i>writeback</i> , this effectively updates the next lower level of memory that holds the data, while completely removing the cached data from the given level of memory. Invalidates combined with writebacks are referred to as <i>writeback-invalidates</i> , and are commonly used for retaining <i>coherence</i> between caches.
Least Recently Used (LRU) allocation	For set-associative and <i>fully-associative caches</i> , least-recently used <i>allocation</i> refers to the method used to choose among <i>line frames</i> in a <i>set</i> when allocating space in the cache. When all of the line frames in the set that the address maps to contain <i>valid</i> data, the line frame in the set that was read or written the least recently (furthest back in time) is selected to hold the newly cached data. The selected line frame is then <i>evicted</i> to make room for the new data.
Line	A cache line is the smallest block of data that the cache operates on. The cache line is typically much larger than the size of data accesses from the CPU or the next higher level of memory. For instance, although the CPU may request single bytes from memory, on a read <i>miss</i> the cache reads an entire line's worth of data to satisfy the request.
Line frame	A location in a cache that holds cached data (one <i>line</i>), an associated <i>tag</i> address, and status information for the line. The status information can include whether the line is <i>valid</i> , <i>dirty</i> , and the current state of that line's <i>LRU</i> .
Line size	The size of a single cache <i>line</i> , in bytes.
Load through	When a CPU request misses both the first-level and second-level caches, the data is fetched from the external memory and stored to both the first-level and second-level cache simultaneously. A cache that stores data and sends that data to the upper-level cache at the same time is a load-through cache. Using a load-through cache reduces the stall time compared to a cache that first stores the data in a lower level and then sends it to the higher-level cache as a second step.
Long-distance access	Accesses made by the CPU to a noncacheable memory. Long-distance accesses are used when accessing external memory that is not marked as cacheable.
Lower-level memory	In a hierarchical memory system, lower-level memories are memories that are further from the CPU. In a C64x system, the lowest level in the hierarchy includes the system memory below L2 and any memory-mapped peripherals.
LRU	Least Recently Used. See <i>least recently used allocation</i> for a description of the LRU replacement policy. When used alone, LRU usually refers to the status information that the cache maintains for identifying the least-recently used <i>line</i> in a <i>set</i> . For example, consider the phrase "accessing a cache line updates the LRU for that line."

Table 2. Terms and Definitions (Continued)

Term	Definition
Memory ordering	Defines what order the effects of memory operations are made visible in memory. (This is sometimes referred to as consistency.) Strong memory ordering at a given level in the memory hierarchy indicates it is not possible to observe the effects of memory accesses in that level of memory in an order different than program order. Relaxed memory ordering allows the memory hierarchy to make the effects of memory operations visible in a different order. Note that strong ordering does not require that the memory system execute memory operations in program order, only that it makes their effects visible to other requestors in an order consistent with program order. Section 8.3 covers the memory ordering assurances that the C64x memory hierarchy provides.
Miss	A cache miss occurs when the data for a requested memory location is not in the cache. A miss may stall the requestor while the <i>line frame</i> is allocated and data is fetched from the next lower level of memory. In some cases, such as a CPU write miss from L1D, it is not strictly necessary to stall the CPU. Cache misses are often divided into three categories: <i>compulsory misses</i> , <i>conflict misses</i> , and <i>capacity misses</i> .
Miss pipelining	The process of servicing a single cache <i>miss</i> is pipelined over several cycles. By pipelining the miss, it is possible to overlap the processing of several misses, should many occur back-to-back. The net result is that much of the overhead for the subsequent misses is hidden, and the incremental stall penalty for the additional misses is much smaller than that for a single miss taken in isolation.
Read allocate	A read-allocate cache only allocates space in the cache on a read <i>miss</i> . A write miss does not cause an <i>allocation</i> to occur unless the cache is also a <i>write-allocate</i> cache. For caches that do not write allocate, the write data would be passed on to the next lower-level cache.
Set	A collection of <i>line frames</i> in a cache that a single address can potentially reside. A <i>direct-mapped cache</i> contains one line frame per set, and an N-way <i>set-associative cache</i> contains N line frames per set. A <i>fully-associative cache</i> has only one set that contains all of the line frames in the cache.
Set-associative cache	A set-associative cache contains multiple <i>line frames</i> that each <i>lower-level memory</i> location can be held in. When allocating room for a new <i>line</i> of data, the selection is made based on the <i>allocation</i> policy for the cache. The C64x devices employ a <i>least recently used allocation</i> policy for its set-associative caches.
Snoop	A method by which a <i>lower-level memory</i> queries a <i>higher-level memory</i> to determine if the higher-level memory contains data for a given address. The primary purpose of snoops is to retain <i>coherency</i> , by allowing a lower-level memory to request updates from a higher-level memory. A snoop operation may trigger a <i>writeback</i> , or more commonly, a <i>writeback-invalidate</i> . Snoops that trigger writeback-invalidates are sometimes called snoop-invalidates.

Table 2. Terms and Definitions (Continued)

Term	Definition
Tag	A storage element containing the most-significant bits of the address stored in a particular <i>line</i> . Tag addresses are stored in special tag memories that are not directly visible to the CPU. The cache queries the tag memories on each access to determine if the access is a <i>hit</i> or a <i>miss</i> .
Thrash	An algorithm is said to thrash the cache when its access pattern causes the performance of the cache to suffer dramatically. Thrashing can occur for multiple reasons. One possible situation is that the algorithm is accessing too much data or program code in a short time frame with little or no reuse. That is, its <i>working set</i> is too large, and thus the algorithm is causing a significant number of <i>capacity misses</i> . Another situation is that the algorithm is repeatedly accessing a small group of different addresses that all map to the same set in the cache, thus causing an artificially high number of <i>conflict misses</i> .
Touch	A memory operation on a given address is said to touch that address. Touch can also refer to reading array elements or other ranges of memory addresses for the sole purpose of allocating them in a particular level of the cache. A CPU-centric loop used for touching a range of memory in order to allocate it into the cache is often referred to as a touch loop. Touching an array is a form of software-controlled prefetch for data.
Valid	When a cache <i>line</i> holds data that has been fetched from the next level memory, that <i>line frame</i> is valid. The invalid state occurs when the line frame holds no data, either because nothing has been cached yet, or because previously cached data has been <i>invalidated</i> for whatever reason (<i>coherence</i> protocol, program request, etc.). The valid state makes no implications as to whether the data has been modified since it was fetched from the <i>lower-level memory</i> ; rather, this is indicated by the <i>dirty</i> or <i>clean</i> state of the line.
Victim	When space is allocated in a <i>set</i> for a new <i>line</i> , and all of the <i>line frames</i> in the set that the address maps to contain <i>valid</i> data, the cache controller must select one of the valid lines to <i>evict</i> in order to make room for the new data. Typically, the <i>least-recently used</i> (LRU) line is selected. The line that is <i>evicted</i> is known as the victim line. If the victim line is <i>dirty</i> , its contents are written to the next lower level of memory using a <i>victim writeback</i> .
Victim Buffer	A special buffer that holds <i>victims</i> until they are <i>written back</i> . Victim <i>lines</i> are moved to the victim buffer to make room in the cache for incoming data.
Victim Writeback	When a <i>dirty line</i> is <i>evicted</i> (that is, a line with updated data is evicted), the updated data is written to the lower levels of memory. This process is referred to as a victim writeback.
Way	In a <i>set-associative cache</i> , each <i>set</i> in the cache contains multiple <i>line frames</i> . The number of line frames in each set is referred to as the number of ways in the cache. The collection of corresponding line frames across all sets in the cache is called a way in the cache. For instance, a 4-way set-associative cache has 4 ways, and each set in the cache has 4 line frames associated with it, one associated with each of the 4 ways. As a result, any given cacheable address in the memory map has 4 possible locations it can map to in a 4-way set-associative cache.

Table 2. Terms and Definitions (Continued)

Term	Definition
Working set	The working set for a program or algorithm is the total set of data and program code that is referenced within a particular period of time. It is often useful to consider the working set on an algorithm-by-algorithm basis when analyzing upper levels of memory, and on a whole-program basis when analyzing lower levels of memory.
Write allocate	A write-allocate cache allocates space in the cache when a write <i>miss</i> occurs. Space is allocated according to the cache's <i>allocation</i> policy (<i>LRU</i> , for example), and the data for the <i>line</i> is read into the cache from the next lower level of memory. Once the data is present in the cache, the write is processed. For a <i>writeback cache</i> , only the current level of memory is updated—the write data is not immediately passed to the next level of memory.
Writeback	The process of writing updated data from a <i>valid</i> but <i>dirty</i> cache <i>line</i> to a <i>lower-level memory</i> . After the writeback occurs, the cache line is considered <i>clean</i> . Unless paired with an <i>invalidate</i> (as in <i>writeback-invalidate</i>), the line remains valid after a writeback.
Writeback cache	A writeback cache will only modify its own data on a write <i>hit</i> . It will not immediately send the update to the next lower-level of memory. The data will be written back at some future point, such as when the cache <i>line</i> is <i>evicted</i> , or when the lower-level memory <i>snoops</i> the address from the <i>higher-level memory</i> . It is also possible to directly initiate a <i>writeback</i> for a range of addresses using cache control registers. A write hit to a writeback cache causes the corresponding line to be marked as <i>dirty</i> —that is, the line contains updates that have yet to be sent to the lower levels of memory.
Writeback-invalidate	A writeback operation followed by an invalidation. See <i>writeback</i> and <i>invalidate</i> . On the C64x devices, a writeback-invalidate on a group of cache <i>lines</i> only writes out data for <i>dirty</i> cache lines, but invalidates the contents of all of the affected cache lines.
Write merging	Write merging combines multiple independent writes into a single, larger write. This improves the performance of the memory system by reducing the number of individual memory accesses it needs to process. For instance, on the C64x device, the L1D write buffer can merge multiple writes under some circumstances if they are to the same double-word address. In this example, the result is a larger effective write-buffer capacity and a lower bandwidth impact on L2.
Write-through cache	A write-through cache passes all writes to the <i>lower-level memory</i> . It never contains updated data that it has not passed on to the lower-level memory. As a result, cache <i>lines</i> can never be <i>dirty</i> in a write-through cache. The C64x devices do not utilize write-through caches.

3 Level 1 Data Cache (L1D)

The level 1 data cache (L1D) services data accesses from the CPU. The following sections describe the parameters and operation of the L1D. The operation of L1D is controlled by various registers, as described in section 7, *Memory System Controls*.

3.1 L1D Parameters

The L1D is a 16K-byte cache. It is a two-way set associative cache with a 64-byte line size and 128 sets. It also features a 64-bit by 4-entry write buffer between L1D and the L2 memory.

Physical addresses map onto the cache in a straightforward manner. The physical address divides into three fields as shown in Figure 3. Bits 5–0 of the address specify an offset within the line. Bits 12–6 of the address select one of the 128 sets within the cache. Bits 31–13 of the address serve as the tag for the line.

Figure 3. L1D Address Allocation

31	13 12	6 5	0
Tag		Set Index	Offset

Because L1D is a two-way cache, each set contains two cache lines, one for each way. On each access, the L1D compares the tag portion of the address for the access to the tag information for both lines in the appropriate set. If the tag matches one of the lines and that line is marked valid, the access is a hit. If these conditions are not met, the access is a miss. Miss penalties are discussed in detail under section 3.2.

The L1D is a read-allocate-only cache. This means that new lines are allocated in L1D for read misses, but not for write misses. For this reason, a 4-entry write buffer exists between the L1D and L2 caches that captures data from write misses. The write buffer is enhanced in comparison to the write buffer on the C621x/C671x devices. The write buffer is described in section 3.2.3.

The L1D implements a least-recently used (LRU) line allocation policy. This means that on an L1D read miss, the L1D evicts the least-recently read or written line within a set in order to make room for the incoming data. Note that invalid lines are always considered least-recently used.

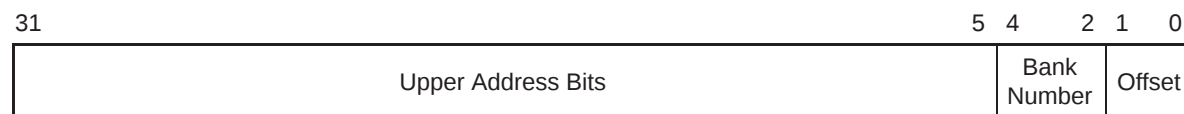
If the selected line is dirty, that is, its contents are updated, then the victim line's data is prepared for writeback to L2 as a victim writeback. The actual victim writeback occurs after the new data is fetched, and then only if the newly fetched data is considered cacheable. If the newly fetched data is noncacheable, the victim writeback is cancelled and the victim line remains in the L1D cache.

3.2 L1D Performance

3.2.1 L1D Memory Banking

The C64x DSP has a least-significant bit (LSB) based memory banking structure that is similar to the structure employed by the C620x/C670x families. The L1D on C64x devices divides memory into eight 32-bit-wide banks. These banks are single-ported, allowing only one access per cycle. This is in contrast to the C621x/C671x devices, which use a single bank of dual-ported memory rather than multiple banks of single-ported memory. In Figure 4, bits 4–2 of the address select the bank and bits 1–0 select the byte within the bank.

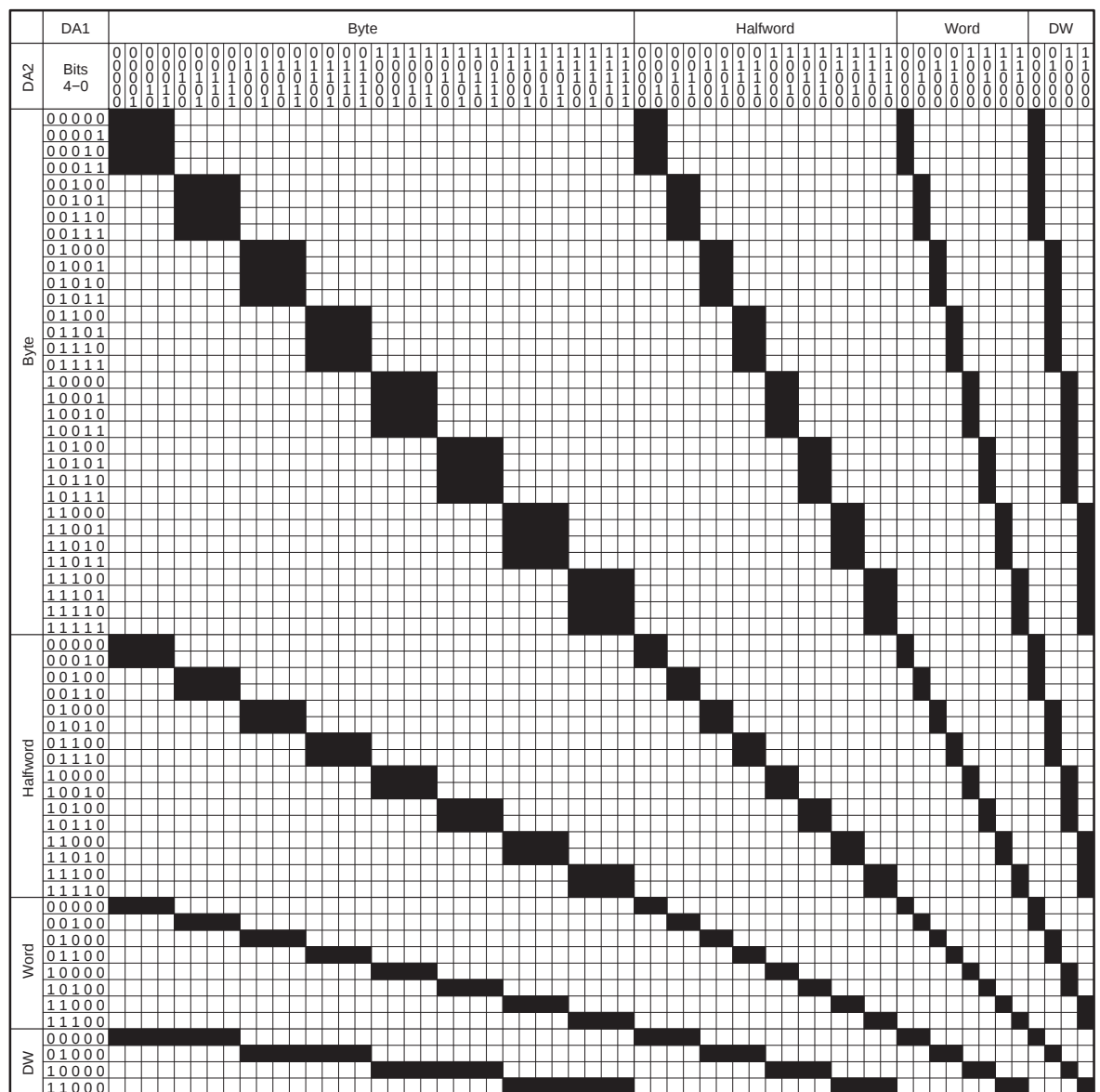
Figure 4. Address to Bank Number Mapping



The shaded areas in Figure 5 show combinations of parallel accesses that may result in bank-conflict stalls according to the LSBs of addresses for the two accesses. Two simultaneous accesses to the same bank incur a one-cycle stall penalty, except under the following special cases:

- ☐ The memory accesses are both writes to nonoverlapping bytes within the same word. That is, bits 31–2 of the address are the same.
- ☐ The memory accesses are both reads that access all or part of the same word. That is, bits 31–2 of the address are the same. In this case, the two accesses may overlap.
- ☐ One or both of the memory accesses is a write that misses L1D and is serviced by the write buffer instead. (See section 3.2.3 for information on the write buffer.)
- ☐ The memory accesses form a single nonaligned access. Nonaligned accesses do not cause bank-conflict stalls, even though the memory system may subdivide them into multiple accesses.

