



ALTIVECPIM/D
6/1999
Rev. 0

AltiVec™ Technology Programming Interface Manual



**Home Page:**

www.freescale.com

email:

support@freescale.com

USA/Europe or Locations Not Listed:

Freescale Semiconductor
Technical Information Center, CH370
1300 N. Alma School Road
Chandler, Arizona 85224
(800) 521-6274
480-768-2130

support@freescale.com

Europe, Middle East, and Africa:

Freescale Halbleiter Deutschland GmbH
Technical Information Center
Schatzbogen 7
81829 Muenchen, Germany
+44 1296 380 456 (English)
+46 8 52200080 (English)
+49 89 92103 559 (German)
+33 1 69 35 48 48 (French)
support@freescale.com

Japan:

Freescale Semiconductor Japan Ltd.
Headquarters
ARCO Tower 15F
1-8-1, Shimo-Meguro, Meguro-ku
Tokyo 153-0064, Japan
0120 191014
+81 2666 8080
support.japan@freescale.com

Asia/Pacific:

Freescale Semiconductor Hong Kong Ltd.
Technical Information Center
2 Dai King Street
Tai Po Industrial Estate,
Tai Po, N.T., Hong Kong
+800 2666 8080
support.asia@freescale.com

For Literature Requests Only:

Freescale Semiconductor
Literature Distribution Center
P.O. Box 5405
Denver, Colorado 80217
(800) 441-2447
303-675-2140
Fax: 303-675-2150
LDCForFreescaleSemiconductor@hibbertgroup.com

Information in this document is provided solely to enable system and software implementers to use Freescale Semiconductor products. There are no express or implied copyright licenses granted hereunder to design or fabricate any integrated circuits or integrated circuits based on the information in this document.

Freescale Semiconductor reserves the right to make changes without further notice to any products herein. Freescale Semiconductor makes no warranty, representation or guarantee regarding the suitability of its products for any particular purpose, nor does Freescale Semiconductor assume any liability arising out of the application or use of any product or circuit, and specifically disclaims any and all liability, including without limitation consequential or incidental damages. "Typical" parameters which may be provided in Freescale Semiconductor data sheets and/or specifications can and do vary in different applications and actual performance may vary over time. All operating parameters, including "Typicals" must be validated for each customer application by customer's technical experts. Freescale Semiconductor does not convey any license under its patent rights nor the rights of others. Freescale Semiconductor products are not designed, intended, or authorized for use as components in systems intended for surgical implant into the body, or other applications intended to support or sustain life, or for any other application in which the failure of the Freescale Semiconductor product could create a situation where personal injury or death may occur. Should Buyer purchase or use Freescale Semiconductor products for any such unintended or unauthorized application, Buyer shall indemnify and hold Freescale Semiconductor and its officers, employees, subsidiaries, affiliates, and distributors harmless against all claims, costs, damages, and expenses, and reasonable attorney fees arising out of, directly or indirectly, any claim of personal injury or death associated with such unintended or unauthorized use, even if such claim alleges that Freescale Semiconductor was negligent regarding the design or manufacture of the part.





Overview **1**

High-Level Language Interface **2**

Application Binary Interface **3**

AltiVec Operations and Predicates **4**

AltiVec Instruction Set/Operations/Predicates Cross-Reference **A**

Glossary of Terms and Abbreviations **GLO**

Index **IND**



1

Overview

2

High-Level Language Interface

3

Application Binary Interface

4

AltiVec Operations and Predicates

A

AltiVec Instruction Set/Operations/Predicates Cross-Reference

GLO

Glossary of Terms and Abbreviations

IND

Index



CONTENTS

Paragraph Number	Title	Page Number
	Audience	xvi
	Organization.....	xvi
	Suggested Reading.....	xvii
	PowerPC Documentation.....	xvii
	General Information.....	xviii

Chapter 1 Overview

1.1	High-Level Language Interface	1-1
1.2	Application Binary Interface (ABI)	1-2

Chapter 2 High-Level Language Interface

2.1	Data Types	2-1
2.2	New Keywords.....	2-2
2.2.1	The Keyword and Predefine Method.....	2-2
2.2.2	The Context Sensitive Keyword Method.....	2-3
2.3	Alignment	2-3
2.3.1	Alignment of Vector Types	2-3
2.3.2	Alignment of Non-Vector Types	2-3
2.3.3	Alignment of Aggregates and Unions Containing Vector Types	2-3
2.4	Extensions of C/C++ Operators for the New Types	2-4
2.4.1	sizeof()	2-4
2.4.2	Assignment	2-4
2.4.3	Address Operator	2-4
2.4.4	Pointer Arithmetic.....	2-4
2.4.5	Pointer Dereferencing	2-4
2.4.6	Type Casting	2-5
2.5	New Operators	2-5
2.5.1	Vector Literals	2-5
2.5.2	Vector Literals and Casts	2-6
2.5.3	Value for Adjusting Pointers	2-7
2.5.4	New Operators Representing AltiVec Operations.....	2-7
2.6	Programming Interface	2-8

Chapter 3 Application Binary Interface (ABI)

3.1	Data Representation.....	3-1
3.2	Register Usage Conventions	3-1



CONTENTS

Paragraph Number	Title	Page Number
3.3	The Stack Frame	3-2
3.3.1	SVR4 ABI and EABI Stack Frame.....	3-3
3.3.2	Apple Macintosh ABI and AIX ABI Stack Frame	3-5
3.3.3	Vector Register Saving and Restoring Functions	3-7
3.4	Function Calls	3-9
3.4.1	SVR4 ABI and EABI Parameter Passing and Varargs.....	3-9
3.4.2	Apple Macintosh ABI and AIX ABI Parameter Passing without Varargs.....	3-9
3.4.3	Apple Macintosh ABI and AIX ABI Parameter Passing with Varargs.....	3-10
3.5	malloc(), vec_malloc(), and new	3-10
3.6	setjmp() and longjmp()	3-11
3.7	Debugging Information.....	3-11
3.8	printf() and scanf() Control Strings.....	3-12
3.8.1	Output Conversion Specifications	3-12
3.8.2	Input Conversion Specifications	3-14

Chapter 4 AltiVec Operations and Predicates

4.1	Vector Status and Control Register.....	4-1
4.2	Byte Ordering.....	4-3
4.3	Notation and Conventions.....	4-4
4.4	Generic and Specific AltiVec Operations.....	4-7
4.5	AltiVec Predicates	4-133

Appendix A AltiVec Instruction Set/Operation/Predicate Cross-Reference

Glossary of Terms and Abbreviations

Index

ILLUSTRATIONS

Figure Number	Title	Page Number
3-1	SVR4 ABI and EABI Stack Frame	3-3
3-2	Apple Macintosh ABI and AIX ABI Stack Frame.....	3-5
4-1	Vector Status and Control Register (VSCR)	4-1
4-2	VSCR Moved to a Vector Register	4-1
4-3	Big-Endian Byte Ordering for a Vector Register	4-3
4-4	Operation Description Format	4-7
4-5	Absolute Value of Sixteen Integer Elements (8-bit)	4-8
4-6	Absolute Value of Eight Integer Elements (16-bit).....	4-9
4-7	Absolute Value of Four Integer Elements (32-bit)	4-9
4-8	Absolute Value of Four Floating-Point Elements (32-bit)	4-9
4-9	Saturated Absolute Value of Sixteen Integer Elements (8-bit)	4-10
4-10	Saturated Absolute Value of Eight Integer Elements (16-bit).....	4-11
4-11	Saturated Absolute Value of Four Integer Elements (32-bit).....	4-11
4-12	Add Sixteen Integer Elements (8-bit)	4-12
4-13	Add Eight Integer Elements (16-bit)	4-13
4-14	Add Four Integer Elements (32-bit)	4-13
4-15	Add Four Floating-Point Elements (32-bit).....	4-14
4-16	Carryout of Four Unsigned Integer Adds (32-bit)	4-15
4-17	Add Saturating Sixteen Integer Elements (8-bit)	4-16
4-18	Add Saturating Eight Integer Elements (16-bit).....	4-17
4-19	Add Saturating Four Integer Elements (32-bit)	4-17
4-20	Logical Bit-Wise AND	4-18
4-21	Logical Bit-Wise AND with Complement	4-19
4-22	Average Sixteen Integer Elements (8-bit)	4-21
4-23	Average Eight Integer Elements (16-bit).....	4-22
4-24	Average Four Integer Elements (32-bit).....	4-22
4-25	Round to Plus Infinity of Four Floating-Point Integer Elements (32-Bit)	4-23
4-26	Compare Bounds of Four Floating-Point Elements (32-Bit).....	4-24
4-27	Compare Equal of Sixteen Integer Elements (8-bits)	4-25
4-28	Compare Equal of Eight Integer Elements (16-Bit)	4-26
4-29	Compare Equal of Four Integer Elements (32-Bit)	4-26
4-30	Compare Equal of Four Floating-Point Elements (32-Bit)	4-26
4-31	Compare Greater-Than-or-Equal of Four Floating-Point Elements (32-Bit)	4-27
4-32	Compare Greater-Than of Sixteen Integer Elements (8-bits).....	4-28
4-33	Compare Greater-Than of Eight Integer Elements (16-Bit)	4-29
4-34	Compare Greater-Than of Four Integer Elements (32-Bit)	4-29
4-35	Compare Greater-Than of Four Floating-Point Elements (32-Bit)	4-29
4-36	Compare Less-Than-or-Equal of Four Floating-Point Elements (32-Bit).....	4-30
4-37	Compare Less-Than of Sixteen Integer Elements (8-bits)	4-31
4-38	Compare Less-Than of Eight Integer Elements (16-Bit).....	4-32
4-39	Compare Less-Than of Four Integer Elements (32-Bit).....	4-32
4-40	Compare Less-Than of Four Floating-Point Elements (32-Bit)	4-32
4-41	Convert Four Integer Elements to Four Floating-Point Elements (32-Bit)	4-33

ILLUSTRATIONS

Figure Number	Title	Page Number
4-42	Convert Four Floating-Point Elements to Four Saturated Signed Integer Elements (32-Bit)	4-34
4-43	Convert Four Floating-Point Elements to Four Saturated Unsigned Integer Elements (32-Bit)	4-35
4-44	Format of b Type (32-bit)	4-38
4-45	Format of b Type (64-bit)	4-38
4-46	Format of b Type (32-bit)	4-40
4-47	Format of b Type (64-bit)	4-40
4-48	Format of b Type (32-bit)	4-42
4-49	Format of b Type (64-bit)	4-42
4-50	Format of b Type (32-bit)	4-44
4-51	Format of b Type (64-bit)	4-44
4-52	2 Raised to the Exponent Estimate Floating-Point for Four Floating-Point Elements (32-Bit)	4-46
4-53	Round to Minus Infinity of Four Floating-Point Integer Elements (32-Bit)	4-47
4-54	Vector Load Indexed Operation	4-48
4-55	Vector Load Element Indexed Operation	4-50
4-56	Vector Load Indexed LRU Operation	4-51
4-57	Log2 Estimate Floating-Point for Four Floating-Point Elements (32-Bit)	4-53
4-58	Multiply-Add Four Floating-Point Elements (32-Bit)	4-56
4-59	Multiply-Add Four Floating-Point Elements (32-Bit)	4-57
4-60	Maximum of Sixteen Integer Elements (8-Bit)	4-58
4-61	Maximum of Eight Integer Elements (16-bit)	4-59
4-62	Maximum of Four Integer Elements (32-bit)	4-59
4-63	Maximum of Four Floating-Point Elements (32-bit)	4-60
4-64	Merge Eight High-Order Elements (8-Bit)	4-61
4-65	Merge Four High-Order Elements (16-bit)	4-62
4-66	Merge Two High-Order Elements (32-bit)	4-62
4-67	Merge Eight Low-Order Elements (8-Bit)	4-63
4-68	Merge Four Low-Order Elements (16-bit)	4-64
4-69	Merge Two Low-Order Elements (32-bit)	4-64
4-70	Vector Move from VSCR	4-65
4-71	Minimum of Sixteen Integer Elements (8-Bit)	4-66
4-72	Minimum of Eight Integer Elements (16-bit)	4-67
4-73	Minimum of Four Integer Elements (32-bit)	4-67
4-74	Minimum of Four Floating-Point Elements (32-bit)	4-68
4-75	Multiply-Add of Eight Integer Elements (16-Bit)	4-69
4-76	Multiply-Add of Eight Integer Elements (16-Bit)	4-70
4-77	Multiply Sum of Sixteen Integer Elements (8-Bit)	4-71
4-78	Multiply Sum of Eight Integer Elements (16-Bit)	4-72
4-79	Multiply-Sum of Integer Elements (16-Bit to 32-Bit)	4-73
4-80	Vector Move to VSCR	4-74
4-81	Even Multiply of Eight Integer Elements (8-Bit)	4-75

ILLUSTRATIONS

Figure Number	Title	Page Number
4-82	Even Multiply of Four Integer Elements (16-Bit)	4-75
4-83	Odd Multiply of Eight Integer Elements (8-Bit)	4-76
4-84	Odd Multiply of Four Integer Elements (16-Bit)	4-76
4-85	Negative Multiply-Subtract of Four Floating-Point Elements (32-Bit)	4-77
4-86	Logical Bit-Wise NOR	4-78
4-87	Logical Bit-Wise OR	4-79
4-88	Pack Sixteen Unsigned Integer Elements (16-Bit) to Sixteen Unsigned Integer Elements (8-Bit)	4-80
4-89	Pack Eight Unsigned Integer Elements (32-Bit) to Eight Unsigned Integer Elements (16-Bit)	4-80
4-90	Pack Eight Pixel Elements (32-Bit) to Eight Elements (16-Bit)	4-81
4-91	Pack Sixteen Integer Elements (16-Bit) to Sixteen Integer Elements (8-Bit)	4-82
4-92	Pack Eight Integer Elements (32-Bit) to Eight Integer Elements (16-Bit)	4-82
4-93	Pack Sixteen Integer Elements (16-Bit) to Sixteen Unsigned Integer Elements (8-Bit)	4-83
4-94	Pack Eight Integer Elements (32-Bit) to Eight Unsigned Integer Elements (16-Bit)	4-83
4-95	Permute Sixteen Integer Elements (8-Bit)	4-84
4-96	Reciprocal Estimate of Four Floating-Point Elements (32-Bit)	4-85
4-97	Left Rotate of Sixteen Integer Elements (8-Bit)	4-86
4-98	Left Rotate of Eight Integer Elements (16-bit)	4-86
4-99	Left Rotate of Four Integer Elements (32-bit)	4-87
4-100	Round to Nearest of Four Floating-Point Integer Elements (32-Bit)	4-88
4-101	Reciprocal Square Root Estimate of Four Floating-Point Elements (32-Bit)	4-89
4-102	Bit-Wise Conditional Select of Vector Contents (128-bit)	4-90
4-103	Shift Bits Left in Sixteen Integer Elements (8-Bit)	4-91
4-104	Shift Bits Left in Eight Integer Elements (16-bit)	4-92
4-105	Shift Bits Left in Four Integer Elements (32-Bit)	4-92
4-106	Bit-Wise Conditional Select of Vector Contents (128-bit)	4-93
4-107	Shift Bits Left in Vector (128-Bit)	4-95
4-108	Left Byte Shift of Vector (128-Bit)	4-96
4-109	Copy Contents to Sixteen Integer Elements (8-Bit)	4-97
4-110	Copy Contents to Eight Elements (16-bit)	4-97
4-111	Copy Contents to Four Integer Elements (32-Bit)	4-98
4-112	Copy Value into Sixteen Signed Integer Elements (8-Bit)	4-99
4-113	Copy Value into Eight Signed Integer Elements (16-Bit)	4-100
4-114	Copy Value into Four Signed Integer Elements (32-Bit)	4-101
4-115	Copy Value into Sixteen Signed Integer Elements (8-Bit)	4-102
4-116	Copy Value into Eight Signed Integer Elements (16-Bit)	4-103
4-117	Copy Value into Four Signed Integer Elements (32-Bit)	4-104
4-118	Shift Bits Right in Sixteen Integer Elements (8-Bit)	4-105
4-119	Shift Bits Right in Eight Integer Elements (16-bit)	4-106
4-120	Shift Bits Right in Four Integer Elements (32-Bit)	4-106

ILLUSTRATIONS

Figure Number	Title	Page Number
4-121	Shift Bits Right in Sixteen Integer Elements (8-Bit).....	4-107
4-122	Shift Bits Right in Eight Integer Elements (16-bit).....	4-108
4-123	Shift Bits Right in Four Integer Elements (32-Bit)	4-108
4-124	Shift Bits Right in Vector (128-Bit)	4-110
4-125	Right Byte Shift of Vector (128-Bit).....	4-111
4-126	Vector Store Indexed	4-112
4-127	Vector Store Element.....	4-115
4-128	Vector Store Indexed LRU	4-116
4-129	Subtract Sixteen Integer Elements (8-bit)	4-118
4-130	Subtract Eight Integer Elements (16-bit).....	4-119
4-131	Subtract Four Integer Elements (32-bit).....	4-119
4-132	Subtract Four Floating-Point Elements (32-bit).....	4-120
4-133	Carryout of Four Unsigned Integer Subtracts (32-bit)	4-121
4-134	Subtract Saturating Sixteen Integer Elements (8-bit)	4-122
4-135	Subtract Saturating Eight Integer Elements (16-bit)	4-123
4-136	Subtract Saturating Four Integer Elements (32-bit)	4-123
4-137	Four Sums in the Integer Elements (32-Bit).....	4-124
4-138	Four Sums in the Integer Elements (32-Bit).....	4-124
4-139	Two Saturated Sums in the Four Signed Integer Elements (32-Bit)	4-125
4-140	Saturated Sum of Five Signed Integer Elements (32-Bit).....	4-126
4-141	Round-to-Zero of Four Floating-Point Integer Elements (32-Bit)	4-127
4-142	Unpack High-Order Elements (8-Bit) to Elements (16-Bit)	4-128
4-143	Unpack High-Order Pixel Elements (16-Bit) to Elements (32-Bit)	4-129
4-144	Unpack High-Order Signed Integer Elements (16-Bit) to Signed Integer Elements (32-Bit)	4-129
4-145	Unpack Low-Order Elements (8-Bit) to Elements (16-Bit)	4-130
4-146	Unpack Low-Order Pixel Elements (16-Bit) to Elements (32-Bit).....	4-130
4-147	Unpack Low-Order Signed Integer Elements (16-Bit) to Signed Integer Elements (32-Bit)	4-131
4-148	Logical Bit-Wise XOR	4-132
4-149	All Equal of Sixteen Integer Elements (8-bits)	4-134
4-150	All Equal of Eight Integer Elements (16-Bit).....	4-135
4-151	All Equal of Four Integer Elements (32-Bit).....	4-135
4-152	All Equal of Four Floating-Point Elements (32-Bit)	4-136
4-153	All Greater Than or Equal of Sixteen Integer Elements (8-bits).....	4-137
4-154	All Greater Than or Equal of Eight Integer Elements (16-Bit)	4-138
4-155	All Greater Than or Equal of Four Integer Elements (32-Bit)	4-138
4-156	All Greater Than or Equal of Four Floating-Point Elements (32-Bit)	4-139
4-157	All Greater Than of Sixteen Integer Elements (8-bits).....	4-140
4-158	All Greater Than of Eight Integer Elements (16-Bit).....	4-141
4-159	All Greater Than of Four Integer Elements (32-Bit).....	4-141
4-160	All Greater Than of Four Floating-Point Elements (32-Bit)	4-142
4-161	All in Bounds of Four Floating-Point Elements (32-Bit)	4-143

ILLUSTRATIONS

Figure Number	Title	Page Number
4-162	All Less Than or Equal of Sixteen Integer Elements (8-bits).....	4-144
4-163	All Less Than or Equal of Eight Integer Elements (16-Bit).....	4-145
4-164	All Less Than or Equal of Four Integer Elements (32-Bit).....	4-145
4-165	All Less Than or Equal of Four Floating-Point Elements (32-Bit)	4-146
4-166	All Less Than of Sixteen Integer Elements (8-bits)	4-147
4-167	All Less Than of Eight Integer Elements (16-Bit)	4-148
4-168	All Less Than of Four Integer Elements (32-Bit).....	4-148
4-169	All Less Than of Four Floating-Point Elements (32-Bit).....	4-149
4-170	All NaN of Four Floating-Point Elements (32-Bit).....	4-150
4-171	All Not Equal of Sixteen Integer Elements (8-bits)	4-151
4-172	All Not Equal of Eight Integer Elements (16-Bit).....	4-152
4-173	All Not Equal of Four Integer Elements (32-Bit).....	4-152
4-174	All Not Equal of Four Floating-Point Elements (32-Bit)	4-153
4-175	All Not Greater Than or Equal of Four Floating-Point Elements (32-Bit)	4-154
4-176	All Not Greater Than of Four Floating-Point Elements (32-Bit)	4-155
4-177	All Not Less Than or Equal of Four Floating-Point Elements (32-Bit)	4-156
4-178	All Not Less Than of Four Floating-Point Elements (32-Bit).....	4-157
4-179	All Numeric of Four Floating-Point Elements (32-Bit)	4-158
4-180	Any Equal of Sixteen Integer Elements (8-bits).....	4-159
4-181	Any Equal of Eight Integer Elements (16-Bit).....	4-160
4-182	Any Equal of Four Integer Elements (32-Bit)	4-160
4-183	Any Equal of Four Floating-Point Elements (32-Bit)	4-161
4-184	Any Greater Than or Equal of Sixteen Integer Elements (8-bits)	4-162
4-185	Any Greater Than or Equal of Eight Integer Elements (16-Bit)	4-163
4-186	Any Greater Than or Equal of Four Integer Elements (32-Bit).....	4-163
4-187	Any Greater Than or Equal of Four Floating-Point Elements (32-Bit).....	4-164
4-188	Any Greater Than of Sixteen Integer Elements (8-bits).....	4-165
4-189	Any Greater Than of Eight Integer Elements (16-Bit)	4-166
4-190	Any Greater Than of Four Integer Elements (32-Bit)	4-166
4-191	Any Greater Than of Four Floating-Point Elements (32-Bit)	4-167
4-192	Any Less Than or Equal of Sixteen Integer Elements (8-bits).....	4-168
4-193	Any Less Than or Equal of Eight Integer Elements (16-Bit).....	4-169
4-194	Any Less Than or Equal of Four Integer Elements (32-Bit)	4-169
4-195	Any Less Than or Equal of Four Floating-Point Elements (32-Bit)	4-170
4-196	Any Less Than of Sixteen Integer Elements (8-bits)	4-171
4-197	Any Less Than of Eight Integer Elements (16-Bit).....	4-172
4-198	Any Less Than of Four Integer Elements (32-Bit).....	4-172
4-199	Any Less Than of Four Floating-Point Elements (32-Bit)	4-173
4-200	Any NaN of Four Floating-Point Elements (32-Bit)	4-174
4-201	Any Not Equal of Sixteen Integer Elements (8-bits).....	4-175
4-202	Any Not Equal of Eight Integer Elements (16-Bit).....	4-176
4-203	Any Not Equal of Four Integer Elements (32-Bit)	4-176
4-204	Any Not Equal of Four Floating-Point Elements (32-Bit)	4-177



ILLUSTRATIONS

Figure Number	Title	Page Number
4-205	Any Not Greater Than or Equal of Four Floating-Point Elements (32-Bit)	4-178
4-206	Any Not Greater Than of Four Floating-Point Elements (32-Bit)	4-179
4-207	Any Not Less Than or Equal of Four Floating-Point Elements (32-Bit)	4-180
4-208	Any Not Less Than of Four Floating-Point Elements (32-Bit)	4-181
4-209	Any Numeric of Four Floating-Point Elements (32-Bit)	4-182
4-210	Any Out of Bounds of Four Floating-Point Elements (32-Bit)	4-183

TABLES

Table Number	Title	Page Number
2-1	AltiVec Data Types	2-1
2-2	Vector Literal Format and Description.....	2-7
2-3	Increment Value for vec_step by Data Type	2-8
3-1	AltiVec Registers.....	3-1
3-2	Vector Registers Valid Tag Format.....	3-3
3-3	ABI Specifications for setjmp() and longjmp()	3-11
4-1	VSCR Field Descriptions.....	4-2
4-2	Notation and Conventions	4-4
4-3	Precedence Rules	4-6
4-4	vec_dss—Vector Data Stream Stop Argument Types.....	4-36
4-5	vec_dst—Vector Data Stream Touch Argument Types	4-39
4-6	vec_dstst—Vector Data Stream for Touch Store Argument Types	4-41
4-7	vec_dststt—Vector Data Stream Touch for Store Transient Argument Types	4-43
4-8	vec_dstt—Vector Data Stream Touch Transient Argument Types	4-45
4-9	vec_ld—Load Vector Indexed Argument Types.....	4-49
4-10	vec_lde(a,b)—Vector Load Element Indexed Argument Types	4-50
4-11	vec_ldl—Vector Load Indexed LRU Argument Types.....	4-52
4-12	vec_lvsl—Load Vector for Shift Left Argument Types.....	4-54
4-13	vec_lvslr—Vector Load for Shift Right Argument Types	4-55
4-14	Vector Move from Vector Status and Control Registers Argument Type and Mapping.....	4-65
4-15	vec_mtvscr—Vector Move to Vector Status and Control Register Argument Types	4-74
4-16	Special Value Results of Reciprocal Estimates	4-85
4-17	Special Value Results of Reciprocal Square Root Estimates	4-89
4-18	vec_st—Vector Store Indexed Argument Types.....	4-113
4-19	vec_stl—Vector Store Index Argument Types.....	4-117
A-1	Instructions to Operations/Predicates Cross-Reference.....	A-1
A-2	Operations to Instructions Cross-Reference	A-7
A-3	Predicate to Instruction Cross-Reference	A-14



TABLES

**Table
Number**

Title

**Page
Number**



About This Book

The primary objective of this manual is to help programmers to provide software that is compatible across the family of PowerPC™ processors using AltiVec™ technology.

To locate any published errata or updates for this document, refer to the website at <http://www.mot.com/SPS/PowerPC/>.

This book is one of two that discuss the AltiVec architecture, the two books are:

- *AltiVec: The Programming Interface Manual (AltiVec PIM)* is used as a reference guide for high-level programmers. The AltiVec PIM provides a mechanism for programmers to access AltiVec functionality from programming languages such as C and C++. The AltiVec PIM defines a programming model for use with the AltiVec instruction set extension to the PowerPC architecture.
- *AltiVec: The Programming Environments Manual (AltiVec PEM)* is used as a reference guide for assembler programmers. The AltiVec PEM provides a description for each instruction that includes the instruction format, an individualized legend that provides such information as the level(s) of the PowerPC architecture in which the instruction may be found, the privilege level of the instruction, and figures to help in understanding how the instruction works.

It is beyond the scope of this manual to describe individual AltiVec technology implementations on PowerPC processors. It must be kept in mind that each PowerPC processor is unique in its implementation of the AltiVec technology.

The information in this book is subject to change without notice, as described in the disclaimers on the title page of this book. As with any technical documentation, it is the readers' responsibility to be sure they are using the most recent version of the documentation. For more information, contact your sales representative or visit our website at: <http://www.mot.com/SPS/PowerPC/>.



Audience

This manual is intended for system software and application programmers who want to develop products using the AltiVec technology extension to the PowerPC processors in general. It is assumed that the reader understands operating systems, microprocessor system design, the basic principles of RISC processing, and the AltiVec Instruction Set.

Organization

Following is a summary and a brief description of the major sections of this manual:

- Chapter 1, “Overview,” is useful for those who want a general understanding of what the programming model defines in the AltiVec technology.
- Chapter 2, “High-Level Language Interface,” is useful for software engineers who need to understand how to access AltiVec functionality from high level languages such as C and C++.
- Chapter 3, “Application Binary Interface (ABI),” describes AltiVec extensions for System V Application Binary Interface PowerPC Processor Supplement (SVR4 ABI), the PowerPC Embedded Application Binary Interface (EABI), Appendix A of The PowerPC Compiler Writer’s Guide (AIX ABI), and the Apple Macintosh ABI.
- Chapter 4, “AltiVec Operations and Predicates,” alphabetically defines the AltiVec operations and predicates. Each AltiVec operation and predicate description includes a pseudocode functional description and figures illustrating that function, a valid set of argument types for that AltiVec operation or predicate, the result type for that set of argument types, and the specific AltiVec instruction generated for that set of arguments.
- Appendix A, “AltiVec Instruction Set/Operation/Predicate Cross-Reference,” cross-references the AltiVec instruction set, operations, and predicates by functionality.
- This manual also includes a glossary and an index.



Suggested Reading

This section lists additional reading that provides background for the information in this manual as well as general information about the AltiVec technology and PowerPC architecture.

PowerPC Documentation

The PowerPC documentation is organized in the following types of documents:

- User's manuals—These books provide details about individual PowerPC implementations and are intended to be used in conjunction with *PowerPC Microprocessor Family: The Programming Environments Manual*.
- *PowerPC Microprocessor Family: The Programming Environments*, Rev. 1 provides information about resources defined by the PowerPC architecture that are common to PowerPC processors. This document describes both the 64- and 32-bit portions of the architecture. MPCFPE/AD (Freescale order #).
- *Implementation Variances Relative to Rev. 1 of The Programming Environments Manual* is available via the world-wide web at <http://www.mot.com/SPS/PowerPC/>.
- Addenda/errata to user's manuals—Because some processors have follow-on parts an addendum is provided that describes the additional features and changes to functionality of the follow-on part. These addenda are intended for use with the corresponding user's manuals.
- Hardware specifications—Hardware specifications provide specific data regarding bus timing, signal behavior, and AC, DC, and thermal characteristics, as well as other design considerations for each PowerPC implementation.
- Technical Summaries—Each PowerPC implementation has a technical summary that provides an overview of its features. This document is roughly the equivalent to the overview (Chapter 1) of an implementation's user's manual.
- *PowerPC Microprocessor Family: The Programmer's Reference Guide: MPCPRG/D* (Freescale order #) is a concise reference that includes the register summary, memory control model, exception vectors, and the PowerPC instruction set.
- *PowerPC Microprocessor Family: The Programmer's Pocket Reference Guide: MPCPRGREF/D* (Freescale order #) This foldout card provides an overview of the PowerPC registers, instructions, and exceptions for 32-bit implementations.
- Application notes—These short documents contain useful information about specific design issues useful to programmers and engineers working with PowerPC processors (available via the worldwide web at <http://www.mot.com/SPS/PowerPC/>).
- Documentation for support chips



Additional literature on AltiVec technology and PowerPC implementations is being released as new processors become available. For a current list of AltiVec technology and PowerPC documentation, refer to the website at <http://www.mot.com/SPS/PowerPC/>.

General Information

The following documentation provides useful information about the PowerPC architecture and computer architecture in general:

- The following books are available from the Morgan-Kaufmann Publishers, 340 Pine Street, Sixth Floor, San Francisco, CA 94104; Tel. (800) 745-7323 (U.S.A.), (415) 392-2665 (International); internet address: mkp@mkp.com.
 - *The PowerPC Architecture: A Specification for a New Family of RISC Processors*, Second Edition, by International Business Machines, Inc.
Updates to the architecture specification are accessible via the world-wide web at <http://www.austin.ibm.com/tech/ppc-chg.html>.
 - *PowerPC Microprocessor Common Hardware Reference Platform: A System Architecture*, by Apple Computer, Inc., International Business Machines, Inc., And Freescale Semiconductor, Inc.
 - *Macintosh Technology in the Common Hardware Reference Platform*, by Apple Computer, Inc.
 - *Computer Organization and Design*, by David A. Patterson and John L. Hennessy.
 - *Computer Architecture: A Quantitative Approach*, Second Edition, by John L. Hennessy and David A. Patterson.
- *PowerPC Programming for Intel Programmers*, by Kip McClanahan; IDG Books Worldwide, Inc., 919 East Hillsdale Boulevard, Suite 400, Foster City, CA, 94404; Tel. (800) 434-3422 (U.S.A.), (415) 655-3022 (International).

Chapter 1

Overview

This document defines a programming model for use with the AltiVec instruction set extension to the PowerPC architecture. There are three types of programming interfaces described in this document:

- A high-level language interface, intended for use within programming languages such as C or C++
- An application binary interface (ABI) defining low-level coding conventions
- An assembly language interface

Although a higher-level application programming interface (API) such as `mediaLib` is intended for use with AltiVec, such a specification is not addressed by this document. For further details on `mediaLib` see the AltiVec website at: <http://www.mot.com/SPS/PowerPC/AltiVec>.

An AltiVec-enabled compiler implementing the model described in this document predefines the value `__VEC__` as the decimal integer 10205.

1.1 High-Level Language Interface

The high-level language interface for AltiVec is a way for programmer to be able to use the AltiVec technology from programming languages such as C and C++. It describes fundamental data type for the AltiVec programming model. Details of this interface are described in Chapter 2, “High-Level Language Interface.”



1.2 Application Binary Interface (ABI)

The AltiVec Programming Model extends the existing PowerPC ABIs and the extension is independent of the endian mode. The ABI reviews what the data types are and what the register usage conventions are for vector register files. The ABI also discusses how to set up the stack frame. The vector register save and restore functions are included in the ABI section to advocate uniformity among compilers on the method used in saving and restoring vector registers.

The Programming Interface Manual provides the valid set of argument types for specific AltiVec operations and predicates as well as the specific AltiVec instruction(s) generated for that set of arguments. The AltiVec operations and predicates are organized alphabetically in Chapter 4, “AltiVec Operations and Predicates.”

Chapter 2

High-Level Language Interface

The AltiVec high-level language interface:

- Provides an efficient and expressive mechanism for programmers to access AltiVec functionality from programming languages such as C and C++.

Note: Access to AltiVec functionality from Java applications is not currently addressed by this specification, but will likely be addressed through a higher level API such as `mediaLib`.

- Defines a minimal set of language extensions that clearly describes the intent of the programmer while minimizing the impact on existing PowerPC compilers and development tools.
- Defines a minimal set of library extensions needed to support AltiVec functionality.

2.1 Data Types

The AltiVec programming model introduces a set of fundamental data types, as described in Table 2-1.

Table 2-1. AltiVec Data Types

New C/C++ Type	Interpretation of Contents	Components Represent Values
vector unsigned char	16 unsigned char	0...255
vector signed char	16 signed char	-128...127
vector bool char	16 unsigned char	0(F), 255 (T)
vector unsigned short	8 unsigned short	0...65536
vector unsigned short int		
vector signed short	8 signed short	-32768...32767
vector signed short int		
vector bool short	8 unsigned short	0 (F), 65535 (T)
vector bool short int		
vector unsigned int	4 unsigned int	$0 \dots 2^{32} - 1$
vector unsigned long*		
vector unsigned long int*		

Table 2-1. AltiVec Data Types (Continued)

New C/C++ Type	Interpretation of Contents	Components Represent Values
vector signed int	4 signed int	$-2^{31} \dots 2^{31}-1$
vector signed long*		
vector signed long int*		
vector bool int	4 unsigned int	0 (F), $2^{32} - 1$ (T)
vector bool long*		
vector bool long int*		
vector float	4 float	IEEE-754 values
vector pixel	8 unsigned short	1/5/5/5 pixel

*The vector types with the long keyword are deprecated and will be eliminated in a future version of this document.

In illustrations where an algorithm could apply to multiple types, `vec_data` represents any one of these types. Introducing fundamental types permits the compiler to provide stronger type checking and supports overloaded operations on vector types.

2.2 New Keywords

The model introduces new uses for the following five identifiers:

- `vector`
- `__vector`
- `pixel`
- `__pixel`
- `bool`

as simple type specifier keywords. Among the type specifiers used in a declaration, the `vector` type specifier must occur first. As in C and C++, the remaining type specifiers may be freely intermixed in any order, possibly with other declaration specifiers. The syntax does not allow the use of a `typedef` name as a type specifier. For example, the following is not allowed:

```
typedef signed short int16;
vector int16 data;
```

These new uses may conflict with their existing use in C and C++. There are two methods that may be used to deal with this conflict. An implementation of the AltiVec programming model may choose either method.