Notice that a read access and a write access in parallel to the same bank will always cause a stall. Two reads or two writes to the same bank may not stall as long as the above conditions are met.



While similar to C620x/C670x devices memory banking, the C64x device banking scheme differs in two ways:

- ❑ Banks are 16-bits wide on the C620x/C670x devices, and 32-bits wide on the C64x device.
- ❑ The C620x/C670x devices always stall on parallel accesses to the same bank, regardless of whether those accesses are to the same word. The C64x device allows many special cases to proceed without stalls, as previously described.

3.2.2 L1D Miss Penalty

The L1D can service up to two data accesses from the CPU every cycle. Accesses that hit L1D complete without stalls, unless a bank conflict occurs as described in section 3.2.1.

Reads that miss L1D stall the CPU while the requested data is fetched. The L1D is a read-allocate cache, and so it will allocate a new line for the requested data, as described in section 3.1. An isolated L1D read miss that hits L2 SRAM stalls the CPU for 6 cycles, and an isolated L1D read miss that hits L2 cache stalls the CPU for 8 cycles. This assumes there is no other memory traffic in L2 that delays the processing of requests from L1D. Section 5.4 discusses interactions between the various requestors that access L2.

An L1D read miss that also misses L2 stalls the CPU while the L2 retrieves the data from external memory. Once the data is retrieved, it is stored in L2 and transferred to the L1D. The external miss penalty varies depending on the type and width of external memory used to hold external data, as well as other aspects of system loading. Section 5.2 describes how L2 handles cache misses on behalf of L1D.

If there are two read misses to the same line in the same cycle, only one miss penalty is incurred. Similarly, if there are two accesses in succession to the same line and the first one is a miss, the second access will not incur any additional miss penalty.

The process of allocating a line in L1D can result in a victim writeback. Victim writebacks move updated data out of L1D to the lower levels of memory. When updated data is evicted from L1D, the cache moves the data to the victim buffer. Once the data is moved to the victim buffer, the L1D resumes processing of the current read miss. Further processing of the victim writeback occurs in the background. Subsequent read and write misses, however, must wait for the victim writeback to be processed. As a result, victim writebacks can noticeably lengthen the time for servicing cache misses.

The L1D pipelines read misses. Consecutive read misses to different lines may be overlapped, reducing the overall stall penalty. The incremental stall penalty can be as small as 2 cycles per miss. Section 3.2.4 discusses miss pipelining.

Write misses do not stall the CPU directly. Rather, write misses are queued in the write buffer that exists between L1D and L2. Although the CPU does not always stall for write misses, the write buffer can stall the CPU under various circumstances. Section 3.2.3 describes the effects of the write buffer.

3.2.3 L1D Write Buffer

The L1D does not write allocate. Rather, write misses are passed directly to L2 without allocating a line in L1D. A write buffer exists between the L1D cache and the L2 memory to capture these write misses. The write buffer provides a 64-bit path for writes from L1D to L2 with room for four outstanding write requests.

Writes that miss L1D do not stall the CPU unless the write buffer is full. If the write buffer is full, a write miss will stall the CPU until there is room in the buffer for the write. The write buffer can also indirectly stall the CPU by extending the time for a read miss. Reads that miss L1D will not be processed as long as the write buffer is not empty. Once the write buffer has emptied, the read miss will be processed. This is necessary as a read miss may overlap an address for which a write is pending in the write buffer.

The L2 can process a new request from the write buffer every cycle, provided that the requested L2 bank is not busy. Section 5.3 describes the L2 banking structure and its impact on performance.

The C64x write buffer allows merging of write requests. It merges two write misses into a single transaction providing all of the following rules are obeyed:

- ☐ The double-word addresses (that is, the upper 29 bits) for the two accesses are the same.
- ☐ The two writes are to locations in L2 SRAM (*not* locations that may be held in L2 cache).
- ☐ The first write has just been placed in the write buffer queue.
- ☐ The second write is presently being placed in the buffer queue.
- ☐ The first write has not yet been presented to the L2 controller.

The previous conditions occur in a number of situations, such as when a program makes a large series of sequential writes or when it makes a burst of small writes to a structure in memory. Write merging increases the effective capacity of the write buffer in these cases by reducing the number of independent stores that are present in the write buffer. This reduces the stall penalty for programs with a large number of write misses.

As a secondary benefit, write merging reduces the number of memory operations executed in L2. This improves the overall performance of the L2 memory by reducing the total number of individual write operations L2 must process. Adjacent accesses are combined into a single access to an L2 bank, rather than multiple accesses to that bank. This allows other requestors to access that bank more quickly, and it allows the CPU to move on to the next bank immediately in the next cycle.

3.2.4 L1D Miss Pipelining

The L1D cache pipelines read misses. A single L1D read miss takes 6 cycles when serviced from L2 SRAM, and 8 cycles when serviced from L2 cache. Miss pipelining can hide much of this overhead by overlapping the processing of several cache misses.

For L1D miss pipelining to be effective, there must be multiple outstanding L1D read misses. Load instructions on the C64x DSP have a 5-cycle-deep pipeline, and the C64x DSP may issue up to two accesses per cycle. In this pipeline, the L1D performs tag comparisons in one pipeline stage (E2), and services cache hits and misses on the following stage (E3). Cache read misses result in a CPU stall.

L1D processes single read misses only when there are no outstanding victim writebacks and when the write buffer is empty. When two cache misses occur in parallel, the L1D processes the misses in program order. (The program order is described in section 8.3.1.) In the case of two write misses, the misses are inserted in the write buffer and the CPU does not stall unless the write buffer is full. (Section 3.2.3 describes the write buffer.) In the case of two read misses or a read and a write miss, the misses are overlapped as long as they are to different sets, that is, their addresses differ in bits 13–6.

Cache misses are processed in the E3 pipeline stage. Once L1D has issued commands to L2 for all of the cache misses in E3, the L1D may decide to advance its state internally by one pipeline stage to consider cache misses due to accesses that were in the E2 pipeline stage. This allows L1D to aggressively overlap requests for cache misses that occur in parallel and cache misses that occur on consecutive cycles. L1D considers the accesses in E2 only if the write buffer and victim writeback buffer are empty. Although the L1D internal state advances, the CPU stall is not released until the data returns for accesses that were in the E3 stage.

Once the CPU stall is released, memory accesses that were in the E2 stage advance to the E3 pipeline stage. This may bring one or two new accesses into the E2 pipeline stage. It also potentially brings one or two unprocessed cache misses from E2 into E3. The L1D first issues commands for any cache misses that are now in E3 but that have not yet been processed. Once the accesses in E3 are processed, the L1D may consider accesses in E2 as previously described. In any case, the L1D stalls the CPU when there are accesses in E3 that have not yet completed.

The net result is that the L1D can generate a continuous stream of requests to L2. Code that issues pairs of memory reads to different cache lines every cycle will maximize this effect. As noted above, this pipelining can result in improved performance, especially in the presence of sustained read misses.

The incremental miss penalty can be as small as 2 cycles per miss when the L1D is able to overlap the processing for a new cache miss with that of prior misses. Therefore, the average miss penalty for a sustained sequence of back-to-back misses approaches 2 cycles per miss in the ideal case. Table 3 and Table 4 illustrate the performance for various numbers of consecutive L1D read misses that hit in L2 cache and L2 SRAM, assuming all misses are able to overlap. These further assume that there is no other memory traffic in L2 that may lengthen the time required for an L1D cache miss, and that all misses are within the same half of the affected L1D cache lines.

Table 3. Cycles Per Miss for Different Numbers of L1D Misses That Hit L2 Cache

Number of Misses	Total Stall Cycles	Mean Cycles Per Miss
1	8	8
2	10	5
3	12	4
4	14	3.5
> 4, even	$6 + (2 * M)$	$2 + (6 / M)$

Note: M = Number of total misses.

Table 4. Cycles Per Miss for Different Numbers of L1D Misses that Hit L2 SRAM

Number of Misses	Total Stall Cycles	Mean Cycles Per Miss
1	6	6
2	8	4
3	10	3.33
4	12	3
> 4, even	$4 + (2 * M)$	$2 + (4 / M)$

Note: M = Number of total misses.

4 Level 1 Program Cache (L1P)

The level 1 program cache (L1P) services program fetches from the CPU. The following sections describe the parameters and operation of the L1P. The operation of L1P is controlled by various registers, as described in section 7, *Memory System Controls*.

4.1 L1P Parameters

The L1P is a 16K-byte cache. It is a direct-mapped cache with a 32-byte line size and 512 sets.

Physical addresses map onto the cache in a fixed manner. The physical address divides into three fields as shown in Figure 6. Bits 4–0 of the address specify an instruction within a set. Bits 13–5 of the address select one of the 512 sets within the cache. Bits 31–14 of the address serve as the tag for the line.

Figure 6. L1P Address Allocation

31	14 13	5 4	0
Tag		Set Index	Offset

Because L1P is direct-mapped cache, each address maps to a fixed location in the cache. That is, each set contains exactly one line frame. On a cache miss, the cache allocates the corresponding line for the incoming data. Because L1P does not support writes from the CPU, the previous contents of the line are discarded.

4.2 L1P Performance

4.2.1 L1P Miss Penalty

A program fetch which hits L1P completes in a single cycle without stalling the CPU. An L1P miss that hits in L2 may stall the CPU for up to 8 cycles, depending on the parallelism of the execute packets in the vicinity of the miss. Section 4.2.2 describes this in more detail.

An L1P miss that misses in L2 cache stalls the CPU until the L2 retrieves the data from external memory and transfers the data to the L1P, which then returns the data to the CPU. This delay depends upon the type of external memory used to hold the program, as well as other aspects of system loading.

The C64x DSP allows an execute packet to span two fetch packets. This spanning does not change the penalty for a single miss. However, if both fetch packets are not present in L1P, two cache misses occur.

4.2.2 L1P Miss Pipelining

The L1P cache pipelines cache misses. A single L1P cache miss requires 8 cycles to retrieve data from L2. Miss pipelining can hide much of this overhead by overlapping the processing for several cache misses. Additionally, some amount of the cache miss overhead can be overlapped with dispatch stalls that occur in the fetch pipeline.

For L1P miss pipelining to be effective, there must be multiple outstanding cache misses. The C64x DSP fetch pipeline accomplishes this by attempting to fetch one new fetch packet every cycle, so long as there is room in the fetch pipeline. To understand how this works, it is necessary to understand the nature of the fetch pipeline itself.

The fetch and decode pipeline is divided into 6 stages leading up to but not including the first execution stage, E1. The stages are:

- ☐ PG – Program Generate
- ☐ PS – Program Send
- ☐ PW – Program Wait
- ☐ PR – Program Read
- ☐ DP – Dispatch
- ☐ DC – Decode

C6000 DSP instructions are grouped into two groupings: fetch packets and execute packets. The CPU fetches instructions from memory in fixed bundles of 8 instructions, known as fetch packets. The instructions are decoded and separated into bundles of parallel-issue instructions known as execute packets. A single execute packet may contain between 1 and 8 instructions. Thus, a single fetch packet may contain multiple execute packets. On the C64x DSP, an execute packet may also span two fetch packets. The Program Read (PR) stage of the pipeline is responsible for identifying a sequence of execute packets within a sequence of fetch packets. The Dispatch (DP) stage is responsible for extracting and dispatching them to functional units.

As a result of the disparity between fetch packets and execute packets, the entire fetch pipeline need not advance every cycle. Rather, the PR pipeline stage only allows the Program Wait (PW) stage to advance its contents into the PR stage when the DP stage has consumed the complete fetch packet held in PR. The stages before PR advance as needed to fill in gaps. Thus, when there are no cache misses, the early stages of the fetch pipeline are stalled while the DP stage pulls the individual execute packets from the current fetch packet. These stalls are referred to as dispatch stalls.

The C64x DSP takes advantage of these dispatch stalls by allowing the earlier stages of the pipeline to advance toward DP while cache misses for those stages are still pending. Cache misses may be pending for the PR, PW, and PS pipeline stages. Because the DP stage stalls the PR stage with a dispatch stall while it consumes the fetch packets in the PR stage of the pipeline, it is not necessary to expose these cache stalls to the CPU. When a fetch packet is consumed completely, however, the contents of the PW stage must advance into the PR stage. At this point, the CPU is stalled if DP requests an execute packet from PR for which there is still an outstanding cache miss.

When a branch is taken, the fetch packet containing the branch target advances through the fetch pipeline every cycle until the branch target reaches the E1 pipeline stage. Branch targets override the dispatch stall described above. As a result, they do not gain as much benefit from miss pipelining as other instructions. The fetch packets that immediately follow a branch target do benefit, however. Although the code in the fetch packets that follows the branch target may not execute immediately, the branch triggers several consecutive fetches for this code, and thus pipelines any misses for that code. In addition, no stalls are registered for fetch packets that were requested prior to the branch being taken, but that never made it to the DP pipeline stage.

The miss penalty for a single L1P miss is 8 cycles. The second miss in a pair of back-to-back misses will see an incremental stall penalty of up to 2 cycles. Sustained back-to-back misses in straight-line (nonbranching) code incurs an average miss penalty based on the average parallelism of the code. The average miss penalty for a long sequence of sustained misses in straight-line code is summarized in Table 5.

Table 5. Average Miss Penalties for Large Numbers of Sequential Execute Packets

Instructions Per Execute Packet	Average Stalls Per Execute Packet
1	0.125
2	0.125
3	0.688
4	1.500
5	1.813
6	2.375
7	2.938
8	4.000

5 Level 2 Unified Memory (L2)

The level 2 unified memory (L2) can operate as SRAM, cache, or both. It services cache misses from both L1P and L1D, as well as DMA accesses using the EDMA controller. The following sections describe the parameters and operation of the L2. The operation of L2 is controlled by various registers, as described in section 7, *Memory System Controls*.

5.1 L2 Cache and L2 SRAM

The L2 can operate as SRAM, as cache, or as both, depending on its mode. L2 SRAM refers to the portion of L2 mapped as SRAM and L2 cache refers to the portion of L2 that acts as cache. For a given device, the total capacity of L2 SRAM and L2 cache together is fixed, regardless of the L2 mode.

The total size of the L2 depends on the specific C64x device. The C6414, C6415, and C6416 devices provide 1024K bytes of L2 memory. The C6411 and C6412 devices provide 256K bytes of L2 memory. For other C64x devices, consult the data sheet to determine the L2 size for the device.

After reset, the entire L2 memory is mapped as SRAM. The L2 SRAM is contiguous and always starts at address 0000 0000h in the memory map. The C6414, C6415, and C6416 devices map L2 SRAM over the address range 0000 0000h to 000F FFFFh. The C6411 and C6412 devices map L2 SRAM over the address range 0000 0000h to 0003 FFFFh.

The L2 cache is a 4-way set associative cache whose capacity varies between 32K bytes and 256K bytes depending on its mode. The L2 cache is enabled by the L2MODE field in the cache configuration register (CCFG). Enabling L2 cache reduces the amount of available L2 SRAM. Section 7.1.3 discusses how the L2 memory map varies according to cache mode and the specific device being used.

External physical memory addresses map onto the L2 cache differently depending on the cache mode. The physical address divides into three fields as shown in Figure 7, Figure 8, Figure 9, and Figure 10. Bits 6–0 of the address specify an offset within a line. The next 6 to 9 bits (depending on mode) of the address specify the set within the cache. The remaining bits of the address serve as the tag for the line.

Figure 7. L2 Address Allocation, 256K Cache (L2MODE = 111b)

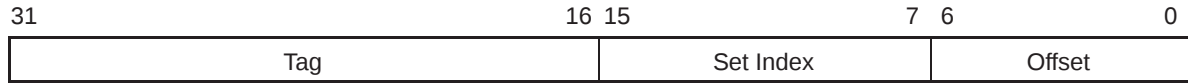


Figure 8. L2 Address Allocation, 128K Cache (L2MODE = 011b)

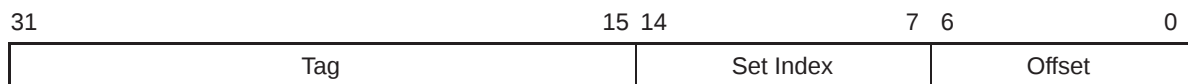


Figure 9. L2 Address Allocation, 64K Cache (L2MODE = 010b)



Figure 10. L2 Address Allocation, 32K Cache (L2MODE = 001b)



5.2 L2 Operation

When L2 cache is enabled, it services requests from L1P and L1D to external addresses. The operation of the L2 cache is similar to the operation of both L1P and L1D. When a request is made, the L2 first determines if the address requested is present in the cache. If the address is present, the access is considered a cache hit and the L2 services the request directly within the cache. If the address is not present, the access results in a cache miss. On a miss, the L2 processes the request according to the cacheability of the affected address.

On a cache hit, the L2 updates the LRU status for the corresponding set in L2 cache. If the access is a read, the L2 returns the requested data. If the access is a write, the L2 updates the contents of the cache line and marks the line as dirty. L2 is a writeback cache, and so write hits in L2 are not immediately forwarded to external memory. The external memory will be updated when this line is later evicted, or is written back using the block-writeback control registers described in section 7.3.2.

The L2 allocates a new line within the L2 cache on a cache miss to a cacheable external memory location. Note that unlike L1D, the L2 allocates a line for both read and write misses. The L2 cache implements a least-recently used policy to select a line within the set to allocate on a cache miss. If this line being replaced is valid, it is evicted as described below. Once space allocated for the new data, an entire L2 line's worth of data is fetched via the EDMA into the allocated line.

Evicting a line from L2 requires several steps, regardless of whether the victim is clean or dirty. For each line in L2, L2 tracks whether the given line is also cached in L1D. If it detects that the victim line is present in L1D, it sends L1D snoop-invalidate requests to remove the affected L1D lines. L1D responds by invalidating the corresponding line. If the line in L1D was dirty, the updated data is passed to L2 and merged with the L2 line that is being evicted. The combined result is written to external memory. If the victim line was not dirty in either L1D or L2, its contents are discarded. These actions ensure that the most recent writes to the affected cache line are written to external memory, but that clean lines are not needlessly written to external memory.

Note that L1P is not consulted when a line is evicted from L2. This allows program code to remain in L1P despite having been evicted from L2. This presumes that program code is never written to. In those rare situations where this is not the case, programs may use the cache controls described in section 7.3.2 to remove cached program code from L1P.

If the cache miss was a read, the data is stored in L2 cache when it arrives. It is also forwarded directly to the requestor, thereby reducing the overall stall time. The portion of the L2 cache line requested by L1 is sent directly to that L1, and is allocated within that cache. The entire L2 line is also stored within the L2 cache as it arrives.

If the cache miss was a write, the incoming data is merged with the write data from L1D, and the merged result is stored in L2. In the case of a write, the line is not immediately allocated in L1D, as L1D does not write allocate.

A cache-miss to a non-cacheable location results in a long-distance access. A long distance read causes the L2 to issue a read transfer for the requested data to the EDMA controller. When the requested data is returned, the L2 forwards the data to the requestor. It does not allocate space for this data in the L2 cache. A long distance read from L1D reads exactly the amount of data requested by the CPU. A long distance read from L1P reads exactly 32 bytes. The data read is not cached by L1D, L1P, or L2.

A long distance write causes the L2 to store a temporary copy of the written data. It then issues a write transfer for the write miss to the EDMA controller. Long distance writes can only originate from L1D using the L1D write buffer. Because the written data is stored in a special holding buffer, it is not necessary to stall the CPU while the long-distance write is being processed. Also, further writes to the L2 SRAM address space or on-chip peripherals may be processed while the long-distance access is being executed.

The L2 cache allows only one long-distance access to be in progress at one time. A long distance write will potentially stall the CPU if a previous long distance write is still in progress. Likewise, a long distance read will not be processed until all previous long distance writes are complete.

The ordering constraints placed on long-distance accesses effectively make accesses to non-cacheable regions of memory strongly ordered. Program order for these accesses is retained, thus making these accesses useful for interfacing to peripherals or synchronizing with other CPUs.

5.3 L2 Bank Structure

The L2 memory is divided into 8 64-bit banks that operate at the CPU's clock rate, but pipeline accesses over two cycles. The L2 begins processing a new request on each cycle, although each bank can process new requests no faster than one every two cycles.

Because each bank requires two cycles to process a request, an access to a bank on one cycle blocks all accesses to that bank on the following cycle. Thus, bank conflicts can occur between accesses on adjacent cycles. Repeated accesses to the same bank are processed at a rate of one every two cycles, regardless of the requestor. This highlights the importance of merging write misses from L1D.

The L2 must service requests from the following sources:

- ☐ L1P read miss that hits in L2.
- ☐ L1D read or write miss that hits in L2.
- ☐ EDMA reads and writes.
- ☐ Internal cache operations (victim writebacks, line fills, snoops).

The internal cache machinery is the fourth source of requests to the L2 memory. These requests represent data movement triggered by normal cache operation. Most of these requests are triggered by cache misses, or by user-initiated cache control operations.

Only one requestor may access the L2 in a given cycle. In the case of an L2 conflict, the L2 prioritizes requests in the above order. That is, L1P read hits have highest priority, followed by L1D, and so on.

L2 contention due simultaneous access or bank conflicts can lengthen the time required for servicing L1P and L1D cache misses. The priority scheme above minimizes this effect by servicing CPU-initiated accesses before EDMA accesses and the background activities performed by the cache.

5.4 L2 Interfaces

The L2 services requests from the L1D cache, L1P cache, and the EDMA. The L2 provides access to its own memory, the peripheral configuration bus (CFGBUS), and the various cache control registers described in section 6, *Registers*. The following sections describe the interaction between L2 and the various requestors that it interfaces.

5.4.1 L1D/L1P-to-L2 Request Servicing

The L2 controller allows the L1P and L1D to access both L2 SRAM and L2 cache. The L2 also acts as intermediary for long-distance accesses to addresses in external memory, and for accesses to on-chip peripherals. For each access, the address and the L2 mode determine the behavior.

Memory accesses to addresses that are mapped as L2 SRAM are serviced directly by L2 and are treated as cacheable by L1P and L1D. Memory accesses to addresses that are outside L2 SRAM but are on-chip are treated as non-cacheable. These accesses include accesses to on-chip peripherals and cache control registers.

L2 SRAM accesses by L1D do not trigger cache operations in either L1P or L2 cache. Likewise, L1P accesses to L2 SRAM do not trigger cache operations in either L1D or L2 cache. This is consistent with the coherence policy outlined in section 8.1. This contrasts with EDMA accesses to L2 SRAM, described in section 5.4.2.

L1P and L1D memory accesses to external addresses may be serviced by the L2 cache, when the cache is enabled and caching is permitted on the given range of addresses. The cacheability of a given external address range is determined by the cache enable bit (CE) in the corresponding memory attribute register (MAR). Section 7.2 describes MAR operation.

When the L2 is in all SRAM mode (L2MODE = 000b), it does not attempt to cache memory accesses directly. However, when servicing read requests for L1P and L1D, the L2 returns the cacheability status of the fetched data, thus allowing L1P and L1D to directly retain copies of memory that is marked as cacheable. Writes and writebacks from L1D for external addresses are sent to external memory directly when in this mode.

Accesses to noncacheable external locations are handled as long-distance accesses. Long distance accesses are described in section 5.2.

Note:

The L2 only supports 32-bit accesses to on-chip peripheral addresses, including its own control registers. Byte, half-word, and double-word accesses may not work as expected.

5.4.2 EDMA-to-L2 Request Servicing

The L2 controller handles EDMA access to L2 SRAM addresses. EDMA read and write requests to L2 SRAM result in specific actions to enforce the memory coherence policy described in section 8.1.

The L2 controller tracks information on what portions of L2 SRAM are held in L1D cache. When the EDMA reads from a location in L2 SRAM, the L2 controller *snoops* L1D for updated data, if necessary, before servicing the read for the EDMA. If L1D contains updated data, the update is written to L2 SRAM before allowing the EDMA read to proceed. The affected line in L1D is left valid but clean. No action is made relative to L1P.

When the EDMA writes to a location in L2 SRAM, the L2 controller sends an invalidate command to L1P and a snoop-invalidate command to L1D, as necessary. For L1P, the invalidate command is always sent since L2 does not track the state of L1P. If the address is present in L1P, the affected line is invalidated. As L1P cannot contain updated data, L2 does not interact further with L1P. For L1D, the snoop-invalidate is sent only if L2 detects that L1D holds a copy of the address being written. If L1D contains updated data, the update is merged with the EDMA write and stored to L2, with the incoming data taking precedence over the data returned by L1D. The affected line in L1D is then marked invalid.

The end result of this system of snoops and invalidates is that coherence is retained between EDMA and CPU accesses to L2 SRAM. The example in section 8.1 illustrates this protocol in action.

5.4.3 L2 Request Servicing Using EDMA

The L2 controller relies on the EDMA to service requests on its behalf. Cache misses, victim writebacks, and long-distance accesses to external addresses cause the L2 controller to issue DMA transfer requests to the EDMA controller.

L2 requests are queued along with other EDMA requests, and are serviced according to the policies set by the EDMA. The L2 requests may be placed on any of the EDMA priority queues. The priority level for requests and the number of outstanding requests permitted can be controlled as described in section 7.4. The EDMA is described in *TMS320C6000 DSP Enhanced Direct Memory Access (EDMA) Peripheral Reference Guide* (SPRU234).

5.4.4 EDMA Access to Cache Controls

The L2 controller manages access to the memory mapped cache control registers. While the L2 controller permits CPU access to these registers, it does not permit EDMA access to any of these registers. EDMA accesses to the cache control registers are dropped. EDMA writes are ignored and EDMA reads return undefined values.

As a result of these restrictions, cache control operations may only be initiated from the CPU. External devices that wish to trigger a cache control operation should do so using an interrupt handler or some other mechanism that triggers CPU writes to the appropriate registers on behalf of the requestor.

5.4.5 HPI and PCI Access to Memory Subsystem

The host port interface (HPI) and peripheral component interconnect (PCI) peripherals are not directly connected to the two-level memory subsystem. HPI and PCI requests are serviced indirectly by the EDMA. Therefore, HPI and PCI accesses follow the same coherence policies and are subject to the same restrictions as EDMA accesses.

Sections 5.4.2 through 5.4.4 describe the interaction between the EDMA and the two-level memory system. Section 8, *Memory System Policies*, describes the coherence and ordering policies for the memory system.

6 Registers

The two-level memory hierarchy is controlled by several memory-mapped control registers listed in Table 6. It is also controlled by the data cache control (DCC) and program cache control (PCC) fields in the CPU control and status register (CSR), see section 7.1. See the device-specific datasheet for the memory address of these registers.

Table 6. Internal Memory Control Registers

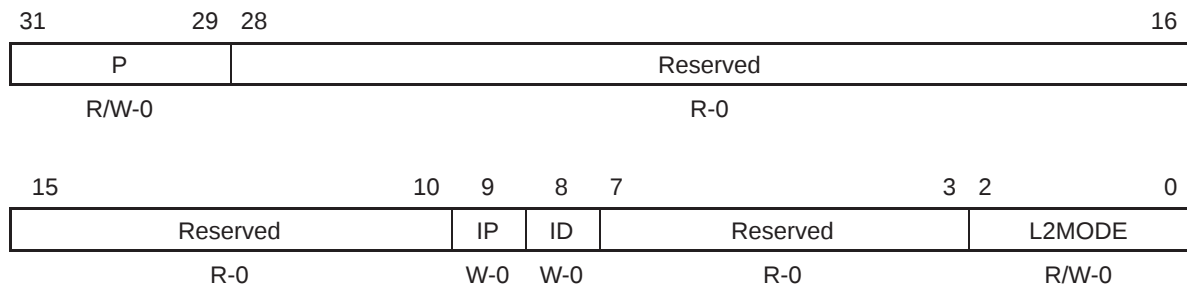
Acronym	Register Name	Section
CCFG	Cache configuration register	6.1
EDMAWEIGHT	L2 EDMA access control register	6.2
L2ALLOC0–3	L2 allocation registers 0–3	6.3
L2WBAR	L2 writeback base address register	6.4
L2WWC	L2 writeback word count register	6.5
L2WIBAR	L2 writeback-invalidate base address register	6.6
L2WIWC	L2 writeback-invalidate word count register	6.7
L2IBAR	L2 invalidate base address register	6.8
L2IWC	L2 invalidate word count register	6.9
L1PIBAR	L1P invalidate base address register	6.10
L1PIWC	L1P invalidate word count register	6.11
L1DWIBAR	L1D writeback-invalidate base address register	6.12
L1DWIWC	L1D writeback-invalidate word count register	6.13
L1DIBAR	L1D invalidate base address register	6.14
L1DIWC	L1D invalidate word count register	6.15
L2WB	L2 writeback all register	6.16
L2WBINV	L2 writeback-invalidate all register	6.17
MAR0–MAR95	Reserved	–
MAR96–MAR111 [†]	L2 memory attribute register 96–111: controls EMIFB CE0–CE3	6.18
MAR112–MAR127	Reserved	–
MAR128–MAR191	L2 memory attribute register 128–191: controls EMIFA CE0–CE3	6.18
MAR192–MAR255	Reserved	–

[†] MAR96–MAR111 only available on the C6414/C6415/C6416 devices; on all other devices, these registers are reserved.

6.1 Cache Configuration Register (CCFG)

The cache configuration register (CCFG) is shown in Figure 11 and described in Table 7.

Figure 11. Cache Configuration Register (CCFG)



Legend: R = Read only; W = Write only; R/W = Read/write; -n = value after reset

Table 7. Cache Configuration Register (CCFG) Field Descriptions

Bit	field [†]	symval [†]	Value	Description
31–29	P	OF(value)	0–7h	L2 requestor priority bits.
		DEFAULT	0	L2 controller requests are placed on urgent priority level.
		URGENT		
		HIGH	1h	L2 controller requests are placed on high priority level.
		MEDIUM	2h	L2 controller requests are placed on medium priority level.
		LOW	3h	L2 controller requests are placed on low priority level.
		–	4h–7h	Reserved
28–10	Reserved	–	0	Reserved. The reserved bit location is always read as 0. Bits in this field should always be written with 0.
9	IP	OF(value)		Invalidate L1P bit.
		DEFAULT	0	Normal L1P operation.
		NORMAL		
		INVALIDATE	1	All L1P lines are invalidated.

[†] For CSL implementation, use the notation `CACHE_CCFG_field_symval`

[‡] C64x devices with less than 256K of L2 memory may not support all L2 cache modes. See the device-specific datasheet to determine the supported cache sizes.