2.2.1 The Keyword and Predefine Method

In this method, `__vector`, `__pixel`, and `bool` are added as keywords while `vector` and `pixel` are predefined macros. `bool` is already a keyword in C++. To allow its use in C as a keyword, it is treated the same as it is in C++. This means that the C language is extended to allow `bool` alone as a set of type specifiers. Typically, this type will map to `int`. To

accommodate a conflict with other uses of the identifiers `vector` and `pixel`, the user can either `#undef` or use a command line option to remove the predefines.

2.2.2 The Context Sensitive Keyword Method

In this method, `__vector` and `__pixel` are added as keywords without regard to context while the new uses of `vector`, `pixel`, and `bool` are keywords only in the context of a type. Since `vector` must be first among the type specifiers, it can be recognized as a type specifier when a type identifier is being scanned. The new uses of `pixel` and `bool` occur after `vector` has been recognized. In all other contexts, `vector`, `pixel`, and `bool` are not reserved. This avoids conflicts such as `class vector`, `typedef int bool`, and allows the use of `vector`, `pixel`, and `bool` as identifiers for other uses.

2.3 Alignment

The following paragraphs described AltiVec alignment requirements. When working with vector data, the programmer must be aware of these alignment issues. Because the AltiVec technology does not generate exceptions, the programmer must determine whether and when vector data becomes unaligned.

2.3.1 Alignment of Vector Types

A defined data item of any vector data type in memory is always aligned on a 16-byte boundary. A pointer to any vector data type always points to a 16-byte boundary. The compiler is responsible for aligning vector data types on 16-byte boundaries. Given that vector data is correctly aligned, a program is incorrect if it attempts to dereference a pointer to a vector type if the pointer does not contain a 16-byte aligned address. In the AltiVec architecture, an unaligned load/store does not cause an alignment exception that might lead to (slow) loading of the bytes at the given address. Instead, the low-order bits of the address are quietly ignored.

2.3.2 Alignment of Non-Vector Types

An array of components to be loaded into vector registers need not be aligned, but will have to be accessed with attention to its alignment. Typically, this is accomplished using either the Load Vector for Shift Right, `vec_lvsr()`, or Load Vector for Shift Left, `vec_lvsl()`, operation and the Vector Permute, `vec_perm()`, operation.

2.3.3 Alignment of Aggregates and Unions Containing Vector Types

Aggregates (structures and arrays) and unions containing vector types must be aligned on 16-byte boundaries and their internal organization padded, if necessary, so that each internal vector type is aligned on a 16-byte boundary. This is an extension to all ABIs (AIX, Apple, SVR4, and EABI).

2.4 Extensions of C/C++ Operators for the New Types

Most C/C++ operators do not permit any of their arguments to be one of the new types. Let *a* and *b* be vector types and *p* be a pointer to a vector type. The normal C/C++ operators are extended to include the following operations.

2.4.1 sizeof()

The operations `sizeof(a)` and `sizeof(*p)` return 16.

2.4.2 Assignment

If either the left hand side or right hand side of an expression has a vector type, then both sides of the expression must be of the same vector type. Thus, the expression `a = b` is valid and represents assignment if *a* and *b* are of the same vector type (or if neither is a vector type). Otherwise, the expression is invalid and must be signaled as an error by the compiler.

2.4.3 Address Operator

The operation `&a` is valid if *a* is a vector type. The result of the operation is a pointer to *a*.

2.4.4 Pointer Arithmetic

The usual pointer arithmetic can be performed on *p*. In particular, `p+1` is a pointer to the next vector after *p*.

2.4.5 Pointer Dereferencing

If *p* is a pointer to a vector type, `*p` implies either a 128-bit vector load from the address obtained by clearing the low order bits of *p*, equivalent to the instruction `vec_ld(0, p)` or a 128-bit vector store to that address equivalent to the instruction `vec_st(0, p)`. If it is desired to mark the data accessed as least-recently-used (LRU), the explicit instruction `vec_ldl(0,p)` or `vec_stl(0, p)` must be used.

Dereferencing a pointer to a non-vector type produces the standard behavior of either a load or a copy of the corresponding type.

Accessing of unaligned memory must be carried out explicitly by a `vec_ld(int, type *)` operation, a `vec_ldl(int, type *)` operation, a `vec_st(int, type *)` operation or a `vec_stl(int, type *)` operation.

2.4.6 Type Casting

Pointers to old and new types may be cast back and forth to each other. Casting a pointer to a new type represents an unchecked assertion that the address is 16-byte aligned. Some new operators are provided to provide the equivalence of casts and data initialization.

Casts from one vector type to another are provided by normal C casts. These should not be needed frequently if the overloaded forms of operators are used. None of the casts performs a conversion; the bit pattern of the result is the same as the bit pattern of the argument that is cast.

- `(vector signed char) vec_data`
- `(vector signed short) vec_data`
- `(vector signed int) vec_data`
- `(vector unsigned char) vec_data`
- `(vector unsigned short) vec_data`
- `(vector unsigned int) vec_data`
- `(vector bool char) vec_data`
- `(vector bool short) vec_data`
- `(vector bool int) vec_data`
- `(vector float) vec_data`
- `(vector pixel) vec_data`

Casts between vector types and scalar types are illegal. To copy data between these types, use the `vec_lde()` or `vec_ste()` operations. An alternative is to use a union consisting of a vector type and an equivalent array of the scalar type and copy the data using the union.

2.5 New Operators

New operators are introduced to construct vector literals, adjust pointers, and allow full access to the functionality provided by the AltiVec architecture.

2.5.1 Vector Literals

A vector literal is written as a parenthesized vector type followed by a parenthesized set of constant expressions. Vector literals may be used either in initialization statements or as constants in executable statements. Table 2-2 lists the formats and descriptions of the vector literals. For each, the compiler generates code that either computes or loads the values into the register.

Table 2-2. Vector Literal Format and Description

Notation	Represents
<code>(vector unsigned char) (unsigned int)</code>	A set of 16 unsigned 8-bit quantities which all have the value specified by the integer.
<code>(vector unsigned char) (unsigned int, ..., unsigned int)</code>	A set of 16 unsigned 8-bit quantities specified by the 16 integers.
<code>(vector signed char) (int)</code>	A set of 16 signed 8-bit quantities that all have the value specified by the integer.
<code>(vector signed char) (int, ..., int)</code>	A set of 16 signed 8-bit quantities specified by the 16 integers.
<code>(vector unsigned short) (unsigned int)</code>	A set of eight unsigned 16-bit quantities which all have the value specified by the unsigned integer.
<code>(vector unsigned short) (unsigned int, ..., unsigned int)</code>	A set of eight unsigned 16-bit quantities specified by the eight unsigned integers.
<code>(vector signed short) (int)</code>	A set of eight signed 16-bit quantities which all have the value specified by the integer.
<code>(vector signed short) (int, ..., int)</code>	A set of eight signed 16-bit quantities specified by the eight integers.
<code>(vector unsigned int) (unsigned int)</code>	A set of four unsigned 32-bit quantities which all have the value specified by the unsigned integer.
<code>(vector unsigned int) (unsigned int, ..., unsigned int)</code>	A set of four unsigned 32-bit quantities specified by the four unsigned integers.
<code>(vector signed int) (int)</code>	A set of four signed 32-bit quantities which all have the value specified by the integer.
<code>(vector signed int) (int, ..., int)</code>	A set of four signed 32-bit quantities specified by the 4 integers.
<code>(vector float) (float)</code>	A set of four floating-point quantities which all have the value specified by the floating-point value.
<code>(vector float) (float, ..., float)</code>	A set of four floating-point quantities which all have the value specified by the four floating-point values.

2.5.2 Vector Literals and Casts

The combination of vector casts and vector literals can complicate some parsers. An implementation is not required to support the cast to a vector type of a vector cast or vector literal when the operand of the cast is not a parenthesized expression. For example, the programmer may write the following:

```
(vector unsigned char)((vector unsigned int)(1, 2, 3, 4))
(vector signed char)((vector unsigned short) variable)
```

The similar expressions below without the parenthesized expression may not be used in a conforming application

```
(vector unsigned char)(vector unsigned int)(1, 2, 3, 4)
(vector signed char)(vector unsigned short) variable
```

2.5.3 Value for Adjusting Pointers

At compile time, the `vec_step(vec_data)` produces the integer value representing the amount by which a pointer to a component of an AltiVec data should increment to cause a pointer increment to increment by 16 bytes. For example, a vector unsigned short data type is considered to contain eight unsigned 2-byte values. A pointer to unsigned 2-byte values used to stream through an array of unsigned 2-byte values by a full vector at a time should increment by `vec_step(vector unsigned short) = 8`. Table 2-3 provides a summary of the values by data type.

Table 2-3. Increment Value for `vec_step` by Data Type

vec_step Expression	Value
<code>vec_step(vector unsigned char)</code> <code>vec_step(vector signed char)</code> <code>vec_step(vector bool char)</code>	16
<code>vec_step(vector unsigned short)</code> <code>vec_step(vector signed short)</code> <code>vec_step(vector bool short)</code>	8
<code>vec_step(vector unsigned int)</code> <code>vec_step(vector signed int)</code> <code>vec_step(vector bool int)</code>	4
<code>vec_step(vector pixel)</code>	8
<code>vec_step(vector float)</code>	4

2.5.4 New Operators Representing AltiVec Operations

New operators are introduced to allow full access to the functionality provided by the AltiVec architecture. The new operators are represented in the programming language by language structures that parse like function calls. The names associated with these operations are all prefixed with `vec_`. The appearance of one of these forms can indicate the following:

- A generic AltiVec operation, like `vec_add()`
- A specific AltiVec operation, like `vec_addubm()`
- A predicate computed from a AltiVec operation like `vec_all_eq()`
- Loading of a vector of components, as discussed in Section 2.5.1, “Vector Literals”

Each AltiVec operator takes a list of arguments that represent the input operands. The order of the operands is prescribed in the architecture specification and includes a returned result (possibly `void`).

The programming model restricts the operand types permitted for each AltiVec operation, whether specific or generic. The programmer may override this constraint by explicitly casting arguments to permissible types.

For a specific operation, the operand types determine whether the operation is acceptable within the programming model and the type of the result. For example, `vec_vaddubm(vector signed char, vector signed char)` is acceptable in the programming model because it represents a reasonable way to do modular addition with signed bytes, while `vec_vaddubs(vector signed char, vector signed char)` and `vec_vaddubh(vector signed char, vector signed char)` are not acceptable. If permitted, the former operation would produce a result in which saturation treats the operands as unsigned; the latter operation would produce a result in which adjacent pairs of signed bytes are treated as signed halfwords.

For a generic operation, the operand types are used to determine whether the operation is acceptable, to select a particular operation according to the types of the arguments, and to determine the type of the result. For example, `vec_add(vector signed char, vector signed char)` will map onto `vec_vaddubm()` and return a result of type `vector signed char`, while `vec_add(vector unsigned short, vector unsigned short)` maps onto `vec_vadduhm()` and return a result of type `vector unsigned short`.

The AltiVec operations that set condition register CR6 (i.e., the compare dot instructions) are treated somewhat differently in the programming model. The programmer can not access specific register names. Instead of directly specifying a compare dot instruction, the programmer makes reference to a predicate that returns an integer value derived from the result of a compare dot instruction. As in C, this value may be used directly as a value (1 is true, 0 is false) or as a condition for branching. It is expected that the compiler will produce the minimum code needed to use the condition. Predicates begin with `vec_all_` or `vec_any_`. Either the true or false state of any bit that can be set by a compare dot instruction has a predicate. For example, `vec_all_gt(x,y)` tests the true value of bit 24 of the CR after executing some `vcmpgt.` instruction. To complete the coverage by predicates, additional predicates exercise compare dot instructions with reversed or duplicated arguments. As examples, `vec_all_lt(x,y)` performs a `vcmpgtx.(y,x)`, and `vec_all_nan(x)` is mapped onto `vcmpeqfp.(x,x)`. If the programmer wishes to have both the result of the compare dot instruction as returned in the vector register and the value of CR6, the programmer specifies two operations. The compiler's job is to determine that these can be merged. The AltiVec operations and predicates are listed in Chapter 4, "AltiVec Operations and Predicates".

2.6 Programming Interface

This document does not prohibit or require an implementation to provide any set of `include` files or `#pragma` preprocessor commands. If an implementation requires that an `include` file be used prior to the use of the syntax described in this document, it is suggested that the `include` file be named `<altivec.h>`. If an implementation supports `#pragma` preprocessor commands, it is suggested that it provide `__ALTIVEC__` as a predefined macro with a nonzero value. A suggested preprocessor command set includes the following:

```
#pragma altivec_codegen on | off
```

When this pragma is on, the compiler may use AltiVec instructions. When you set this pragma off, the `altivec_model` pragma is also set to off.

```
#pragma altivec_model on | off
```

When this pragma is on, the compiler accepts the syntax specified in this document, and the `altivec_codegen` pragma is also set to on.

```
#pragma altivec_vrsave on | off | allon
```

When this pragma is on, the compiler maintains the VRSAVE register. With `allon` selected, the compiler changes the VRSAVE register to have all bits set. It is combined with `#pragma altivec_vrsave off` by having a parent function do the work once of setting the value of the VRSAVE register with `#pragma altivec_vrsave allon` and the function it calls uses the setting `#pragma altivec_vrsave off`.



Chapter 3

Application Binary Interface (ABI)

Note: The ABI extensions described herein for embedded applications are still under review by the PowerPC EABI industry working group, and may be subject to change. Modifications, if any, will be highlighted in future revisions of this document.

The AltiVec programming model extends the existing PowerPC ABIs. This chapter specifies extensions to the System V Application Binary Interface PowerPC Processor Supplement (SVR4 ABI), the PowerPC Embedded Application Binary Interface (EABI), Appendix A of The PowerPC Compiler Writer's Guide (AIX ABI), and the Apple Macintosh ABI. The SVR4 ABI and EABI specifications define both a Big-Endian ABI and a Little-Endian ABI. This extension is independent of the endian mode.

3.1 Data Representation

The vector data types are 16-bytes long and 16-byte aligned. All ABIs are extended similarly. Aggregates (structures and arrays) and unions containing vector types must be aligned on 16-byte boundaries and their internal organization padded, if necessary, so that each internal vector type is aligned on a 16-byte boundary. The Apple ABI and AIX ABI specify a maximum alignment for aggregates and unions of 4-bytes; the EABI specifies a maximum alignment of 8-bytes. Increasing the alignment to 16-bytes creates the opportunity for padding or holes in the parameter lists involving these aggregates described in Section 3.4.2, "Apple Macintosh ABI and AIX ABI Parameter Passing without Varargs."

3.2 Register Usage Conventions

The register usage conventions for the vector register file are defined as follows:

Table 3-1. AltiVec Registers

Register	Intended use	Behavior across call sites
v0–v1	General use	Volatile (Caller save)
v2–v13	Parameters, general	Volatile (Caller save)
v14–v19	General	Volatile (Caller save)
v20–v31	General	Non-volatile (Callee save)

Table 3-1. AltiVec Registers

Register	Intended use	Behavior across call sites
VRSAVE	Special, see Section 3.3, "The Stack Frame"	Non-volatile (Callee save)

The VRSAVE special purpose register (SPR256, named `vrsave` in assembly instructions) is used to inform the operating system which vector registers (VRs) need to be saved and reloaded across context switches. Bit n of this register is set to 1 if vector register vn needs to be saved and restored across a context switch. Otherwise, the operating system may return that register with any value that does not violate security after a context switch. The most significant bit in the 32-bit word is bit 0.

The EABI does not use VRSAVE for any special purpose, but VRSAVE is a non-volatile register.

3.3 The Stack Frame

The stack pointer maintains 16-byte alignment in the SVR4 ABI and the AIX ABI and 8-byte alignment in the EABI and the Apple Macintosh ABI and AIX ABI. It is not necessary to align the stack dynamically in either the SVR4 ABI or the AIX ABI, however, the alignment padding space is specified for both. The additions to the stack frame are the vector register save area, the `vrsave` word, and the alignment padding space to dynamically align the stack to a quadword boundary.

The following additional requirements apply to the stack frame:

- Before a function changes the value of `vrsave`, it shall save the value of VRSAVE at the time of entry to the function in the `vrsave` word.
- The alignment padding space shall be either 0, 4, 8, or 12 bytes long so that the address of the vector register save area (and subsequent stack locations) are quadword aligned.
- If the code establishing the stack frame dynamically aligns the stack pointer, it shall update the stack pointer atomically with an `stwux` instruction. The code may assume the stack pointer on entry is aligned on an 8-byte boundary.
- Before a function changes the value in any non-volatile vector register, vn , it shall save the value in vn in the word in the vector register save area $16 \cdot (32 - n)$ bytes before the low-addressed end of the alignment padding space.
- Local variables of a vector data type which need to be saved to memory will be placed on the stack frame on a 16-byte alignment boundary in the same stack frame region used for local variables of other types.

SP in the figures denotes the stack pointer (general purpose register `r1`) of the called function after it has executed code establishing its stack frame.

3.3.1 SVR4 ABI and EABI Stack Frame

The size of the vector register save area and the presence of the VRSAVE word may vary within a function and are determined by a new registers valid tag. Note: In the SVR4 ABI, the registers valid tag is the most general way to describe a stack frame. It is associated with a frame or frame valid tag. Figure 3-1 shows an SVR4 and EABI stack frame.

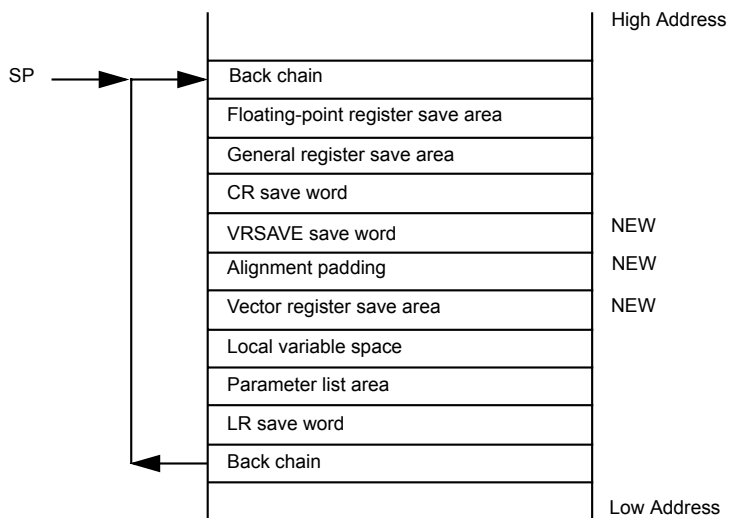


Figure 3-1. SVR4 ABI and EABI Stack Frame

Table 3-2. Vector Registers Valid Tag Format

Word	Bits	Name	Description
1	0–17	RESERVED	0
1	18–29	START_OFFSET	The number of words between the BASE of the nearest preceding Frame or Frame Valid tag and the first instruction to which this tag applies.
1	30–31	TYPE	2
2	0–11	VECTOR_REGS	One bit for each non-volatile vector register, bit 0 for v31,..., bit 11 for v20, with a 1 signifying that the register is saved in the vector register save area.
2	12	VRSAVE_AREA ¹	1 if and only if the VRSAVE word is allocated in the register save area.
1.If more than one Vector Registers Valid Tag applies to the same Frame or Frame Valid tag, they shall all have the same values for VRSAVE_AREA and VR.			

Table 3-2. Vector Registers Valid Tag Format

Word	Bits	Name	Description
2	13-17	VR ¹	Size in quadwords of the vector register save area.
2	18-29	RANGE	The number of words between the first and the last instruction to which this tag applies.
2	30	VRSAVE_REG	1 if and only if VRSAVE is saved in the VRSAVE word.
2	31	SUBTYPE	1
1. If more than one Vector Registers Valid Tag applies to the same Frame or Frame Valid tag, they shall all have the same values for VRSAVE_AREA and VR.			

The code example below shows sample prologue and epilogue code with full saves of all the non-volatile floating-point (FPRs), general (GPRs), and VRs for a stack frame of less than 32 Kbytes. The example aligns the stack pointer dynamically, addresses incoming arguments via r30, uses volatile VRs v0–v10, maintains VRSAVE, does not alter the non-volatile fields of the CR and does no dynamic stack allocation. Saving and restoring the VRs and updating vrsave can occur in either order. A function that does not need to address incoming arguments but does align the stack pointer dynamically can recover the address of the original stack pointer with an instruction such as `lwz r11,0(sp)`. The computation of `len` in the example and whether to use `subfic` or `addi` to align the stack dynamically is based on the size of the components of the frame. Starting with the components at higher addresses, the value of `len` is computed by adding the size of the FPR save area, the GPR save area, the CR save word, and the VRSAVE word.

The size of the alignment padding space is then computed as the smallest number of bytes needed to make `len` a multiple of 16. In the example below, the alignment padding space is 4 bytes. Consequently, `subfic` is used to dynamically align the stack by increasing the size of the alignment padding space by either 0 or 8 bytes. Had the alignment padding space been 8 or 12 bytes, `addi` would be used to align the stack dynamically by decreasing the size of the alignment padding space by either 0 or 8 bytes. Continuing, the value of `len` is updated by adding the size of the vector register save area, the local variable space, the outgoing parameter list area, and the LR save word. The size of the local variable space is adjusted so that the overall value of `len` is a multiple of 16. The following is SVR4 ABI and EABI prologue and epilogue sample code.

```

function: mflr      r0                # Save return address ...
          stw        r0,4(sp)         # ... in caller's frame.
          ori        r11,sp,0        # Save end of fpr save area
          rlwinm     r12,sp,0,28,28  # 0 or 8 based on SP alignment
          subfic     r12,r12,-len    # Add in stack length
          stwux      sp,sp,r12       # Establish new aligned frame
          bl         _savefpr_14     # Save floating-point registers
          addi       r11,r11,-144    # Compute end of gpr save area
          bl         _savegpr_14_g   # Save gprs and fetch GOT ptr
          mflr      r31              # Place GOT ptr in r31
          # Save CR here if necessary
          addi       r30,r11,144     # Save pointer to incoming

```

```

# arguments
mfspr    r0,vrsave      # Save VRSAVE ...
stw      r0,-220(r30)   # ... in caller's frame.
oris     r0,r0,0xff70   # Use v0-v10 and ...
ori      r0,r0,0x0fff   # v20-v31 (for example)
mtspr    vrsave,r0      # Update VRSAVE
addi     r0,sp,len-224  # Compute end of vr save area
bl       _savevr20      # Save VRs
# Body of function
addi     r0,sp,len-224  # Address of vr save area to r0
bl       _restvr20      # Restore VRs
lwz      r0,-220(r30)   # Fetch prior value of VRSAVE
mtspr    vrsave,r0      # Restore VRSAVE
addi     r11,r30,-144   # Address of gpr save area to r11
bl       _restgpr_14    # Restore gprs
addi     r11,r11,144    # Address of fpr save area to r11
bl       _restfpr_14_x  # Restore fprs and return

```

3.3.2 Apple Macintosh ABI and AIX ABI Stack Frame

Figure 3-2 shows how the Apple Macintosh ABI and AIX ABI stack frame is set up.

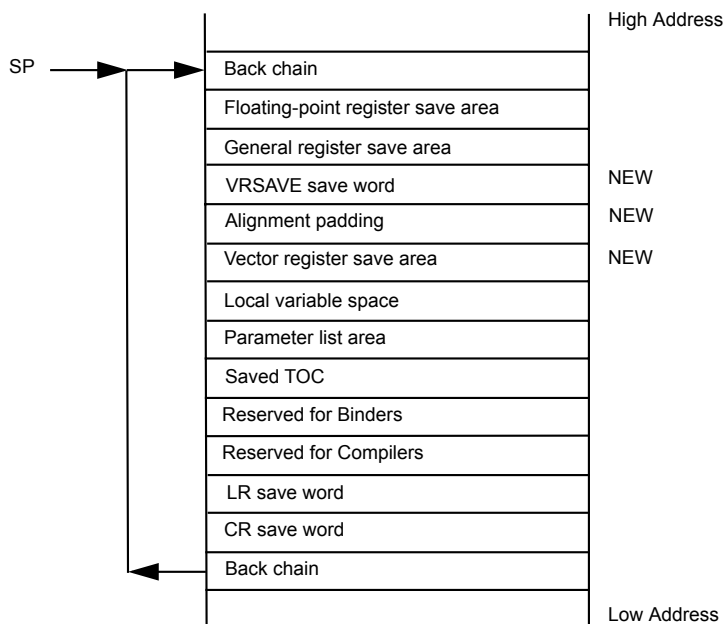


Figure 3-2. Apple Macintosh ABI and AIX ABI Stack Frame

The Apple Macintosh ABI and AIX ABI stack frame allow the use of a 220-byte area at a negative offset from the stack pointer. This area can be used to save non-volatile registers before the stack pointer has been updated. This size of this area is not changed. Depending

on the number of non-volatile registers saved, it may be necessary to update the stack pointer before saving the VRs. However, it remains unnecessary to update the stack pointer before saving the GPRs or FPRs.

The size of the VR save area and the presence of the VRSAVE word are determined by a traceback table entry. The `spare3` 2-bit field in the fixed portion of the traceback table is changed to the following:

`has_vec_info` This 1-bit field is set if the procedure saves non-volatile VRs in the vector register save area, saves `vrsave` in the VRSAVE word, specifies the number of vector parameters, or uses Altivec instructions.

`spare4` Reserved 1-bit field.

When the `has_vec_info` bit is set, all the following optional fields of the traceback table are present following the position of the `alloca_reg` field.

`vr_saved` This 6-bit field represents the number of non-volatile VRs saved by this procedure. Because the last register saved is always `v31`, a value of 2 in `vr_saved` indicates that `v30` and `v31` are saved.

`saves_vrsave` If this routine saves `vrsave`, this 1-bit field is set. If so, the VRSAVE word in the register save area must be used to restore the prior value before returning from this procedure.

`has_varargs` If this function has a variable argument list, this 1-bit field is set. Otherwise, it is set to 0.

`vectorparms` This 7-bit field records the number of vector parameters. The field may be set to a non-zero value for a procedure with vector parameters that does not have a variable argument list. Otherwise, `parmsonstk` must be set.

`vec_present` This 1-bit field is set if Altivec instructions are performed within the procedure.

The following code shows sample prologue and epilogue code with full saves of all the non-volatile floating-point, general, and VRs for a stack frame of less than 32 Kbytes. The code example dynamically aligns the stack pointer, addresses incoming arguments via `r31`, uses volatile VRs `v0–v10`, maintains VRSAVE, does not alter the non-volatile fields of the CR and does no dynamic stack allocation. Saving and restoring the VRs and updating the `vrsave` register can occur in either order. A function that does not need to address incoming arguments but does align the stack pointer dynamically can recover the address of the original stack pointer with an instruction such as `lwz r11,0(sp)`.

The computation of `len` in the example and whether to use `subfic` or `addi` to align the stack dynamically are based on the size of the components of the frame. Starting with the components at higher addresses, the value of `len` is computed by adding the size of the floating-point register save area, the general register save area, and the VRSAVE word. The size of the alignment padding space is then computed as the smallest number of bytes

needed to make `len` a multiple of 16. In the example below, the alignment padding space is 0 bytes. Consequently, `subfic` is used to align the stack dynamically by increasing the size of the alignment padding space by either 0 or 8 bytes. Had the alignment padding space been 8 or 12 bytes, `addi` is used to align the stack dynamically by decreasing the size of the alignment padding space by either 0 or 8 bytes. Continuing, the value of `len` is updated by adding the size of the vector register save area, the local variable space, the outgoing parameter list area, and 24 for the size of the link area. The size of the local variable space is adjusted so that the overall value of `len` is a multiple of 16.

The following is Apple Macintosh ABI and AIX ABI prologue and epilogue sample code.

```
function: mflr      r0          # Save return address ...
          stw       r0,8(sp)    # ... in the caller's frame.
          bl       _savef14     # Save floating-point registers.
          stmw      r13,-220(sp) # Save gprs in gpr save area
                                   # Save CR here if necessary
          ori       r31,sp,0     # Save pointer to incoming
                                   # arguments
          rlwinm    r12,sp,0,28,28 # 0 or 8 based on SP alignment
          subfic    r12,r12,-len  # Add in stack length
          stwux     sp,sp,r12     # Establish new aligned frame
          mfspr     r0,vrsave     # Save VRSAVE ...
          stw       r0,-224(r31)  # ... in caller's frame.
          oris      r0,r0,0xff70  # Use v0-v10 v20-v31 and ...
          ori       r0,r0,0x0fff  # v20-v31 (for example)
          mtspr     vrsave,r0     # Update VRSAVE
          addi      r0,sp,len-224  # Compute end of VRSAVE area
          bl       _savev20      # Save VRs
                                   # Body of function
          addi      r0,sp,len-224  # Address of VRSAVE area to r0
          bl       _restv20      # Restore VRs
          lwz       r0,-224(r31)  # Fetch prior value of VRSAVE
          mtspr     vrsave,r0     # Restore Vrsave
          ori       sp,r31       # Restore SP
          lmw       r13,-220(sp)  # Restore gprs
          lwz       r0,8(sp)     # Restore return address ...
          mtlr      r0          # ... and return from _restf14
          b        _restf14      # Restore fprs and return
```

3.3.3 Vector Register Saving and Restoring Functions

The vector register saving and restoring functions described in this section are not part of the ABI. They are defined here only to encourage uniformity among compilers in the code used to save and restore VRs.

On entry to the functions described in this section, `r0` contains the address of the word just beyond the end of the vector register save area, and they leave `r0` undisturbed. They modify the value of `r12`. The following code is an example of saving a vector register.

```
_savev20: addi      r12,r0,-192
          stvx      v20,r12,r0    # save v20
_savev21: addi      r12,r0,-176
```

```

                stvx          v21,r12,r0          # save v21
_savev22:      addi          r12,r0,-160
                stvx          v22,r12,r0          # save v22
_savev23:      addi          r12,r0,-144
                stvx          v23,r12,r0          # save v23
_savev24:      addi          r12,r0,-128
                stvx          v24,r12,r0          # save v24
_savev25:      addi          r12,r0,-112
                stvx          v25,r12,r0          # save v25
_savev26:      addi          r12,r0,-96
                stvx          v26,r12,r0          # save v26
_savev27:      addi          r12,r0,-80
                stvx          v27,r12,r0          # save v27
_savev28:      addi          r12,r0,-64
                stvx          v28,r12,r0          # save v28
_savev29:      addi          r12,r0,-48
                stvx          v29,r12,r0          # save v29
_savev30:      addi          r12,r0,-32
                stvx          v30,r12,r0          # save v30
_savev31:      addi          r12,r0,-16
                stvx          v31,r12,r0          # save v31
                blr              # return to prologue

```

The following code shows how to restore a vector register.

```

_restv20:      addi          r12,r0,-192
                lvx           v20,r12,r0          # restore v20
_restv21:      addi          r12,r0,-176
                lvx           v21,r12,r0          # restore v21
_restv22:      addi          r12,r0,-160
                lvx           v22,r12,r0          # restore v22
_restv23:      addi          r12,r0,-144
                lvx           v23,r12,r0          # restore v23
_restv24:      addi          r12,r0,-128
                lvx           v24,r12,r0          # restore v24
_restv25:      addi          r12,r0,-112
                lvx           v25,r12,r0          # restore v25
_restv26:      addi          r12,r0,-96
                lvx           v26,r12,r0          # restore v26
_restv27:      addi          r12,r0,-80
                lvx           v27,r12,r0          # restore v27
_restv28:      addi          r12,r0,-64
                lvx           v28,r12,r0          # restore v28
_restv29:      addi          r12,r0,-48
                lvx           v29,r12,r0          # restore v29
_restv30:      addi          r12,r0,-32
                lvx           v30,r12,r0          # restore v30
_restv31:      addi          r12,r0,-16
                lvx           v31,r12,r0          # restore v31
                blr              # return to prologue

```

3.4 Function Calls

This section applies to all user functions. Note that the intrinsic AltiVec operations are not treated as function calls, so these comments don't apply to those operations.

The first twelve vector parameters are placed in VRs v2–v13. If fewer (or no) vector type arguments are passed, the unneeded registers are not loaded and contain undefined values upon entry to the called function.

Functions that declare a vector data type as a return value place that return value in v2.

Any function that returns a vector type or has a vector parameter requires a prototype. This requirement enables the compiler to avoid shadowing VRs in GPRs.

3.4.1 SVR4 ABI and EABI Parameter Passing and Varargs

The SVR4 ABI algorithm for passing parameters considers the arguments as ordered from left (first argument) to right, although the order of evaluation of the arguments is unspecified. The vector arguments maintain their ordering. The algorithm is modified to add `vr` to contain the number of the next available vector register. In the INITIALIZE step, set `vr=2`. In the SCAN loop, add a case for the next argument `VECTOR_ARG` as follows:

- If the next argument is in the variable portion of a parameter list, set `vr=14`. This leaves the fixed portion of a variable argument list in VRs and places the variable portion in memory.
- If `vr>13` (that is, there are no more available VRs), go to OTHER. Otherwise, load the argument value into vector register `vr`, set `vr` to `vr+1`, and go to SCAN.

The OTHER case is modified only to understand that vector arguments have 16-byte size and alignment.

Aggregates are passed by reference (i.e., converted to a pointer to the object), so no change is needed to deal with 16-byte aligned aggregates.

The `va_list` type is unchanged, but an additional `va_arg_type` value of 4 named `arg_VECTOR` is defined for the `__va_arg()` interface. Since vector parameters in the variable portion of a parameter list are passed in memory, the `__va_arg()` routine can access the vector value from the `overflow_arg_area` value in the `va_list` type.

3.4.2 Apple Macintosh ABI and AIX ABI Parameter Passing without Varargs

If the function does not take a variable argument list, the non-vector parameters are passed in the same registers and stack locations as they would be if the vector parameters were not present. The only change is that aggregates and unions may be 16-byte aligned instead of 4-byte aligned. This can result in words in the parameter list being skipped for alignment (padding) and left with undefined value.

The first twelve vector parameters are placed in v2–v13. These parameters are not shadowed in GPRs. They are not allocated space in the memory argument list. Any additional vector parameters are passed through memory on the program stack. They appear together, 16-byte aligned, and after any non-vector parameters.

3.4.3 Apple Macintosh ABI and AIX ABI Parameter Passing with Varargs

The `va_list` type continues to be a pointer to the memory location of the next parameter. If `va_arg()` accesses a vector type, the `va_list` value must first be aligned to a 16-byte boundary.

A function that takes a variable argument list has all parameters, including vector parameters, mapped in the argument area as ordered and aligned according to their type. The first 8 words of the argument area are shadowed in the GPRs only if they correspond to the variable portion of the parameter list. The first parameter word is named PW0 and is at stack offset 0x24. A vector parameter must be aligned on a 16-byte boundary. This means there are two cases where vector parameters are passed in GPRs. If a vector parameter is passed in PW2:PW5 (stack offset 0x32), its value is placed in GPR5–GPR8. If a vector parameter is passed in PW6:PW9 (stack offset 0x48), its value PW6:PW7 is placed in GPR9 and GPR10 and the value PW8:PW9 is placed on the stack. All parameters after the first 8 words of the argument area that correspond to the variable portion of the parameter list are passed in memory.

In the fixed portion of the parameter list, vector parameters are placed in v2–v13, but are provided a stack location corresponding to their position in the parameter list.

3.5 `malloc()`, `vec_malloc()`, and `new`

In the interest of saving space, `malloc()`, `calloc()`, and `realloc()` are not required to return a 16-byte aligned address. Instead, a new set of memory management functions is introduced that return a 16-byte aligned address. The new functions are named `vec_malloc()`, `vec_calloc()`, `vec_realloc()`, and `vec_free()`. The two sets of memory management functions may not be interchanged: memory allocated with `malloc()`, `calloc()`, or `realloc()` may only be freed with `free()` and reallocated with `realloc()`; memory allocated with `vec_malloc()`, `vec_calloc()`, or `vec_realloc()` may only be freed with `vec_free()` and reallocated with `vec_realloc()`.

The user must use the appropriate set of functions based on the alignment requirement of the type involved. In the case of the C++ operator `new`, the implementation of `new` is required to use the appropriate set of functions based on the alignment requirement of the type.