Table 7. Cache Configuration Register (CCFG) Field Descriptions (Continued)

Bit	field [†]	symval [†]	Value	Description
8	ID	OF(value)		Invalidate L1D bit.
		DEFAULT	0	Normal L1D operation.
		NORMAL		
		INVALIDATE	1	All L1D lines are invalidated.
7–3	Reserved	–	0	Reserved. The reserved bit location is always read as 0. Bits in this field should always be written with 0.
2–0	L2MODE [‡]	OF(value)	0–7h	L2 operation mode bits. See the device-specific datasheet to determine the available cache options for a device and how the various L2 cache modes affect the memory map for different sizes of L2.
		DEFAULT	0	L2 cache disabled (all SRAM mode)
		0KC		
		32KC	1h	4-way cache (32K L2 cache)
		64KC	2h	4-way cache (64K L2 cache)
		128KC	3h	4-way cache (128K L2 cache)
		–	4h–6h	Reserved
		256KC	7h	4-way cache (256K L2 cache). Not supported on devices with L2 cache sizes of less than 256-Kbytes memory, such as the C6410 device.

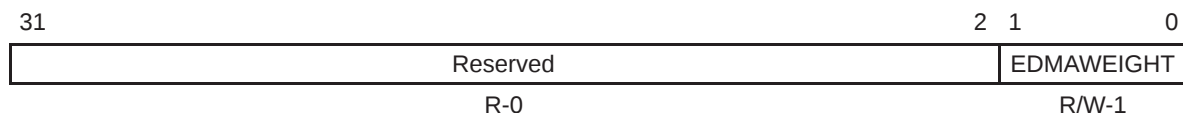
[†] For CSL implementation, use the notation `CACHE_CCFG_field_symval`

[‡] C64x devices with less than 256K of L2 memory may not support all L2 cache modes. See the device-specific datasheet to determine the supported cache sizes.

6.2 L2 EDMA Access Control Register (EDMAWEIGHT)

The L2 EDMA access control register (EDMAWEIGHT) is shown in Figure 12 and described in Table 8.

Figure 12. L2 EDMA Access Control Register (EDMAWEIGHT)



Legend: R = Read only; R/W = Read/Write; -n = value after reset

Table 8. L2 EDMA Access Control Register (EDMAWEIGHT) Field Descriptions

Bit	Field	Value	Description
31-2	Reserved	0	Reserved. The reserved bit location is always read as 0. A value written to this field has no effect.
1-0	EDMAWEIGHT	0-3h	EDMA weight limits the amount of time L1D blocks EDMA access to L2.
		0	L1D access always a higher priority than EDMA access to L2. EDMA never receives priority.
		1h	EDMA receives priority after 16 L1D priority cycles.
		2h	EDMA receives priority after 4 L1D priority cycles.
		3h	EDMA receives priority after 1 L1D priority cycle.

6.3 L2 Allocation Registers (L2ALLOC0–L2ALLOC3)

The L2 allocation registers (L2ALLOC n) are shown in Figure 13 and described in Table 9.

Figure 13. L2 Allocation Registers (L2ALLOC0–L2ALLOC3)

(a) L2ALLOC0

31		3	2	0
Reserved				Q0CNT
R-0				R/W-110

(b) L2ALLOC1

31		3	2	0
Reserved				Q1CNT
R-0				R/W-010

(c) L2ALLOC2

31		3	2	0
Reserved				Q2CNT
R-0				R/W-010

(d) L2ALLOC3

31		3	2	0
Reserved				Q3CNT
R-0				R/W-010

Legend: R = Read only; R/W = Read/write; - n = value after reset

Table 9. L2 Allocation Registers (L2ALLOC0–L2ALLOC3) Field Descriptions

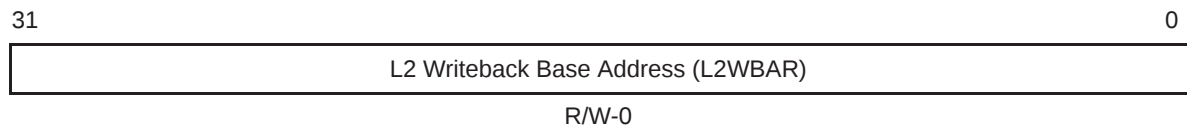
Bit	field [†]	symval [†]	Value	Description
31–3	Reserved	–	0	Reserved. The reserved bit location is always read as 0. Bits in this field should always be written with 0.
2–0	QnCNT	OF(value)	0–7h	The the total number of outstanding L2 and QDMA requests permitted on the corresponding EDMA priority level. Further requests on that priority level are stalled until the number of outstanding requests falls below the QnCNT setting.
		DEFAULT	2h	For L2ALLOC1, L2ALLOC2, and L2ALLOC3.
		DEFAULT	6h	For L2ALLOC0.

[†] For CSL implementation, use the notation `CACHE_L2ALLOCn_field_symval`

6.4 L2 Writeback Base Address Register (L2WBAR)

The L2 writeback base address register (L2WBAR) is shown in Figure 14 and described in Table 10.

Figure 14. L2 Writeback Base Address Register (L2WBAR)



Legend: R/W = Read/write; -n = value after reset

Table 10. L2 Writeback Base Address Register (L2WBAR) Field Descriptions

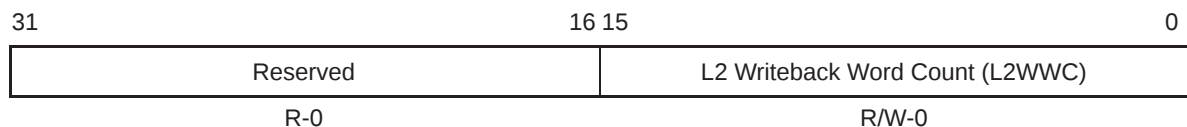
Bit	Field	symval [†]	Value	Description
31–0	L2WBAR	OF(value)	0–FFFF FFFFh	L2 writeback base address.
		DEFAULT	0	

[†] For CSL implementation, use the notation `CACHE_L2WBAR_L2WBAR_symval`

6.5 L2 Writeback Word Count Register (L2WWC)

The L2 writeback word count register (L2WWC) is shown in Figure 15 and described in Table 11.

Figure 15. L2 Writeback Word Count Register (L2WWC)



Legend: R = Read only; R/W = Read/write; -n = value after reset

Table 11. L2 Writeback Word Count Register (L2WWC) Field Descriptions

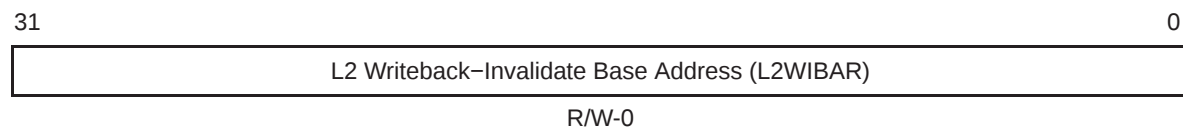
Bit	Field	symval [†]	Value	Description
31–16	Reserved	–	0	Reserved. The reserved bit location is always read as 0. Bits in this field should always be written with 0.
15–0	L2WWC	OF(value)	0–FFFFh	L2 writeback word count.
		DEFAULT	0	

[†] For CSL implementation, use the notation `CACHE_L2WWC_L2WWC_symval`

6.6 L2 Writeback-Invalidate Base Address Register (L2WIBAR)

The L2 writeback-invalidate base address register (L2WIBAR) is shown in Figure 16 and described in Table 12.

Figure 16. L2 Writeback-Invalidate Base Address Register (L2WIBAR)



Legend: R/W = Read/write; -n = value after reset

Table 12. L2 Writeback-Invalidate Base Address Register (L2WIBAR) Field Descriptions

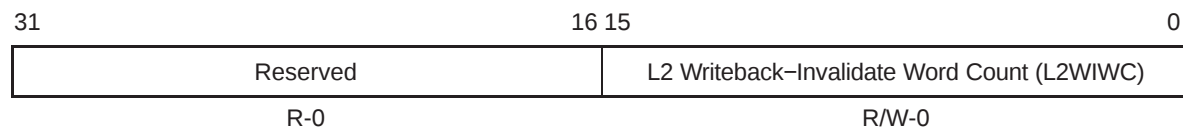
Bit	Field	symval [†]	Value	Description
31–0	L2WIBAR	OF(value)	0–FFFF FFFFh	L2 writeback-invalidate base address.
		DEFAULT	0	

[†] For CSL implementation, use the notation `CACHE_L2WIBAR_L2WIBAR_symval`

6.7 L2 Writeback-Invalidate Word Count Register (L2WIWC)

The L2 writeback-invalidate word count register (L2WIWC) is shown in Figure 17 and described in Table 13.

Figure 17. L2 Writeback-Invalidate Word Count Register (L2WIWC)



Legend: R = Read only; R/W = Read/write; -n = value after reset

Table 13. L2 Writeback-Invalidate Word Count Register (L2WIWC) Field Descriptions

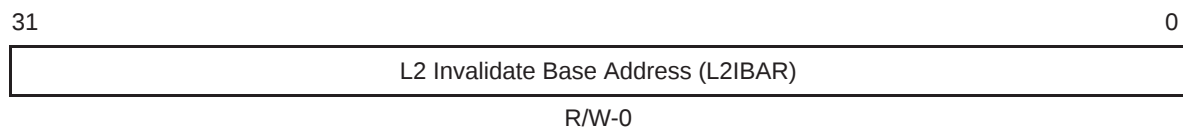
Bit	Field	symval [†]	Value	Description
31–16	Reserved	–	0	Reserved. The reserved bit location is always read as 0. Bits in this field should always be written with 0.
15–0	L2WIWC	OF(value)	0–FFFFh	L2 writeback-invalidate word count.
		DEFAULT	0	

[†] For CSL implementation, use the notation `CACHE_L2WIWC_L2WIWC_symval`

6.8 L2 Invalidate Base Address Register (L2IBAR)

The L2 invalidate base address register (L2IBAR) is shown in Figure 18 and described in Table 12.

Figure 18. L2 Invalidate Base Address Register (L2IBAR)



Legend: R/W = Read/write; -n = value after reset

Table 14. L2 Invalidate Base Address Register (L2IBAR) Field Descriptions

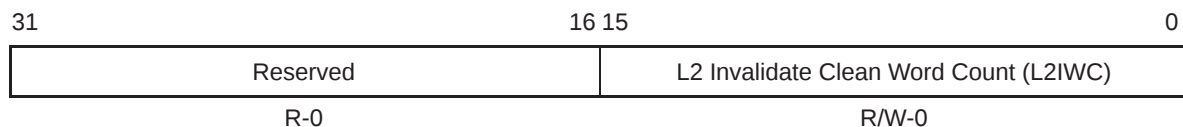
Bit	Field	symval [†]	Value	Description
31–0	L2IBAR	OF(value)	0–FFFF FFFFh	L2 invalidate base address.
		DEFAULT	0	

[†] For CSL implementation, use the notation `CACHE_L2IBAR_L2IBAR_symval`

6.9 L2 Invalidate Word Count Register (L2IWC)

The L2 invalidate word count register (L2IWC) is shown in Figure 19 and described in Table 15.

Figure 19. L2 Invalidate Word Count Register (L2IWC)



Legend: R = Read only; R/W = Read/write; -n = value after reset

Table 15. L2 Invalidate Word Count Register (L2IWC) Field Descriptions

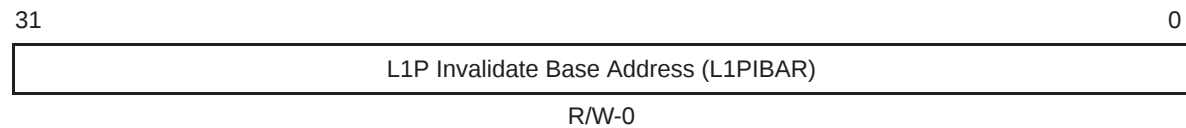
Bit	Field	symval [†]	Value	Description
31–16	Reserved	–	0	Reserved. The reserved bit location is always read as 0. Bits in this field should always be written with 0.
15–0	L2IWC	OF(value)	0–FFFFh	L2 invalidate word count.
		DEFAULT	0	

[†] For CSL implementation, use the notation `CACHE_L2IWC_L2IWC_symval`

6.10 L1P Invalidate Base Address Register (L1PIBAR)

The L1P invalidate base address register (L1PIBAR) is shown in Figure 20 and described in Table 16.

Figure 20. L1P Invalidate Base Address Register (L1PIBAR)



Legend: R/W = Read/write; -n = value after reset

Table 16. L1P Invalidate Base Address Register (L1PIBAR) Field Descriptions

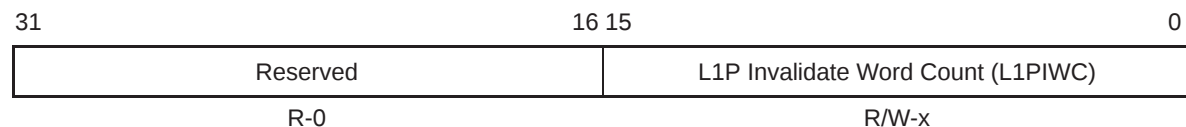
Bit	Field	symval [†]	Value	Description
31–0	L1PIBAR	OF(value)	0–FFFF FFFFh	L1P invalidate base address.
		DEFAULT	0	

[†] For CSL implementation, use the notation `CACHE_L1PIBAR_L1PIBAR_symval`

6.11 L1P Invalidate Word Count Register (L1PIWC)

The L1P invalidate word count register (L1PIWC) is shown in Figure 21 and described in Table 17.

Figure 21. L1P Invalidate Word Count Register (L1PIWC)



Legend: R = Read only; R/W = Read/write; -n = value after reset

Table 17. L1P Invalidate Word Count Register (L1PIWC) Field Descriptions

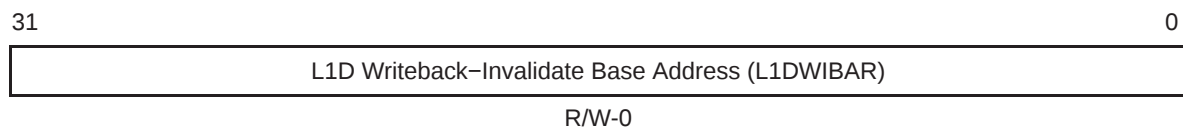
Bit	Field	symval [†]	Value	Description
31–16	Reserved	–	0	Reserved. The reserved bit location is always read as 0. Bits in this field should always be written with 0.
15–0	L1PIWC	OF(value)	0–FFFFh	L1P invalidate word count.
		DEFAULT	0	

[†] For CSL implementation, use the notation `CACHE_L1PIWC_L1PIWC_symval`

6.12 L1D Writeback-Invalidate Base Address Register (L1DWIBAR)

The L1D writeback-invalidate base address register (L1DWIBAR) is shown in Figure 22 and described in Table 18.

Figure 22. L1D Writeback-Invalidate Base Address Register (L1DWIBAR)



Legend: R/W = Read/write; -n = value after reset

Table 18. L1D Writeback-Invalidate Base Address Register (L1DWIBAR) Field Descriptions

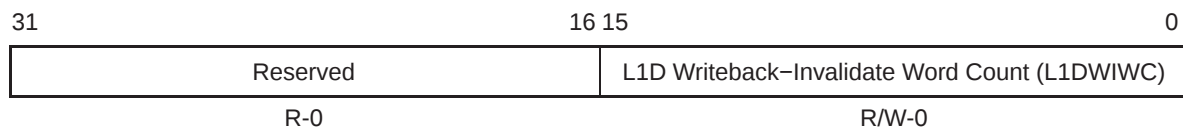
Bit	Field	symval [†]	Value	Description
31-0	L1DWIBAR	OF(value)	0-FFFF FFFFh	L1D writeback-invalidate base address.
		DEFAULT	0	

[†] For CSL implementation, use the notation `CACHE_L1DWIBAR_L1DWIBAR_symval`

6.13 L1D Writeback-Invalidate Word Count Register (L1DWIWC)

The L1D writeback-invalidate word count register (L1DWIWC) is shown in Figure 23 and described in Table 19.

Figure 23. L1D Writeback-Invalidate Word Count Register (L1DWIWC)



Legend: R = Read only; R/W = Read/write; -n = value after reset

Table 19. L1D Writeback-Invalidate Word Count Register (L1DWIWC) Field Descriptions

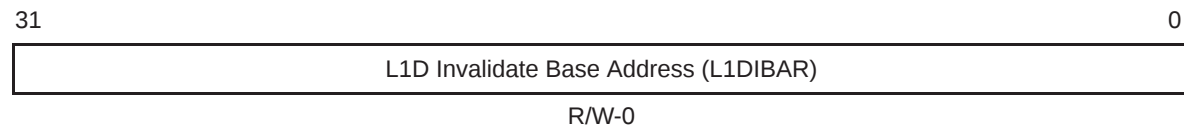
Bit	Field	symval [†]	Value	Description
31-16	Reserved	–	0	Reserved. The reserved bit location is always read as 0. Bits in this field should always be written with 0.
15-0	L1DWIWC	OF(value)	0-FFFFh	L1D writeback-invalidate word count.
		DEFAULT	0	

[†] For CSL implementation, use the notation `CACHE_L1DWIWC_L1DWIWC_symval`

6.14 L1D Invalidate Base Address Register (L1DIBAR)

The L1D invalidate base address register (L1DIBAR) is shown in Figure 24 and described in Table 20.

Figure 24. L1D Invalidate Base Address Register (L1DIBAR)



Legend: R/W = Read/write; -n = value after reset

Table 20. L1D Invalidate Base Address Register (L1DIBAR) Field Descriptions

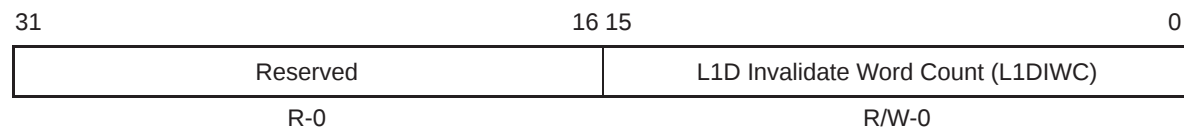
Bit	Field	symval [†]	Value	Description
31–0	L1DIBAR	OF(value)	0–FFFF FFFFh	L1D invalidate base address.
		DEFAULT	0	

[†] For CSL implementation, use the notation `CACHE_L1DIBAR_L1DIBAR_symval`

6.15 L1D Invalidate Word Count Register (L1DIWC)

The L1D invalidate word count register (L1DIWC) is shown in Figure 25 and described in Table 21.

Figure 25. L1D Invalidate Word Count Register (L1DIWC)



Legend: R = Read only; R/W = Read/write; -n = value after reset

Table 21. L1D Invalidate Word Count Register (L1DIWC) Field Descriptions

Bit	Field	symval [†]	Value	Description
31–16	Reserved	–	0	Reserved. The reserved bit location is always read as 0. Bits in this field should always be written with 0.
15–0	L1DIWC	OF(value)	0–FFFFh	L1D invalidate word count.
		DEFAULT	0	

[†] For CSL implementation, use the notation `CACHE_L1DIWC_L1DIWC_symval`

6.16 L2 Writeback All Register (L2WB)

The L2 writeback all register (L2WB) is shown in Figure 26 and described in Table 22.

Figure 26. L2 Writeback All Register (L2WB)

31		1	0
Reserved		C	
R-0		R/W-0	

Legend: R = Read only; R/W = Read/write; -n = value after reset

Table 22. L2 Writeback All Register (L2WB) Field Descriptions

Bit	Field	<i>symval</i> [†]	Value	Description
31–1	Reserved	–	0	Reserved. The reserved bit location is always read as 0. Bits in this field should always be written with 0.
0	C	OF(<i>value</i>)		L2 writeback all.
		DEFAULT	0	Normal L2 operation
		NORMAL		
		FLUSH	1	All L2 lines writeback all

[†] For CSL implementation, use the notation `CACHE_L2WB_C_symval`

6.17 L2 Writeback-Invalidate All Register (L2WBINV)

The L2 writeback-invalidate all register (L2WBINV) is shown in Figure 27 and described in Table 23.

Figure 27. L2 Writeback-Invalidate All Register (L2WBINV)

31		1	0
Reserved			C
R-0			R/W-0

Legend: R = Read only; R/W = Read/write; -n = value after reset

Table 23. L2 Writeback-Invalidate All Register (L2WBINV) Field Descriptions

Bit	Field	symval [†]	Value	Description
31–1	Reserved	–	0	Reserved. The reserved bit location is always read as 0. Bits in this field should always be written with 0.
0	C	OF(value)		L2 writeback-invalidate all.
		DEFAULT	0	Normal L2 operation
		NORMAL		
		CLEAN	1	All L2 lines writeback-invalidate all

[†] For CSL implementation, use the notation `CACHE_L2WBINV_C_symval`

6.18 L2 Memory Attribute Registers (MAR0–MAR255)

The L2 memory attribute register (MAR) is shown in Figure 28 and described in Table 24. Each MAR controls the cacheability of a 16-Mbyte address range, see section 7.2.

Figure 28. L2 Memory Attribute Register (MAR)

31		1	0
Reserved			CE
R-0			R/W-0

Legend: R = Read only; R/W = Read/write; -n = value after reset

Table 24. Memory Attribute Register (MAR) Field Descriptions

Bit	Field	symval [†]	Value	Description
31–1	Reserved	–	0	Reserved. The reserved bit location is always read as 0. Bits in this field should always be written with 0.
0	CE	OF(value)		Cache enable bit determines whether the L1D, L1P, and L2 are allowed to cache the corresponding address range.
		DEFAULT DISABLE	0	Memory range is not cacheable.
		ENABLE	1	Memory range is cacheable.

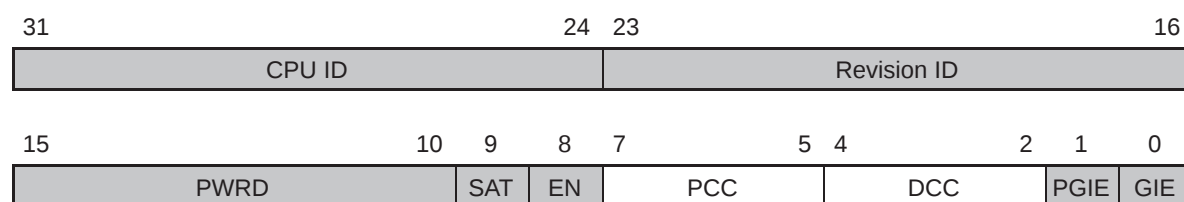
[†] For CSL implementation, use the notation `CACHE_MAR_CE_symval`

7 Memory System Control

7.1 Cache Mode Selection

The cache mode for the two-level memory hierarchy is determined by the cache configuration register (CCFG), and the data cache control (DCC) and program cache control (PCC) fields of the CPU control and status register (CSR). The CSR is shown in Figure 29.

Figure 29. CPU Control and Status Register (CSR)



7.1.1 L1D Mode Selection Using DCC Field in CSR

The L1D only operates as a cache and cannot be memory mapped. The L1D does not support freeze or bypass modes. The only values allowed for the data cache control (DCC) field are 000b and 010b. All other values for DCC are reserved, as shown in Table 25.

Table 25. L1D Mode Setting Using DCC Field

Bit	Field	Value	Description
4–2	DCC		Data cache control bit.
		000	2-way cache enabled
		001	Reserved
		010	2-way cache enabled
		011–111	Reserved

7.1.2 L1P Mode Selection Using PCC Field in CSR

The L1P only operates as a cache and cannot be memory mapped. The L1P does not support freeze or bypass modes. The only values allowed for the program cache control (PCC) field are 000b and 010b. All other values for PCC are reserved, as shown in Table 26.

Table 26. L1P Mode Setting Using PCC Field

Bit	Field	Value	Description
7–5	PCC		Program cache control bit.
		000	Direct-mapped cache enabled
		001	Reserved
		010	Direct-mapped cache enabled
		011–111	Reserved

7.1.3 L2 Mode Selection Using L2MODE Field in CCFG

The L2 memory can function as mapped SRAM, as cache, or as some combination of both. The L2MODE field in the cache configuration register (CCFG) determines what portion of L2 is mapped as SRAM and what portion acts as cache. For some settings of L2MODE, either L2 SRAM or L2 cache may not be present. The mode with no L2 cache is referred to as all SRAM mode. The CCFG is shown in Figure 11 (page 39) and described in Table 7.

The reset value of the L2MODE field is 000b, so the L2 is configured in all SRAM mode after reset. That is, there is no L2 cache enabled and the entire L2 SRAM address range is available. The L2 still services external memory requests for the L1P and L1D in this mode. It will not, however, retain copies of the data itself.

By default, external memory is not marked as cacheable. Cacheability is controlled separately of the cache mode. Section 7.2 discusses the memory attribute registers (MAR) that control the cacheability of external memory.

As various amounts of L2 cache are enabled, the amount of L2 SRAM available decreases. Addresses at the top of the L2 SRAM address space are converted to cache. For example, consider the C6414 DSP, which offers 1024K bytes of L2 memory at addresses from 0000 0000h–000F FFFFh. When the L2 is in 256K cache mode (L2MODE = 111b), the uppermost 256K-byte range of addresses, 000C 0000h–000F FFFFh, is not available to programs.

Note:

Reads or writes to L2 address ranges that are configured as L2 cache may result in undesired operation of the cache hierarchy. Programs must confine L2 accesses to L2 addresses that are mapped as L2 SRAM to ensure correct program operation.

The L2 controller processes reads to the address range 0000 0000h–0010 0FFFh without undesired cache operation, even if some of these addresses are configured as L2 cache. This address range represents the entire allotted L2 address range, plus some additional space to allow for certain program optimizations. Therefore, the restriction above does not apply to reads; however, programs should not interpret values returned by reads nor should programs perform writes to L2 addresses that are not configured as L2 SRAM.

To ensure correct operation when switching L2 cache modes, you must perform a series of operations. Table 27 specifies the required operations when either adding or removing L2 memory as mapped SRAM. Failure to follow these guidelines may result in data loss and undefined L2 operation.

Table 27. L2 Mode Switch Procedure

To Switch From...	To...	Perform the Following Steps
All SRAM Mode	Mode with L2 cache	1) Use EDMA to transfer any data needed out of the L2 SRAM space to be converted into cache. 2) Perform a block writeback-invalidate in L1D of L2 SRAM addresses that are about to become L2 cache. 3) Wait for block writeback-invalidate to complete. 4) Perform block writeback-invalidates for any external address ranges that may be cached in L1D. (This step is unnecessary if no CE bit has been set to 1 in any MAR.) 5) Wait for the block writeback-invalidates from step 4 to complete. 6) Write to CCFG to change mode. 7) Force CPU to wait for CCFG modification by reading CCFG. 8) Execute 8 cycles of NOP.
Mode with mixed L2 SRAM and L2 cache	Mode with less L2 mapped SRAM	1) Use EDMA to transfer any data needed out of the L2 SRAM space to be converted into cache. 2) Perform a block writeback-invalidate in L1D of L2 SRAM addresses that are about to become L2 cache. 3) Wait for block writeback-invalidate to complete. 4) Perform global writeback-invalidate of L2 (L2WBINV). 5) Wait for L2WBINV to complete. 6) Write to CCFG to change mode. 7) Force CPU to wait for CCFG modification by reading CCFG. 8) Execute 8 cycles of NOP.
Any L2 mode	Mode with more L2 mapped SRAM	1) Perform global writeback-invalidate of L2 (L2WBINV). 2) Wait for L2WBINV to complete. 3) Write to CCFG to change mode. 4) Force CPU to wait for CCFG modification by reading CCFG. 5) Execute 8 cycles of NOP.

7.2 Cacheability Controls

The cacheability of external address ranges is controlled by the memory attribute registers (MAR0–MAR255). Each MAR controls the cacheability of a 16-Mbyte address range as listed in Table 28. MAR is shown in Figure 28 (page 51) and described in Table 24. See the device-specific datasheet for the memory address of these registers.

The cache enable (CE) bit in each MAR determines whether the L1D, L1P, and L2 are allowed to cache the corresponding address range. After reset, the CE bit in each MAR is cleared to 0, thereby disabling caching of external memory by default. This is in contrast to L2 SRAM, which is always considered cacheable.

To enable caching on a particular external address range, an application should set the CE bit in the appropriate MAR to 1. No special procedure is necessary. Subsequent accesses to the affected address range are cached by the two-level memory system.

To disable caching for a given address range, programs should follow the following sequence to ensure that all future accesses to the particular address range are not cached.

- 1) Ensure that all addresses within the affected range are removed from the L1 and L2 caches. This is accomplished in one of the following ways. Any one of the following operations should be sufficient.
 - a) If L2 cache is enabled, invoke a global writeback-invalidate using L2WBINV. Wait for the C bit in L2WBINV to read as 0. Alternately, invoke a block writeback-invalidate of the affected range using L2WIBAR/L2WIWC. Wait for L2WIWC to read as 0.
 - b) If L2 is in all SRAM mode, invoke a block writeback-invalidate of the affected range using L1DWIBAR/L1DWIWC. Wait for L1DWIWC to read as 0.

Note that the block-oriented cache controls can only operate on a 256K-byte address range at a time, so multiple block writeback-invalidate operations may be necessary to remove the entire affected address range from the cache. These cache controls are discussed in section 7.3.

- 2) Clear the CE bit in the appropriate MAR to 0.

Table 28. Memory Attribute Registers

Acronym	Register Name
MAR0–MAR95	Reserved
MAR96 [†]	L2 memory attribute register 96: controls EMIFB CE0 range 6000 0000h–60FF FFFFh
MAR97 [†]	L2 memory attribute register 97: controls EMIFB CE0 range 6100 0000h–61FF FFFFh
MAR98 [†]	L2 memory attribute register 98: controls EMIFB CE0 range 6200 0000h–62FF FFFFh
MAR99 [†]	L2 memory attribute register 99: controls EMIFB CE0 range 6300 0000h–63FF FFFFh
MAR100 [†]	L2 memory attribute register 100: controls EMIFB CE1 range 6400 0000h–64FF FFFFh
MAR101 [†]	L2 memory attribute register 101: controls EMIFB CE1 range 6500 0000h–65FF FFFFh
MAR102 [†]	L2 memory attribute register 102: controls EMIFB CE1 range 6600 0000h–66FF FFFFh
MAR103 [†]	L2 memory attribute register 103: controls EMIFB CE1 range 6700 0000h–67FF FFFFh
MAR104 [†]	L2 memory attribute register 104: controls EMIFB CE2 range 6800 0000h–68FF FFFFh
MAR105 [†]	L2 memory attribute register 105: controls EMIFB CE2 range 6900 0000h–69FF FFFFh
MAR106 [†]	L2 memory attribute register 106: controls EMIFB CE2 range 6A00 0000h–6AFF FFFFh
MAR107 [†]	L2 memory attribute register 107: controls EMIFB CE2 range 6B00 0000h–6BFF FFFFh
MAR108 [†]	L2 memory attribute register 108: controls EMIFB CE3 range 6C00 0000h–6CFF FFFFh
MAR109 [†]	L2 memory attribute register 109: controls EMIFB CE3 range 6D00 0000h–6DFF FFFFh
MAR110 [†]	L2 memory attribute register 110: controls EMIFB CE3 range 6E00 0000h–6EFF FFFFh
MAR111 [†]	L2 memory attribute register 111: controls EMIFB CE3 range 6F00 0000h–6FFF FFFFh
MAR112–MAR127	Reserved
MAR128	L2 memory attribute register 128: controls EMIFA CE0 range 8000 0000h–80FF FFFFh
MAR129	L2 memory attribute register 129: controls EMIFA CE0 range 8100 0000h–81FF FFFFh
MAR130	L2 memory attribute register 130: controls EMIFA CE0 range 8200 0000h–82FF FFFFh
MAR131	L2 memory attribute register 131: controls EMIFA CE0 range 8300 0000h–83FF FFFFh
MAR132	L2 memory attribute register 132: controls EMIFA CE0 range 8400 0000h–84FF FFFFh
MAR133	L2 memory attribute register 133: controls EMIFA CE0 range 8500 0000h–85FF FFFFh
MAR134	L2 memory attribute register 134: controls EMIFA CE0 range 8600 0000h–86FF FFFFh
MAR135	L2 memory attribute register 135: controls EMIFA CE0 range 8700 0000h–87FF FFFFh
MAR136	L2 memory attribute register 136: controls EMIFA CE0 range 8800 0000h–88FF FFFFh
MAR137	L2 memory attribute register 137: controls EMIFA CE0 range 8900 0000h–89FF FFFFh

[†] MAR96–MAR111 only available on the C6414/C6415/C6416 devices; on all other devices, these registers are reserved.