3.6 setjmp() and longjmp()

The context required to be saved and restored by `setjmp()`, `longjmp()`, and related functions now includes the 12 non-volatile VRs and `vrsave`. The user types `sigjmp_buf` and `jmp_buf` are extended by 48 words. An unused word in the existing `jmp_buf` is used to save `VRSAVE`.

Table 3-3. ABI Specifications for `setjmp()` and `longjmp()`

ABI	jmp_buf Size	VRSAVE Offset	v20–v31 Offset
AIX ABI	448	100	256
Apple Macintosh ABI	448	16	256
SVR4 ABI and EABI	448	248	256

There are complications in implementing `setjmp()` and `longjmp()`:

- The user types must be enlarged. Existing applications that use these interfaces will have to be recompiled even though they make no use of the AltiVec instruction set.
- The implementation that saves and restores the VRs can only assume that the v20–v31 offset is aligned on a 4-byte boundary. A method where the VRs are saved at the first aligned location in the `jmp_buf` was rejected because the user types are only 4-byte aligned and may be copied by value to a location with different alignment.
- The implementation that saves and restores the VRs and `vrsave` uses instructions that do not exist on a non-AltiVec enabled PowerPC implementation. The method for testing whether the AltiVec instructions operate is privileged. One solution is to define an O/S interface that saves and restores the VRs and `vrsave` if and only if the AltiVec instructions exist and are enabled.

A simple solution to these complications is to define `setjmp()`, `longjmp()` and the user types `sigjmp_buf` and `jmp_buf` differently when compiled with an AltiVec-enabled compiler (i.e., when `__VEC__` is defined). These bindings result in a larger `jmp_buf` with 16-byte alignment. The bindings for `setjmp()` and `longjmp()` unconditionally save and restore the vector state. Such an implementation does not save and restore the vector state when these interfaces are compiled without an AltiVec-enabled compiler. The application must ensure that these two sets of bindings are not mixed.

3.7 Debugging Information

Extensions to the debugging information format are required to describe vector types and vector register locations. While vector types can be described as fixed length arrays of existing C types, the implementation should describe these as new fundamental types. Doing so allows a debugger to provide mechanisms to display vector values, assign vector values, and create vector literals.

This section is subject to change. It is intended to describe the extensions to the standard debugging formats: xcoff stabstrings, DWARF version 1.1.0, and DWARF version 2.0.0.

Xcoff stabstrings used in the AIX ABI and adopted by the Apple Macintosh ABI support the location of objects in GPRs and FPRs. The stabstring code “R” describes a parameter passed by value in the given GPR; “r” describes a local variable residing in the given GPR. The stabstring code “X” describes a parameter passed by value in the given vector register; “x” describes a local variable residing in the given vector register.

DWARF 2.0 debugging DIEs support the location of objects in any machine register. The SVR4 ABI specifies the DWARF register number mapping. The VRs v0–v31 are assigned register numbers 1124–1155. The VRSAVE SPR is SPR256 and is assigned the register number 356.

3.8 printf() and scanf() Control Strings

The conversion specifications in control strings for input functions (fscanf, scanf, sscanf) and output functions (fprintf, printf, sprintf, vsprintf, vprintf, vsprintf) are extended to support vector types.

3.8.1 Output Conversion Specifications

The output conversion specifications have the following general form:

```
%[<flags>][<width>][<precision>][<size>]<conversion>
```

where,

<flags>	::=<flag-char> <flags><flag-char>
<flag-char>	::=<std-flag-char> <c-sep>
<std-flag-char>	::= '−' '+' '0' '#' ' ' '−'
<c-sep>	::= ',' ';' ':' '_'
<width>	::= <decimal-integer> '*'
<precision>	::= '.' <width>
<size>	::= 'll' 'L' 'l' 'h' <vector-size>
<vector-size>	::= 'vl' 'vh' 'lv' 'hv' 'v'
<conversion>	::= <char-conv> <str-conv> <fp-conv> <int-conv> <misc-conv>
<char-conv>	::= 'c'
<str-conv>	::= 's' 'P'
<fp-conv>	::= 'e' 'E' 'f' 'g' 'G'
<int-conv>	::= 'd' 'i' 'u' 'o' 'p' 'x' 'X'
<misc-conv>	::= 'n' '%'

The extensions to the output conversion specification for vector types are shown in bold.

The <vector-size> indicates that a single vector value is to be converted. The vector value is displayed in the following general form:

```
value1 C value2 C ... C valuen
```

where C is a separator character defined by <c-sep> and there are 4, 8, or 16 output values depending on the <vector-size> each formatted according to the <conversion>, as follows:

- A <vector-size> of 'vl' or 'lv' consumes one argument and modifies the <int-conv> conversion; it should be of type `vector signed int`, `vector unsigned int`, or `vector bool int`; it is treated as a series of four 4-byte components.
- A <vector-size> of 'vh' or 'hv' consumes one argument and modifies the <int-conv> conversion; it should be of type `vector signed short`, `vector unsigned short`, `vector bool short`, or `vector pixel`; it is treated as a series of eight 2-byte components.
- A <vector-size> of 'v' with <int-conv> or <char-conv> consumes one argument; it should be of type `vector signed char`, `vector unsigned char`, or `vector bool char`; it is treated as a series of sixteen 1-byte components.
- A <vector-size> of 'v' with <fp-conv> consumes one argument; it should be of type `vector float`; it is treated as a series of four 4-byte floating-point components.
- All other combinations of <vector-size> and <conversion> are undefined.

The default value for the separator character is a space unless the 'c' conversion is being used. For the 'c' conversion the default separator character is null. Only one separator character may be specified in <flags>.

Examples:

```
vector signed char s8 = vector signed char('a','b',' ','d','e','f',
                                           'g','h','i','j','k','l',
                                           'm','','','o','p');
vector unsigned short u16 = vector unsigned short(1,2,3,4,5,6,7,8);
vector signed int s32 = vector signed int(1, 2, 3, 99);
vector float f32 = vector float(1.1, 2.2, 3.3, 4.39501);
printf("s8 = %vc\n", s8);
printf("s8 = %,vc\n", s8);
printf("u16 = %vhu\n", u16);
printf("s32 = %,2lvd\n", s32);
printf("f32 = %,5.2vf\n", f32);
```

This code produces the following output:

```
s8 = ab defghijklm,op
s8 = a,b, ,d,e,f,g,h,i,j,k,l,m,,,o,p
u16 = 1 2 3 4 5 6 7 8
s32 = 1, 2, 3,99
f32 = 1.10 ,2.20 ,3.30 ,4.40
```

3.8.2 Input Conversion Specifications

The input conversion specifications have the following general form:

```
%[<flags>][<width>][<size>]<conversion>
```

where,

```
<flags>      ::=      '*' | <b>c-sep</b> ['*'] | ['*'] <b>c-sep</b>
<c-sep>      ::=      ', ' | ';' | ':' | '_'
<width>      ::=      <decimal-integer>
<size>       ::=      'll' | 'L' | 'l' | 'h' | <vector-size>
<vector-size> ::=      'vl' | 'vh' | 'lv' | 'hv' | 'v'
<conversion> ::=      <char-conv> | <str-conv> | <fp-conv> |
                        <int-conv> | <misc-conv>
<char-conv>  ::=      'c'
<str-conv>   ::=      's' | 'p'
<fp-conv>    ::=      'e' | 'E' | 'f' | 'g' | 'G'
<int-conv>   ::=      'd' | 'i' | 'u' | 'o' | 'p' | 'x' | 'X'
<misc-conv>  ::=      'n' | '%' | '['
```

The extensions to the input conversion specification for vector types are shown in bold.

The **<vector-size>** indicates that a single vector value is to be scanned and converted. The vector value to be scanned is in the following general form:

```
value1 C value2 C ... C valuen
```

where C is a separator sequence defined by **<c-sep>** (the separator character optionally preceded by whitespace characters) and 4, 8, or 16 values are scanned depending on the **<vector-size>** each value scanned according to the **<conversion>**, as follows:

- A **<vector-size>** of 'vl' or 'lv' consumes one argument and modifies the **<int-conv>** conversion; it should be of type vector signed int * or vector unsigned int * depending on the **<int-conv>** specification; four values are scanned.
- A **<vector-size>** of 'vh' or 'hv' consumes one argument and modifies the **<int-conv>** conversion; it should be of type vector signed * or vector unsigned short * depending on the **<int-conv>** specification; 8 values are scanned.
- A **<vector-size>** of 'v' with **<int-conv>** or **<char-conv>** consumes one argument; it should be of type vector signed char * or vector unsigned char * depending on the **<int-conv>** or **<char-conv>** specification; 16 values are scanned.
- A **<vector-size>** of 'v' with **<fp-conv>** consumes one argument; it should be of type vector float *; four floating-point values are scanned.
- All other combinations of **<vector-size>** and **<conversion>** are undefined.

For the 'c' conversion the default separator character is null, and the separator sequence does not include whitespace characters preceding the separator character. For other than the

'c' conversions, the default separator character is a space, and the separator sequence does include whitespace characters preceding the separator character.

If the input stream reaches end-of-file or there is a conflict between the control string and a character read from the input stream, the input functions return EOF and do not assign to their vector argument.

When a conflict occurs, the character causing the conflict remains unread and is processed by the next input operation.

Examples:

```
sscanf("ab defghijklm,op", "%vc", &s8);
sscanf("a,b, ,d,e,f,g,h,i,j,k,l,m,,,o,p", "%,vc", &s8);
sscanf("1 2 3 4 5 6 7 8", "%vhu", &u16);
sscanf("1, 2, 3,99", "%,2lvd", &s32);
sscanf("1.10 ,2.20 ,3.30 ,4.40", "%,5vf", &f32);
```

This is equivalent to:

```
vector signed char s8 = vector signed char('a','b',' ','d','e','f',
                                           'g','h','i','j','k','l',
                                           'm',' ',' ','o','p');
vector unsigned short u16 = vector unsigned short(1,2,3,4,5,6,7,8);
vector signed int s32 = vector signed int(1, 2, 3, 99);
vector float f32 = vector float(1.1, 2.2, 3.3, 4.4);
```



Chapter 4

AltiVec Operations and Predicates

The following three subsections provide some background information that is helpful in understanding the descriptions provided for each operation and predicate. This is followed by a detailed listing of AltiVec operations followed by a separate section describing the AltiVec predicates. The final subsection contains compiler notes for handling predicates.

4.1 Vector Status and Control Register

The vector status and control register (VSCR) is a special 32-bit vector register shown in Figure 4-1.

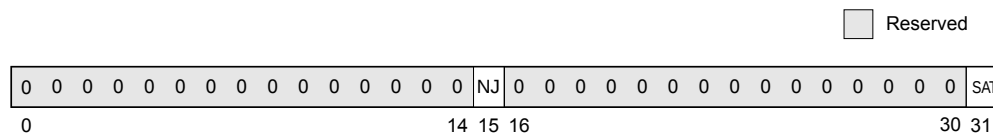


Figure 4-1. Vector Status and Control Register (VSCR)

The VSCR has two defined bits, the AltiVec non-Java mode (NJ) bit (VSCR[15]) and the AltiVec saturation (SAT) bit (VSCR[31]); the remaining bits are reserved. The `vec_mfvscr` operation moves the VSCR to a vector register. When moved, the 32-bit VSCR is right-justified in the 128-bit vector register, and the upper 96 bits `VRx[0–95]` of the vector register are cleared, so the VSCR in a vector register looks as shown in Figure 4-2.

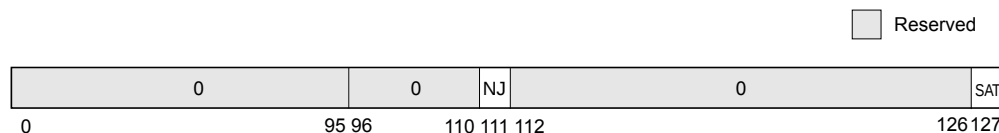


Figure 4-2. VSCR Moved to a Vector Register

VSCR bit settings are shown in Table 4-1.

Table 4-1. VSCR Field Descriptions

Bits	Name	Description
0–14	—	Reserved. Software is permitted to write any value to such a bit. A subsequent reading of the bit returns 0 if the value last written to the bit was 0 and returns an undefined value (0 or 1) otherwise.
15	NJ	Non-Java. A mode control bit that determines whether AltiVec floating-point operations will be performed in a Java-IEEE-C9X-compliant mode or a possibly faster non-Java/non-IEEE mode. 0 The Java-IEEE-C9X-compliant mode is selected. Denormalized values are handled as specified by Java, IEEE, and C9X standard. 1 The non-Java/non-IEEE-compliant mode is selected. If an element in a source vector register contains a denormalized value, the value 0 is used instead. If an instruction causes an underflow exception, the corresponding element in the target VR is cleared to 0. In both cases the 0 has the same sign as the denormalized or underflowing value. This mode is described in detail in the AltiVec Programming Environments Manual.
16–30	—	Reserved. Software is permitted to write any value to such a bit. A subsequent reading of the bit returns 0 if the value last written to the bit was 0 and returns an undefined value (0 or 1) otherwise.
31	SAT	Saturation. A sticky status bit indicating that some field in a saturating instruction saturated since the last time SAT was cleared. In other words, when SAT = 1 it remains set until it is cleared by an explicit instruction. 0 Indicates no saturation occurred, an instruction can explicitly clear this bit. 1 The AltiVec saturate instruction implicitly sets the SAT field when saturation has occurred on the results one of the AltiVec instructions or vector operations having saturate in its name.

After `vec_mfvscr` executes, the result in the target vector register is architecturally precise. That is, it reflects all updates to the SAT bit that could have been made by vector instructions logically preceding it in the program flow, and further, it does not reflect any SAT updates that may be made to it by vector instructions logically following it in the program flow. Reading the VSCR can be much slower than typical AltiVec instructions, and therefore care must be taken in reading it to avoid performance problems.

The first six 16-bit elements of the result are 0. The seventh element of the result contains the high-order 16 bits of the VSCR (including NJ). The eighth element of the result contains the low-order 16 bits of the VSCR (including SAT).

The setting of the Non-Java mode (NJ) bit (VSCR[15]) affects some vector floating-point operations. The other special bit (VSCR[31]) is the AltiVec Saturation (SAT) bit that is set when an operation generates a saturated result. Saturation is defined with respect to the type of resulting element. The result *d* of saturating a value *x* with respect to a type *t* means:

$$d = \max(\text{minimum}(t), \min(\text{maximum}(t), x))$$

where `minimum(t)` is the algebraically smallest value representable by a number of type *t* and `maximum(t)` is the algebraically largest value by a number of type *t*.

For each operation, where applicable, the effects of the NJ bit setting and/or the effects on the SAT bit are described in the operation description.

4.2 Byte Ordering

The default mapping for AltiVec ISA is PowerPC big-endian. The endian support of the PowerPC architecture does not address any data element larger than a double word; the basic memory unit for vectors is a quad word. Big-endian byte ordering is shown in Figure 4-3.

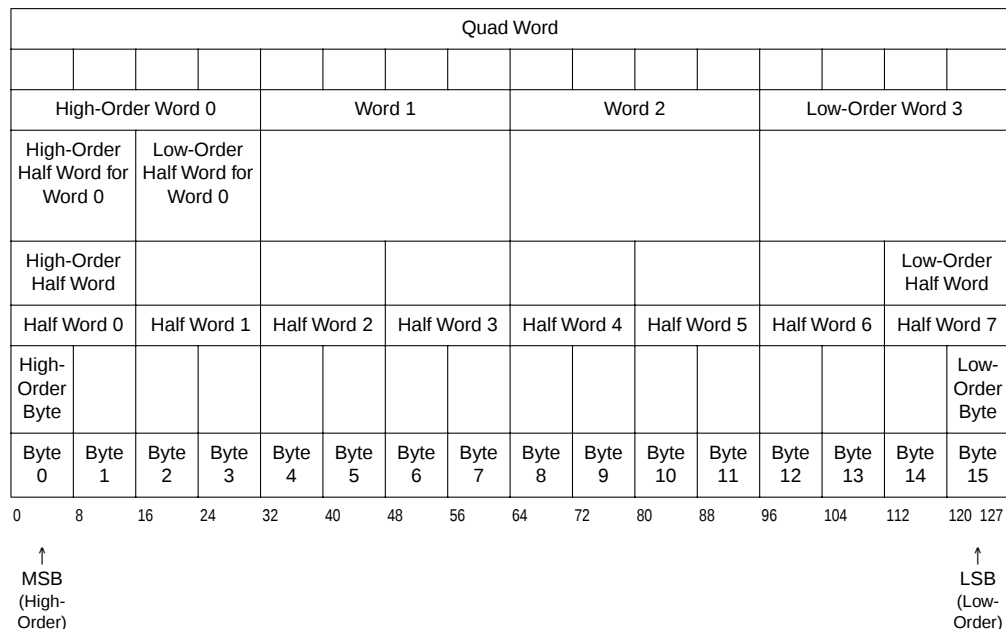


Figure 4-3. Big-Endian Byte Ordering for a Vector Register

As shown in Figure 4-3, the vector register elements are numbered using big-endian byte ordering. For example, the high-order (or most significant) byte element is numbered 0 and the low-order (or least significant) byte element is numbered 15.

When defining high-order and low-order for elements in a vector register, be careful not to confuse its meaning based on the bit numbering. For example, in Figure 4-3 the high-order half word for word 0 would be half word 0 (bits 0–7), and the low-order half word for word 0 would be half word 1 (bits 8–15).

4.3 Notation and Conventions

Operation and predicate functionality is described in this section by a semiformal pseudocode language. Table 4-2 lists the pseudocode notation and conventions used throughout the section.

Table 4-2. Notation and Conventions

Notation/Convention	Meaning
$\leftarrow$	Assignment
$+$, $+_{fp}$	Add, single-precision floating-point add
$-$, $-_{fp}$	Subtract, single-precision floating-point subtract
$*$, $*_{fp}$	Multiply, single-precision floating-point multiply
$/$	Integer division with non-negative remainder
$<$, $<_{fp}$	Less than, single-precision floating-point less than
$\leq$, $\leq_{fp}$	Less than or equal, single-precision floating-point less than or equal
$>$, $>_{fp}$	Greater than, single-precision floating-point greater than
$\geq$, $\geq_{fp}$	Greater than or equal, single-precision floating-point greater than or equal
$\neq$, $\neq_{fp}$	Not equal, floating-point not equal
$=$, $=_{fp}$	Equal, floating-point equal
$+\infty$, $-\infty$	Positive infinity, negative infinity
$\ $	Concatenation of two bit strings (e.g., 010 $\ $ 111 is the same as 010111)
$\&$	AND bit-wise operator
$ $	OR bit-wise operator
$\oplus$	Exclusive-OR bit-wise operator
$\neg$	NOT logical operator (one's complement)
0bnnnn	A number expressed in binary format
0xn timer	A number expressed in hexadecimal format
a,b,c,d	These symbols represent whole operands in an AltiVec operation or predicate. This is typically a vector, but in some operations it can represent a specific length literal value.
a_i, b_i, c_i, d_i	These symbols represent the i^{th} component elements of a vector a, b, c, or d, respectively.
ABS(x)	Absolute value of x
BorrowOut(x - y)	Borrow out of the difference of x and y
BoundAlign(x,y)	Align x to a y-byte boundary.
CarryOut(x + y)	Carry out of the sum of x and y
Ceil(x)	The smallest single-precision floating-point integer that is greater than or equal to x
do i=x to y	Do loop. • Do the following starting at x and iterating to y • Indenting shows range. • "To" and/or "by" clauses specify incrementing an iteration variable. • "While" clauses give termination conditions.
end	Indicates the end of a do loop

Table 4-2. Notation and Conventions (Continued)

Notation/Convention	Meaning
Floor(x)	The largest single-precision floating-point integer that is less than or equal to x
FP2 ^X Est(x)	3-bit-accurate floating-point estimate of 2 ^x
FPLog ₂ Est(x)	3-bit-accurate floating-point estimate of log ₂ (x)
FPRecipEst(x)	12-bit-accurate floating-point estimate of 1/x
if...then...else...	Conditional execution, indenting shows range, else is optional.
ISNaN(x)	Result is 1 if x is a not a number (NaN) and 0 if x is a number
ISNUM(x)	Result is 1 if x is a number and 0 if x is not a number (NaN)
MAX(x,y)	Returns the larger of x or y. For floating-point values, the following applies: - the maximum of +0.0 and -0.0 is +0.0 - the maximum of any value and a NaN is a QNaN
MEM(x,y)	Value at memory location x of size y bytes
MIN(x,y)	Returns the smaller of x or y. For floating-point values, the following applies: - the minimum of +0.0 and -0.0 is -0.0 - the minimum of any value and a NaN is a QNaN
mod(x,y)	Remainder of x/y
NaN	Not a Number, non-numeric
NEG(x)	Result is -x
NGE(x,y)	Result is 1 if x or y is a NaN or if x < y, and 0 otherwise
NGT(x,y)	Result is 1 if x or y is a NaN or x ≤ y, and 0 otherwise
NLE(x,y)	Result is 1 if x or y is a NaN or x > y, and 0 otherwise
NLT(x,y)	Result is 1 if x or y is a NaN or x ≥ y, and 0 otherwise
QNaN	NaN that propagates through most arithmetic operations without signalling an exception
RecipSQRTTest(x)	Result is a 12-bit accurate single-precision floating-point estimate of the reciprocal of the square root of x
RndToFPINear(x)	The single-precision floating-point integer that is nearest in value to x (in case of a tie, the even single-precision floating-point value is used).
RndToFPITrunc(x)	The largest single-precision floating-point integer that is less than or equal to x if x ≥ 0, or the smallest single-precision floating-point integer that is greater than or equal to x if x < 0
RndToFPNearest(x)	IEEE rounding to nearest floating-point number
ROTL(x,y)	Result of rotating x left by y bits
S	Represents a propagated sign bit in a figure
Saturate(x)	y ← Saturate(x) means saturate x to the type of y
ShiftRight(x,y) ShiftLeft(x,y)	Shift the contents of x right or left y bits, clearing vacated bits (logical shift). This operation is used for shift instructions.
ShiftRightA(x,y)	Shift the contents of x right y bits, copying the sign bit to the vacated bits (algebraic shift)
SignExtend(x,y)	Sign-extend x on the left with sign bits (that is, with copies of bit 0 of x) to produce y-bit value; represented in figures by a single S
SIToFP(x,y)	Result of converting the signed integer x to a y-bit floating-point value using Round-to-Nearest mode

Table 4-2. Notation and Conventions (Continued)

Notation/Convention	Meaning
UIToUImod(x,y)	Truncate an unsigned integer x to y-bit unsigned integer
Undefined	An undefined value. The value may vary from one implementation to another, and from one execution to another on the same implementation.
x_i	The i^{th} element of vector x where the size and type of the element are determined by the type of x
$x[i]$	The i^{th} byte of vector x
$x[y:x]$	Bits i through j of vector x, where i can equal j if referring to a single bit
x_0	A bit string of x zeros
x_1	A bit string of x ones
x_y	A bit string of x copies of y, for example, $^31 = 111$
x^n	x raised to the nth power

Precedence rules for pseudocode operators are summarized in Table 4-3.

Table 4-3. Precedence Rules

Operators	Associativity
$x[i]$, $x[y]$, $x[y:z]$ function evaluation	Left to right
x^y or replication, x^y or exponentiation	Right to left
unary $-$, $\neg$	Right to left
$*$, * _{fp} , $/$	Left to right
$+$, $+$ _{fp} , $-$, $-$ _{fp}	Left to right
$ $	Left to right
$=$, $=$ _{fp} , $!=$, $!=$ _{fp} , $<$, $<$ _{fp} , $\leq$, $\leq$ _{fp} , $>$, $>$ _{fp} , $\geq$, $\geq$ _{fp}	Left to right
$\&$, $\oplus$	Left to right
$ $	Left to right
$\leftarrow$	None

Operators higher in Table 4-3 are applied before those lower in the table. Operators at the same level in the table associate from left to right, from right to left, or not at all, as shown. For example, ‘ $-$ ’ (unary minus) associates from left to right, so $a - b - c = (a - b) - c$. Parentheses are used to override the evaluation order implied by Table 4-3, or to increase clarity; parenthesized expressions are evaluated before serving as operands.

4.4 Generic and Specific AltiVec Operations

The AltiVec operations are organized alphabetically by generic operation name with a definition of the permitted generic and specific AltiVec operations. The operations are listed in alphabetical order by mnemonic. Figure 4-4 shows the format for each operation description page.

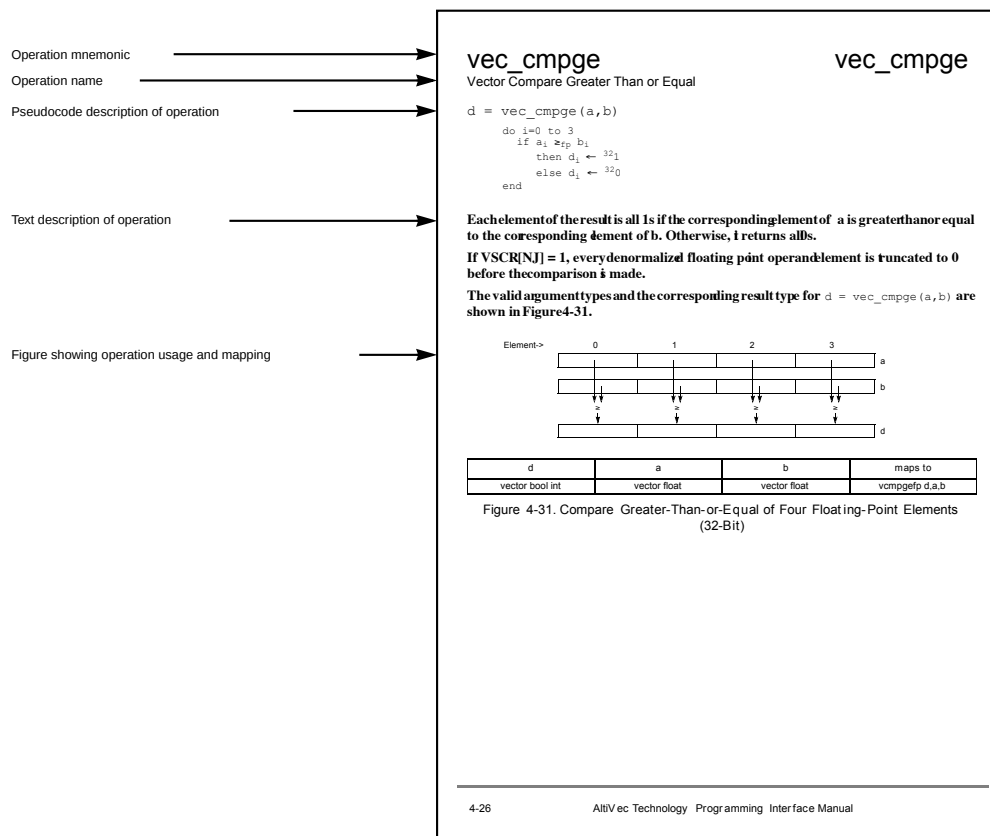


Figure 4-4. Operation Description Format

Where possible, each description is supported by reference figures indicating data modifications and including a table that lists:

- the valid set of argument types for that generic AltiVec operation,
- the result type for each set of argument types, and
- the specific AltiVec instruction(s) generated for that set of arguments.

Any operation not explicitly permitted in this section is prohibited.

vec_abs

vec_abs

Vector Absolute Value

```
d = vec_abs(a)

n ← number of elements
do i=0 to n-1
  di ← ABS(ai)
end
```

Each element of the result is the absolute value of the corresponding element of a. The arithmetic is modular for integer types.

For vector float argument types, the operation is independent of VSCR[NJ].

Programming note: Unlike other operations, `vec_abs` maps to multiple instructions. The programmer should consider alternatives. For example, to compute the absolute difference of two vectors `a` and `b`, the expression `vec_abs(vec_sub(a,b))` expands to four instructions. A simpler method uses the expression `vec_sub(vec_max(a,b), vec_min(a,b))` that expands to three instructions.

The valid combinations of argument types and the corresponding result types for `d = vec_abs(a)` are shown in Figure 4-5, Figure 4-6, Figure 4-7, and Figure 4-8. It is necessary to use the generic name since there is no specific operation for `vec_abs`.

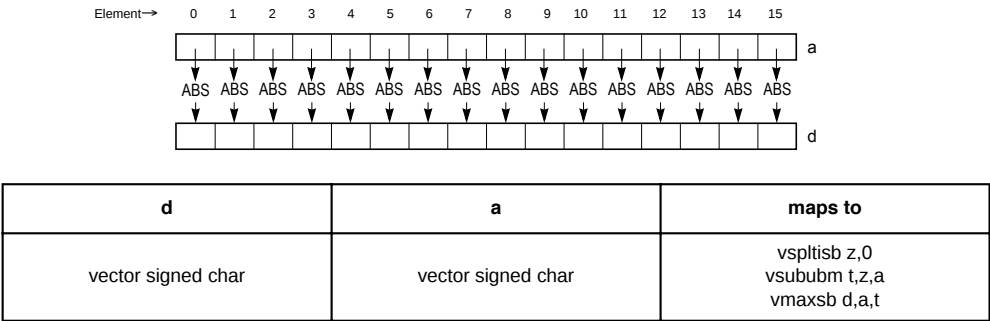


Figure 4-5. Absolute Value of Sixteen Integer Elements (8-bit)

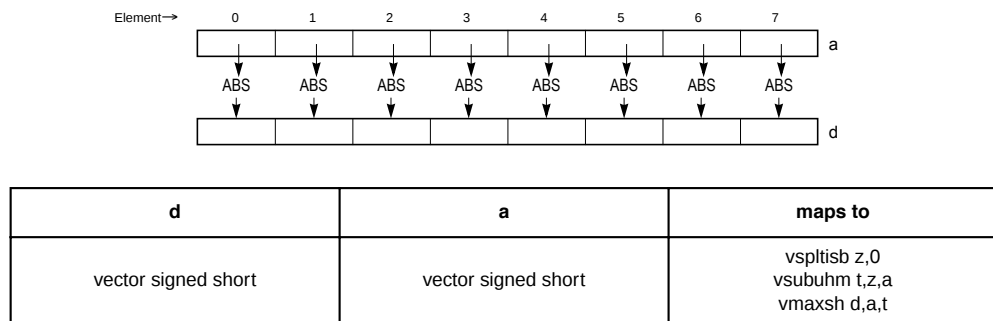


Figure 4-6. Absolute Value of Eight Integer Elements (16-bit)

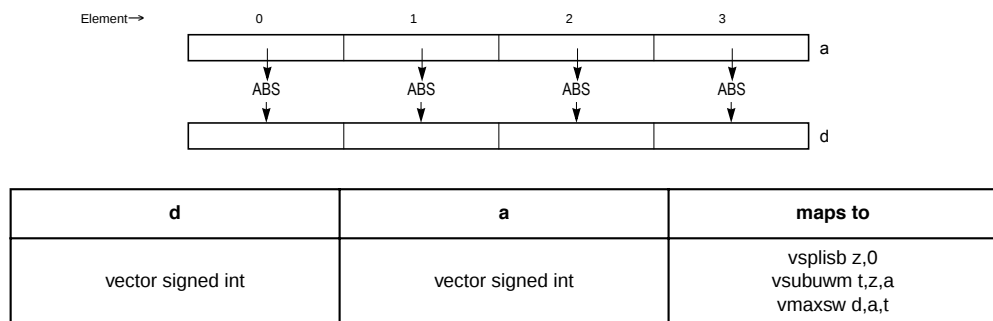


Figure 4-7. Absolute Value of Four Integer Elements (32-bit)

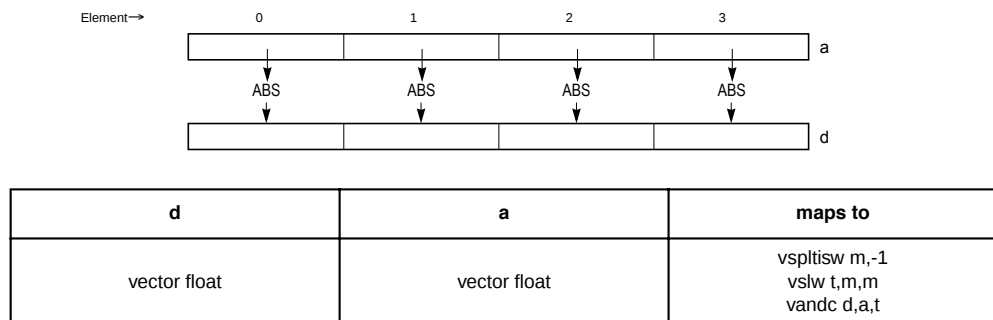


Figure 4-8. Absolute Value of Four Floating-Point Elements (32-bit)

vec_abss

vec_abss

Vector Absolute Value Saturated

```
d = vec_abss(a)

n ← number of elements
do i=0 to n-1
  di ← Saturate(ABS(ai))
end
```

Each element of the result is the absolute value of the corresponding element of a. The arithmetic is saturated for integer types. If saturation occurs, VSCR[SAT] is set (see Table 4-1).

Programming note: Unlike other operations, `vec_abss` maps to multiple instructions. The programmer should consider alternatives. For example, to compute the absolute difference of two vectors a and b, the expression `vec_abss(vec_subs(a,b))` expands to four instructions. A simpler method uses the expression `vec_subs(vec_max(a,b),vec_min(a,b))` that expands to three instructions.

The valid combinations of argument types and the corresponding result types for `d = vec_abss(a)` are shown in Figure 4-9, Figure 4-10, and Figure 4-11. It is necessary to use the generic name since there is no specific operation for `vec_abss`.

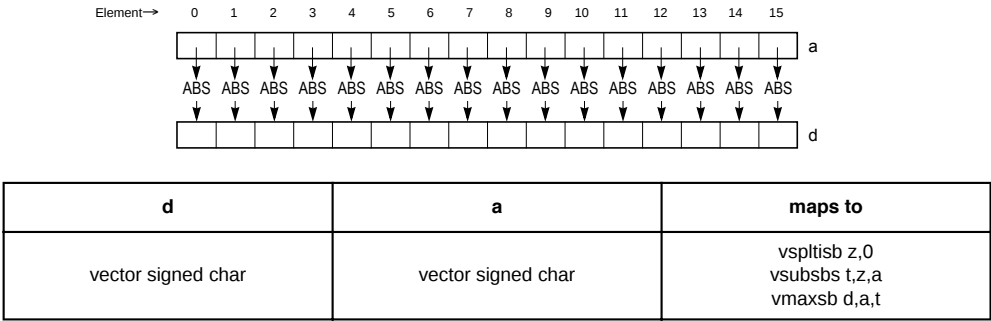


Figure 4-9. Saturated Absolute Value of Sixteen Integer Elements (8-bit)

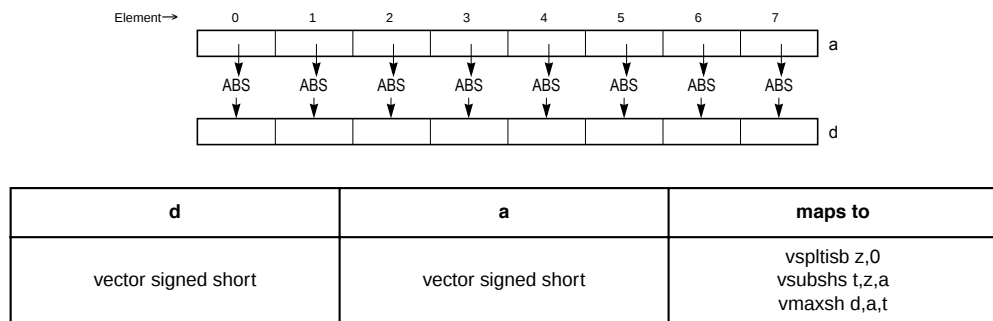


Figure 4-10. Saturated Absolute Value of Eight Integer Elements (16-bit)

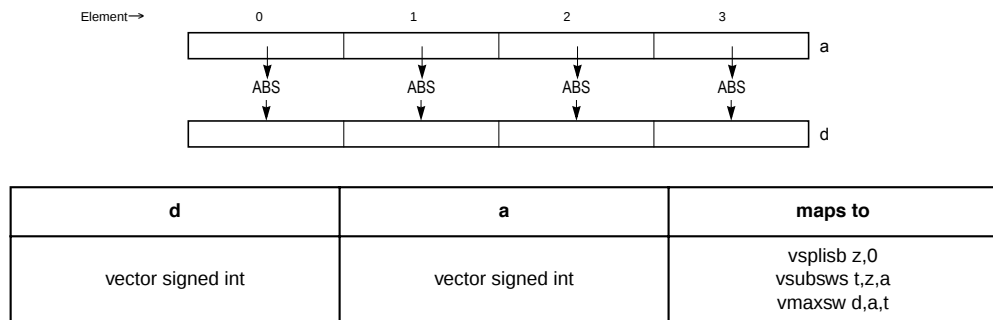


Figure 4-11. Saturated Absolute Value of Four Integer Elements (32-bit)

vec_add

Vector Add

vec_add

d = vec_add(**a**,**b**)

- Integer add:

```
n ← number of elements
do i=0 to n-1
  di ← ai + bi
end
```

- Floating-point add:

```
do i=0 to 3
  di ← ai +fp bi
end
```

Each element of **a** is added to the corresponding element of **b**. Each sum is placed in the corresponding element of **d**.