Table 28. Memory Attribute Registers (Continued)

Acronym	Register Name
MAR138	L2 memory attribute register 138: controls EMIFA CE0 range 8A00 0000h–8AFF FFFFh
MAR139	L2 memory attribute register 139: controls EMIFA CE0 range 8B00 0000h–8BFF FFFFh
MAR140	L2 memory attribute register 140: controls EMIFA CE0 range 8C00 0000h–8CFF FFFFh
MAR141	L2 memory attribute register 141: controls EMIFA CE0 range 8D00 0000h–8DFF FFFFh
MAR142	L2 memory attribute register 142: controls EMIFA CE0 range 8E00 0000h–8EFF FFFFh
MAR143	L2 memory attribute register 143: controls EMIFA CE0 range 8F00 0000h–8FFF FFFFh
MAR144	L2 memory attribute register 144: controls EMIFA CE1 range 9000 0000h–90FF FFFFh
MAR145	L2 memory attribute register 145: controls EMIFA CE1 range 9100 0000h–91FF FFFFh
MAR146	L2 memory attribute register 146: controls EMIFA CE1 range 9200 0000h–92FF FFFFh
MAR147	L2 memory attribute register 147: controls EMIFA CE1 range 9300 0000h–93FF FFFFh
MAR148	L2 memory attribute register 148: controls EMIFA CE1 range 9400 0000h–94FF FFFFh
MAR149	L2 memory attribute register 149: controls EMIFA CE1 range 9500 0000h–95FF FFFFh
MAR150	L2 memory attribute register 150: controls EMIFA CE1 range 9600 0000h–96FF FFFFh
MAR151	L2 memory attribute register 151: controls EMIFA CE1 range 9700 0000h–97FF FFFFh
MAR152	L2 memory attribute register 152: controls EMIFA CE1 range 9800 0000h–98FF FFFFh
MAR153	L2 memory attribute register 153: controls EMIFA CE1 range 9900 0000h–99FF FFFFh
MAR154	L2 memory attribute register 154: controls EMIFA CE1 range 9A00 0000h–9AFF FFFFh
MAR155	L2 memory attribute register 155: controls EMIFA CE1 range 9B00 0000h–9BFF FFFFh
MAR156	L2 memory attribute register 156: controls EMIFA CE1 range 9C00 0000h–9CFF FFFFh
MAR157	L2 memory attribute register 157: controls EMIFA CE1 range 9D00 0000h–9DFF FFFFh
MAR158	L2 memory attribute register 158: controls EMIFA CE1 range 9E00 0000h–9EFF FFFFh
MAR159	L2 memory attribute register 159: controls EMIFA CE1 range 9F00 0000h–9FFF FFFFh
MAR160	L2 memory attribute register 160: controls EMIFA CE2 range A000 0000h–A0FF FFFFh
MAR161	L2 memory attribute register 161: controls EMIFA CE2 range A100 0000h–A1FF FFFFh
MAR162	L2 memory attribute register 162: controls EMIFA CE2 range A200 0000h–A2FF FFFFh
MAR163	L2 memory attribute register 163: controls EMIFA CE2 range A300 0000h–A3FF FFFFh
MAR164	L2 memory attribute register 164: controls EMIFA CE2 range A400 0000h–A4FF FFFFh
MAR165	L2 memory attribute register 165: controls EMIFA CE2 range A500 0000h–A5FF FFFFh

† MAR96–MAR111 only available on the C6414/C6415/C6416 devices; on all other devices, these registers are reserved.

Table 28. Memory Attribute Registers (Continued)

Acronym	Register Name
MAR166	L2 memory attribute register 166: controls EMIFA CE2 range A600 0000h–A6FF FFFFh
MAR167	L2 memory attribute register 167: controls EMIFA CE2 range A700 0000h–A7FF FFFFh
MAR168	L2 memory attribute register 168: controls EMIFA CE2 range A800 0000h–A8FF FFFFh
MAR169	L2 memory attribute register 169: controls EMIFA CE2 range A900 0000h–A9FF FFFFh
MAR170	L2 memory attribute register 170: controls EMIFA CE2 range AA00 0000h–AAFF FFFFh
MAR171	L2 memory attribute register 171: controls EMIFA CE2 range AB00 0000h–ABFF FFFFh
MAR172	L2 memory attribute register 172: controls EMIFA CE2 range AC00 0000h–ACFF FFFFh
MAR173	L2 memory attribute register 173: controls EMIFA CE2 range AD00 0000h–ADFF FFFFh
MAR174	L2 memory attribute register 174: controls EMIFA CE2 range AE00 0000h–AEFF FFFFh
MAR175	L2 memory attribute register 175: controls EMIFA CE2 range AF00 0000h–AFF FFFFh
MAR176	L2 memory attribute register 176: controls EMIFA CE3 range B000 0000h–B0FF FFFFh
MAR177	L2 memory attribute register 177: controls EMIFA CE3 range B100 0000h–B1FF FFFFh
MAR178	L2 memory attribute register 178: controls EMIFA CE3 range B200 0000h–B2FF FFFFh
MAR179	L2 memory attribute register 179: controls EMIFA CE3 range B300 0000h–B3FF FFFFh
MAR180	L2 memory attribute register 180: controls EMIFA CE3 range B400 0000h–B4FF FFFFh
MAR181	L2 memory attribute register 181: controls EMIFA CE3 range B500 0000h–B5FF FFFFh
MAR182	L2 memory attribute register 182: controls EMIFA CE3 range B600 0000h–B6FF FFFFh
MAR183	L2 memory attribute register 183: controls EMIFA CE3 range B700 0000h–B7FF FFFFh
MAR184	L2 memory attribute register 184: controls EMIFA CE3 range B800 0000h–B8FF FFFFh
MAR185	L2 memory attribute register 185: controls EMIFA CE3 range B900 0000h–B9FF FFFFh
MAR186	L2 memory attribute register 186: controls EMIFA CE3 range BA00 0000h–BAFF FFFFh
MAR187	L2 memory attribute register 187: controls EMIFA CE3 range BB00 0000h–BBFF FFFFh
MAR188	L2 memory attribute register 188: controls EMIFA CE3 range BC00 0000h–BCFF FFFFh
MAR189	L2 memory attribute register 189: controls EMIFA CE3 range BD00 0000h–BDFF FFFFh
MAR190	L2 memory attribute register 190: controls EMIFA CE3 range BE00 0000h–BEFF FFFFh
MAR191	L2 memory attribute register 191: controls EMIFA CE3 range BF00 0000h–BFFF FFFFh
MAR192–MAR255	Reserved

† MAR96–MAR111 only available on the C6414/C6415/C6416 devices; on all other devices, these registers are reserved.

7.3 Program-Initiated Cache Operations

The memory system provides a set of cache control operations. These allow programs to specifically request that certain data be written back or invalidated. The cache operations fall into two categories, global operations that operate on the entire cache, and block operations that operate on a specific range of addresses. The global operations are discussed in section 7.3.1 and the block operations are discussed in section 7.3.2.

The memory system can only perform one program-initiated cache operation at a time. This includes global operations, block operations, and mode changes. For this reason, the memory system may stall accesses to cache control registers while a cache control operation is in progress. Table 29 gives a summary of the available operations and their impact on the memory system.

Table 29. Summary of Program-Initiated Cache Operations

Type of Operation	Cache Operation	Register Usage	L1P Cache Effect	L1D Cache Effect	L2 Cache Effect
Global Operation	L2 Writeback All	L2WB	No effect.	Updated lines holding addresses also held in L2 cache are written back. All lines corresponding to addresses held in L2 cache are invalidated. [†]	Updated data is written back. All lines kept valid.
	L2 Writeback and Invalidate All	L2WBINV	Entire contents are discarded. [†]	Updated lines holding addresses also held in L2 cache are written back. All lines corresponding to addresses held in L2 cache are invalidated. [†]	Updated lines are written back. All lines are invalidated.
	L1P Invalidate All	IP bit in CCFG	Entire contents are discarded.	No effect.	No effect.
	L1D Invalidate All	ID bit in CCFG	No effect.	Entire contents are discarded. No updated data is written back.	No effect.

[†] As described in section 7.3.3, these operations only affect L1D and L1P when L2 cache is enabled. When L2 cache is enabled, these operations only affect those portions of L1P and L1D that are also held in L2 cache. Future C6000 devices may cause these operands to act on the entire contents of L1P and L1D, regardless of whether the same addresses are cached in L2, or whether L2 cache is enabled.

[‡] In contrast to its behavior on C621x/C671x devices, L1DWIBAR/L1DWIWC on C64x devices do not cause the corresponding block to be invalidated in L1P.

Table 29. Summary of Program-Initiated Cache Operations (Continued)

Type of Operation	Cache Operation	Register Usage	L1P Cache Effect	L1D Cache Effect	L2 Cache Effect
Block Operation	L2 Block Writeback	L2WBAR, L2WWC	No effect	Updated lines in block are written to L2. All lines in block are invalidated in L1D. [†]	Updated lines in block are written to external memory. Lines in block are kept valid in L2.
	L2 Block Writeback with Invalidate	L2WIBAR, L2WIWC	All lines in block are invalidated. [†]	Updated lines in block are written to L2. All lines in block are invalidated in L1D. [†]	Updated lines in block are written to external memory. All lines in block are invalidated in L2.
	L2 Block Invalidate	L2IBAR, L2IWC	All lines in block are invalidated. [†]	All lines in block are invalidated in L1D. Updated data in block is discarded. [†]	All lines in block are invalidated in L2. Updated data in block is discarded.
	L1P Block Invalidate	L1PIBAR, L1PIWC	All lines in block are invalidated.	No effect.	No effect.
	L1D Block Writeback with Invalidate	L1DWIBAR, L1DWIWC	No effect. [‡]	Updated lines in block are written to L2. All lines in block are invalidated in L1D.	No effect.
	L1D Block Invalidate	L1DIBAR, L1DIWC	No effect.	All lines in block are invalidated in L1D. Updated data in block is discarded.	No effect.

[†] As described in section 7.3.3, these operations only affect L1D and L1P when L2 cache is enabled. When L2 cache is enabled, these operations only affect those portions of L1P and L1D that are also held in L2 cache. Future C6000 devices may cause these operands to act on the entire contents of L1P and L1D, regardless of whether the same addresses are cached in L2, or whether L2 cache is enabled.

[‡] In contrast to its behavior on C621x/C671x devices, L1DWIBAR/L1DWIWC on C64x devices do not cause the corresponding block to be invalidated in L1P.

7.3.1 Global Cache Operations

The global cache operations execute on the entire contents of a particular cache. Global operations take precedence over program accesses to the cache. Program accesses (either data or program fetches) to a particular cache are stalled while a global cache operation is active on that cache.

Global operations in L1D and L1P are initiated by the ID and IP bits in CCFG (Figure 11). The L1D and L1P only offer global invalidate. By writing a 1 to the ID or IP bit in CCFG, a program can invalidate the entire contents of the corresponding L1 cache. Upon initiation of the global invalidate, the entire contents of the corresponding cache is discarded — no updated data is written back. Reading CCFG after the write will stall the program until the invalidate operation is complete.

The L1D global-invalidate causes all updated data in L1D to be discarded, rather than written back to the lower levels of memory. This can cause incorrect operation in programs that expect the updates to be written to the lower levels of memory. Therefore, most programs use either the L1D block writeback-invalidate (described in section 7.3.2) or the global L2 operations, rather than the L1D global-invalidate.

The L2 offers both global writeback and global writeback-invalidate operations. Global cache operations in L2 are initiated by writing a 1 to the C bit in either L2WB shown in Figure 26 (page 49) and described in Table 22, or L2WBINV shown in Figure 27 (page 50) and described in Table 23. Writing a 1 to the C bit of L2WB initiates a global writeback of L2; writing a 1 to the C bit of L2WBINV initiates a global writeback-invalidate of L2. The C bit continues to read as 1 until the cache operation is complete. Programs can poll to determine when a cache operation is complete.

Global operations on L2 have indirect effects on the L1 caches, as discussed in section 7.3.3.

7.3.2 Block Cache Operations

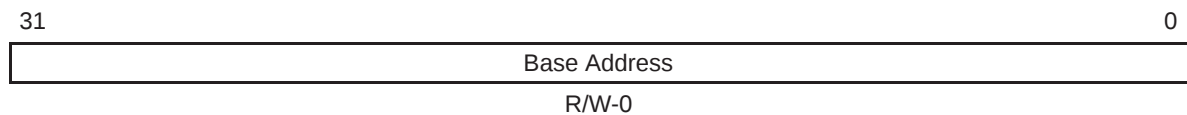
The block cache operations execute on a range of addresses that may be held in the cache. Block operations execute in the background, allowing other program accesses to interleave with the block cache operation.

Programs initiate block cache operations with two writes. The program first writes the starting address to one of the base address registers (BAR), shown in Figure 30. Next, the program writes the total number of words to operate on to the corresponding word count register (WC), shown in Figure 31. The cache operation begins as soon as the word count register is written with a non-zero value. The cache provides a set of BAR/WC pseudo-register pairs, one for each block operation the cache supports. The complete list of supported operations is shown in Table 29.

Notice that the word count field in WC is only 16-bits wide. This limits the block size to 65535 words (approximately 256K bytes). Larger ranges require multiple block commands to be issued to cover the entire range.

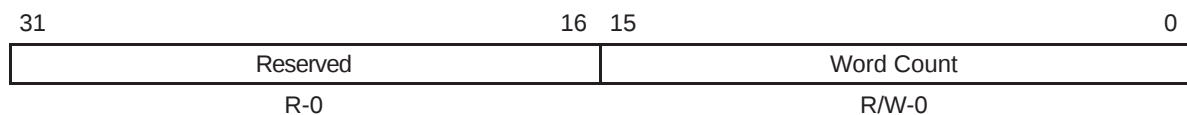
Although block operations specify the block in terms of a word address and word count, the block operations always operate on whole cache lines. Whole lines are always be written back and/or invalidated in each affected cache. For this reason, programs should be careful to align arrays on cache-line boundaries, and to pad arrays to be a multiple of the cache line size. This is especially true when invoking the invalidate-only commands with respect to these arrays.

Figure 30. Block Cache Operation Base Address Register (BAR)



Legend: R/W = Read/write; -n = value after reset

Figure 31. Block Cache Operation Word Count Register (WC)



Legend: R = Read only; R/W = Read/write; -n = value after reset

In current device implementations, the block command registers (BAR and WC) are implemented as a single register pair in the underlying hardware. They appear as pseudo-registers at multiple addresses in the address map. Therefore, programs **must not** interleave writes to different command register pairs. Writes to one BAR/WC pair should be performed atomically with respect to other BAR/WC writes. Depending on the nature of the program, it may be necessary to disable interrupts while writing to BAR and WC.

Despite this restriction, programs may attempt to initiate a new block command while a previous command is in progress. The cache controller will stall the write to BAR until the previous cache operation completes, thus preserving correct operation. To avoid stalling, programs may determine if a block cache operation has completed by polling WC. WC returns a non-zero value while the cache operation is in progress, and zero when the operation has completed. The non-zero value returned may vary between device implementations, so programs should only rely on the register being non-zero.

Programs should not assume that the value of BAR is retained between block cache operations. Programs should always write an address to BAR and a word count to WC for each cache operation. Also, programs should not assume that the various BAR and WC map to the same physical register. For each cache operation, programs should write to BAR and WC for that operation.