For `vector float` argument types, if `VSCR[NJ] = 1`, every denormalized operand element is truncated to a 0 of the same sign before the operation is carried out, and each denormalized result element is truncated to a 0 of the same sign.

The valid combinations of argument types and the corresponding result types for **d** = vec_add(**a**,**b**) are shown in Figure 4-12, Figure 4-13, Figure 4-14, and Figure 4-15.

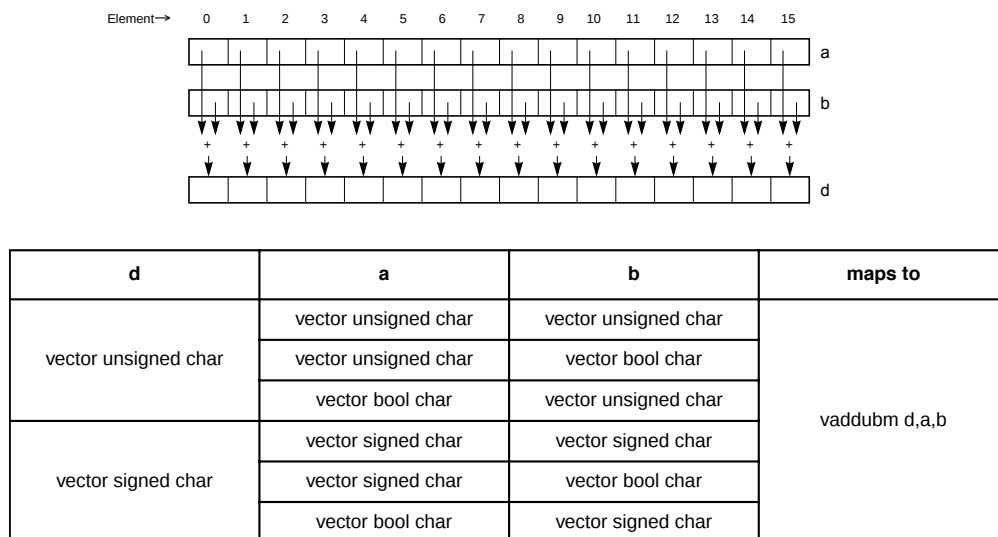
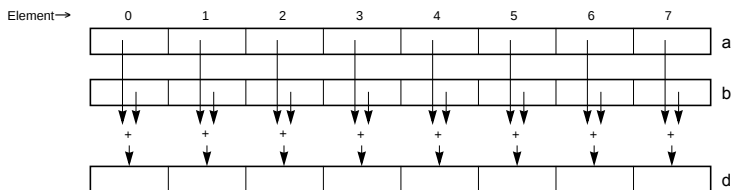
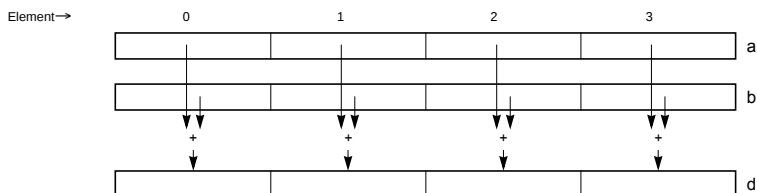


Figure 4-12. Add Sixteen Integer Elements (8-bit)



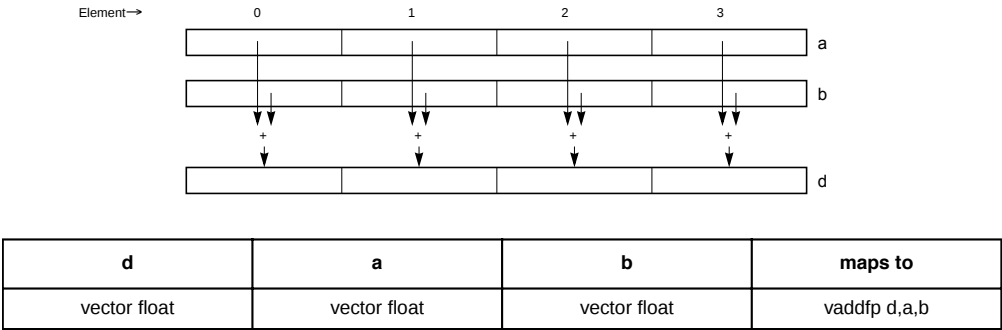
d	a	b	maps to
vector unsigned short	vector unsigned short	vector unsigned short	vadduhm d,a,b
	vector unsigned short	vector bool short	
	vector bool short	vector unsigned short	
vector signed short	vector signed short	vector signed short	
	vector signed short	vector bool short	
	vector bool short	vector signed short	

Figure 4-13. Add Eight Integer Elements (16-bit)



d	a	b	maps to
vector unsigned int	vector unsigned int	vector unsigned int	vadduwm d,a,b
	vector unsigned int	vector bool int	
	vector bool int	vector unsigned int	
vector signed int	vector signed int	vector signed int	
	vector signed int	vector bool int	
	vector bool int	vector signed int	

Figure 4-14. Add Four Integer Elements (32-bit)



vec_addc

Vector Add Carryout Unsigned Word

vec_addc

```
d = vec_addc(a,b)
```

```
do i=0 to 3
  di = CarryOut(ai + bi)
end
```

Each element of **a** is added to the corresponding element in **b**. The carry from each sum is zero-extended and placed into the corresponding element of **d**. CarryOut (**a** + **b**) is 1 if there is a carry, and otherwise 0. The valid argument types and the corresponding result type for **d** = vec_addc(**a**, **b**) are shown in Figure 4-16.

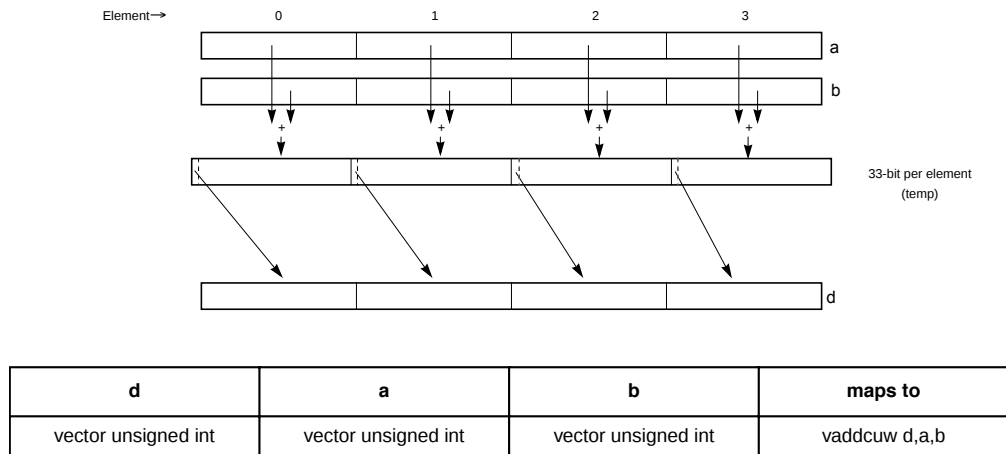


Figure 4-16. Carryout of Four Unsigned Integer Adds (32-bit)

vec_adds

Vector Add Saturated

vec_adds

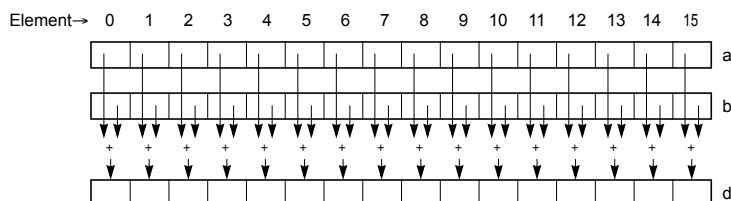
d = vec_adds(**a**,**b**)

```

n ← number of elements
do i=0 to n-1
  di ← Saturate(ai + bi)
end

```

Each element of **a** is added to the corresponding element of **b**. If saturation occurs, VSCR[SAT] is set (see Table 4-1). The signed-integer result is placed into the corresponding element of **d**. The valid combinations of argument types and the corresponding result types for **d** = vec_adds(**a**,**b**) are shown in Figure 4-17, Figure 4-18, and Figure 4-19.



d	a	b	maps to
vector unsigned char	vector unsigned char	vector unsigned char	vaddubs d,a,b
	vector unsigned char	vector bool char	
	vector bool char	vector unsigned char	
vector signed char	vector signed char	vector signed char	vaddubs d,a,b
	vector signed char	vector bool char	
	vector bool char	vector signed char	

Figure 4-17. Add Saturating Sixteen Integer Elements (8-bit)

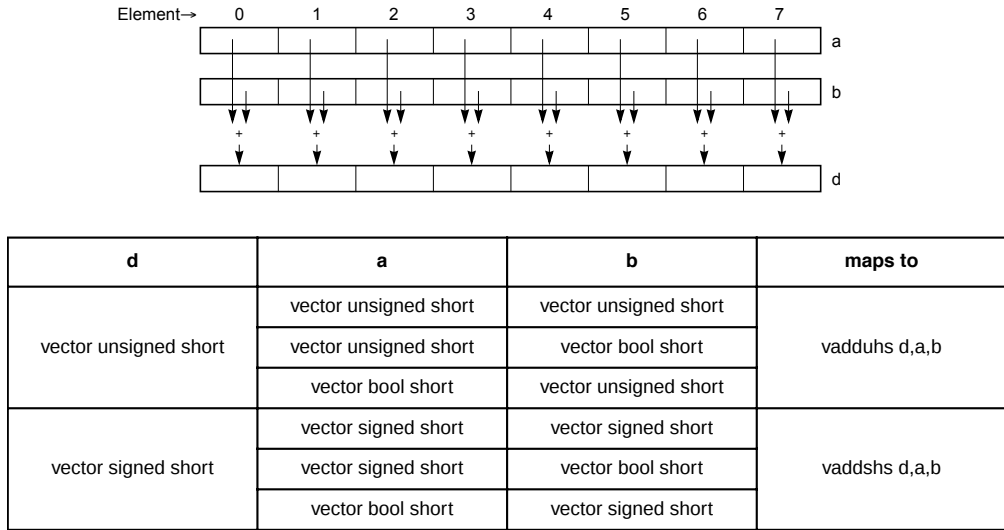


Figure 4-18. Add Saturating Eight Integer Elements (16-bit)

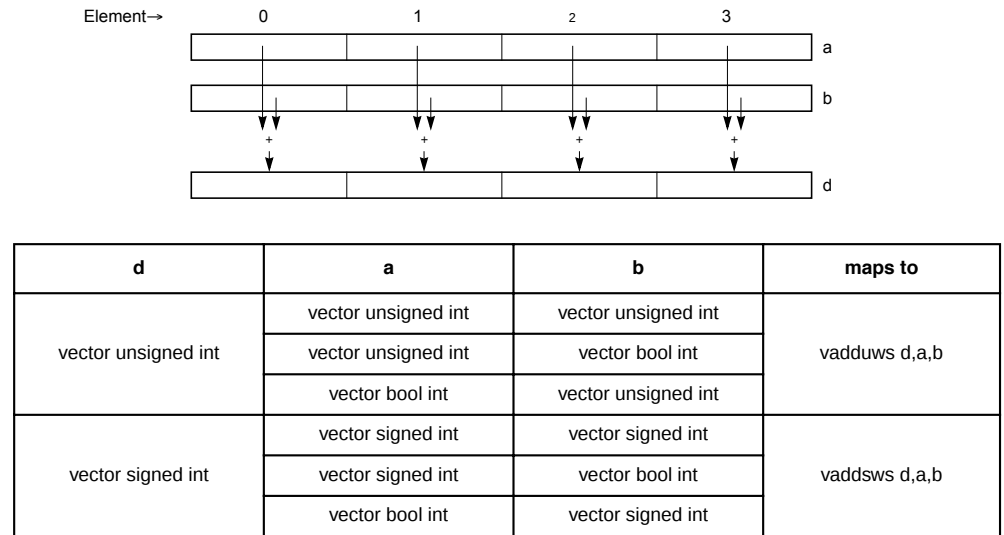


Figure 4-19. Add Saturating Four Integer Elements (32-bit)

vec_and

Vector Logical AND

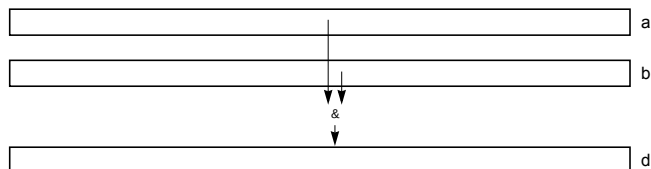
vec_and

d = vec_and(**a**,**b**)

d ← **a** & **b**

Each bit of the result is the logical AND of the corresponding bits of **a** and **b**. The valid combinations of argument types and the corresponding result types for

d = vec_and(**a**,**b**) are shown in Figure 4-20.



d	a	b	maps to
vector unsigned char	vector unsigned char	vector unsigned char	vand d,a,b
	vector unsigned char	vector bool char	
	vector bool char	vector unsigned char	
vector signed char	vector signed char	vector signed char	
	vector signed char	vector bool char	
	vector bool char	vector signed char	
vector bool char	vector bool char	vector bool char	
vector unsigned short	vector unsigned short	vector unsigned short	
	vector unsigned short	vector bool short	
	vector bool short	vector unsigned short	
vector signed short	vector signed short	vector signed short	
	vector signed short	vector bool short	
	vector bool short	vector signed short	
vector bool short	vector bool short	vector bool short	
vector unsigned int	vector unsigned int	vector unsigned int	
	vector unsigned int	vector bool int	
	vector bool int	vector unsigned int	
vector signed int	vector signed int	vector signed int	
	vector signed int	vector bool int	
	vector bool int	vector signed int	
vector bool int	vector bool int	vector bool int	
vector float	vector bool int	vector float	
	vector float	vector bool int	
	vector float	vector float	

Figure 4-20. Logical Bit-Wise AND

vec_andc

Vector Logical AND with Complement

d = vec_andc(**a**,**b**)

d ← **a** & ¬**b**

Each bit of the result is the logical AND of the corresponding bit of **a** and the one's complement of the corresponding bit of **b**. The valid combinations of argument types and the corresponding result types for **d** = vec_andc(**a**,**b**) are shown in Figure 4-21.

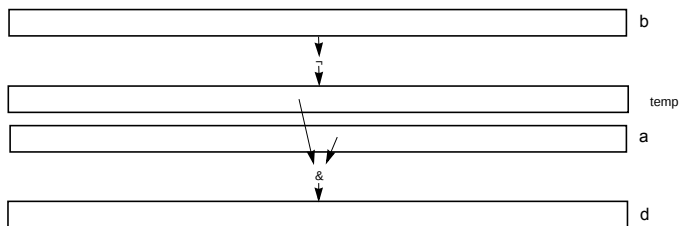


Figure 4-21. Logical Bit-Wise AND with Complement

d	a	b	maps to
vector unsigned char	vector unsigned char	vector unsigned char	vandc d,a,b
	vector unsigned char	vector bool char	
	vector bool char	vector unsigned char	
vector signed char	vector signed char	vector signed char	
	vector signed char	vector bool char	
	vector bool char	vector signed char	
vector bool char	vector bool char	vector bool char	
vector unsigned short	vector unsigned short	vector unsigned short	
	vector unsigned short	vector bool short	
	vector bool short	vector unsigned short	
vector signed short	vector signed short	vector signed short	
	vector signed short	vector bool short	
	vector bool short	vector signed short	
vector bool short	vector bool short	vector bool short	
vector unsigned int	vector unsigned int	vector unsigned int	
	vector unsigned int	vector bool int	
	vector bool int	vector unsigned int	
vector signed int	vector signed int	vector signed int	
	vector signed int	vector bool int	
	vector bool int	vector signed int	
vector bool int	vector bool int	vector bool int	
vector float	vector bool int	vector float	
	vector float	vector bool int	
	vector float	vector float	

Figure 4-21. Logical Bit-Wise AND with Complement

vec_avg

Vector Average

vec_avg

d = vec_avg(**a**,**b**)

```

n ← number of elements
do i=0 to n-1
  di ← (ai + bi + 1) / 2
end

```

Each element of the result is a rounded average of the corresponding elements of **a** and **b**. Intermediate calculations are not limited by the element size. The value 1 is added to the sum of elements in **a** and **b** to ensure the result is rounded up. The valid combinations of argument types and the corresponding result types for **d** = vec_avg(**a**,**b**) are shown in Figure 4-22, Figure 4-23, and Figure 4-24.

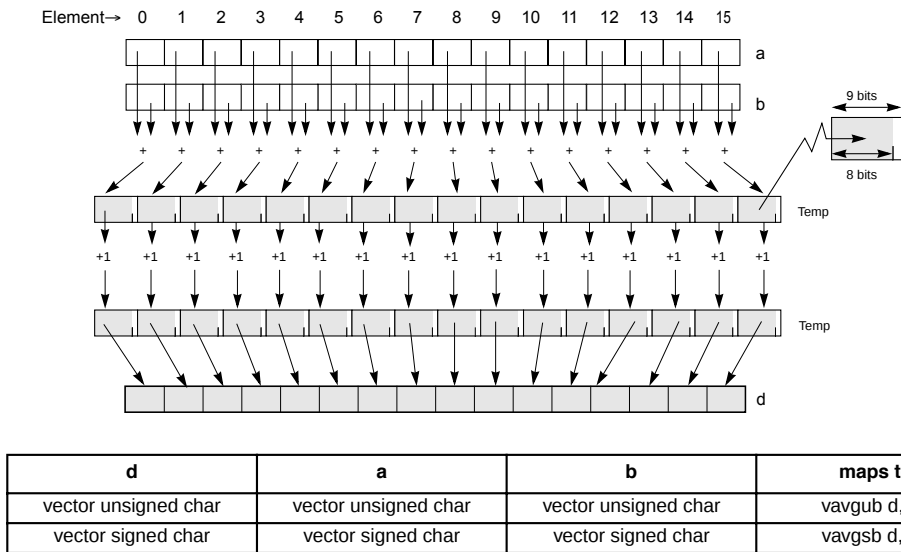


Figure 4-22. Average Sixteen Integer Elements (8-bit)

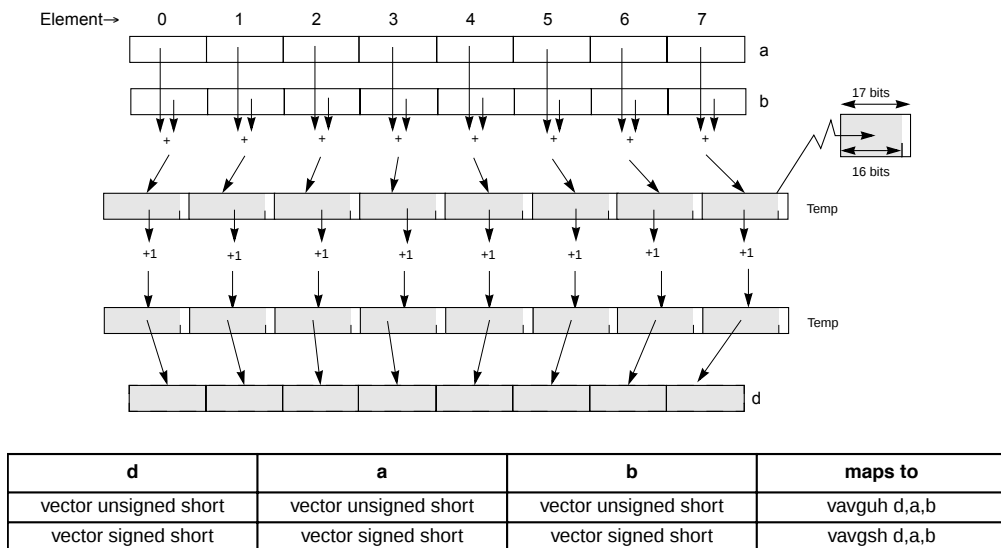


Figure 4-23. Average Eight Integer Elements (16-bit)

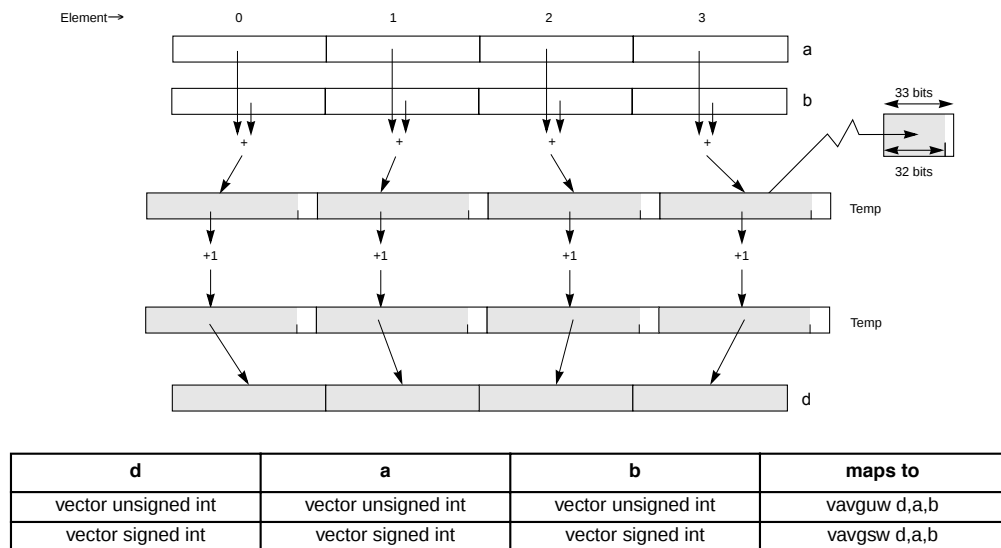


Figure 4-24. Average Four Integer Elements (32-bit)

vec_ceil

Vector Ceiling

d = vec_ceil(**a**)

```
do i=0 to 3
  di ← Ceil(ai)
end
```

Each single-precision floating-point element in **a** is rounded to a single-precision floating-point integer using the rounding mode Round toward +Infinity, and placed into the corresponding word element of **d**. If an element a_i is infinite, the corresponding element d_i equals a_i . If an element a_i is finite, the corresponding element d_i is the smallest represented floating-point value $\geq a_i$. For example, if the floating-point element was 123.45, the resulting integer would be 124.

If VSCR[NJ] = 1, every denormalized operand element is truncated to 0 before the operation.

The valid argument types and the corresponding result type for **d** = vec_ceil(**a**,**b**) are shown in Figure 4-25.

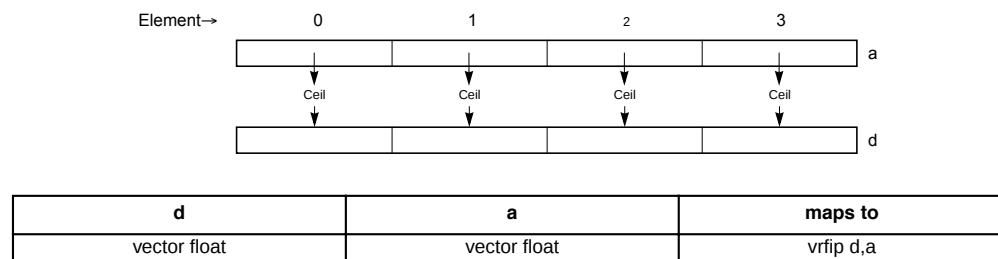


Figure 4-25. Round to Plus Infinity of Four Floating-Point Integer Elements (32-Bit)

vec_cmpb

Vector Compare Bounds Floating-Point

vec_cmpb

d = vec_cmpb(**a**,**b**)

```

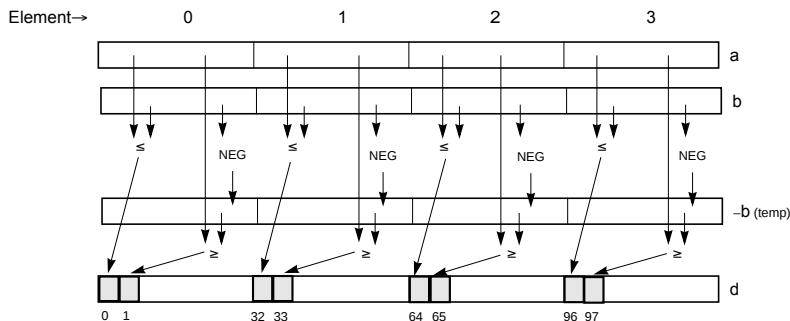
do i=0 to 3
  di ← 0
  if ai ≤fp bi
    then di[0] ← 0
    else di[0] ← 1
  if ai ≥fp -bi
    then di[1] ← 0
    else di[1] ← 1
end

```

Each element in **a** is compared to the corresponding element in **b**. The 2-bit result indicates whether the element in **a** is within the bounds specified by the element in **b**. Bit 0 of each result is 0 if the element in **a** is less than or equal to the element in **b** (i.e., in bounds high), and is 1 otherwise (i.e., out of bounds high). Bit 1 of the 2-bit value is 0 if the element in **a** is greater than or equal to the negative of the element in **b** (i.e., in bounds low), and is 1 otherwise (i.e., out of bounds low). The 2-bit result is placed into the high-order two bits (bit 0 and 1) of the corresponding element in **d** (which correspond to bits 0–1, 32–33, 64–65, and 96–97 of **d**, respectively) and the remaining bits are cleared. If any single-precision floating-point word element in **b** is negative; the corresponding element in **a** is out of bounds. If an element in **a** or **b** element is a NaN, the two high-order bits of the corresponding result are both 1.

If VSCR[NJ] = 1, every denormalized operand element is truncated to 0 before the comparison.

The valid argument types and the corresponding result type for **d** = vec_cmpb(**a**,**b**) are shown in Figure 4-26.



d	a	b	maps to
vector signed int	vector float	vector float	vcmpbfp d,a,b

Figure 4-26. Compare Bounds of Four Floating-Point Elements (32-Bit)

vec_cmpeq

Vector Compare Equal

vec_cmpeq

d = vec_cmpeq(**a**,**b**)

- Integer compare equal:

```

n ← number of elements
m ← number of bits in an element (128/n)
do i=0 to n-1
  if ai = bi
    then di ← m1
    else di ← m0
end

```

- Floating-point compare equal:

```

do i=0 to 3
  if ai =fp bi
    then di ← 321
    else di ← 320
end

```

Each element of the result is all ones if the corresponding element of **a** is equal to the corresponding element of **b**. Otherwise, it returns all zeros.

For `vector float` argument types, if `VSCR[NJ] = 1`, every denormalized floating-point operand element is truncated to 0 before the comparison.

The valid combinations of argument types and the corresponding result types for **d** = vec_cmpeq(**a**,**b**) are shown in Figure 4-27, Figure 4-28, Figure 4-29, and Figure 4-30.

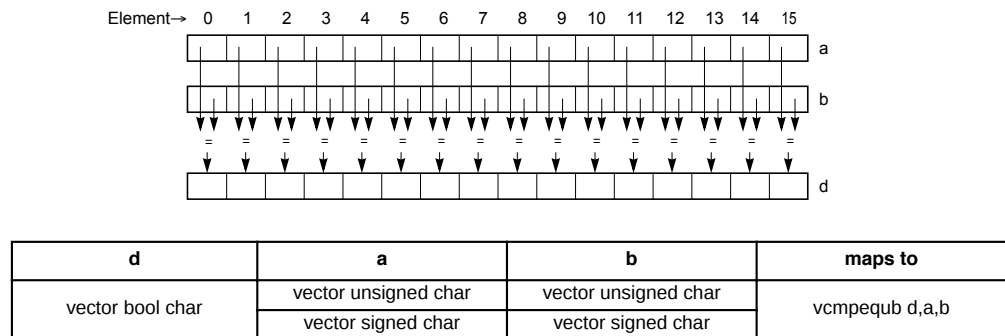


Figure 4-27. Compare Equal of Sixteen Integer Elements (8-bits)

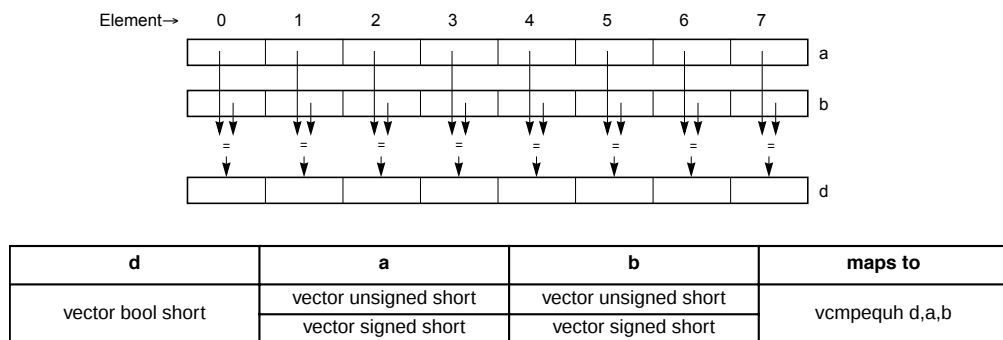


Figure 4-28. Compare Equal of Eight Integer Elements (16-Bit)

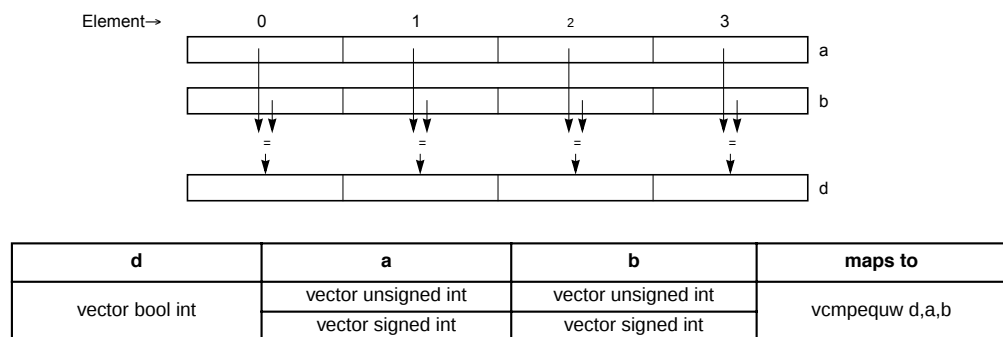


Figure 4-29. Compare Equal of Four Integer Elements (32-Bit)

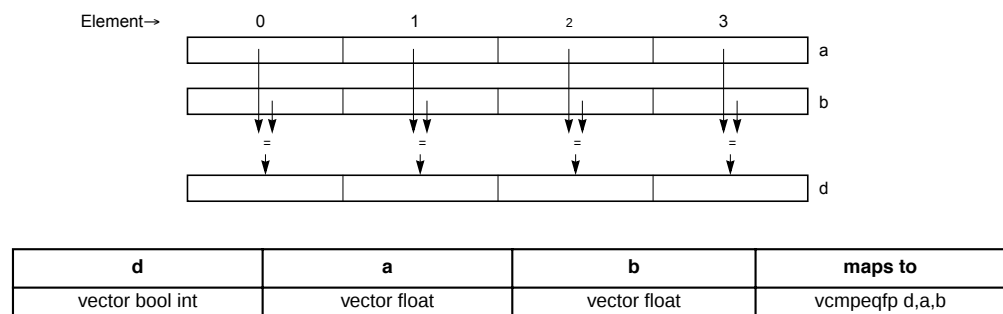


Figure 4-30. Compare Equal of Four Floating-Point Elements (32-Bit)

vec_cmpge

Vector Compare Greater Than or Equal

vec_cmpge

d = vec_cmpge(**a**,**b**)

```
do i=0 to 3
  if ai ≥fp bi
    then di ← 321
    else di ← 320
end
```

Each element of the result is all ones if the corresponding element of **a** is greater than or equal to the corresponding element of **b**. Otherwise, it returns all zeros.

If VSCR[NJ] = 1, every denormalized floating-point operand element is truncated to 0 before the comparison.

The valid argument types and the corresponding result type for **d** = vec_cmpge(**a**,**b**) are shown in Figure 4-31.

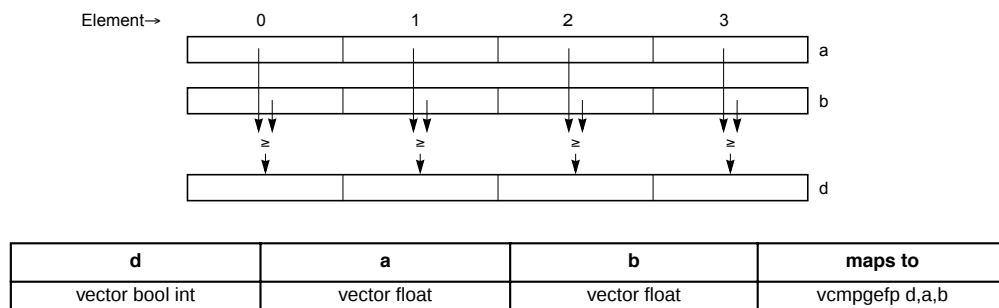


Figure 4-31. Compare Greater-Than-or-Equal of Four Floating-Point Elements (32-Bit)

vec_cmpgt

Vector Compare Greater Than

vec_cmpgt

d = vec_cmpgt(**a**,**b**)

- Integer compare greater than:

```

n ← number of elements
m ← number of bits in an element (128/n)
do i=0 to n-1
  if ai > bi
    then di ← m1
    else di ← m0
end

```

- Floating-point compare greater than:

```

do i=0 to 3
  if ai >fp bi
    then di ← 321
    else di ← 320
end

```

Each element of the result is all ones if the corresponding element of a is greater than the corresponding element of b. Otherwise, it returns all zeros.

For vector float types, if VSCR[NJ] = 1, every denormalized floating-point operand element is truncated to 0 before the comparison.

The valid combinations of argument types and the corresponding result types for **d** = vec_cmpgt(**a**,**b**) are shown in Figure 4-32, Figure 4-33, Figure 4-34, and Figure 4-35.

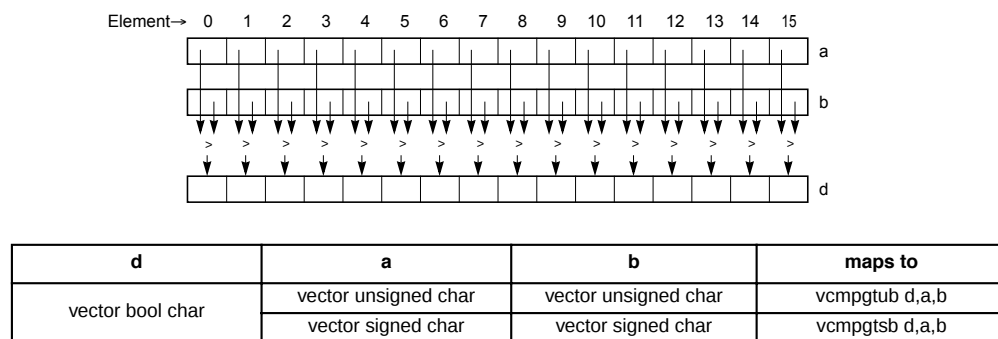


Figure 4-32. Compare Greater-Than of Sixteen Integer Elements (8-bits)

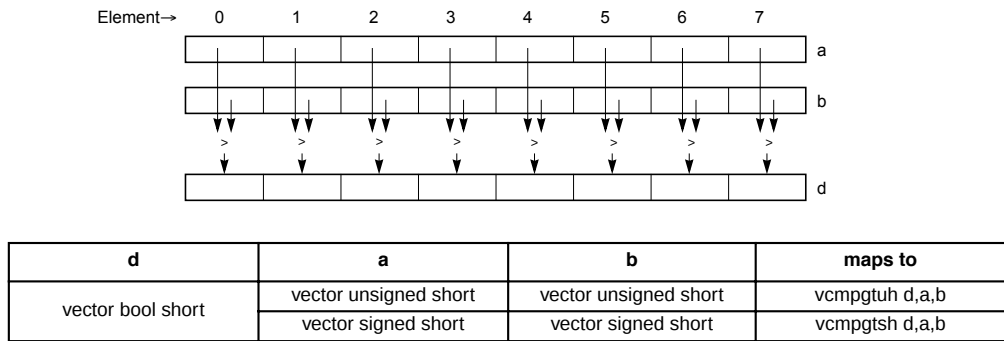


Figure 4-33. Compare Greater-Than of Eight Integer Elements (16-Bit)

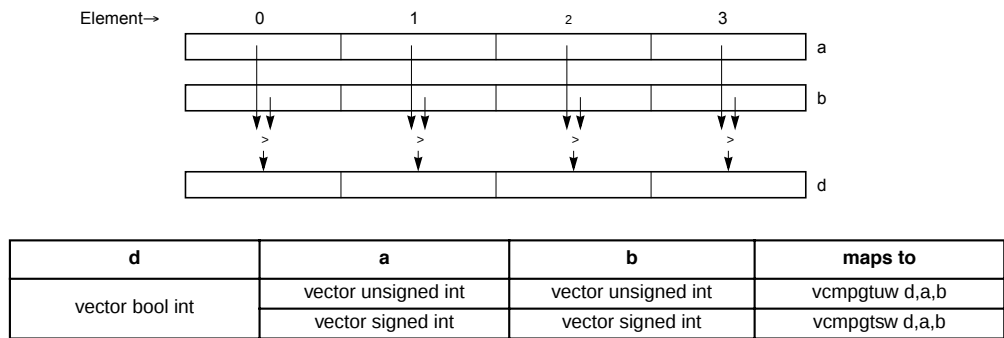


Figure 4-34. Compare Greater-Than of Four Integer Elements (32-Bit)

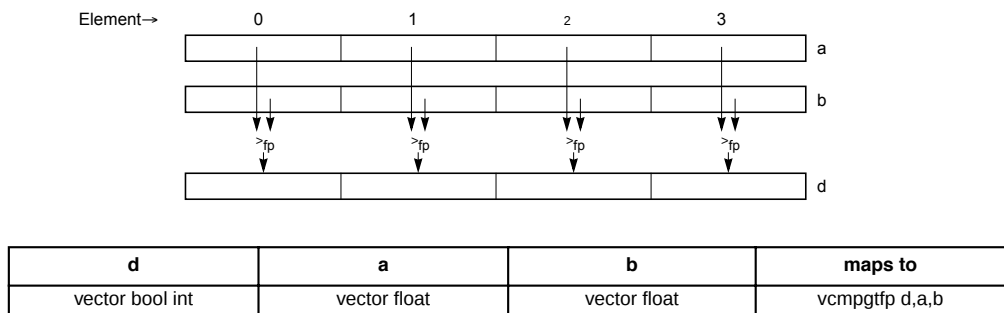


Figure 4-35. Compare Greater-Than of Four Floating-Point Elements (32-Bit)

vec_cmple

vec_cmple

Vector Compare Less Than or Equal

```
d = vec_cmple(a,b)

do i=0 to 3
  if ai ≤fp bi
    then di ← 321
    else di ← 320
end
```

Each element of the result is all ones if the corresponding element of a is less than or equal to the corresponding element of b. Otherwise, it returns all zeros.

If VSCR[NJ] = 1, every denormalized floating-point operand element is truncated to 0 before the comparison.

The valid argument types and the corresponding result type for **d = vec_cmple(a,b)** are shown in Figure 4-36. It is necessary to use the generic name, since the specific operation `vec_vcmpgefp` does not reverse its operands.

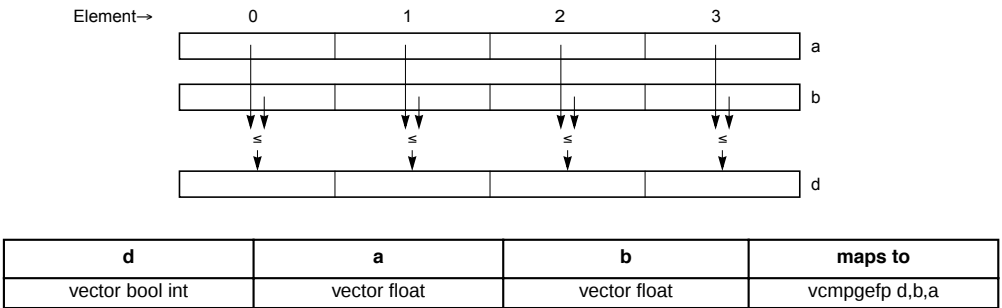


Figure 4-36. Compare Less-Than-or-Equal of Four Floating-Point Elements (32-Bit)

vec_cmplt

Vector Compare Less Than

vec_cmplt

d = vec_cmplt(**a**,**b**)

- Integer compare less than:

```

n ← number of elements
m ← number of bits in an element (128/n)
do i=0 to n-1
  if ai < bi
    then di ← m1
    else di ← m0
end

```

- Floating-point compare less than:

```

do i=0 to 3
  if ai <fp bi
    then di ← 321
    else di ← 320
end

```

Each element of the result is all ones if the corresponding element of a is less than the corresponding element of b. Otherwise, it returns all zeros.

For vector float types, if VSCR[NJ] = 1, every denormalized floating-point operand element is truncated to 0 before the comparison.

The valid combinations of argument types and the corresponding result types for **d** = vec_cmplt(**a**,**b**) are shown in Figure 4-37, Figure 4-38, Figure 4-39, and Figure 4-40. It is necessary to use the generic name, since the specific operations do not reverse their operands.

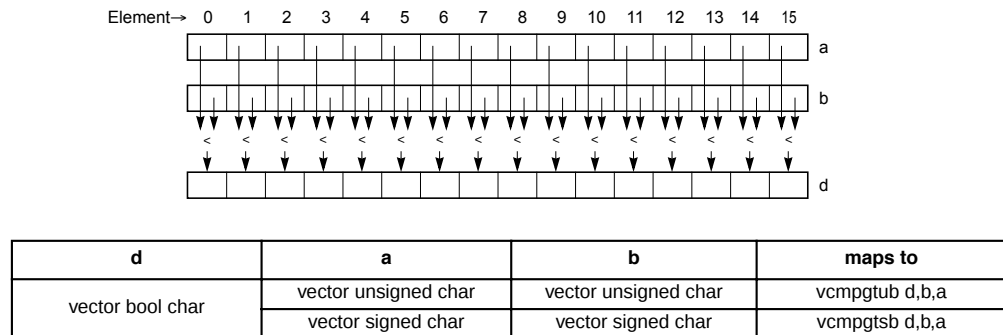


Figure 4-37. Compare Less-Than of Sixteen Integer Elements (8-bits)

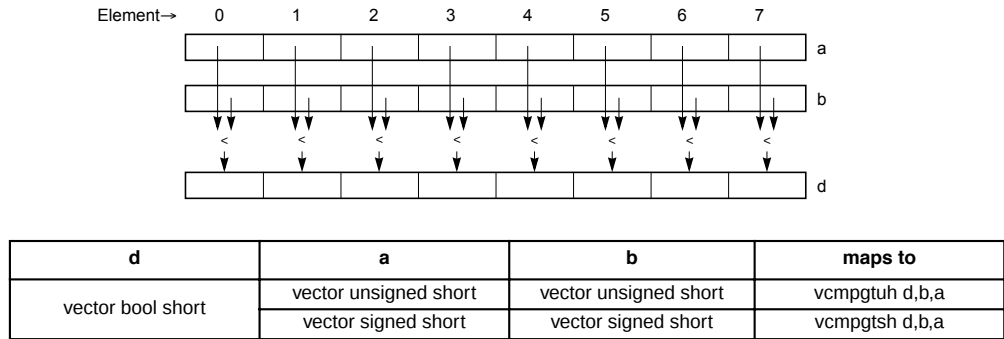


Figure 4-38. Compare Less-Than of Eight Integer Elements (16-Bit)

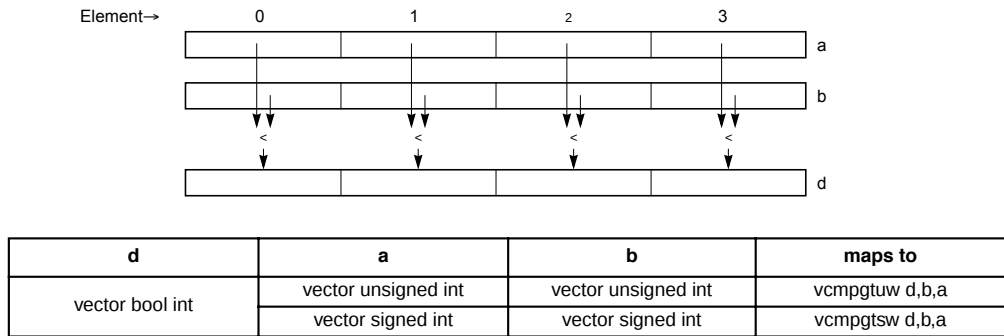


Figure 4-39. Compare Less-Than of Four Integer Elements (32-Bit)

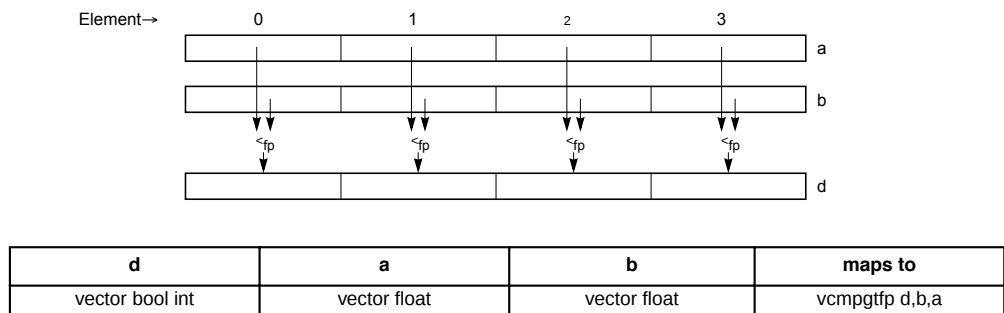


Figure 4-40. Compare Less-Than of Four Floating-Point Elements (32-Bit)

vec_ctf

Vector Convert from Fixed-Point Word

vec_ctf

d = vec_ctf(**a**,**b**)

```
do i=0 to 3
  di ← SIToFP(ai) * 2-b
end
```

Each element of the result is the closest floating-point representation of the number obtained by dividing the corresponding element of **a** by 2 to the power of **b**.

The operation is independent of VSCR[NJ].

The valid argument types and the corresponding result type for **d** = vec_ctf(**a**,**b**) are shown in Figure 4-41.

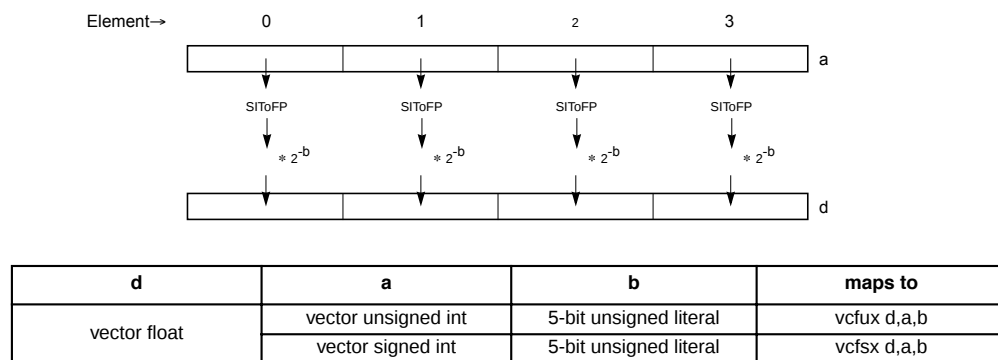


Figure 4-41. Convert Four Integer Elements to Four Floating-Point Elements (32-Bit)

vec_cts

Vector Convert to Signed Fixed-Point Word Saturated

vec_cts

d = vec_cts(**a**,**b**)

```
do i=0 to 3
  di ← Saturate(ai * 2b)
end
```

Each element of the result is the saturated signed value obtained after truncating the product of the corresponding element of **a** and 2 to the power of **b**.

If VSCR[NJ] = 1, every denormalized floating-point operand element is truncated to 0 before the operation.

If saturation occurs, VSCR[SAT] is set (see Table 4-1).

The valid argument types and the corresponding result type for **d** = vec_cts(**a**,**b**) are shown in Figure 4-42.

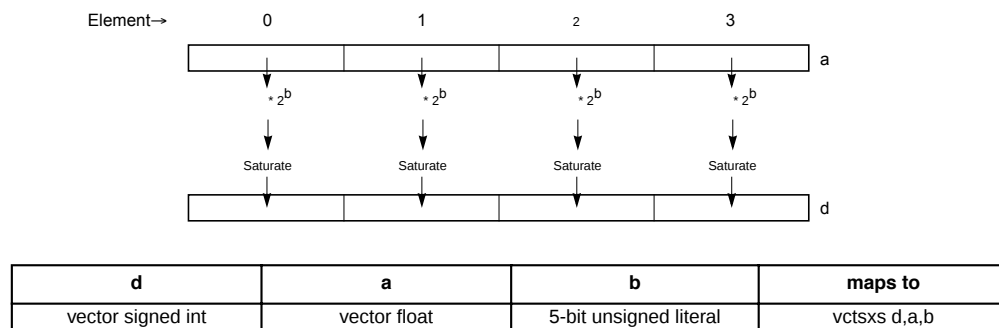


Figure 4-42. Convert Four Floating-Point Elements to Four Saturated Signed Integer Elements (32-Bit)

vec_ctu

Vector Convert to Unsigned Fixed-Point Word Saturated

vec_ctu

d = vec_ctu(**a**,**b**)

```
do i=0 to 3
  di ← Saturate (ai * 2b)
end
```

Each element of the result is the saturated unsigned value obtained after truncating the number obtained by multiplying the corresponding element of **a** by 2 to the power of **b**.

If VSCR[NJ] = 1, every denormalized floating-point operand element is truncated to 0 before the operation.

If saturation occurs, VSCR[SAT] is set (see Table 4-1).

The valid argument types and the corresponding result type for **d** = vec_ctu(**a**,**b**) are shown in Figure 4-43.

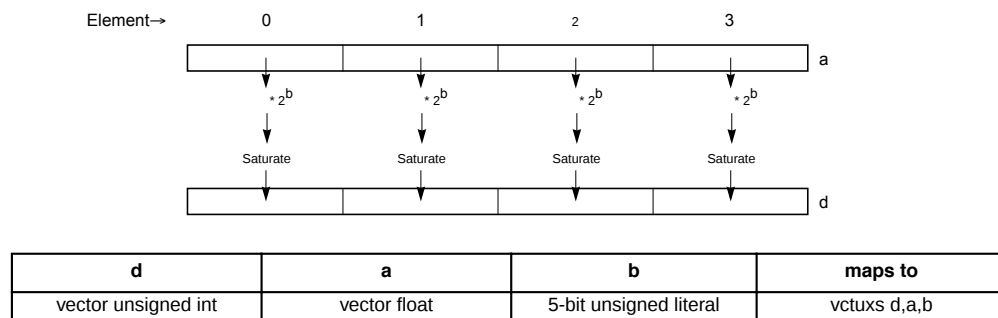


Figure 4-43. Convert Four Floating-Point Elements to Four Saturated Unsigned Integer Elements (32-Bit)

vec_dss

Vector Data Stream Stop

vec_dss

```
vec_dss(a)

DataStreamPrefetchControl ← "stop" || a
```

Each operation stops cache touches for the data stream associated with tag a. The result is void. The valid argument type for `vec_dss(a)` is shown in Table 4-4. The result type is void.

Table 4-4. vec_dss—Vector Data Stream Stop Argument Types

a	maps to
2-bit unsigned literal	dss a

vec_dssall

Vector Stream Stop All

vec_dssall()

```
DataStreamPrefetchControl ← "stop"
```

The operation stops cache touches for all data streams. All argument and result types for `vec_dssall()` are void. `vec_dssall` maps to the `dssall` instruction.

vec_dssall

vec_dst

vec_dst

Vector Data Stream Touch

vec_dst(a,b,c)

```
addr[0:63] ← a
DataStreamPrefetchControl ← "start" || c || 0 || b || addr
```

Each operation initiates cache touches for loads for the data stream associated with tag c at the address a using the data block in b. The result type is void.

The a type may also be a pointer to a const-qualified type. Plain char * is excluded in the mapping for a.

The b type is encoded for 32-bit as follows:

- Block size: b[3:7] if b[3:7] != 0; otherwise 32
- Block count: b[8:15] if b[8:15] != 0; otherwise 256
- Block stride: b[16:31] if b[16:31] != 0; otherwise 32768

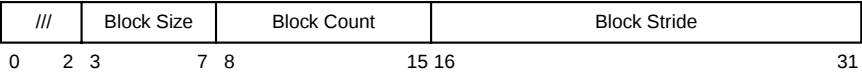


Figure 4-44. Format of b Type (32-bit)

The b type is encoded for 64-bit as follows:

- Block size: b[35:39] if b[35:39] != 0; otherwise 32
- Block count: b[40:47] if b[40:47] != 0; otherwise 256
- Block stride: b[48:63] if b[48:63] != 0; otherwise 32768

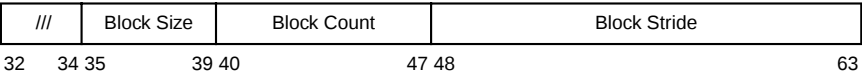


Figure 4-45. Format of b Type (64-bit)

The c type is a 2-bit unsigned literal tag used to identify a specific data stream. Up to four streams can be set up with this mechanism.

The valid combinations of argument types for vec_dst(a,b,c) are shown in Table 4-5. The result type is void.

Table 4-5. vec_dst—Vector Data Stream Touch Argument Types

a	b	c	maps to
vector unsigned char *	any integral type	2-bit unsigned literal	dst a,b,c
vector signed char *	any integral type	2-bit unsigned literal	
vector bool char *	any integral type	2-bit unsigned literal	
vector unsigned short *	any integral type	2-bit unsigned literal	
vector signed short *	any integral type	2-bit unsigned literal	
vector bool short *	any integral type	2-bit unsigned literal	
vector pixel *	any integral type	2-bit unsigned literal	
vector unsigned int *	any integral type	2-bit unsigned literal	
vector signed int *	any integral type	2-bit unsigned literal	
vector bool int *	any integral type	2-bit unsigned literal	
vector float *	any integral type	2-bit unsigned literal	
unsigned char *	any integral type	2-bit unsigned literal	
signed char *	any integral type	2-bit unsigned literal	
unsigned short *	any integral type	2-bit unsigned literal	
short *	any integral type	2-bit unsigned literal	
unsigned int *	any integral type	2-bit unsigned literal	
int *	any integral type	2-bit unsigned literal	
unsigned int *	any integral type	2-bit unsigned literal	
float *	any integral type	2-bit unsigned literal	

vec_dstst

vec_dstst

Vector Data Stream Touch for Store

vec_dstst(a,b,c)

```
addr[0:63] ← a
DataStreamPrefetchControl ← "start" || 0 || static || b || addr
```

Each operation initiates cache touches for stores for the data stream associated with tag c at the address a using the data block in b. The result type is void.

The a type may also be a pointer to a const-qualified type. Plain char * is excluded in the mapping for a.

The b type is encoded for 32-bit as follows:

- Block size: b[3:7] if b[3:7] != 0; otherwise 32
- Block count: b[8:15] if b[8:15] != 0; otherwise 256
- Block stride: b[16:31] if b[16:31] != 0; otherwise 32768

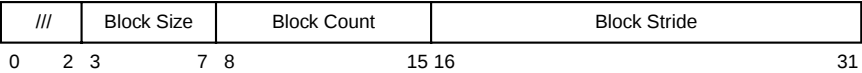


Figure 4-46. Format of b Type (32-bit)

The b type is encoded for 64-bit as follows:

- Block size: b[35:39] if b[35:39] != 0; otherwise 32
- Block count: b[40:47] if b[40:47] != 0; otherwise 256
- Block stride: b[48:63] if b[48:63] != 0; otherwise 32768

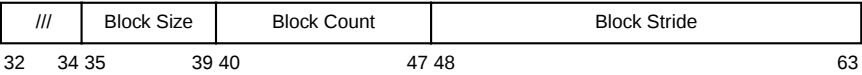


Figure 4-47. Format of b Type (64-bit)

The c type is a 2-bit unsigned literal tag used to identify a specific data stream. Up to four streams can be set up with this mechanism.

The valid combinations of argument types for vec_dstst(a,b,c) are shown in Table 4-6. The result type is void.

Table 4-6. vec_dstst—Vector Data Stream for Touch Store Argument Types

a	b	c	maps to
vector unsigned char *	any integral type	2-bit unsigned literal	dstst a,b,c
vector signed char *	any integral type	2-bit unsigned literal	
vector bool char *	any integral type	2-bit unsigned literal	
vector unsigned short *	any integral type	2-bit unsigned literal	
vector signed short *	any integral type	2-bit unsigned literal	
vector bool short *	any integral type	2-bit unsigned literal	
vector pixel *	any integral type	2-bit unsigned literal	
vector unsigned int *	any integral type	2-bit unsigned literal	
vector signed int *	any integral type	2-bit unsigned literal	
vector bool int *	any integral type	2-bit unsigned literal	
vector float *	any integral type	2-bit unsigned literal	
unsigned char *	any integral type	2-bit unsigned literal	
signed char *	any integral type	2-bit unsigned literal	
unsigned short *	any integral type	2-bit unsigned literal	
short *	any integral type	2-bit unsigned literal	
unsigned int *	any integral type	2-bit unsigned literal	
int *	any integral type	2-bit unsigned literal	
unsigned int *	any integral type	2-bit unsigned literal	
float *	any integral type	2-bit unsigned literal	

vec_dststt

Vector Data Stream Touch for Store Transient

vec_dststt

vec_dststt(a,b,c)

```
addr[0:63] ← a
DataStreamPrefetchControl ← "start" || 1 || static || b || addr
```

Each operation initiates cache touches for transient stores for the data stream associated with tag c at the address a using the data block in b. The result type is void.

The a type may also be a pointer to a const-qualified type. Plain char * is excluded in the mapping for a.

The b type is encoded for 32-bit as follows:

- Block size: b[3:7] if b[3:7] != 0; otherwise 32
- Block count: b[8:15] if b[8:15] != 0; otherwise 256
- Block stride: b[16:31] if b[16:31] != 0; otherwise 32768

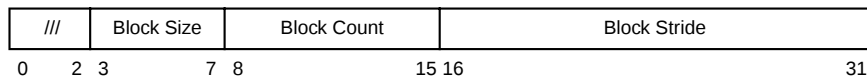


Figure 4-48. Format of b Type (32-bit)

The b type is encoded for 64-bit as follows:

- Block size: b[35:39] if b[35:39] != 0; otherwise 32
- Block count: b[40:47] if b[40:47] != 0; otherwise 256
- Block stride: b[48:63] if b[48:63] != 0; otherwise 32768

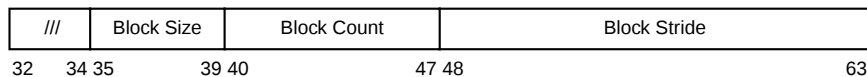


Figure 4-49. Format of b Type (64-bit)

The c type is a 2-bit unsigned literal tag used to identify a specific data stream. Up to four streams can be set up with this mechanism.

The valid combinations of argument types for vec_dststt(a,b,c) are shown in Table 4-7. The result type is void.

Table 4-7. vec_dststt—Vector Data Stream Touch for Store Transient Argument Types

a	b	c	maps to
vector unsigned char *	any integral type	2-bit unsigned literal	dststt a,b,c
vector signed char *	any integral type	2-bit unsigned literal	
vector bool char *	any integral type	2-bit unsigned literal	
vector unsigned short *	any integral type	2-bit unsigned literal	
vector signed short *	any integral type	2-bit unsigned literal	
vector bool short *	any integral type	2-bit unsigned literal	
vector pixel *	any integral type	2-bit unsigned literal	
vector unsigned int *	any integral type	2-bit unsigned literal	
vector signed int *	any integral type	2-bit unsigned literal	
vector bool int *	any integral type	2-bit unsigned literal	
vector float *	any integral type	2-bit unsigned literal	
unsigned char *	any integral type	2-bit unsigned literal	
signed char *	any integral type	2-bit unsigned literal	
unsigned short *	any integral type	2-bit unsigned literal	
short *	any integral type	2-bit unsigned literal	
unsigned int *	any integral type	2-bit unsigned literal	
int *	any integral type	2-bit unsigned literal	
unsigned int *	any integral type	2-bit unsigned literal	
float *	any integral type	2-bit unsigned literal	

vec_dstt

Vector Data Stream Touch Transient

vec_dstt

vec_dstt(a,b,c)

```
addr[0:63] ← a
DataStreamPrefetchControl ← "start" || c || 1 || b || addr
```

Each operation initiates cache touches for transient loads for the data stream associated with tag c at the address a using the data block in b. The result type is void.

The a type may also be a pointer to a const-qualified type. Plain char * is excluded in the mapping for a.

The b type is encoded for 32-bit as follows:

- Block size: b[3:7] if b[3:7] != 0; otherwise 32
- Block count: b[8:15] if b[8:15] != 0; otherwise 256
- Block stride: b[16:31] if b[16:31] != 0; otherwise 32768

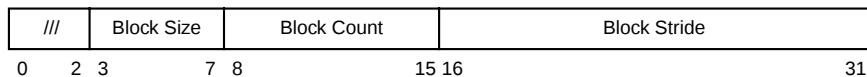


Figure 4-50. Format of b Type (32-bit)

The b type is encoded for 64-bit as follows:

- Block size: b[35:39] if b[35:39] != 0; otherwise 32
- Block count: b[40:47] if b[40:47] != 0; otherwise 256
- Block stride: b[48:63] if b[48:63] != 0; otherwise 32768

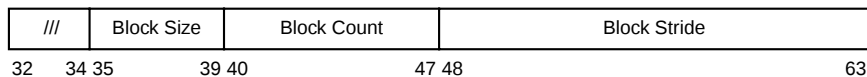


Figure 4-51. Format of b Type (64-bit)

The c type is a 2-bit unsigned literal tag used to identify a specific data stream. Up to four streams can be set up with this mechanism.

The valid combinations of argument types for vec_dstt(a,b,c) are shown in Table 4-8. The result type is void.

Table 4-8. vec_dstt—Vector Data Stream Touch Transient Argument Types

a	b	c	maps to
vector unsigned char *	any integral type	2-bit unsigned literal	dst a,b,c
vector signed char *	any integral type	2-bit unsigned literal	
vector bool char *	any integral type	2-bit unsigned literal	
vector unsigned short *	any integral type	2-bit unsigned literal	
vector signed short *	any integral type	2-bit unsigned literal	
vector bool short *	any integral type	2-bit unsigned literal	
vector pixel *	any integral type	2-bit unsigned literal	
vector unsigned int *	any integral type	2-bit unsigned literal	
vector signed int *	any integral type	2-bit unsigned literal	
vector bool int *	any integral type	2-bit unsigned literal	
vector float *	any integral type	2-bit unsigned literal	
unsigned char *	any integral type	2-bit unsigned literal	
signed char *	any integral type	2-bit unsigned literal	
unsigned short *	any integral type	2-bit unsigned literal	
short *	any integral type	2-bit unsigned literal	
unsigned int *	any integral type	2-bit unsigned literal	
int *	any integral type	2-bit unsigned literal	
unsigned int *	any integral type	2-bit unsigned literal	
float *	any integral type	2-bit unsigned literal	

vec_expte

vec_expte

Vector Is 2 Raised to the Exponent Estimate Floating-Point

```
d = vec_expte(a)
do i=0 to 3
  di ← FP2XEst(ai)
end
```

Each element of the result is an estimate of 2 raised to the corresponding element of a.

If VSCR[NJ] = 1, every denormalized operand element is truncated to a 0 of the same sign before the operation is carried out, and each denormalized result element is truncated to a 0 of the same sign.

The valid argument type and corresponding result type for **d** = vec_expte(**a**) are shown in Figure 4-52.

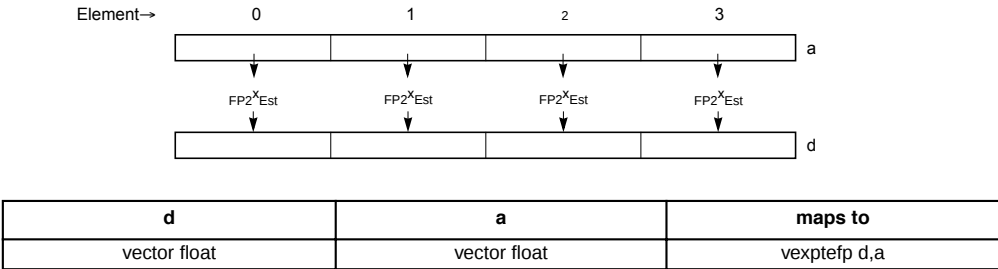


Figure 4-52. 2 Raised to the Exponent Estimate Floating-Point for Four Floating-Point Elements (32-Bit)

vec_floor

Vector Floor

d = vec_floor(**a**)

```
do i=0 to 3
  di ← Floor(ai)
end
```

Each single-precision floating-point word element in **a** is rounded to a single-precision floating-point integer using the rounding mode Round towards $-\infty$, and placed into the corresponding word element of **d**. Each element of the result is thus the largest representable floating-point integer not greater than **a**. For example, if the floating-point element was 123.85, the resulting integer would be 123.

If VSCR[NJ] = 1, every denormalized operand element is truncated to 0 before rounding.

The valid argument type and corresponding result type for **d** = vec_floor(**a**) are shown in Figure 4-53.

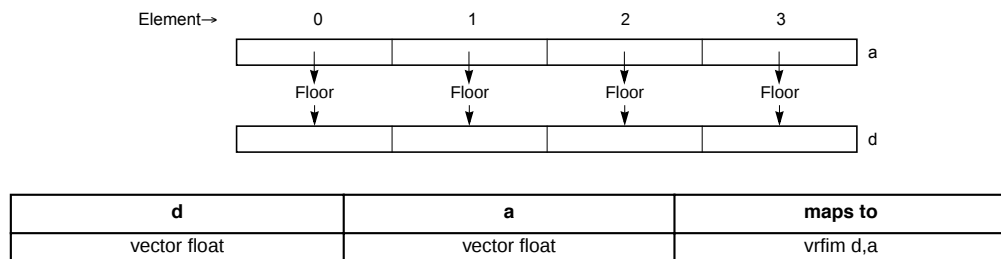


Figure 4-53. Round to Minus Infinity of Four Floating-Point Integer Elements (32-Bit)

vec_ld

Vector Load Indexed

vec_ld

d = vec_ld(**a**,**b**)

```
EA ← BoundAlign(a+b, 16)
d ← MEM(EA, 16)
```

Each operation performs a 16-byte load at a 16-byte aligned address. The **a** is taken to be an integer value, while **b** is a pointer. BoundAlign(**a**+**b**,16) is the largest value less than or equal to **a** + **b** that is a multiple of 16. This load is the one that is generated for a loading dereference of a pointer to a vector type. The **b** type may also be a pointer to a const-qualified type. Plain `char *` is excluded in the mapping for **b**. The valid combinations of argument types and the corresponding result types for **d** = vec_ld(**a**,**b**) are shown in Table 4-9.

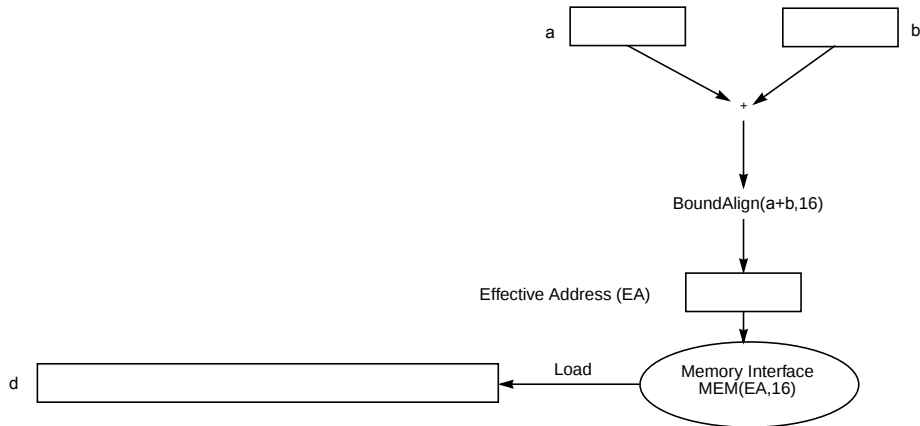


Figure 4-54. Vector Load Indexed Operation

Table 4-9. vec_id—Load Vector Indexed Argument Types

d	a	b	maps to
vector unsigned char	any integral type	vector unsigned char *	lvs d,a,b
	any integral type	unsigned char *	
vector signed char	any integral type	vector signed char *	
	any integral type	signed char *	
vector bool char	any integral type	vector bool char *	
vector unsigned short	any integral type	vector unsigned short *	
	any integral type	unsigned short *	
vector signed short	any integral type	vector signed short *	
	any integral type	short *	
vector bool short	any integral type	vector bool short *	
vector pixel	any integral type	vector pixel *	
vector unsigned int	any integral type	vector unsigned int *	
	any integral type	unsigned int*	
	any integral type	unsigned int *	
vector signed int	any integral type	vector signed int *	
	any integral type	int *	
	any integral type	int *	
vector bool int	any integral type	vector bool int *	
vector float	any integral type	vector float *	
	any integral type	float *	

vec_lde

Vector Load Element Indexed

vec_lde

d = vec_lde(**a**,**b**)

```

s ← 16/(number of elements)
EA ← BoundAlign(a+b,s)
i ← mod(EA,16)/s
di ← MEM(EA,s)

```

Each operation loads a single element into the position in the vector register corresponding to its address, leaving the remaining elements of the register undefined. The **a** is taken to be an integer value, while **b** is a pointer. BoundAlign(**a**+**b**,**s**) is the largest value less than or equal to **a** + **b** that is a multiple of **s**, where **s** is 1 for char pointers, 2 for short pointers, and 4 for int or float pointers. The **b** type may also be a pointer to a const-qualified type. Plain char * is excluded in the mapping for **b**. The valid combinations of argument types and the corresponding result types for **d** = vec_lde(**a**,**b**) are shown in Table 4-10.