The L1P block invalidate can be used in conjunction with the L1D block writeback-invalidate to provide software controlled coherence between L1D and L1P. (Section 8.1 discusses the memory system coherence policies.) To execute code that was previously written to the CPU, the program should use L1PIBAR/L1PIWC to invalidate the block in L1P, and L1DWIBAR/L1DWIWC to writeback-invalidate the block in L1D. These operations can be performed in either order. The specific timing of these operations relative to program fetches is not well defined. Therefore, programs should wait for L1DWIWC and L1PIWC to read as zero prior to branching to an address range that has been invalidated in this manner. (Note that the behavior of L1DIBAR/L1DIWC differs on C621x/C671x devices. See *TMS320C621x/C671x DSP Two-Level Internal Memory Reference Guide*, SPRU609, for details.)

Block cache operations in L2 can indirectly affect L1D and L1P, as noted in Table 29. Section 7.3.3 discusses these interactions in detail.

Note:

Reads or writes to the addresses within the block being operated on while a block cache operation is in progress may cause those addresses to not be written back or invalidated as requested. To avoid this, programs should not access addresses within the range of cache lines affected by a block cache operation while the operation is in progress. Programs may consult the appropriate WC to determine when the block operation is complete.

7.3.3 Effect of L2 Commands on L1 Caches

Note:

The behaviors described in this section may change on future C6000 devices that implement the two-level memory subsystem. Future devices may remove the inclusivity of L1D within L2, and may cause the L2 cache operations to act on L1, regardless of the current L2 cache mode. Forward-compatible programs should not rely on these specific memory system behaviors.

Cache operations in L2 indirectly operate on the L1D and L1P cache. As a result, the L2 cache operations have *no effect on any of the caches* when L2 is in all SRAM mode. This is true for both the global and block commands. Otherwise, when L2 cache is enabled, program-initiated cache operations in L2 may operate on the corresponding contents of the L1D and L1P.

Under normal circumstances, the L1D cache is *inclusive* in L2, and L1P is not. Inclusive implies that the entire contents of L1D are also held either in L2 SRAM or L2 cache. The L2 cache operations are designed with these properties in mind.

Because L1P is not inclusive in L2, the L2 cache operations that invalidate lines in L2 send explicit invalidation commands to L1P. A global writeback-invalidate of L2 (L2WBINV) triggers a complete invalidation of L1P. Block invalidate and writeback-invalidate operations in L2 blindly send invalidate commands to L1P for the corresponding L1P cache lines. This ensures that L1P always fetches the most recent contents of memory after the cache operation is complete.

Because L1D is normally inclusive in L2, the L2 relies on normal cache protocol operation to operate on L1D indirectly. Writebacks and invalidates triggered for lines in L2 result in snoop-invalidate commands sent to L1D when L2 detects that L1D is also holding a copy of the corresponding addresses. Therefore, the L2 global writeback (L2WB) and L2 global writeback-invalidate (L2WBINV) causes all external addresses held in L1D to be written back and invalidated from L1D. Likewise, block operations in L2 cause the corresponding range in L1D to be written to L2 and invalidated from L1D using the indirect snoop-invalidates.

One result of this is that L2 SRAM addresses cached in L1D are not affected by program-initiated cache operations in L2, as L2 cache never holds copies of L2 SRAM. To remove L2 SRAM addresses from L1D, programs must use the L1D block cache operations directly. Ordinarily, direct removal of L2 SRAM addresses from L1D is required only when changing L2 cache modes. The coherence policy described in section 8.1 makes unnecessary most of the need for programs to manually write back portions of L1D to L2 SRAM.

Another result is nonintuitive behavior when L1D is *not* inclusive in L2. L1D is inclusive in L2 under normal circumstances, and so most programs do not need to be concerned about this situation. Indeed, the recommended L2 cache mode-change procedure in section 7.1.3 ensures that the memory system is never in this state. When *not* following the procedure precisely, it is possible for L1D to hold copies of external memory that are not held in L2. This noninclusive state is achieved in the following rare sequence:

- ☐ The program enabled caching for an external address range while L2 was in all SRAM mode.
- ☐ The program read directly from this external address range.
- ☐ The program then enabled L2 cache without first removing the external address range from L1D with a block cache operation.

To prevent this situation, programs should ensure no external addresses are cached in L1D when enabling L2 cache. This is accomplished by not enabling caching for external addresses prior to enabling L2 cache, or by removing external addresses from L1D using the block cache operations described in section 7.3.2.

7.4 L2-to-EDMA Request Control

As described in section 5.4.3, the L2 controller relies on the EDMA to service requests on its behalf. L2 requests may be placed on any of the four EDMA priority queues. The priority queue selection for cache-service requests is made using the P field in CCFG (Figure 11, page 39).

Programs should take care when changing the priority level for L2 requests in order to ensure proper operation of the cache. The following sequence should be followed:

- 1) Poll the EDMA priority queue status register (PQSR) and wait for the PQ bit that corresponds to the current priority level to read as 1. PQSR is described in *TMS320C6000 DSP Enhanced Direct Memory Access (EDMA) Peripheral Reference Guide* (SPRU234).

This step may require that other transfers using this same priority queue, such as externally-triggered EDMA transfers, be disabled. Otherwise, in a heavily loaded system, the PQ bit in PQSR may not read as 1 for an arbitrarily long period of time.

- 2) Write a new value to the P field in CCFG.
- 3) Read CCFG back to ensure the write has completed.

To ensure fairness between L2 requests and other EDMA requests in the system, the L2 provides four priority queue allocation registers, L2ALLOC0, L2ALLOC1, L2ALLOC2, and L2ALLOC3 (Figure 13, page 42). These registers control the rate at which L2 requests are issued to the EDMA. Because the L2 also contains the QDMA control registers, the limits set by L2ALLOC0–3 also apply to QDMA requests. QDMAs are described in *TMS320C6000 DSP Enhanced Direct Memory Access (EDMA) Peripheral Reference Guide* (SPRU234).

The setting of a given L2ALLOC determines the the total number of outstanding L2 and QDMA requests permitted on the corresponding EDMA priority level. Further requests on that priority level are stalled until the number of outstanding requests falls below the L2ALLOC setting. An outstanding transfer in this context refers to a transfer that has been submitted to a given priority queue, but which has not been extracted from the queue into the EDMA channel registers. L2ALLOC0–3 behave similarly to the priority queue allocation registers (PQAR) described in *TMS320C6000 DSP Enhanced Direct Memory Access (EDMA) Peripheral Reference Guide* (SPRU234).

Table 30 lists the default queue allocation. The L2ALLOC settings must also take into account the current settings for PQAR so that no transfers are lost. The correct procedure for modifying the L2ALLOC and PQAR settings is described in *TMS320C6000 DSP Enhanced Direct Memory Access (EDMA) Peripheral Reference Guide* (SPRU234).

Table 30. L2ALLOC Default Queue Allocations

Priority Level	Allocation Register	Default L2/QDMA Allocation
Urgent	L2ALLOC0	6
High	L2ALLOC1	2
Medium	L2ALLOC2	2
Low	L2ALLOC3	2

7.5 EDMA Access Into L2 Control

The C64x DSP incorporates an L2 EDMA access control register (EDMAWEIGHT), located in the L2 cache register memory map, that controls the relative priority weighting of EDMA versus L1D access into L2. EDMAWEIGHT gives EDMA accesses a temporary boost in priority by limiting the amount of time L1D blocks EDMA access to L2. This priority boost only applies when competing with write data from the CPU that misses in L1D, but hits in L2 cache or L2 SRAM. Normal line allocate and eviction operation result in gaps in the L1D-to-L2 access pattern that can be used by EDMA accesses. EDMAWEIGHT lets you control how often this priority boost is given. When the EDMA priority is raised, it is allowed to complete one access before priority is returned to the CPU data. The EDMAWEIGHT is shown in Figure 12, page 41, and described in Table 8.

If L1D blocks L2 for n consecutive cycles, then EDMA is given priority for a single cycle.

8 Memory System Policies

This section discusses the various policies of the memory system, such as coherence between CPU and EDMA or host accesses and the order in which memory updates are made.

8.1 Memory System Coherence

Cache memories work by retaining copies of data from lower levels of the memory hierarchy in the hierarchy's higher levels. This provides a performance benefit as higher levels of the memory hierarchy may be accessed more quickly than the lower levels, and so the lower levels need not be consulted on every access. Because many accesses to memory are captured at higher levels of the hierarchy, the opportunity exists for the CPU and other devices to see a different picture of what is in memory.

A memory system is coherent if all requestors into that memory see updates to individual memory locations occur in the same order. A requestor is a device such as a CPU or EDMA. A coherent memory system ensures that all writes to a given memory location are visible to future reads of that location from any requestor, so long as no intervening write to that location overwrites the value. If the same requestor is writing and reading, the results of a write are immediately visible. If one requestor writes and different requestor reads, the reader may not see the updated value immediately, but it will be able to see updates after a sufficient period of time. Coherence also implies that all writes to a given memory location appear to be serialized. That is, that all requestors see the same order of writes to a memory location, even when multiple requestors are writing to one location.

The two-level cache memory system is coherent with various qualifications. The C64x devices ensure coherence in a limited manner for the programmer. This limited coherence model is designed to simplify and improve the performance of the hardware, while still providing a reasonable programming model.

Memory system coherence only applies to cacheable data. For noncacheable data, the memory system merely passes the access on to its final destination without retaining any intermediate copies, and thus coherence is not an issue. Therefore, this discussion focuses only on cacheable memory, which is generally confined to various forms of RAM and ROM.

The coherence model of the C64x memory system is expressed in terms of requestors and where the memory is located. The two-level memory system must support requests from three sources: CPU program fetches, CPU data accesses, and EDMA accesses. (HPI and PCI accesses to the two-level memory system are handled using the EDMA.) Cacheable areas of interest include both internal and external RAM. Since ROM is typically not writeable, we consider only RAM. Table 31 shows where the memory system ensures coherency.

Table 31. Coherence Assurances in the Two-Level Memory System

Requestor	CPU Program Fetch	CPU Data Access	EDMA Access
CPU program fetch	Coherent	Software-managed, L2 level only	On-chip SRAM only
CPU data access	Software-managed, L2 level only	Coherent	On-chip SRAM only
EDMA access	On-chip SRAM only	On-chip SRAM only	Coherent

Notice that the hardware ensures that accesses by the CPU and EDMA to internal SRAM addresses are coherent, but external addresses are not. Software must ensure external addresses accessed by the EDMA are not held in cache when the EDMA accesses them. Failure to do so can result in data corruption on the affected range of addresses. See section 7.3 for the steps required to ensure particular ranges of addresses are not held in cache.

Also notice that CPU data and program accesses are not coherent. The reason that these are not considered coherent is that the L1P does not query the L1D when it makes a program fetch, and thus CPU writes captured in L1D may not be visible to L1P for an arbitrarily long period of time. Therefore, programs that write to locations that are subsequently executed from must ensure that the updated data is written from L1D to at least L2 before execution. The L1DWIBAR/L1DWIWC and L1PIBAR/L1PIWC described in section 7.3 can be used for this purpose.

Although the memory system is coherent, it does not ensure that all requestors see updates to cacheable memory occurring in the same order. This is of primary importance when a given buffer of memory is accessed by the CPU and an EDMA or host port access. Unless care is taken, it is possible for a EDMA or host port access to see a mixture of old and new data if the access occurs while the CPU updates are being processed by the memory system. Section 8.3 discusses the order in which CPU accesses are made visible to the memory system.

8.2 EDMA Coherence in L2 SRAM Example

As indicated in section 8.1, the memory system ensures coherence between CPU and EDMA accesses to L2 SRAM. It does not provide coherence between CPU and EDMA accesses to external memory. This section illustrates how the L2 SRAM coherence assurances are enforced by the hardware within the context of an example.

A common DSP programming technique is to process streams of data with a double-buffering scheme. In this setup, a large set of data external to the device is processed by transferring in small chunks and operating on them. The results are then transferred to external memory. These transfers are generally executed by the EDMA controller. In pseudo-code, a typical processing sequence might as shown in Figure 32.

Because the EDMA operates independently of the CPU, it is possible to overlap the EDMA transfers with data processing. Double-buffering allows this by having two sets of input and output buffers—one for processing and one for transfers-in-progress. The CPU and EDMA alternate sets of buffers, “ping-ponging” between them. (For this reason, double-buffering is sometimes referred to as ping-pong buffering.) In pseudo-code, a double-buffered processing sequence might be as shown in Figure 33. Here, PING and PONG are pairs pointers to the double-buffers. PING points to the buffers that the CPU will process, and PONG points to the buffers the EDMA is filling.

Graphically, the time sequence for the internal input and output buffers would look as shown in Figure 34. EDMA reads are one step ahead of the processing and EDMA writes are one step behind. In Figure 34, step 2 operates in parallel with steps 3 and 4; steps 5 and 6 overlap with steps 7 and 8.

Figure 35 shows the read/process/write pipeline in software, with the EDMA transfers overlapping processing.

Figure 32. Streaming Data Pseudo-Code

```
while (data left to process)
{
    Read next block to internal input buffer

    Process internal input buffer, placing result in internal output buffer

    Write the internal output buffer to the external output
}
```

Figure 33. Double Buffering Pseudo-Code

```
Read first block to internal PING input buffer.
while (data left to process)
{
    Start reading next block to internal "PONG" input buffer.

    Wait for PING input and output buffers to be ready.

    Process internal PING input buffer, placing result in internal
    PING output buffer.

    Swap PING and PONG buffer pointers.

    Start writing out PONG output buffer.
}
```

Figure 34. Double-Buffering Time Sequence

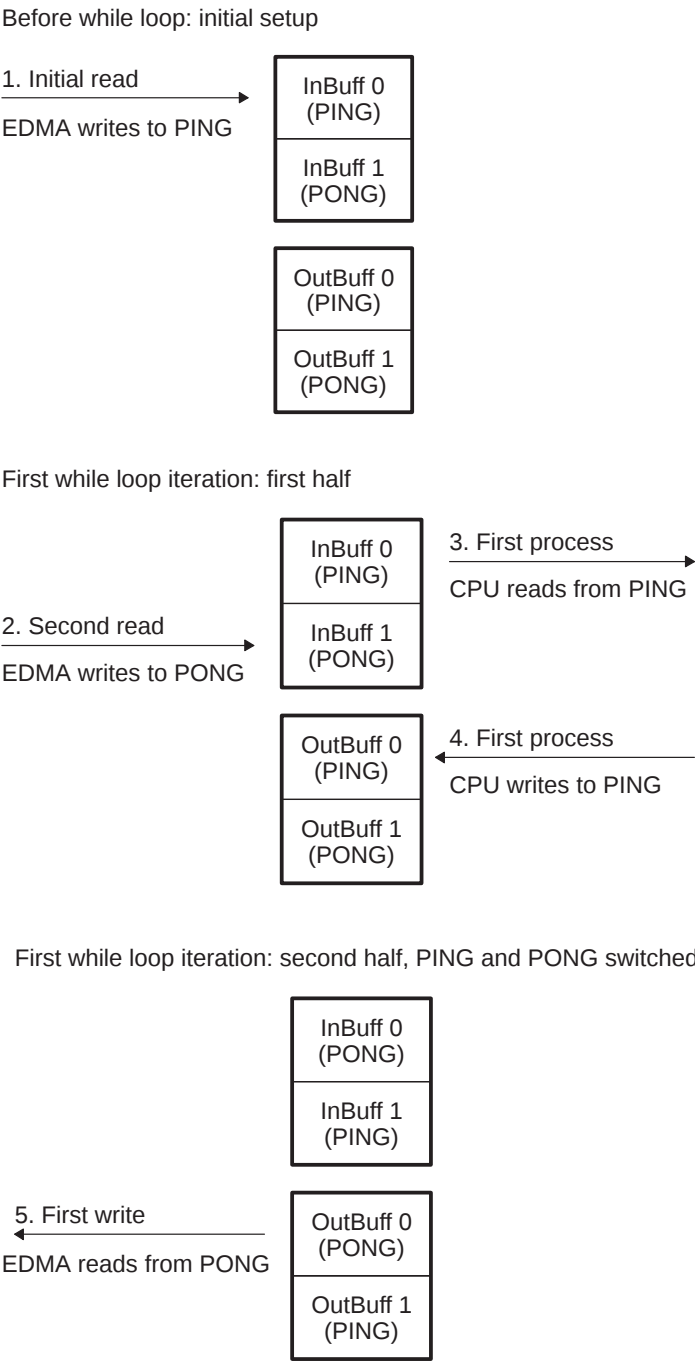


Figure 34. Double-Buffering Time Sequence (Continued)