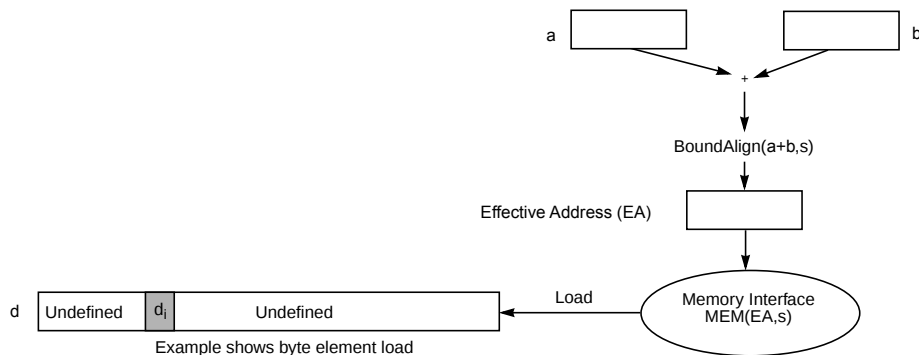


Figure 4-55. Vector Load Element Indexed Operation

Table 4-10. vec_lde(a,b)—Vector Load Element Indexed Argument Types

d	a	b	Maps to
vector unsigned char	any integral type	unsigned char *	lvebx d,a,b
vector signed char	any integral type	signed char *	
vector unsigned short	any integral type	unsigned short *	lvehx d,a,b
vector signed short	any integral type	short *	
vector unsigned int	any integral type	unsigned int *	lvewx d,a,b
	any integral type	unsigned int *	
vector signed int	any integral type	int *	
vector float	any integral type	float *	

vec_ldl

Vector Load Indexed LRU

d = vec_ldl(**a**,**b**)

```
EA ← BoundAlign(a+b,16)
d ← MEM(EA,16)
```

Each operation performs a 16-byte load at a 16-byte aligned address. The **a** is taken to be an integer value, while **b** is a pointer. BoundAlign(**a**+**b**,16) is the largest value less than or equal to **a** + **b** that is a multiple of 16. These operations mark the cache line as least-recently-used. The **b** type may also be a pointer to a `const`-qualified type. Plain `char *` is excluded in the mapping for **b**. The valid combinations of argument types and the corresponding result types for **d** = vec_ldl(**a**,**b**) are shown in Table 4-11.

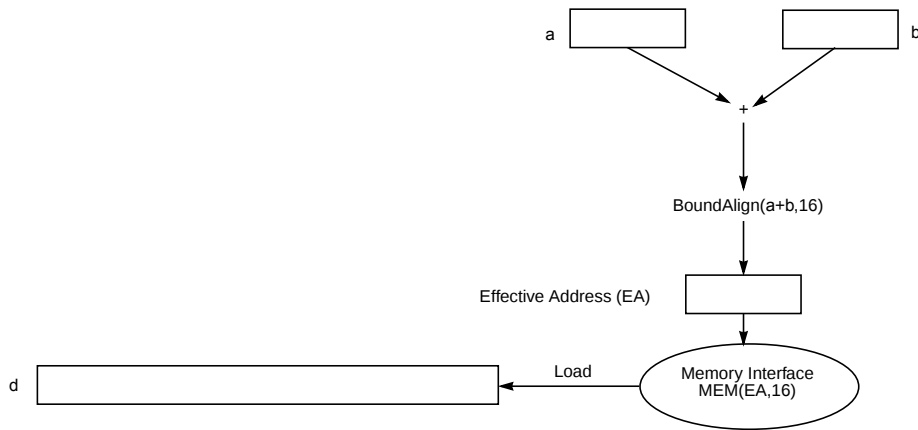


Figure 4-56. Vector Load Indexed LRU Operation

Table 4-11. vec_idl—Vector Load Indexed LRU Argument Types

d	a	b	Maps to
vector unsigned char	any integral type	vector unsigned char *	lxl d,a,b
	any integral type	unsigned char *	
vector signed char	any integral type	vector signed char *	
	any integral type	signed char *	
vector bool char	any integral type	vector bool char *	
vector unsigned short	any integral type	vector unsigned short *	
	any integral type	unsigned short *	
vector signed short	any integral type	vector signed short *	
	any integral type	short *	
vector bool short	any integral type	vector bool short *	
vector pixel	any integral type	vector pixel *	
vector unsigned int	any integral type	vector unsigned int *	
	any integral type	unsigned int *	
vector signed int	any integral type	vector signed int *	
	any integral type	int *	
vector bool int	any integral type	vector bool int *	
vector float	any integral type	vector float *	
	any integral type	float *	

vec_loge

Vector Log₂ Estimate Floating-Point

vec_loge

d = vec_loge(**a**)

```
do i=0 to 3
  di ← FPLog2Est(ai)
end
```

Each element of the result is an estimate of the logarithm to base 2 of the corresponding element of **a**.

If VSCR[NJ] = 1, every denormalized operand element is truncated to a 0 of the same sign before the operation is carried out.

The valid argument type and corresponding result type for **d** = vec_loge(**a**) are shown in Figure 4-57

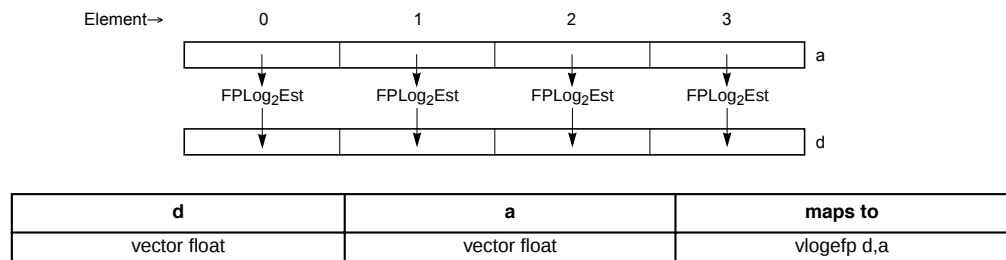


Figure 4-57. Log₂ Estimate Floating-Point for Four Floating-Point Elements (32-Bit)



vec_lvsl

Vector Load for Shift Left

vec_lvsl

d = vec_lvsl(**a**,**b**)

```

EA ← a + b
sh ← EA[28:31]
if sh = 0x0 then d ← 0x000102030405060708090A0B0C0D0E0F
if sh = 0x1 then d ← 0x0102030405060708090A0B0C0D0E0F10
if sh = 0x2 then d ← 0x02030405060708090A0B0C0D0E0F1011
if sh = 0x3 then d ← 0x030405060708090A0B0C0D0E0F101112
if sh = 0x4 then d ← 0x0405060708090A0B0C0D0E0F10111213
if sh = 0x5 then d ← 0x05060708090A0B0C0D0E0F1011121314
if sh = 0x6 then d ← 0x060708090A0B0C0D0E0F101112131415
if sh = 0x7 then d ← 0x0708090A0B0C0D0E0F10111213141516
if sh = 0x8 then d ← 0x08090A0B0C0D0E0F1011121314151617
if sh = 0x9 then d ← 0x090A0B0C0D0E0F101112131415161718
if sh = 0xA then d ← 0x0A0B0C0D0E0F10111213141516171819
if sh = 0xB then d ← 0x0B0C0D0E0F101112131415161718191A
if sh = 0xC then d ← 0x0C0D0E0F101112131415161718191A1B
if sh = 0xD then d ← 0x0D0E0F101112131415161718191A1B1C
if sh = 0xE then d ← 0x0E0F101112131415161718191A1B1C1D
if sh = 0xF then d ← 0x0F101112131415161718191A1B1C1D1E

```

Each operation generates a permutation useful for aligning data from an unaligned address. The **b** type may also be a pointer to a `const-` or `volatile-`qualified type. Plain `char *` is excluded in the mapping for **b**. The valid combination of argument types and the corresponding result type for **d** = vec_lvsl(**a**,**b**) are shown in Table 4-12.

Table 4-12. vec_lvsl—Load Vector for Shift Left Argument Types

d	a	b	maps to
vector unsigned char	any integral type	unsigned char *	lvsl d,a,b
	any integral type	signed char *	
	any integral type	unsigned short *	
	any integral type	short *	
	any integral type	unsigned int *	
	any integral type	int *	
	any integral type	float *	

vec_lvsr

Vector Load Shift Right

vec_lvsr

d = vec_lvsr(**a**,**b**)

```

EA ← a + b
sh ← EA[28:31]
if sh=0x0 then d ← 0x101112131415161718191A1B1C1D1E1F
if sh=0x1 then d ← 0x0F101112131415161718191A1B1C1D1E
if sh=0x2 then d ← 0x0E0F101112131415161718191A1B1C1D
if sh=0x3 then d ← 0x0D0E0F101112131415161718191A1B1C
if sh=0x4 then d ← 0x0C0D0E0F101112131415161718191A1B
if sh=0x5 then d ← 0x0B0C0D0E0F101112131415161718191A
if sh=0x6 then d ← 0x0A0B0C0D0E0F10111213141516171819
if sh=0x7 then d ← 0x090A0B0C0D0E0F101112131415161718
if sh=0x8 then d ← 0x08090A0B0C0D0E0F1011121314151617
if sh=0x9 then d ← 0x0708090A0B0C0D0E0F10111213141516
if sh=0xA then d ← 0x060708090A0B0C0D0E0F101112131415
if sh=0xB then d ← 0x05060708090A0B0C0D0E0F1011121314
if sh=0xC then d ← 0x0405060708090A0B0C0D0E0F10111213
if sh=0xD then d ← 0x030405060708090A0B0C0D0E0F101112
if sh=0xE then d ← 0x02030405060708090A0B0C0D0E0F1011
if sh=0xF then d ← 0x0102030405060708090A0B0C0D0E0F10

```

Each operation generates a permutation useful for aligning data from an unaligned address. The **b** type may also be a pointer to a `const-` or `volatile-`qualified type. Plain `char *` is excluded in the mapping for **b**. The valid combinations of argument types and the corresponding result type for **d** = vec_lvsr(**a**,**b**) are shown in Table 4-13.

Table 4-13. vec_lvsr—Vector Load for Shift Right Argument Types

d	a	b	Maps to
vector unsigned char	any integral type	unsigned char *	lvsr d,a,b
	any integral type	signed char *	
	any integral type	unsigned short *	
	any integral type	short *	
	any integral type	unsigned int *	
	any integral type	int *	
	any integral type	float *	

vec_madd

Vector Multiply Add

vec_madd

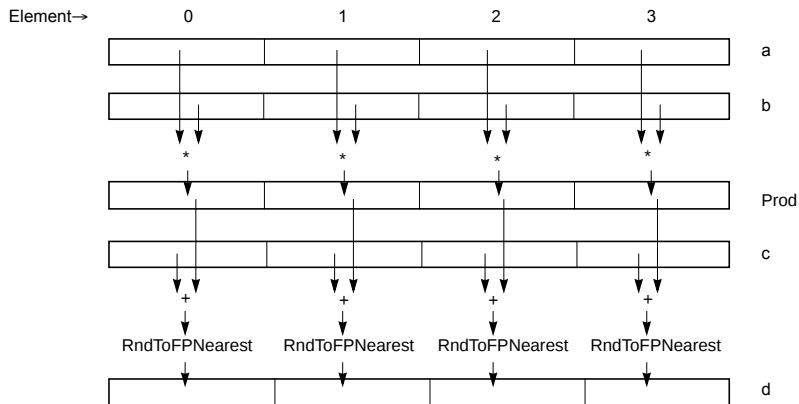
d = vec_madd(**a**,**b**,**c**)

```
do i=0 to 3
  di ← RndToFPNearest(ai * bi + ci)
end
```

Each element of the result is the sum of the corresponding element of **c** and the product of the corresponding elements of **a** and **b**.

If VSCR[NJ] = 1, every denormalized operand element is truncated to a 0 of the same sign before the operation is carried out, and each denormalized result element truncates to a 0 of the same sign.

The valid argument types and the corresponding result type for **d** = vec_madd(**a**,**b**,**c**) are shown in Figure 4-58



d	a	b	c	maps to
vector float	vector float	vector float	vector float	vmaddfp d,a,b,c

Figure 4-58. Multiply-Add Four Floating-Point Elements (32-Bit)

vec_madds

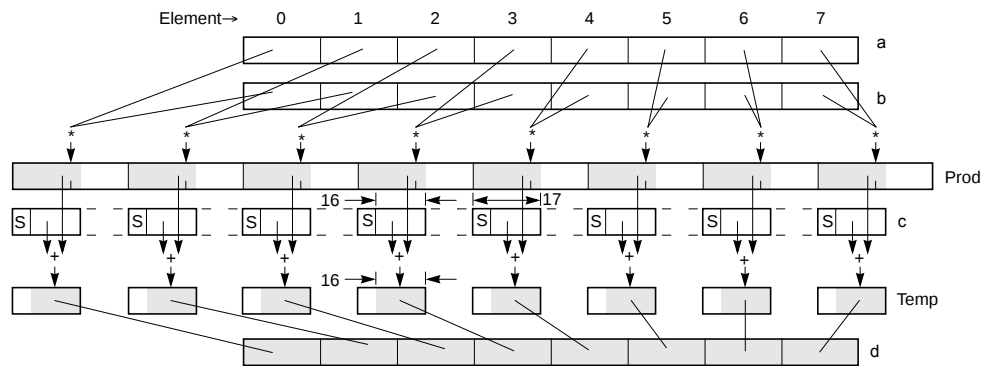
Vector Multiply Add Saturated

vec_madds

d = vec_madds(**a**,**b**,**c**)

```
do i=0 to 7
  di ← Saturate((ai * bi)/215 + ci)
end
```

Each element of the result is the 16-bit saturated sum of the corresponding element of **c** and the high-order 17 bits of the product of the corresponding elements of **a** and **b**. If saturation occurs, VSCR[SAT] is set (see Table 4-1). The valid argument types and the corresponding result type for **d** = vec_madds(**a**,**b**,**c**) are shown in Figure 4-59.



d	a	b	c	maps to
vector signed short	vector signed short	vector signed short	vector signed short	vmhaddshs d,a,b,c

Figure 4-59. Multiply-Add Four Floating-Point Elements (32-Bit)



vec_max

Vector Maximum

vec_max

d = vec_max(**a**,**b**)

```
n ← number of elements
do i=0 to n-1
  di ← MAX(ai,bi)
end
```

Each element of the result is the larger of the corresponding elements of **a** and **b**.

For `vector float` argument types, if VSCR[NJ] is set, every denormalized operand element is truncated to a 0 of the same sign before the operation is carried out, and each denormalized result element truncates to a 0 of the same sign. The maximum of +0.0 and -0.0 is +0.0. The maximum of any value and a NaN is a QNaN.

The valid combinations of argument types and the corresponding result types for **d** = vec_max(**a**,**b**) are shown in Figure 4-60, Figure 4-61, Figure 4-62, and Figure 4-63.

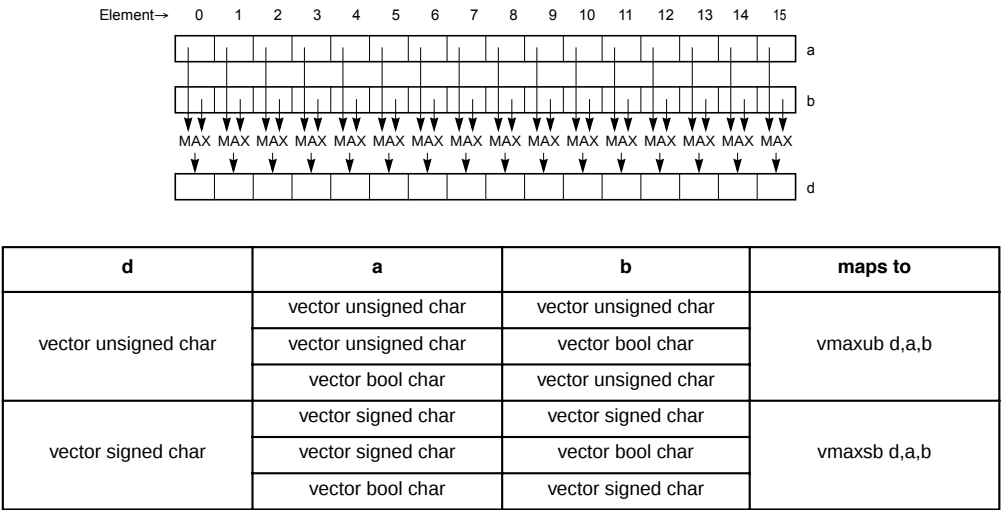


Figure 4-60. Maximum of Sixteen Integer Elements (8-Bit)

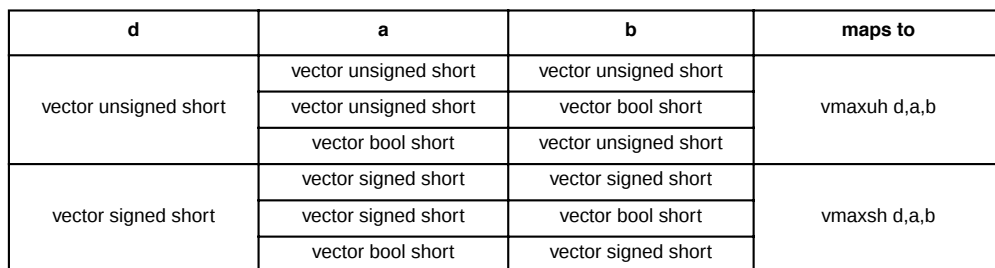


Figure 4-61. Maximum of Eight Integer Elements (16-bit)

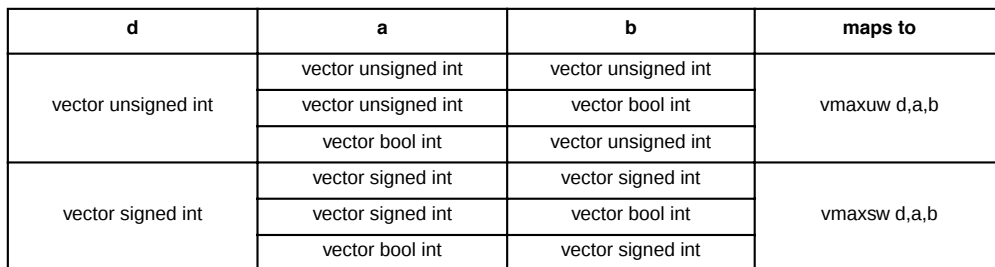


Figure 4-62. Maximum of Four Integer Elements (32-bit)

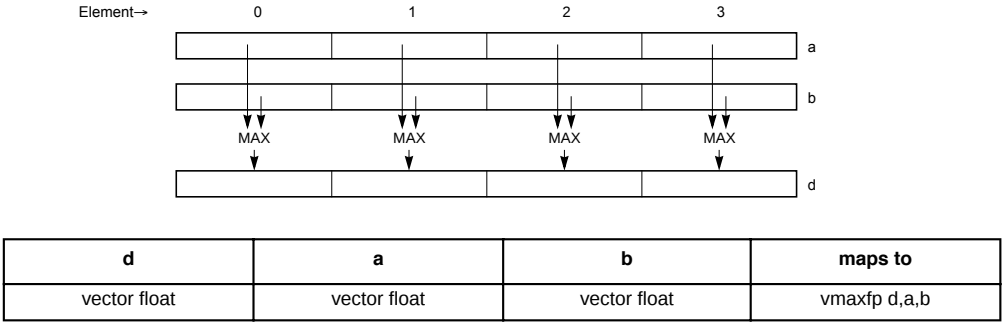


Figure 4-63. Maximum of Four Floating-Point Elements (32-bit)

vec_mergeh

Vector Merge High

vec_mergeh

d = vec_mergeh(**a**,**b**)

```

m ← (number of elements)/2
do i=0 to m-1
  d2i ← ai
  d2i+1 ← bi
end

```

The even elements of the result are obtained left-to-right from the high elements of **a**. The odd elements of the result are obtained left-to-right from the high elements of **b**. The valid combinations of argument types and the corresponding result types for **d** = vec_mergeh(**a**,**b**) are shown in Figure 4-64, Figure 4-65, and Figure 4-66.

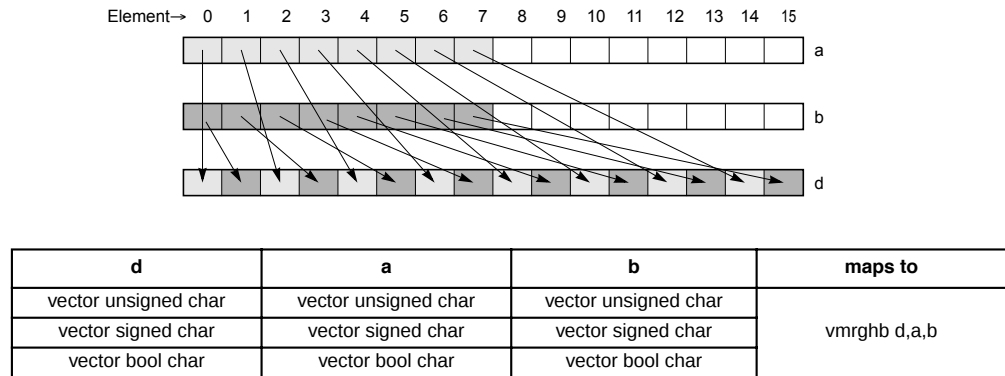


Figure 4-64. Merge Eight High-Order Elements (8-Bit)

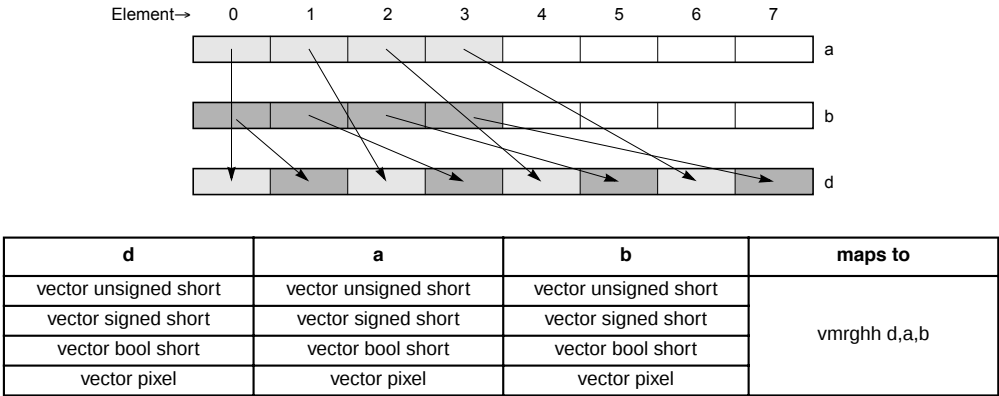


Figure 4-65. Merge Four High-Order Elements (16-bit)

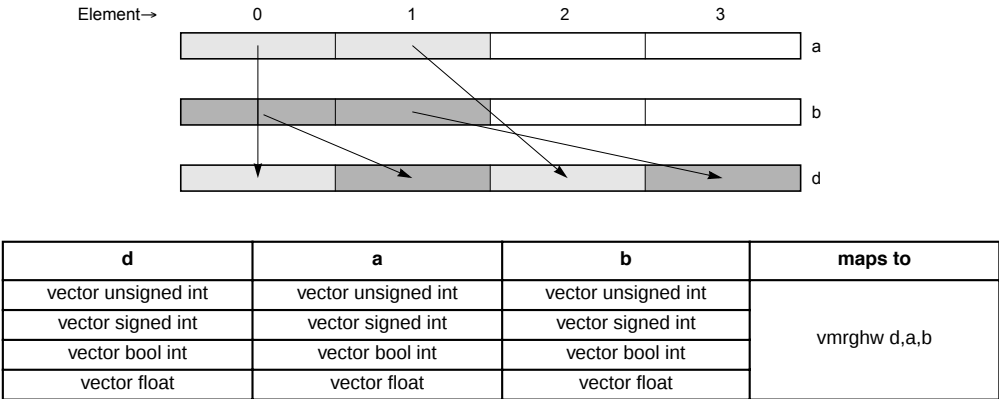


Figure 4-66. Merge Two High-Order Elements (32-bit)

vec_mergel

Vector Merge Low

vec_mergel

d = vec_mergel(**a**,**b**)

```

m ← (number of elements)/2
do i=0 to m-1
  d2i ← ai+m
  d2i+1 ← bi+m
end

```

The even elements of the result are obtained left-to-right from the low elements of **a**. The odd elements of the result are obtained left-to-right from the low elements of **b**. The valid combinations of argument types and the corresponding result types for **d** = vec_mergel(**a**,**b**) are shown in Figure 4-67, Figure 4-68, and Figure 4-69.

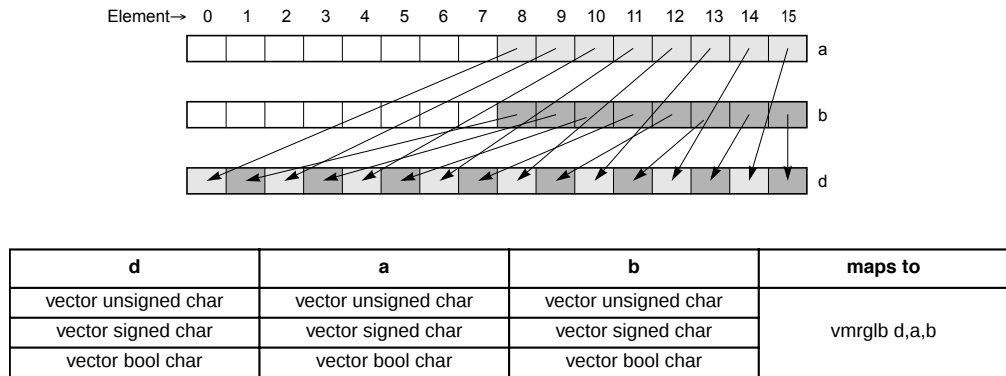


Figure 4-67. Merge Eight Low-Order Elements (8-Bit)

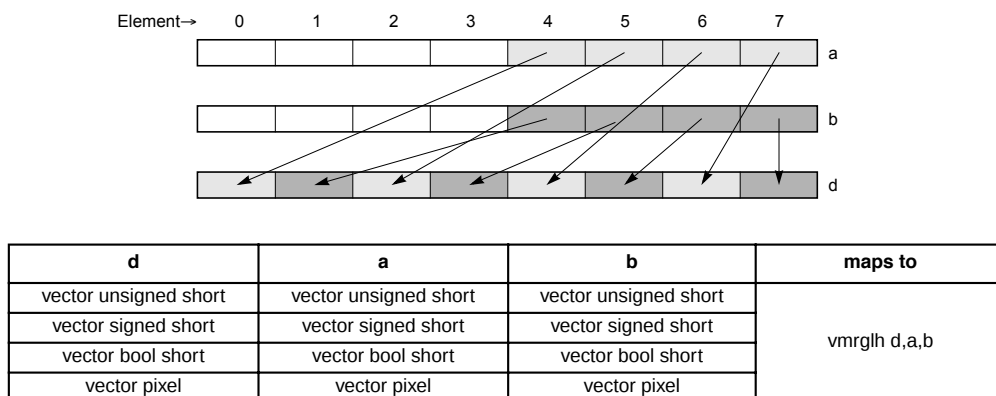


Figure 4-68. Merge Four Low-Order Elements (16-bit)

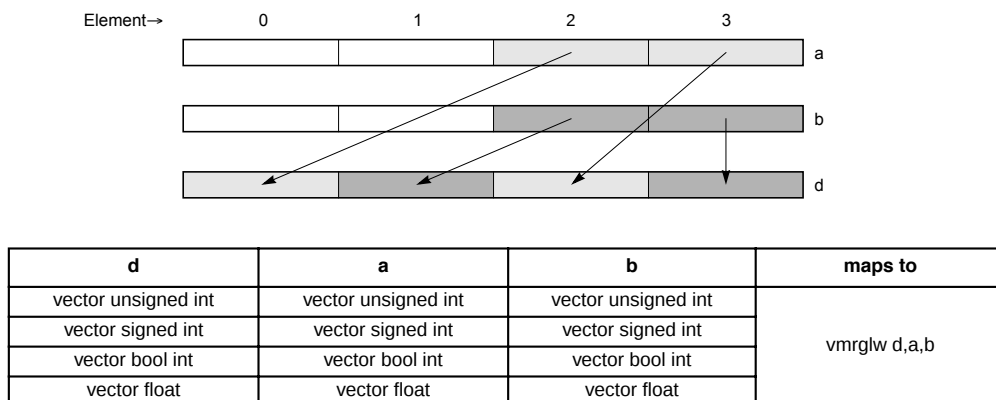


Figure 4-69. Merge Two Low-Order Elements (32-bit)

vec_mfvscr

Vector Move from Vector Status and Control Register

vec_mfvscr

d = vec_mfvscr

$$d \leftarrow {}^{96}0 \parallel (VSCR)$$

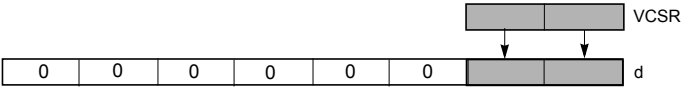


Figure 4-70. Vector Move from VSCR

Table 4-14. Vector Move from Vector Status and Control Registers Argument Type and Mapping

d	Maps to
vector unsigned short	mfvscr



vec_min

Vector Minimum

vec_min

d = vec_min(**a**,**b**)

```
n ← number of elements
do i=0 to n-1
  di ← MIN(ai,bi)
end
```

Each element of the result is the smaller of the corresponding elements of a and b.

For `vector float` argument types, if VSCR[NJ] is set, every denormalized operand element is truncated to a 0 of the same sign before the operation is carried out, and each denormalized result element truncates to a 0 of the same sign. The minimum of +0.0 and -0.0 is -0.0. The minimum of any value and a NaN is a QNaN.

The valid combinations of argument types and the corresponding result types for **d** = vec_min(**a**,**b**) are shown in Figure 4-71, Figure 4-72, Figure 4-73, and Figure 4-74.

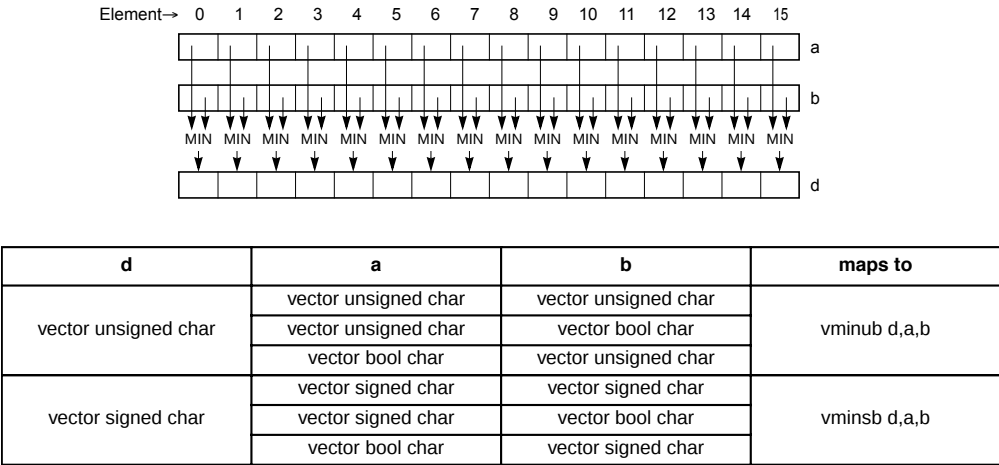


Figure 4-71. Minimum of Sixteen Integer Elements (8-Bit)

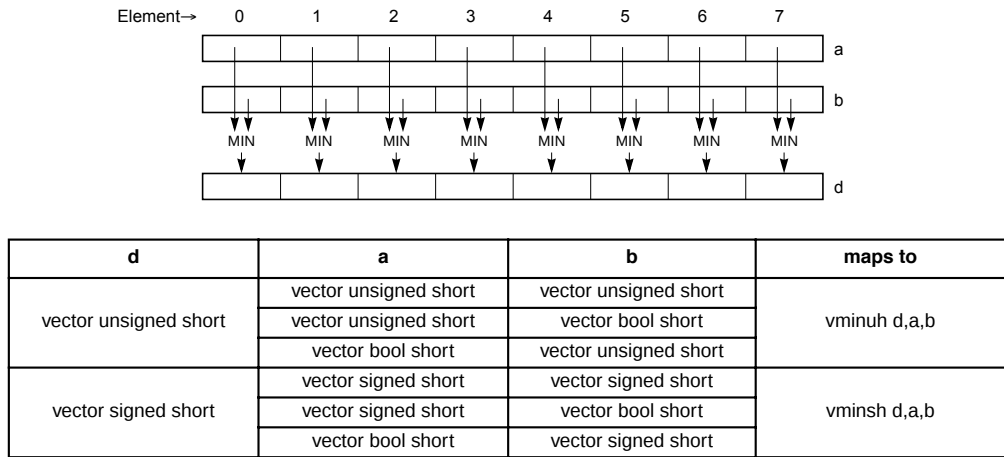


Figure 4-72. Minimum of Eight Integer Elements (16-bit)

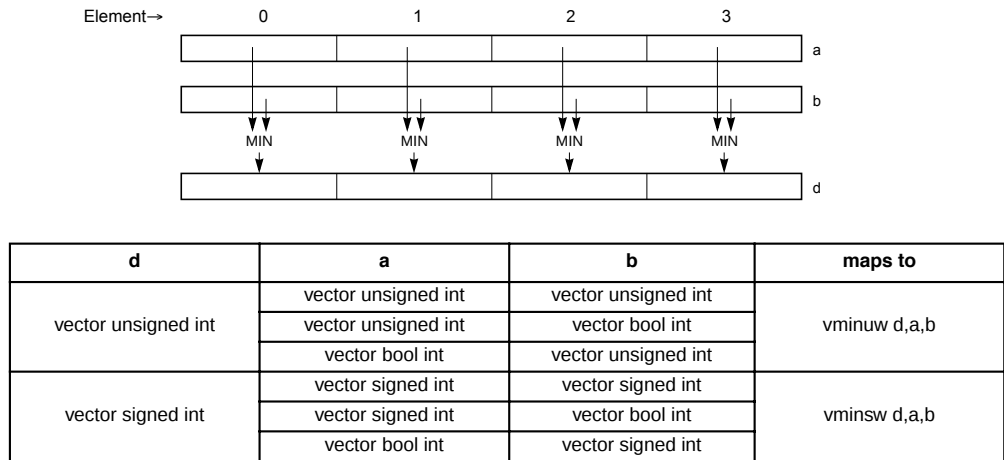


Figure 4-73. Minimum of Four Integer Elements (32-bit)

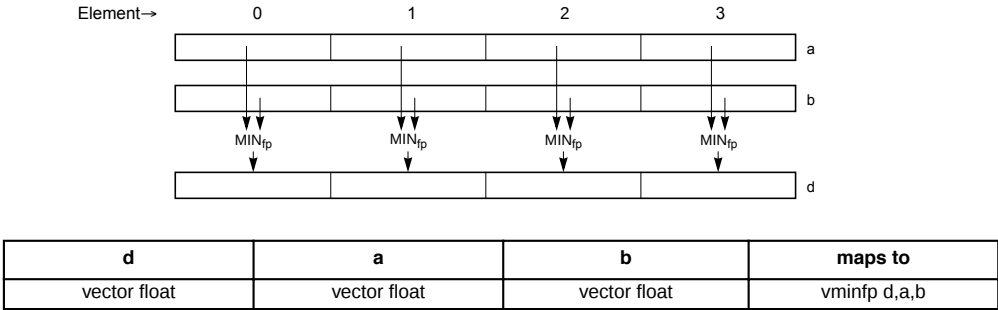


Figure 4-74. Minimum of Four Floating-Point Elements (32-bit)

vec_mladd

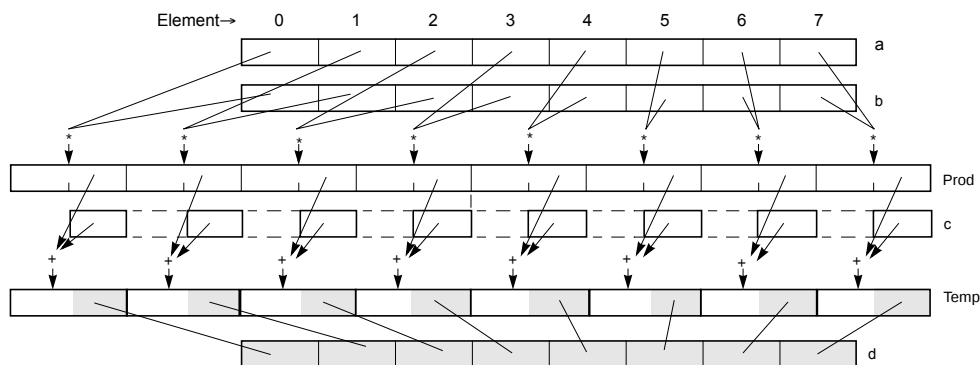
Vector Multiply Low and Add Unsigned Half Word

vec_mladd

d = vec_mladd(**a**,**b**,**c**)

```
do i=0 to 7
  di ← (ai * bi) + ci
end
```

Each element of the result is the low-order 16 bits of the sum of the corresponding element of **c** and the product of the corresponding elements of **a** and **b**. The valid combinations of argument types and the corresponding result types for **d** = vec_mladd(**a**,**b**) are shown in Figure 4-75.