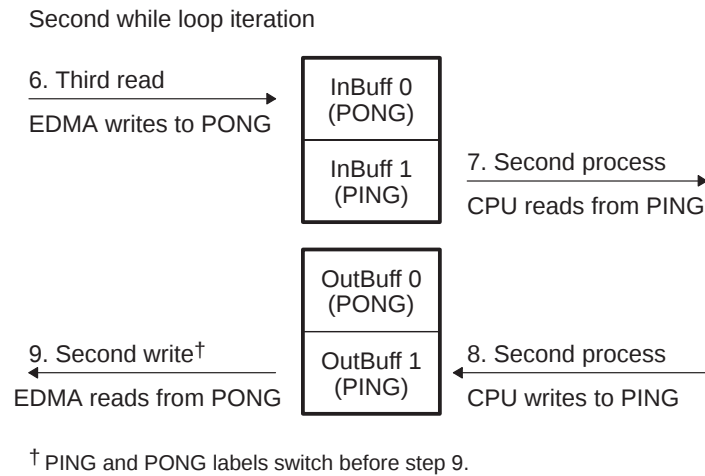
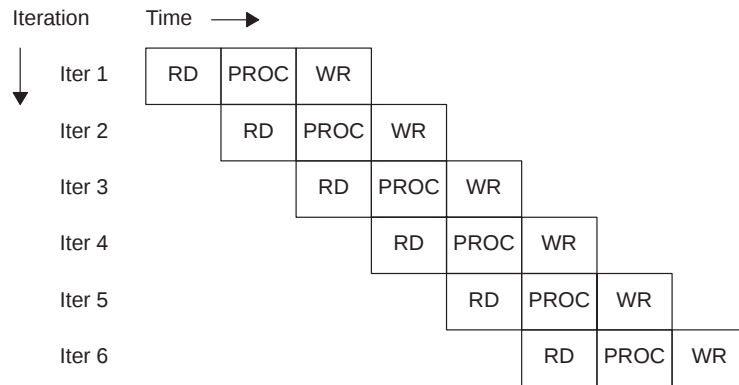


Figure 35. Double Buffering as a Pipelined Process



The memory system's coherence directly supports this programming technique. A system of snoops and invalidates (as discussed in section 5.4.2) keeps L2 SRAM and L1D in sync without programmer intervention. In the context of this example, the memory system performs the following operations in response to each of the numbered steps in Figure 34.

- 1) The EDMA writes to InBuff 0. As the EDMA writes to InBuff 0, the L2 controller sends snoop-invalidates for the corresponding lines in L1D, if it detects that L1D may be caching these addresses. Updated data written back from L1D is merged with the data written by the EDMA, so that L2 contains the most up-to-date data. The corresponding lines are invalidated in L1D.
- 2) This step is similar to step 1, EDMA writes to InBuff 1.
- 3) The CPU reads InBuff 0. These lines were snoop-invalidated from L1D in step 1. Therefore, these accesses miss L1D, forcing L1D to read the new data from L2 SRAM.
- 4) The CPU writes to OutBuff 0. In practice, this step may overlap or interleave with step 3. Suppose initially OutBuff 0 is present in L1D. In this case, the writes hit L1D and mark the lines as dirty. If OutBuff 0 were not present in L1D, the writes would go directly to L2 SRAM.
- 5) The EDMA reads OutBuff 0. The EDMA reads in L2 trigger a snoop for each cache line held in L1D. Any dirty data for the line is written to L2 before the EDMA's read is processed. Thus, the EDMA sees the most up-to-date data. The line is marked clean and is left valid. (On C621x/C671x devices, the line is subsequently invalidated in L1D.)
- 6) The EDMA writes to InBuff 0. This step proceeds identically to step 1. That is, InBuff 0 is snoop-invalidated from L1D as needed. The EDMA writes are processed after any dirty data is written back to L2 SRAM.
- 7) The CPU reads InBuff 1. These lines were snoop-invalidated in step 6. Therefore, these accesses miss L1D, forcing L1D to read the new data from L2 SRAM.
- 8) The CPU writes to OutBuff 1. This step proceeds identically to step 4.
- 9) The EDMA reads OutBuff 1. This step proceeds identically to step 5.

Notice that the system of snoops and snoop-invalidates automatically keeps the EDMA and CPU synchronized for the input and output buffers. Further, double-buffering allows EDMA and CPU accesses to occur in parallel.

8.3 Memory Access Ordering

8.3.1 Program Order of Memory Accesses

To optimize throughput, power, and ease of programming, the C6000 DSP architecture supports a combination of strongly ordered and relaxed ordered memory models. These terms are defined relative to the order of memory operations implied by a particular program sequence. This ordering is referred to as the program order of memory accesses.

The C6000 DSP cores may initiate up to two parallel memory operations per cycle. The program order of memory accesses defines the outcome of memory accesses in terms of a hypothetical serial implementation of the architecture. That is, it describes the order that the parallel memory operations are processed such that time-sequence terms such as earlier and later are used precisely with respect to a particular sequence of operations.

Program order is defined with respect to instructions within an execute packet and with respect to a sequence of execute packets. Memory operations (including those that are issued in parallel) are described as being earlier or later with respect to each other. The terms earlier and later are strictly opposites: if X is not earlier than Y, then X is later than Y.

Memory accesses initiated from different execute packets have the same temporal ordering as the execute packets themselves. That is, in the defined program order, memory operations issued on cycle i are always earlier than memory accesses issued on cycle $i + 1$, and are always later than those issued on cycle $i - 1$.

For accesses issued in parallel, the type of operations (reads or writes), and the data address ports that execute the operations determine the ordering. Table 32 describes the ordering rules.

Table 32. Program Order for Memory Operations Issued From a Single Execute Packet

Data Address 1 (DA1) Operation	Data Address 2 (DA2) Operation	Program Order of Accesses
Load	Load	DA1 is earlier than DA2
Load	Store	DA1 is earlier than DA2
Store	Load	DA2 is earlier than DA1
Store	Store	DA1 is earlier than DA2

The data address port for a load or store instruction is determined by the datapath that provides the data (as opposed to the address) for the memory operation. Load and store instructions that operate on data in the A datapath use DA1. Load and store instructions that operate on data in the B datapath use DA2. Note that the datapath that provides the data to be operated on determines whether DA1 or DA2 is used. The datapath that provides the address for the access is irrelevant.

The C64x DSP supports nonaligned memory accesses to memory using the LDNW, STNW, LDNDW, and STNDW instructions. The memory system does not assure that these memory accesses will be atomic. Rather, it may divide the accesses for these instructions into multiple operations. The program order of memory accesses does not define the order of the individual memory operations that comprise a single nonaligned access. The program order only defines how the entire nonaligned access is ordered relative to earlier and later accesses. So, although the complete nonaligned access does follow the program order defined above with respect to the CPU itself, other requestors may see the nonaligned memory access occur in pieces.

The previous definition describes the memory system semantics. The memory system assures that the semantics of the program order of memory accesses will be retained for CPU accesses relative to themselves. The memory system may, however, relax the ordering of operations as they are executed within the memory hierarchy so long as the correct semantics are retained. It may also allow other requestors to the memory system to see the accesses occur in an order other than the original program order. Section 8.3.2 describes this in detail.

8.3.2 Strong and Relaxed Memory Ordering

The program order of memory accesses (as described in section 8.3.1) describes the desired semantics of a particular sequence of memory accesses. In most circumstances, it is not necessary for the memory system to execute the memory accesses in this exact order at all levels of memory hierarchy in order to retain these semantics. Indeed, for performance and power reasons, it is extremely advantageous for the memory hierarchy to relax the order of accesses.

When communicating with peripherals and when coordinating with other devices (the EDMA, other CPUs) that are accessing the same memory, retaining program order potentially becomes more important. Given these diverse needs, the memory hierarchy supports both strongly ordered and relaxed orderings for memory operations.

Memory system coherence implies that writes to a single memory location are serialized for all requestors, and that all requestors see the same sequence of writes to that location. Coherence does not make any implications about the ordering of accesses to different locations, or the ordering of reads with respect to other reads of the same location. Rather, the memory system ordering rules (strong or relaxed) describe the ordering assurances applied to accesses to different locations.

A sequence of memory operations are said to be strongly ordered if it is not possible to observe a different sequence of memory operations from some place in the system. For instance, if one requestor writes to location X earlier than it writes to a different location Y, a second requestor must not see Y updated before it sees X updated for the accesses to be considered strongly ordered. In this example, regardless of the order in which the memory system processes writes from the first requestor and reads from the second requestor, the second requestor must not be able to observe the write to Y occurring before the write to X. Strong ordering does not require the writes to be serialized, however. The writes may occur in parallel or even in opposite order so long as it is not possible for the second CPU to observe Y being updated before X.

Nonaligned accesses (those issued by the LDNW, STNW, LDNDW, and STNDW instructions) may be strongly ordered relative to other accesses. Because nonaligned accesses are not assured to be atomic, strong ordering assurances apply only to the nonaligned access as a whole. No ordering assurances are made between the individual operations that comprise a single nonaligned access, should the memory system divide the nonaligned access into multiple memory operations.

The memory hierarchy provides strong ordering for all noncacheable long-distance accesses. Such accesses are typically to peripherals and external memory for which the corresponding MAR cache enable bit is not set.

For cacheable locations in external memory, the memory hierarchy provides only a relaxed ordering. Thus, it is possible for other requestors to observe updates occurring in memory in a different order than the original program order of memory accesses. Despite this relaxed ordering, the CPU still sees the desired memory system semantics as described in section 8.3.1.

For cacheable locations in L2 SRAM, the C64x device provides strong ordering for CPU accesses to locations that are within the same L1D cache line, and for locations that are not present in L1D. (On the C621x/C671x devices, this means that data accesses to addresses whose upper 27 bits are equal are strongly ordered.) On the C64x device, this means that data accesses to addresses whose upper 26 bits are equal are strongly ordered.

For locations in L2 SRAM that are not within the same cache line, strong ordering is provided only on writes and only as long as addresses involved are not present in L1D. This can be ensured by using the L1DWIBAR and L1DWIWC control registers described in section 7.3.2. In all other cases, a relaxed ordering is provided for CPU accesses to L2 SRAM.

The L2 provides limited ordering assurances for EDMA access to L2 SRAM. EDMA reads to L2 are not reordered relative to other EDMA reads. EDMA writes to L2 are not reordered relative to other EDMA writes. Reads and writes, however, may be reordered relative to each other.

This page is intentionally left blank.

Revision History

Table 33 lists the changes made since the previous version of this document.

Table 33. Document Revision History

Page	Additions/Modifications/Deletions
35	Changed the text, section 7, Memory System Controls, to section 6, Registers, in section 5.4.
72	Changed Figure 33.

This page is intentionally left blank.

Index

B

- block cache operation base address register (BAR) 63
- block cache operation word count register (WC) 63
- block cache operations 63
- block diagram
 - C64x DSP 9
 - two-level internal memory 12

C

- C bit
 - in L2WB 49
 - in L2WBINV 50
- cache configuration register (CCFG) 39
- cacheability controls 56
- CCFG 39
- CE bit 51

D

- DCC field in CSR 52

E

- EDMA access into L2 control 68
- EDMA access to cache controls 37
- EDMA coherence in L2 SRAM example 71
- EDMA-to-L2 request servicing 36
- EDMAWEIGHT 41
- effect of L2 commands on L1 caches 65

G

- global cache operations 62

H

- HPI and PCI access to memory subsystem 37

I

- ID bit 39
- IP bit 39

L

- L1D invalidate base address register (L1DIBAR) 48
- L1D invalidate word count register (L1DIWC) 48
- L1D memory banking 20
- L1D miss penalty 22
- L1D miss pipelining 24
- L1D parameters 19
- L1D performance 20
- L1D write buffer 23
- L1D writeback-invalidate base address register (L1DWIBAR) 47
- L1D writeback-invalidate word count register (L1DWIWC) 47
- L1D/L1P-to-L2 request servicing 35
- L1DIBAR 48
- L1DIWC 48
- L1DWIBAR 47
- L1DWIWC 47
- L1P invalidate base address register (L1PIBAR) 46
- L1P invalidate word count register (L1PIWC) 46
- L1P miss penalty 27
- L1P miss pipelining 28

- L1P parameters 27
- L1P performance 27
- L1PIBAR 46
- L1PIWC 46
- L2 allocation registers (L2ALLOC0–L2ALLOC3) 42
- L2 bank structure 34
- L2 cache and L2 SRAM 30
- L2 EDMA access control register (EDMAWEIGHT) 41
- L2 interfaces 35
- L2 invalidate base address register (L2IBAR) 45
- L2 invalidate word count register (L2IWC) 45
- L2 memory attribute registers (MAR0–MAR255) 51
 - address ranges 57
- L2 operation 32
- L2 request servicing using EDMA 36
- L2 writeback all register (L2WB) 49
- L2 writeback base address register (L2WBAR) 43
- L2 writeback word count register (L2WWC) 43
- L2 writeback–invalidate all register (L2WBINV) 50
- L2 writeback–invalidate base address register (L2WIBAR) 44
- L2 writeback–invalidate word count register (L2WIWC) 44
- L2-to-EDMA request control 67
- L2ALLOC0–L2ALLOC3 42
- L2IBAR 45
- L2IWC 45
- L2MODE bits 39
- L2WB 49
- L2WBAR 43
- L2WBINV 50
- L2WIBAR 44
- L2WIWC 44
- L2WWC 43
- level 1 data cache (L1D) 19
 - mode selection using DCC field in CSR 52
 - parameters 19
 - performance 20

- level 1 program cache (L1P) 27
 - mode selection using PCC field in CSR 53
 - parameters 27
 - performance 27
- level 2 unified memory (L2) 30
 - bank structure 34
 - cache and SRAM 30
 - interfaces 35
 - mode selection using L2MODE field in CCFG 53
 - operation 32

M

- MAR 51
- memory access ordering 76
 - program order of memory accesses 76
 - strong and relaxed memory ordering 77
- memory hierarchy overview 9
- memory system coherence 69
- memory system control 52
 - cache mode selection 52
 - cacheability controls 56
 - EDMA access into L2 control 68
 - L2-to-EDMA request control 67
 - program-initiated cache operations 60
- memory system policies 69
 - coherence 69
 - EDMA coherence in L2 SRAM example 71
 - memory access ordering 76

N

- notational conventions 3

P

- P bits 39
- PCC field in CSR 53
- program order of memory accesses 76
- program-initiated cache operations 60

R

registers 38

cache configuration register (CCFG) 39

L1D invalidate base address register
(L1DIBAR) 48

L1D invalidate word count register
(L1DIWC) 48

L1D writeback-invalidate base address register
(L1DWIBAR) 47

L1D writeback-invalidate word count register
(L1DWIWC) 47

L1P invalidate base address register
(L1PIBAR) 46

L1P invalidate word count register (L1PIWC) 46

L2 allocation registers
(L2ALLOC0–L2ALLOC3) 42

L2 EDMA access control register
(EDMAWEIGHT) 41

L2 invalidate base address register
(L2IBAR) 45

L2 invalidate word count register (L2IWC) 45

L2 memory attribute registers
(MAR0–MAR255) 51

L2 writeback all register (L2WB) 49

L2 writeback base address register
(L2WBAR) 43

L2 writeback word count register (L2WWC) 43

L2 writeback-invalidate all register
(L2WBINV) 50

L2 writeback-invalidate base address register
(L2WIBAR) 44

L2 writeback-invalidate word count register
(L2WIWC) 44

related documentation from Texas Instruments 3

revision history 81

S

strong and relaxed memory ordering 77

T

terms and definitions 13

trademarks 4