d	a	b	c	maps to
vector unsigned short	vector unsigned short	vector unsigned short	vector unsigned short	vmladduhm d,a,b,c
vector signed short	vector unsigned short	vector signed short	vector signed short	
	vector signed short	vector unsigned short	vector unsigned short	
	vector signed short	vector signed short	vector signed short	

Figure 4-75. Multiply-Add of Eight Integer Elements (16-Bit)

vec_mradds

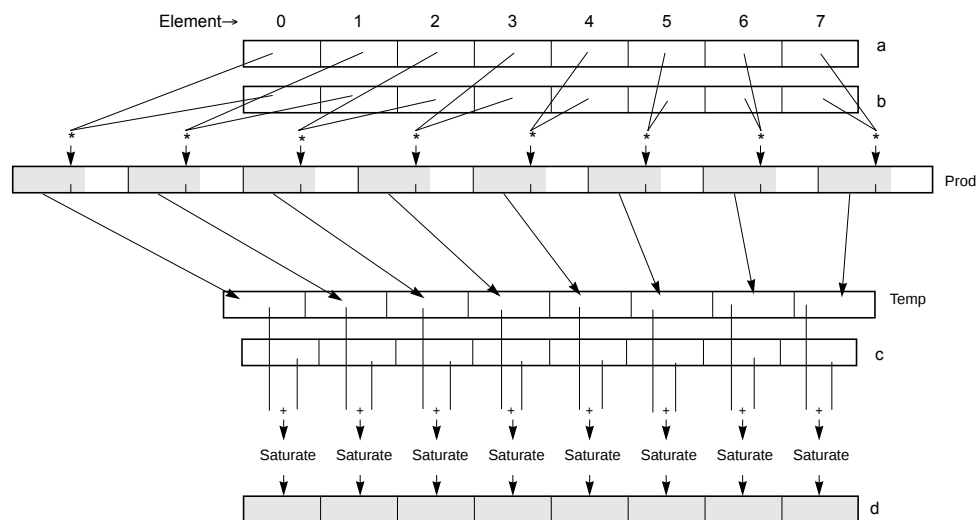
Vector Multiply Round and Add Saturated

vec_mradds

$\mathbf{d} = \text{vec_mradds}(\mathbf{a}, \mathbf{b}, \mathbf{c})$

```
do i=0 to 7
   $d_i \leftarrow \text{Saturate}((a_i * b_i + 2^{14})/2^{15} + c_i)$ 
end
```

Each element of the result is the 16-bit saturated sum of the corresponding element of $\mathbf{c}$ and the high-order 17 bits of the rounded product of the corresponding elements of $\mathbf{a}$ and $\mathbf{b}$. If saturation occurs, VSCR[SAT] is set (see Table 4-1). The valid argument types and the corresponding result type for $\mathbf{d} = \text{vec_mradds}(\mathbf{a}, \mathbf{b}, \mathbf{c})$ are shown in Figure 4-76.



d	a	b	c	maps to
vector signed short	vector signed short	vector signed short	vector signed short	<code>vmhraddshs d,a,b,c</code>

Figure 4-76. Multiply-Add of Eight Integer Elements (16-Bit)

vec_msum

Vector Multiply Sum

vec_msum

d = vec_msum(**a**,**b**,**c**)

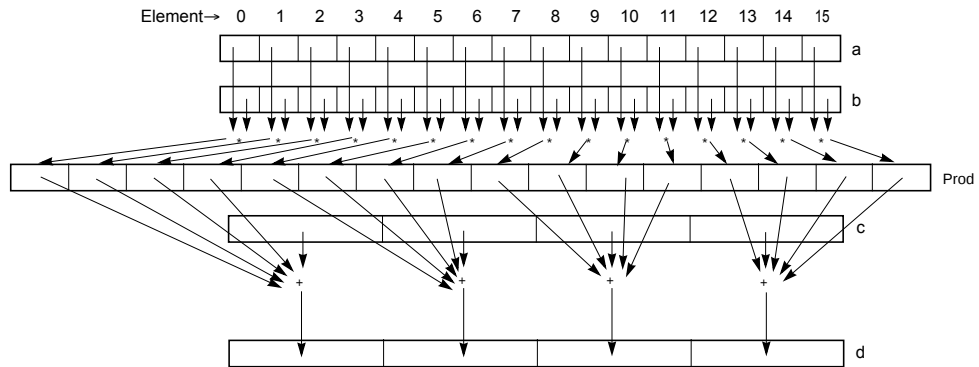
- For Multiply Sum of Sixteen 8-bit elements

```
do i=0 to 3
  di ← (a4i * b4i) + (a4i+1 * b4i+1) + (a4i+2 * b4i+2) + (a4i+3 * b4i+3) + ci
end
```

- For Multiply Sum of Eight 16-bit elements

```
do i=0 to 3
  di ← (a2i * b2i) + (a2i+1 * b2i+1) + ci
end
```

Each element of the result is the sum of the corresponding element of **c** and the products of the elements of **a** and **b** which overlap the positions of that element of **c**. For **vec_msum**, the sum is performed with 32-bit modular addition. The valid combinations of argument types and the corresponding result types for **d** = **vec_msum(a,b,c)** are shown in Figure 4-77 and Figure 4-78.



d	a	b	c	maps to
vector unsigned int	vector unsigned char	vector unsigned char	vector unsigned int	vmsumubm d,a,b,c
vector signed int	vector signed char	vector unsigned char	vector signed int	vmsummbm d,a,b,c

Figure 4-77. Multiply Sum of Sixteen Integer Elements (8-Bit)

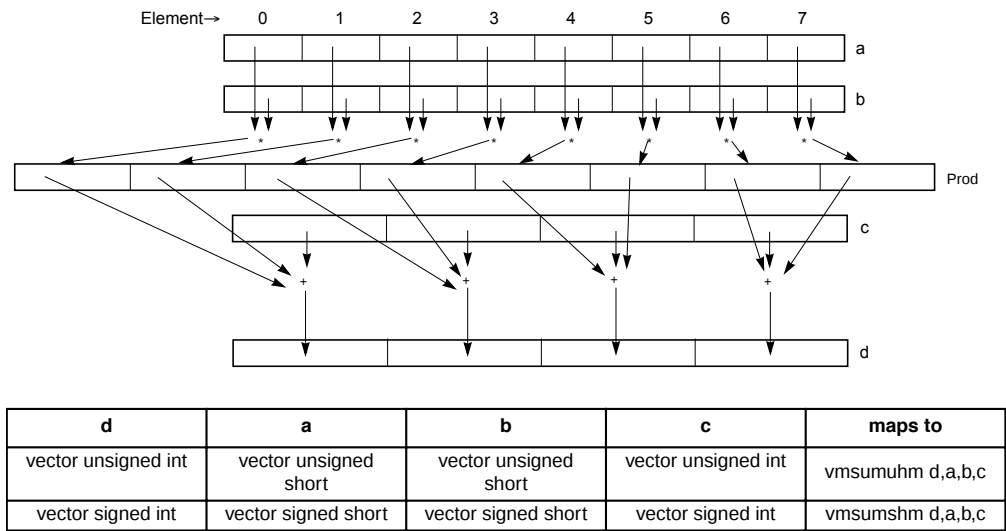


Figure 4-78. Multiply Sum of Eight Integer Elements (16-Bit)

vec_msums

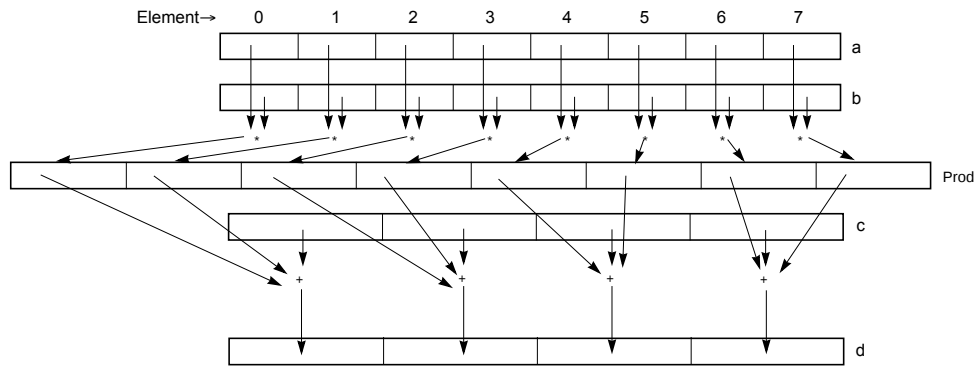
Vector Multiply Sum Saturated

vec_msums

d = vec_msums(**a**,**b**,**c**)

```
do i=0 to 3
  di ← Saturate((a2i * b2i) + (a2i+1 * b2i+1) + ci)
end
```

Each element of the result is the sum of the corresponding element of **c** and the products of the elements of **a** and **b** which overlap the positions of that element of **c**. The sum is performed with 32-bit saturating addition. If saturation occurs, VSCR[SAT] is set (see Table 4-1). The valid combinations of argument types and the corresponding result types for **d** = vec_msums(**a**,**b**,**c**) are shown in Figure 4-79.



d	a	b	c	maps to
vector unsigned int	vector unsigned short	vector unsigned short	vector unsigned int	vmsumuhs d,a,b,c
vector signed int	vector signed short	vector signed short	vector signed int	vmsumshs d,a,b,c

Figure 4-79. Multiply-Sum of Integer Elements (16-Bit to 32-Bit)

vec_mtvscr

vec_mtvscr

Vector Move to Vector Status and Control Register

vec_mtvscr(a)

VSCR ← a[96:127]

The VSCR is set by the elements in a which occupy the last 32 bits. The result is void.

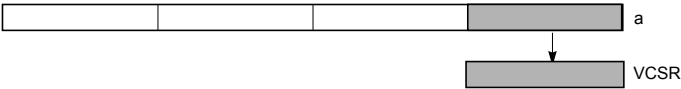


Figure 4-80. Vector Move to VSCR

Refer to the description of `vec_mfvscr` for a detailed description of the VSCR (see Figure 4-1). The valid argument types for `vec_mtvscr(a)` are shown in Table 4-15. The result type is `void`.

Table 4-15. `vec_mtvscr`—Vector Move to Vector Status and Control Register Argument Types

a	Maps to
vector unsigned char	mtvscr a
vector signed char	
vector bool char	
vector unsigned short	
vector signed short	
vector bool short	
vector pixel	
vector unsigned int	
vector signed int	
vector bool int	

vec_mule

Vector Multiply Even

vec_mule

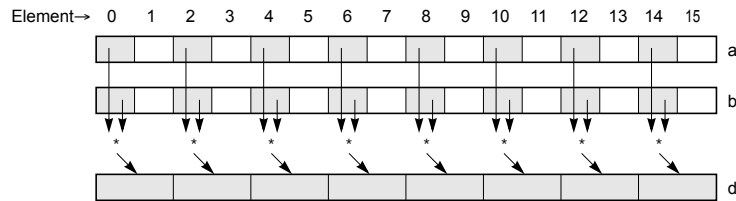
d = vec_mule(**a**,**b**)

```

n ← number of elements in d
do i=0 to n-1
  di ← a2i * b2i
end

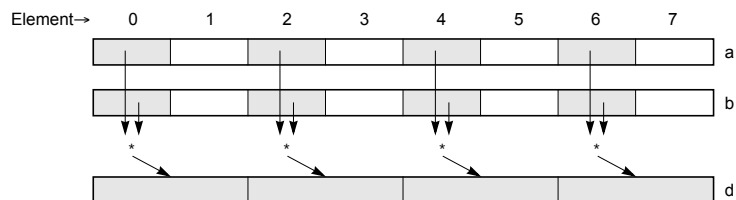
```

Each element of the result is the product of the corresponding high half-width elements of **a** and **b**. The odd elements of **a** and **b** are ignored. The valid combinations of argument types and the corresponding result types for **d** = vec_mule(**a**,**b**) are shown in Figure 4-81 and Figure 4-82.



d	a	b	maps to
vector unsigned short	vector unsigned char	vector unsigned char	vmuleub d,a,b
vector signed short	vector signed char	vector signed char	vmulesb d,a,b

Figure 4-81. Even Multiply of Eight Integer Elements (8-Bit)



d	a	b	maps to
vector unsigned int	vector unsigned short	vector unsigned short	vmuleuh d,a,b
vector signed int	vector signed short	vector signed short	vmulesh d,a,b

Figure 4-82. Even Multiply of Four Integer Elements (16-Bit)

vec_mulo

Vector Multiply Odd

vec_mulo

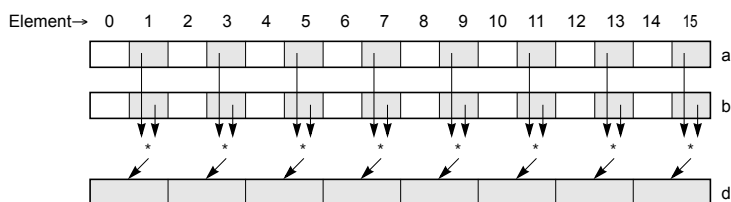
d = vec_mulo(**a**,**b**)

```

n ← number of elements in d
do i=0 to n-1
  di ← a2i+1 * b2i+1
end

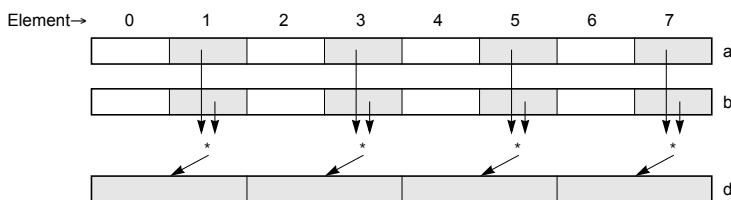
```

Each element of the result is the product of the corresponding low half-width elements of **a** and **b**. The even elements of **a** and **b** are ignored. The valid combinations of argument types and the corresponding result types for **d** = vec_mulo(**a**,**b**) are shown in Figure 4-83 and Figure 4-84.



d	a	b	maps to
vector unsigned short	vector unsigned char	vector unsigned char	vmuloub d,a,b
vector signed short	vector signed char	vector signed char	vmulosb d,a,b

Figure 4-83. Odd Multiply of Eight Integer Elements (8-Bit)



d	a	b	maps to
vector unsigned int	vector unsigned short	vector unsigned short	vmulouh d,a,b
vector signed int	vector signed short	vector signed short	vmulosh d,a,b

Figure 4-84. Odd Multiply of Four Integer Elements (16-Bit)

vec_nmsub

Vector Negative Multiply Subtract

vec_nmsub

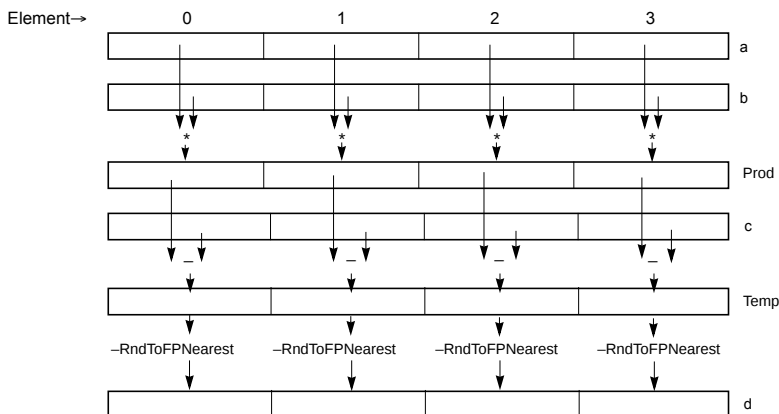
d = vec_nmsub(**a**,**b**,**c**)

```
do i=0 to 3
  di ← -RndToFPNearest(ai * bi - ci)
end
```

Each element of the result is the negative of the difference of the corresponding element of **c** and the product of the corresponding elements of **a** and **b**.

For vector float argument types, if VSCR[NJ] is set, every denormalized operand element is truncated to a 0 of the same sign before the operation is carried out, and each denormalized result element truncates to a 0 of the same sign.

The valid argument types and the corresponding result type for **d** = vec_nmsub(**a**,**b**,**c**) are shown in Figure 4-85.



d	a	b	c	maps to
vector float	vector float	vector float	vector float	vnmsubfp d,a,b,c

Figure 4-85. Negative Multiply-Subtract of Four Floating-Point Elements (32-Bit)

vec_nor

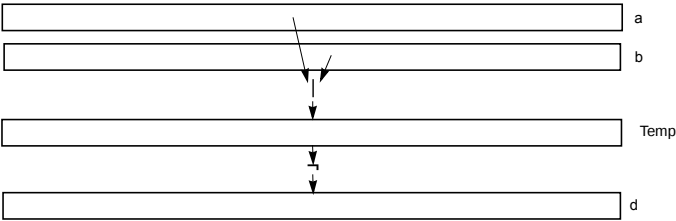
Vector Logical NOR

d = vec_nor(**a**,**b**)

$$d \leftarrow \neg(a \mid b)$$

Each bit of the result is the logical NOR of the corresponding bits of **a** and **b**.
The valid combinations of argument types and the corresponding result types for **d** = vec_nor(**a**,**b**) are shown in Figure 4-86.

vec_nor



d	a	b	maps to
vector unsigned char	vector unsigned char	vector unsigned char	vnor d,a,b
vector signed char	vector signed char	vector signed char	
vector bool char	vector bool char	vector bool char	
vector unsigned short	vector unsigned short	vector unsigned short	
vector signed short	vector signed short	vector signed short	
vector bool short	vector bool short	vector bool short	
vector unsigned int	vector unsigned int	vector unsigned int	
vector signed int	vector signed int	vector signed int	
vector bool int	vector bool int	vector bool int	
vector float	vector float	vector float	

Figure 4-86. Logical Bit-Wise NOR

vec_or

Vector Logical OR

vec_or

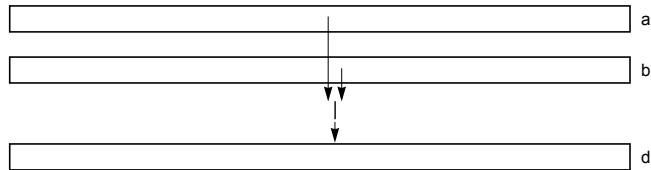
d = vec_or(**a**,**b**)

d ← **a** | **b**

Each bit of the result is the logical OR of the corresponding bits of **a** and **b**.

The valid combinations of argument types and the corresponding result types for

d = vec_or(**a**,**b**) are shown in Figure 4-87.



d	a	b	maps to
vector unsigned char	vector unsigned char	vector unsigned char	vor d,a,b
	vector unsigned char	vector bool char	
	vector bool char	vector unsigned char	
vector signed char	vector signed char	vector signed char	
	vector signed char	vector bool char	
	vector bool char	vector signed char	
vector bool char	vector bool char	vector bool char	
vector unsigned short	vector unsigned short	vector unsigned short	
	vector unsigned short	vector bool short	
	vector bool short	vector unsigned short	
vector signed short	vector signed short	vector signed short	
	vector signed short	vector bool short	
	vector bool short	vector signed short	
vector bool short	vector bool short	vector bool short	
vector unsigned int	vector unsigned int	vector unsigned int	
	vector unsigned int	vector bool int	
	vector bool int	vector unsigned int	
vector signed int	vector signed int	vector signed int	
	vector signed int	vector bool int	
	vector bool int	vector signed int	
vector bool int	vector bool int	vector bool int	
vector float	vector bool int	vector float	
	vector float	vector bool int	
	vector float	vector float	

Figure 4-87. Logical Bit-Wise OR

vec_pack

Vector Pack

vec_pack

d = vec_pack(**a**,**b**)

```

n ← number of elements in a
s ← element size in d (64/n)
do i=0 to n-1
  di ← UIToUImod(ai,s)
  di+n ← UIToUImod(bi,s)
end
  
```

Each high element of the result is the truncation of the corresponding wider element of a. Each low element of the result is the truncation of the corresponding wider element of b. The valid combinations of argument types and the corresponding result types for **d** = vec_pack(**a**,**b**) are shown in Figure 4-88 and Figure 4-89.

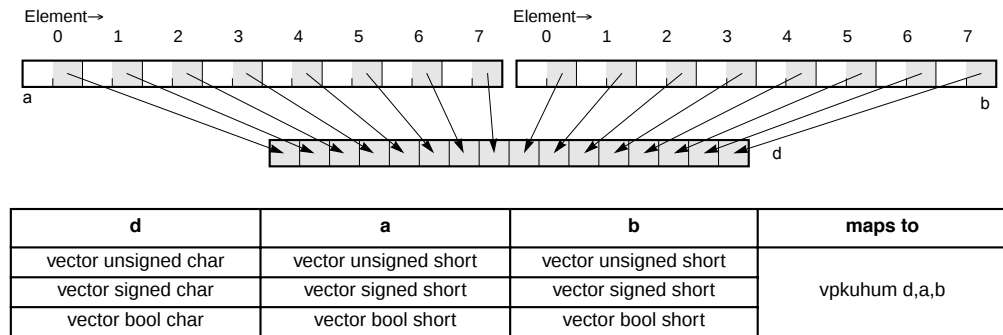


Figure 4-88. Pack Sixteen Unsigned Integer Elements (16-Bit) to Sixteen Unsigned Integer Elements (8-Bit)

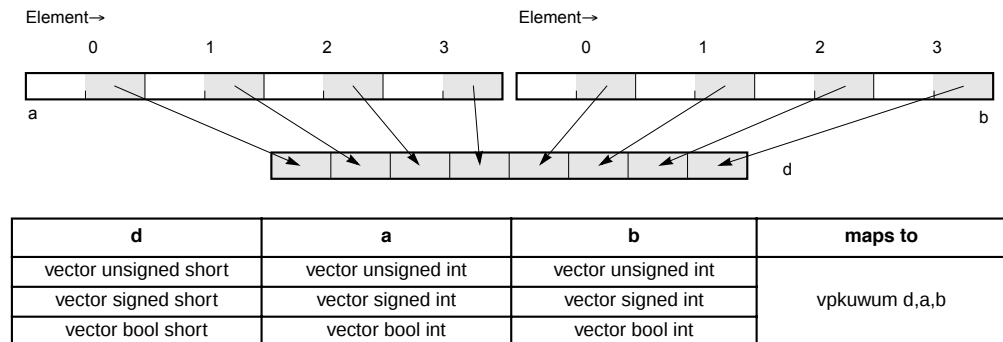


Figure 4-89. Pack Eight Unsigned Integer Elements (32-Bit) to Eight Unsigned Integer Elements (16-Bit)

vec_packpx

Vector Pack Pixel

vec_packpx

d = vec_packpx(**a**,**b**)

```
do i=0 to 3
  di      ← ai[7] || ai[8:12] || ai[16:20] || ai[24:28]
  di+4    ← bi[7] || bi[8:12] || bi[16:20] || bi[24:28]
end
```

Each high element of the result is the packed pixel from the corresponding wider element of **a**. Each low element of the result is the packed pixel from the corresponding wider element of **b**.

Programming note: Each source word can be considered to be a 32-bit pixel consisting of four 8-bit channels. Each target half-word can be considered to be a 16-bit pixel consisting of one 1-bit channel and three 5-bit channels. A channel can be used to specify the intensity of a particular color, such as red, green, or blue, or to provide other information needed by the application.

The usual transformation from a 32-bit pixel to a 16-bit pixel uses the most significant bit of the 8-bit intensity channel. This operation uses the least significant bit. To use the most significant bit, first perform the following operation:

```
(vector unsigned int) vec_rl ((vector unsigned char) a,
                             (vector unsigned char) (1,0,0,0,1,0,0,0,
                                                       1,0,0,0,1,0,0,0))
```

on each input **a** and **b**.

The valid argument types and the corresponding result type for **d** = vec_packpx(**a**,**b**) are shown in Figure 4-90.

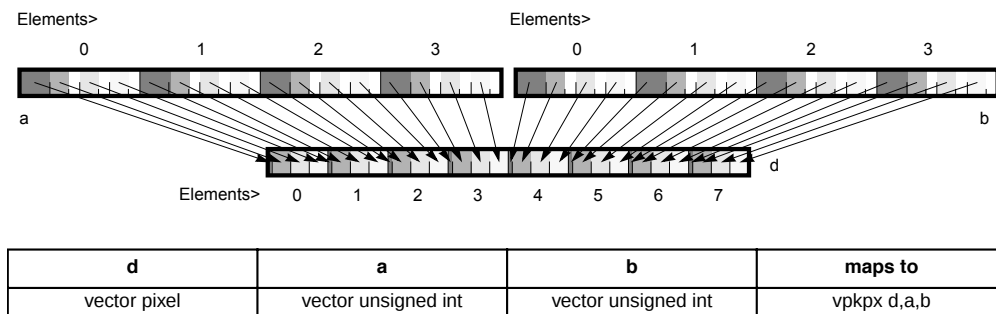


Figure 4-90. Pack Eight Pixel Elements (32-Bit) to Eight Elements (16-Bit)

vec_packs

Vector Pack Saturated

vec_packs

d = vec_packs(**a**,**b**)

```

n ← number of elements in a
do i=0 to n-1
  di ← Saturate(ai)
  di+n ← Saturate(bi)
end

```

Each high element of the result is the saturated value of the corresponding wider element of a. Each low element of the result is the saturated value of the corresponding wider element of b. If saturation occurs, VSCR[SAT] is set (see Table 4-1).

The valid combinations of argument types and the corresponding result types for **d** = vec_packs(**a**,**b**) are shown in Figure 4-91 and Figure 4-92.

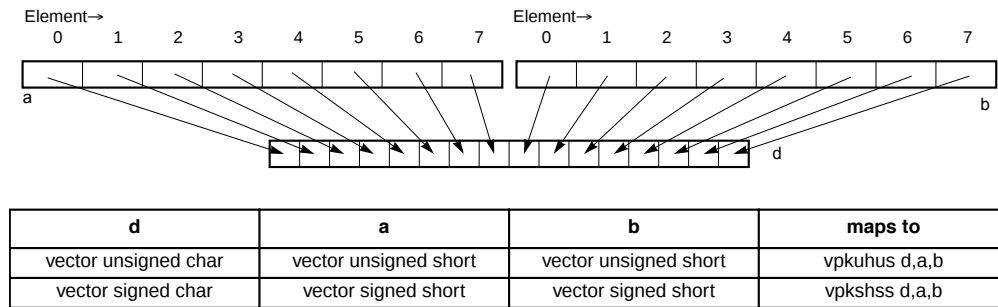


Figure 4-91. Pack Sixteen Integer Elements (16-Bit) to Sixteen Integer Elements (8-Bit)

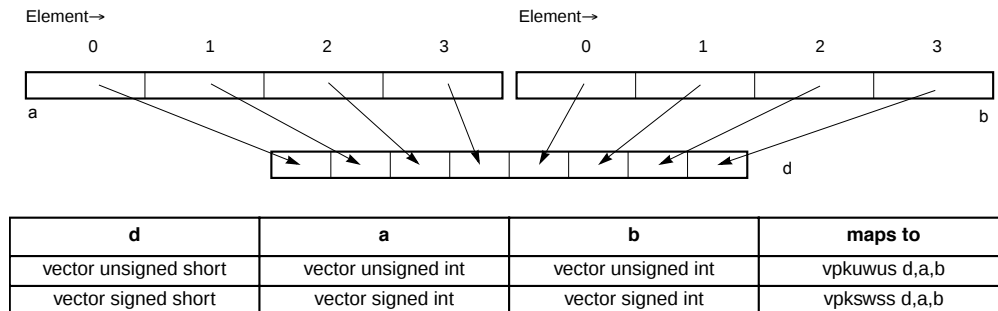


Figure 4-92. Pack Eight Integer Elements (32-Bit) to Eight Integer Elements (16-Bit)

vec_packsu

Vector Pack Saturated Unsigned

vec_packsu

d = vec_packsu(**a**,**b**)

```

n ← number of elements in a
do i=0 to n-1
  di ← Saturate(ai)
  di+n ← Saturate(bi)
end

```

Each high element of the result is the saturated value of the corresponding wider element of a. Each low element of the result is the saturated value of the corresponding wider element of b. If saturation occurs, VSCR[SAT] is set (see Table 4-1). The result elements are all unsigned. The valid combinations of argument types and the corresponding result types for **d** = vec_packsu(**a**,**b**) are shown in Figure 4-93 and Figure 4-94.

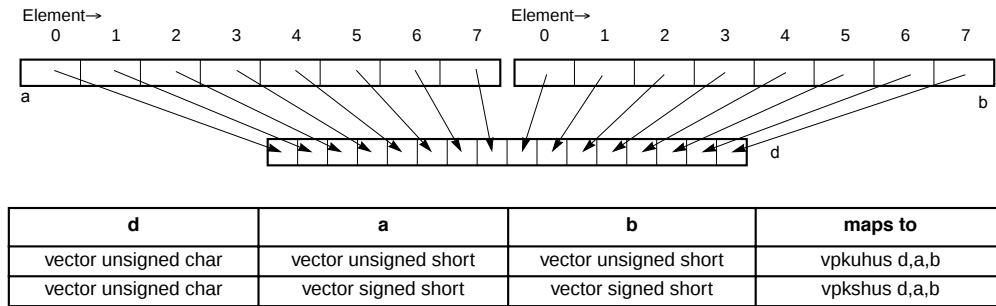


Figure 4-93. Pack Sixteen Integer Elements (16-Bit) to Sixteen Unsigned Integer Elements (8-Bit)

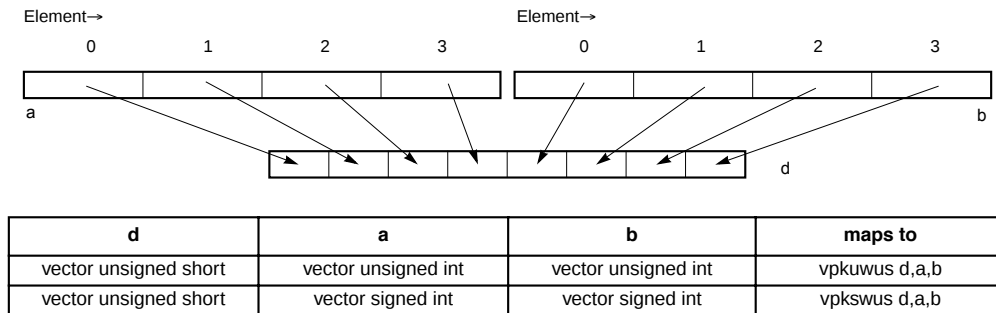


Figure 4-94. Pack Eight Integer Elements (32-Bit) to Eight Unsigned Integer Elements (16-Bit)



vec_perm

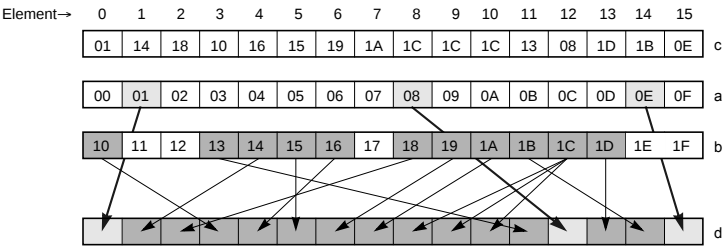
Vector Permute

vec_perm

d = vec_perm(**a**,**b**,**c**)

```
do i=0 to 15
  j ← c{i}{4:7}
  if c{i}{3} = 0
    then d{i} ← a{j}
    else d{i} ← b{j}
end
```

Each element of the result is selected independently by indexing the byte elements of **a** and **b** by the value of the corresponding element of **c**. For example, 0x1C in **c** selects byte 12 in **b**. The value 0x0C selects byte 12 in **a**. The valid combinations of argument types and the corresponding result types for **d** = vec_perm(**a**,**b**,**c**) are shown in Figure 4-95.



d	a	b	c	maps to
vector unsigned char	vector unsigned char	vector unsigned char	vector unsigned char	vperm d,a,b,c
vector signed char	vector signed char	vector signed char	vector unsigned char	
vector bool char	vector bool char	vector bool char	vector unsigned char	
vector unsigned short	vector unsigned short	vector unsigned short	vector unsigned char	
vector signed short	vector signed short	vector signed short	vector unsigned char	
vector bool short	vector bool short	vector bool short	vector unsigned char	
vector pixel	vector pixel	vector pixel	vector unsigned char	
vector unsigned int	vector unsigned int	vector unsigned int	vector unsigned char	
vector signed int	vector signed int	vector signed int	vector unsigned char	
vector bool int	vector bool int	vector bool int	vector unsigned char	
vector float	vector float	vector float	vector unsigned char	

Figure 4-95. Permute Sixteen Integer Elements (8-Bit)

vec_re

Vector Reciprocal Estimate

vec_re

d = vec_re(**a**)

```
do i=0 to 3
  di ← FPREcipEst(ai)
end
```

Each element of the result **d** is an estimate of the reciprocal to the corresponding element of **a**. For results that are not a +0, −0, +∞, −∞, or QNaN, the estimate has a relative error in precision no greater than one part in 4096, that is:

$$\left| \frac{\text{estimate} - 1/x}{1/x} \right| \leq \frac{1}{4096}$$

where *x* is the value of the element in **a**. Note that the value placed into the element of **d** may vary between implementations, and between different executions on the same implementation.

Operation with various special values of the element in **a** is summarized below.

Table 4-16. Special Value Results of Reciprocal Estimates

a	d
−∞	−0
−0	−∞
+0	+∞
+∞	+0
NaN	QNaN

If VSCR[NJ] = 1, every denormalized operand element is truncated to a 0 of the same sign before the operation is carried out, and each denormalized result element truncates to a 0 of the same sign.

The valid argument type and corresponding result type for **d** = vec_re(**a**) are shown in Figure 4-96.

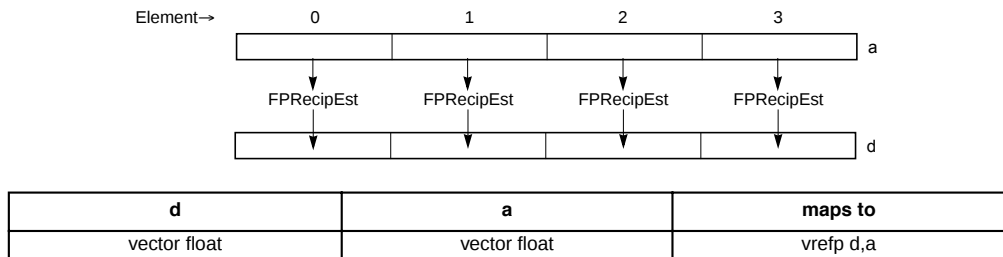


Figure 4-96. Reciprocal Estimate of Four Floating-Point Elements (32-Bit)



vec_rl

Vector Rotate Left

vec_rl

d = vec_rl(**a**,**b**)

```
n ← number of elements
do i=0 to n-1
  di ← ROTL(ai, bi)
end
```

Each element of the result is the result of rotating left the corresponding element of **a** by the number of bits indicated by the corresponding element of **b**. The valid combinations of argument types and the corresponding result types for **d** = vec_rl(**a**,**b**) are shown in Figure 4-97, Figure 4-98, and Figure 4-99.

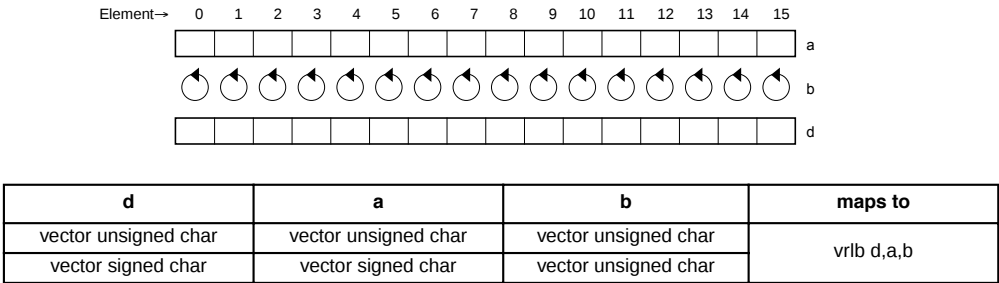


Figure 4-97. Left Rotate of Sixteen Integer Elements (8-Bit)

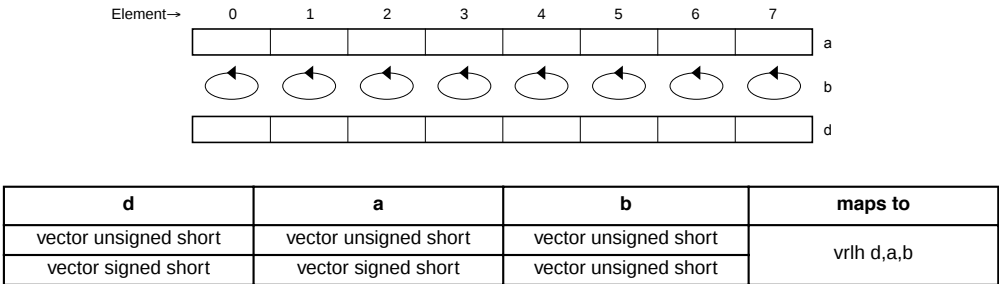


Figure 4-98. Left Rotate of Eight Integer Elements (16-bit)

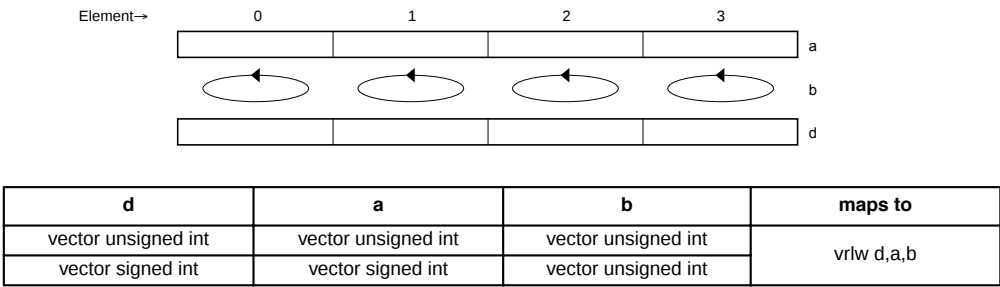


Figure 4-99. Left Rotate of Four Integer Elements (32-bit)

vec_round

Vector Round

$$\mathbf{d} = \text{vec_round}(\mathbf{a})$$

```
do i=0 to 3
di ← RndToFPINear(ai)
end
```

Each element of the result is the nearest representable single-precision floating-point integer to the corresponding element of `a`, using IEEE Round-to-Nearest mode. If the integers are equally near, rounding is to the even integer.

The operation is independent of VSCR[NJ].

The valid argument type and corresponding result type for `a = vec_round(a)` are shown in Figure 4-100.

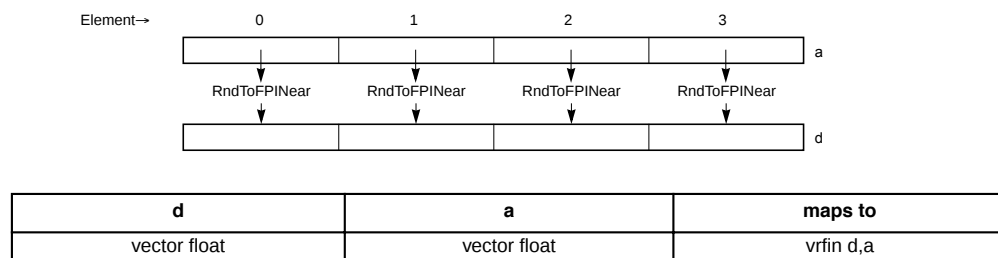


Figure 4-100. Round to Nearest of Four Floating-Point Integer Elements (32-Bit)

vec_rsqрте

Vector Reciprocal Square Root Estimate

vec_rsqрте

d = vec_rsqрте(**a**)

```
do i=0 to 3
  di ← RecipSQRTes(ai)
end
```

Each element of the result is an estimate of the reciprocal square root of the corresponding element of **a**. The single-precision estimate of the reciprocal of the square root of each single-precision element in **a** is placed into the corresponding word element of **d**. The estimate has a relative error in precision no greater than one part in 4096, that is:

$$\left| \frac{\text{estimate} - 1/\sqrt{x}}{1/\sqrt{x}} \right| \leq \frac{1}{4096}$$

where **x** is the value of the element in **a**. The value placed into the element of **d** may vary between implementations and between different executions on the same implementation. If VSCR[NJ] = 1, every denormalized operand element is truncated to a 0 of the same sign before the operation is carried out, and each denormalized result element truncates to a 0 of the same sign. Operation with various special values of the element in **a** is summarized below.

Table 4-17. Special Value Results of Reciprocal Square Root Estimates

a	d
$-\infty$	QNaN
less than 0	QNaN
-0	$-\infty$
+0	$+\infty$
$+\infty$	+0
NaN	QNaN

The valid argument type and corresponding result type for **d** = vec_rsqрте(**a**) are shown in Figure 4-101.

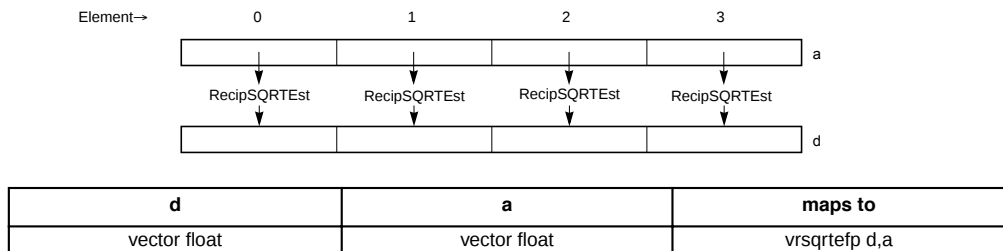


Figure 4-101. Reciprocal Square Root Estimate of Four Floating-Point Elements (32-Bit)



vec_sel

Vector Select

vec_sel

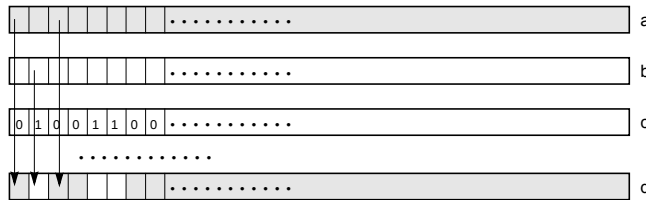
d = vec_sel(**a**,**b**,**c**)

```

do i=0 to 127
  if ci=0
    then d[i] ← a[i]
    else d[i] ← b[i]
end

```

Each bit of the result is the corresponding bit of **a** if the corresponding bit of **c** is 0. Otherwise, it is the corresponding bit of **b**. The valid combinations of argument types and the corresponding result types for **d** = vec_sel(**a**,**b**,**c**) are shown in Figure 4-102.



d	a	b	c	maps to
vector unsigned char	vector unsigned char	vector unsigned char	vector unsigned char	vsel d,a,b,c
	vector unsigned char	vector unsigned char	vector bool char	
vector signed char	vector signed char	vector signed char	vector unsigned char	
	vector signed char	vector signed char	vector bool char	
vector bool char	vector bool char	vector bool char	vector unsigned char	
	vector bool char	vector bool char	vector bool char	
vector unsigned short	vector unsigned short	vector unsigned short	vector unsigned short	
	vector unsigned short	vector unsigned short	vector bool short	
vector signed short	vector signed short	vector signed short	vector unsigned short	
	vector signed short	vector signed short	vector bool short	
vector bool short	vector bool short	vector bool short	vector unsigned short	
	vector bool short	vector bool short	vector bool short	
vector unsigned int	vector unsigned int	vector unsigned int	vector unsigned int	
	vector unsigned int	vector unsigned int	vector bool int	
vector signed int	vector signed int	vector signed int	vector unsigned int	
	vector signed int	vector signed int	vector bool int	
vector bool int	vector bool int	vector bool int	vector unsigned int	
	vector bool int	vector bool int	vector bool int	
vector float	vector float	vector float	vector unsigned int	
	vector float	vector float	vector bool int	

Figure 4-102. Bit-Wise Conditional Select of Vector Contents (128-bit)

vec_sl

Vector Shift Left

vec_sl

$$d = \text{vec_sl}(a,b)$$

```

n ← number of elements
s ← 128/n
do i=0 to n-1
  di ← ShiftLeft(ai,mod(bi,s))
end
  
```

Each element in **d** is the result of shifting the corresponding element of **a** left by the number of bits of the corresponding element of **b**. The valid combinations of argument types and the corresponding result types for **d = vec_sl(a,b)** are shown in Figure 4-103, Figure 4-104, and Figure 4-105.

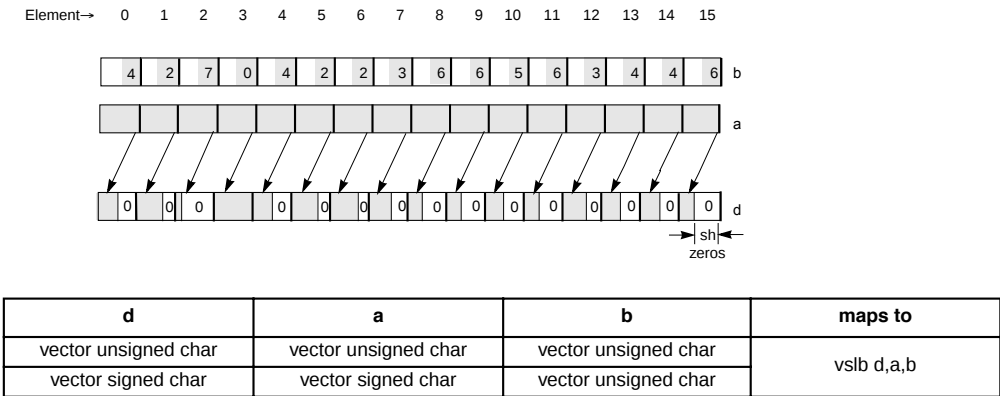


Figure 4-103. Shift Bits Left in Sixteen Integer Elements (8-Bit)

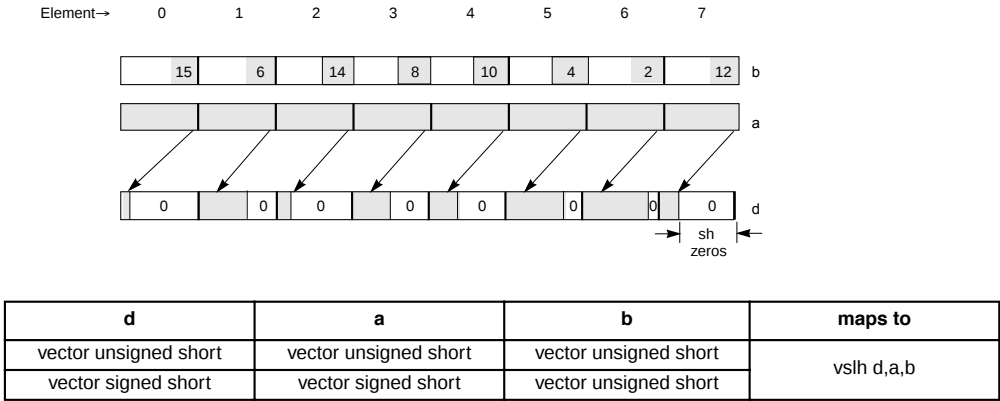


Figure 4-104. Shift Bits Left in Eight Integer Elements (16-Bit)

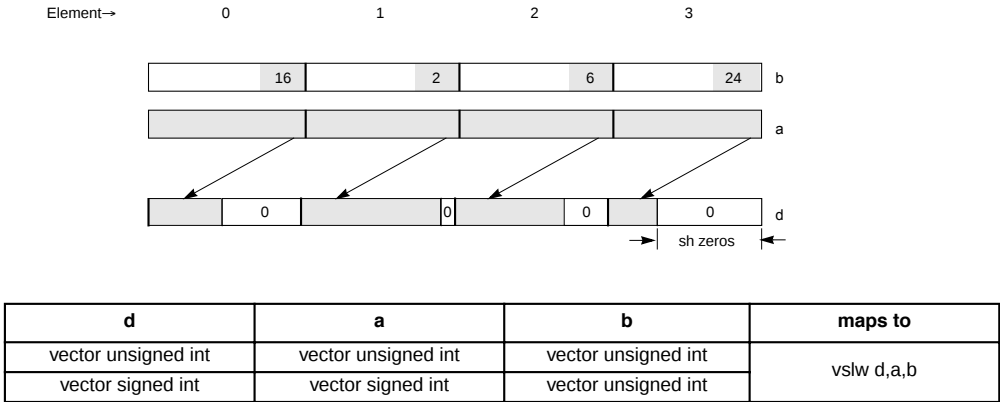


Figure 4-105. Shift Bits Left in Four Integer Elements (32-Bit)

vec_sld

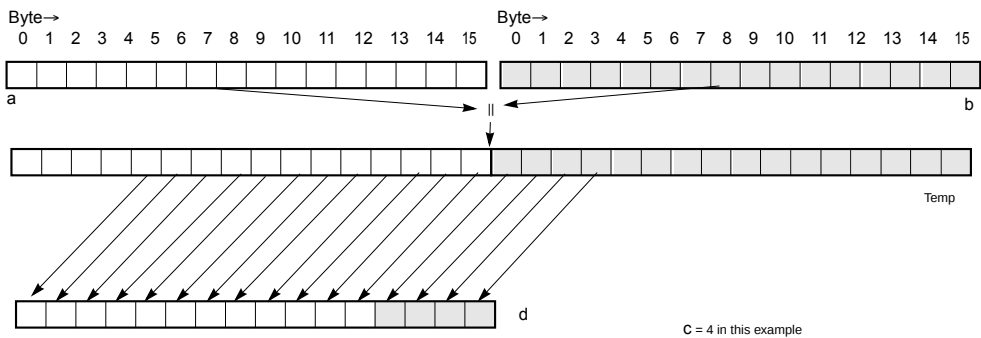
Vector Shift Left Double

vec_sld

d = vec_sld(**a**,**b**,**c**)

```
do i=0 to 15
  if (i+c) < 16
    then d{i} ← a{i+c}
    else d{i} ← b{i+c-16}
end
```

The result is obtained by selecting the top 16 bytes obtained by shifting left (unsigned) by the value of **c** bytes a 32-byte quantity formed by concatenating **a** with **b**. The valid combinations of argument types and the corresponding result types for **d** = vec_sld(**a**,**b**,**c**) are shown in Figure 4-106.



d	a	b	c	maps to
vector unsigned char	vector unsigned char	vector unsigned char	4-bit unsigned literal	vsldoi d,a,b,c
vector signed char	vector signed char	vector signed char	4-bit unsigned literal	
vector unsigned short	vector unsigned short	vector unsigned short	4-bit unsigned literal	
vector signed short	vector signed short	vector signed short	4-bit unsigned literal	
vector pixel	vector pixel	vector pixel	4-bit unsigned literal	
vector unsigned int	vector unsigned int	vector unsigned int	4-bit unsigned literal	
vector signed int	vector signed int	vector signed int	4-bit unsigned literal	
vector float	vector float	vector float	4-bit unsigned literal	

Figure 4-106. Bit-Wise Conditional Select of Vector Contents (128-bit)



vec_sll

Vector Shift Left Long

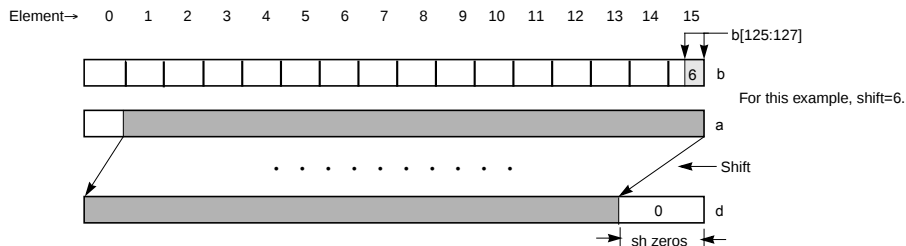
d = vec_sll(**a**,**b**)

```
m ← b[125:127]
If each bi[5:7] = m, where i ranges from 0 to 14
then d ← ShiftLeft(a,m)
else d ← Undefined
```

The result is obtained by shifting **a** left by a number of bits specified by the last 3 bits of the last element of **b**. The valid combinations of argument types and the corresponding result types for **d** = vec_sll(**a**,**b**) are shown in Figure 4-107.

Note that the three low-order bits of all byte elements in **b** must be the same; otherwise the value placed into **d** is undefined.

vec_sll



d	a	b	maps to
vector unsigned char	vector unsigned char	vector unsigned char	vsl d,a,b
	vector unsigned char	vector unsigned short	
	vector unsigned char	vector unsigned int	
vector signed char	vector signed char	vector unsigned char	
	vector signed char	vector unsigned short	
	vector signed char	vector unsigned int	
vector bool char	vector bool char	vector unsigned char	
	vector bool char	vector unsigned short	
	vector bool char	vector unsigned int	
vector unsigned short	vector unsigned short	vector unsigned char	
	vector unsigned short	vector unsigned short	
	vector unsigned short	vector unsigned int	
vector signed short	vector signed short	vector unsigned char	
	vector signed short	vector unsigned short	
	vector signed short	vector unsigned int	
vector bool short	vector bool short	vector unsigned char	
	vector bool short	vector unsigned short	
	vector bool short	vector unsigned int	
vector pixel	vector pixel	vector unsigned char	
	vector pixel	vector unsigned short	
	vector pixel	vector unsigned int	
vector unsigned int	vector unsigned int	vector unsigned char	
	vector unsigned int	vector unsigned short	
	vector unsigned int	vector unsigned int	
vector signed int	vector signed int	vector unsigned char	
	vector signed int	vector unsigned short	
	vector signed int	vector unsigned int	
vector bool int	vector bool int	vector unsigned char	
	vector bool int	vector unsigned short	
	vector bool int	vector unsigned int	

Figure 4-107. Shift Bits Left in Vector (128-Bit)

vec_slo

Vector Shift Left by Octet

vec_slo

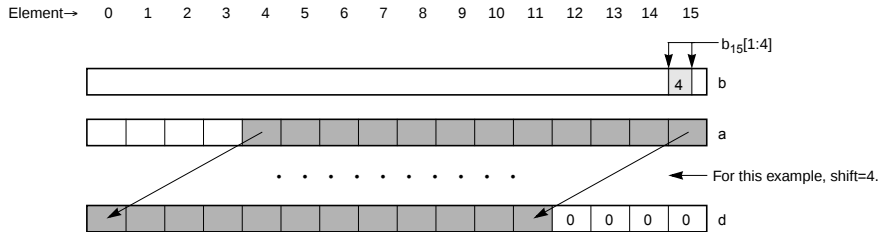
d = vec_slo(**a**,**b**)

```

m ← b15[1:4]
do i=0 to 15
  j ← i + m
  if j < 16
    then d{i} ← a{j}
    else d{i} ← 0
  end
end

```

The contents of **a** are shifted left by the number of bytes specified by bits **b₁₅[1:4]**; only these 4 bits in **b** are significant for the shift value. Bytes shifted out of byte 0 are lost. Zeros are supplied to the vacated bytes on the right. The result is placed into **d**. The valid combinations of argument types and the corresponding result types for **d = vec_slo(a, b)** are shown in Figure 4-108.



d	a	b	maps to
vector unsigned char	vector unsigned char	vector unsigned char	vslo d,a,b
	vector unsigned char	vector signed char	
vector signed char	vector signed char	vector unsigned char	
	vector signed char	vector signed char	
vector unsigned short	vector unsigned short	vector unsigned char	
	vector unsigned short	vector signed char	
vector signed short	vector signed short	vector unsigned char	
	vector signed short	vector signed char	
vector pixel	vector pixel	vector unsigned char	
	vector pixel	vector signed char	
vector unsigned int	vector unsigned int	vector unsigned char	
	vector unsigned int	vector signed char	
vector signed int	vector signed int	vector unsigned char	
	vector signed int	vector signed char	
vector float	vector float	vector unsigned char	
	vector float	vector signed char	

Figure 4-108. Left Byte Shift of Vector (128-Bit)

vec_splat

Vector Splat

vec_splat

d = vec_splat(**a**,**b**)

```

n ← number of elements
do i=0 to n-1
  j ← mod(b,n)
  di ← aj
end

```

Each element of the result is component **b** of **a**. The valid combinations of argument types and the corresponding result types for **d** = vec_splat(**a**,**b**) are shown in Figure 4-109, Figure 4-110, and Figure 4-111.

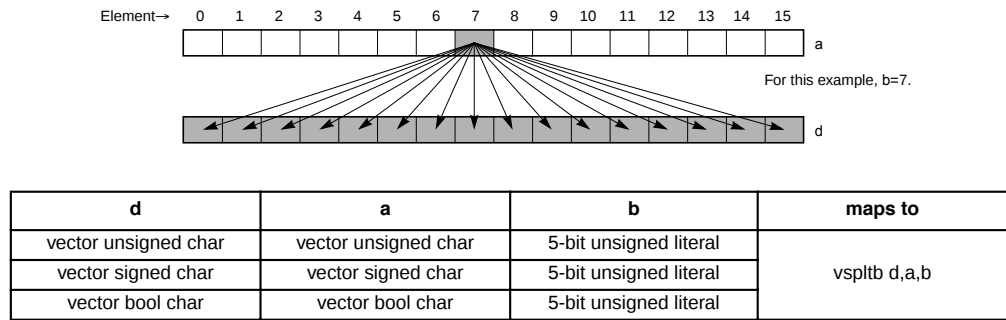


Figure 4-109. Copy Contents to Sixteen Integer Elements (8-Bit)

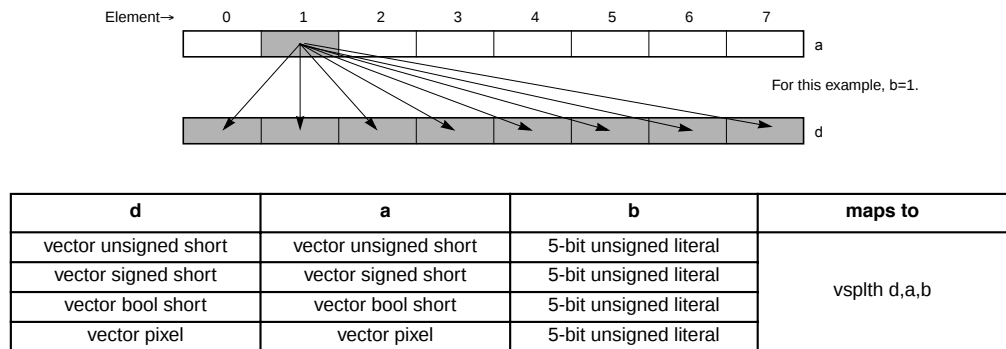


Figure 4-110. Copy Contents to Eight Elements (16-Bit)

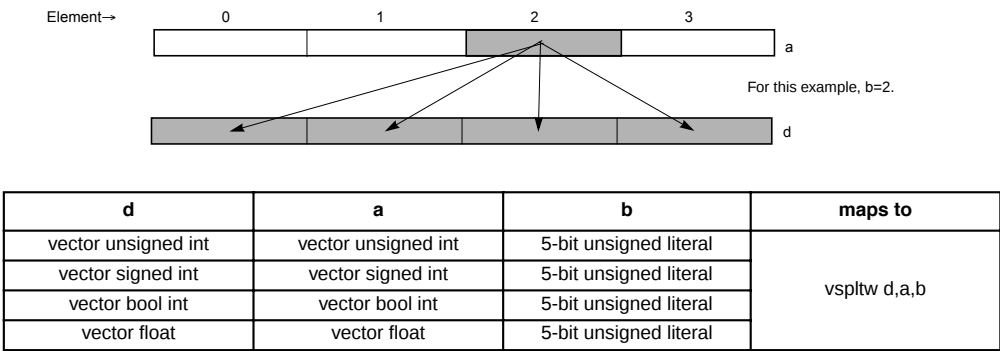


Figure 4-111. Copy Contents to Four Integer Elements (32-Bit)

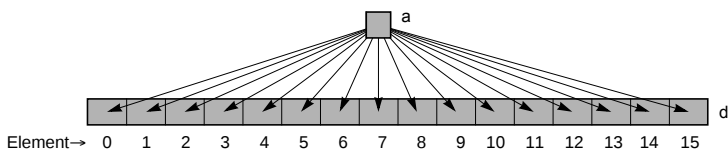
vec_splat_s8

Vector Splat Signed Byte

d = vec_splat_s8(**a**)

```
do i=0 to 15
  di ← SignExtend(a)
end
```

Each element of the result is the value obtained by sign-extending **a**. This permits values ranging from -16 to 15 only. The valid argument type and corresponding result type for **d** = vec_splat_s8(**a**) are shown in Figure 4-112.



d	a	maps to
vector signed char	5-bit signed literal	vspltisb d,a

Figure 4-112. Copy Value into Sixteen Signed Integer Elements (8-Bit)

vec_splat_s16

Vector Splat Signed Half-Word

```
d = vec_splat_s16(a)
do i=0 to 7
  di ← SignExtend(a)
end
```

Each element of the result is the value obtained by sign-extending a. This permits values ranging from -16 to 15 only. The valid argument type and corresponding result type for **d = vec_splat_s16(a)**, are shown in Figure 4-113.

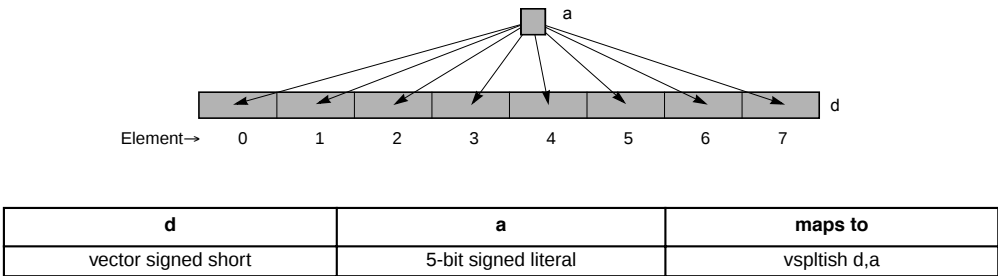


Figure 4-113. Copy Value into Eight Signed Integer Elements (16-Bit)

vec_splat_s32

Vector Splat Signed Word

d = vec_splat_s32(**a**)

```
do i=0 to 3
  di ← SignExtend(a)
end
```

Each element of the result is the value obtained by sign-extending **a**. This permits values ranging from -16 to 15 only. The valid argument type are corresponding result type for **d** = vec_splat_s32(**a**) are shown in Figure 4-114.

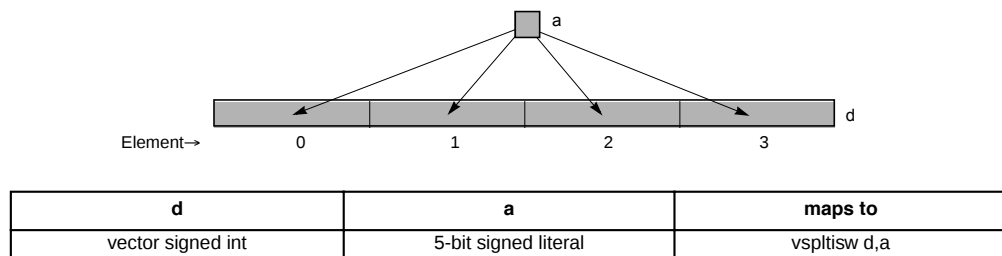


Figure 4-114. Copy Value into Four Signed Integer Elements (32-Bit)



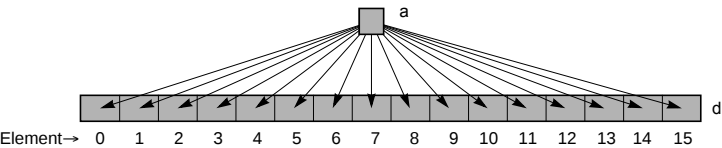
vec_splat_u8

Vector Splat Unsigned Byte

vec_splat_u8

```
d = vec_splat_u8(a)
do i=0 to 15
  di ← SignExtend(a)
end
```

Each element of the result is the value obtained by sign-extending a and casting it to an unsigned char value. Each element of **d** is set to 256*sign(a) + a, where sign(a) is 0 for non-negative a and 1 for negative a. The valid argument type and corresponding result type for **d = vec_splat_u8(a)** are shown in Figure 4-115. It is necessary to use the generic name, since the specific operation `vec_vspltisb` returns a `vector signed char` value.



d	a	maps to
vector unsigned char	5-bit signed literal	vspltisb d,a

Figure 4-115. Copy Value into Sixteen Signed Integer Elements (8-Bit)

vec_splat_u16

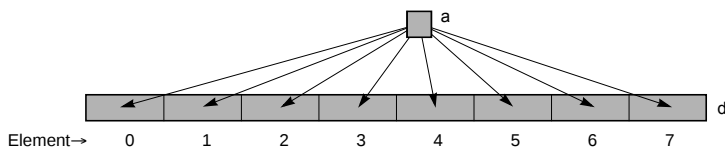
Vector Splat Unsigned Half-Word

vec_splat_u16

d = vec_splat_u16(**a**)

```
do i=0 to 7
  di ← SignExtend(a)
end
```

Each element of the result is the value obtained by sign-extending **a** and casting it to an unsigned short value. Each element of **d** is set to $65536 * \text{sign}(\mathbf{a}) + \mathbf{a}$, where $\text{sign}(\mathbf{a})$ is 0 for non-negative **a** and 1 for negative **a**. The valid argument type and corresponding result type for **d** = vec_splat_u16(**a**) are shown in Figure 4-116. It is necessary to use the generic name, since the specific operation vec_vspltish returns a vector signed short value.



d	a	maps to
vector unsigned short	5-bit signed literal	vspltish d,a

Figure 4-116. Copy Value into Eight Signed Integer Elements (16-Bit)

vec_splat_u32

Vector Splat Unsigned Word

vec_splat_u32

```
d = vec_splat_u32(a)

do i=0 to 3
  di ← SignExtend(a)
end
```

Each element of the result is the value obtained by sign-extending *a*, and casting it to an unsigned int value. Each element of *d* is set to 4294967296*sign(*a*) + *a*, where sign(*a*) is 0 for non-negative *a* and 1 for negative *a*. The valid argument type and corresponding result type for **d = vec_splat_u32(a)** are shown in Figure 4-117. It is necessary to use the generic name, since the specific operation *vec_vspltisw* returns a vector signed int value.

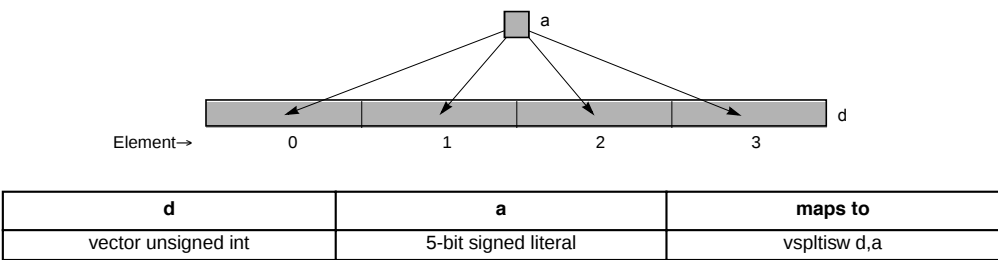


Figure 4-117. Copy Value into Four Signed Integer Elements (32-Bit)

vec_sr

Vector Shift Right

vec_sr

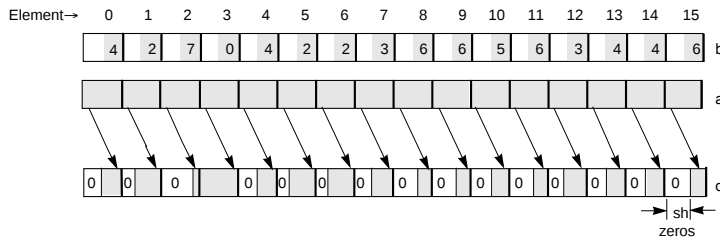
d = vec_sr(**a**,**b**)

```

n ← number of elements
s ← 128/n
do i=0 to n-1
  di ← ShiftRight(ai,mod(bi,s))
end

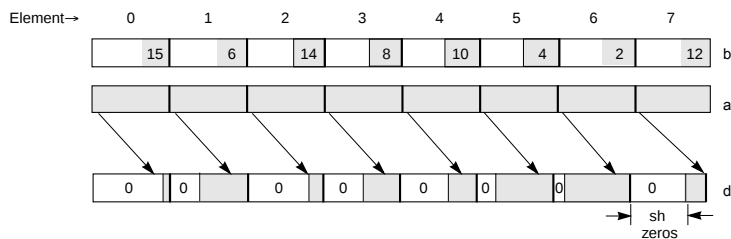
```

Each element of the result is the result of shifting the corresponding element of **a** right by the number of bits of the corresponding element of **b**. Zero bits are shifted in from the left for both signed and unsigned argument types. The valid combinations of argument types and the corresponding result types for **d** = vec_sr(**a**,**b**) are shown in Figure 4-118, Figure 4-119, and Figure 4-120.



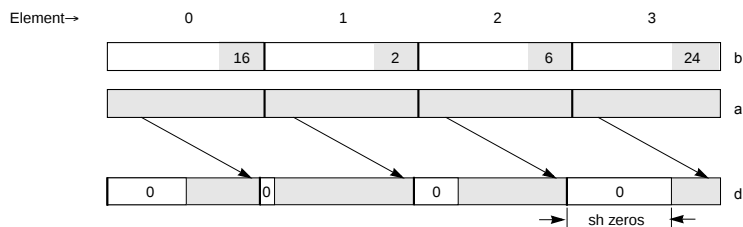
d	a	b	maps to
vector unsigned char	vector unsigned char	vector unsigned char	vsrb d,a,b
vector signed char	vector signed char	vector unsigned char	

Figure 4-118. Shift Bits Right in Sixteen Integer Elements (8-Bit)



d	a	b	maps to
vector unsigned short	vector unsigned short	vector unsigned short	vsrh d,a,b
vector signed short	vector signed short	vector unsigned short	

Figure 4-119. Shift Bits Right in Eight Integer Elements (16-bit)



d	a	b	maps to
vector unsigned int	vector unsigned int	vector unsigned int	vsrw d,a,b
vector signed int	vector signed int	vector unsigned int	

Figure 4-120. Shift Bits Right in Four Integer Elements (32-Bit)

vec_sra

Vector Shift Right Algebraic

vec_sra

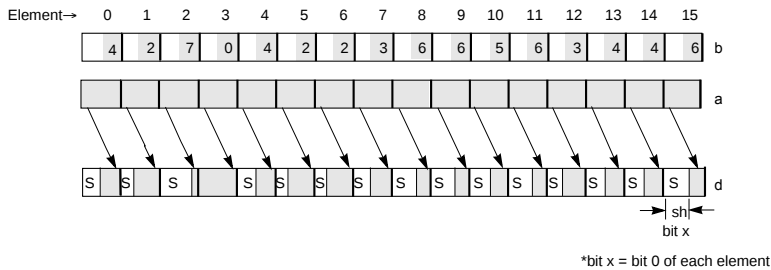
d = vec_sra(**a**,**b**)

```

n ← number of elements
s ← 128/n
do i=0 to n-1
  di ← ShiftRightA(ai, mod(bi, s))
end

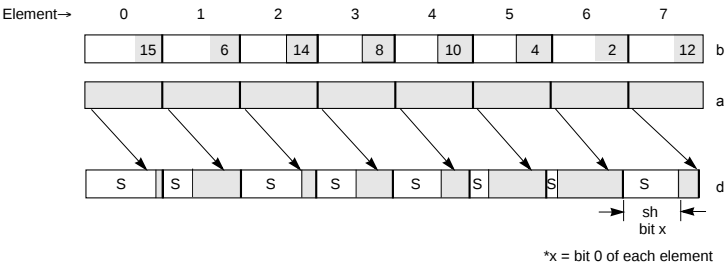
```

Each element of the result is the result of shifting the corresponding element of **a** right by the number of bits of the corresponding element of **b**. Copies of the sign bit are shifted in from the left for both signed and unsigned argument types. The valid combinations of argument types and the corresponding result types for **d** = vec_sra(**a**,**b**) are shown in Figure 4-121, Figure 4-122, and Figure 4-123.



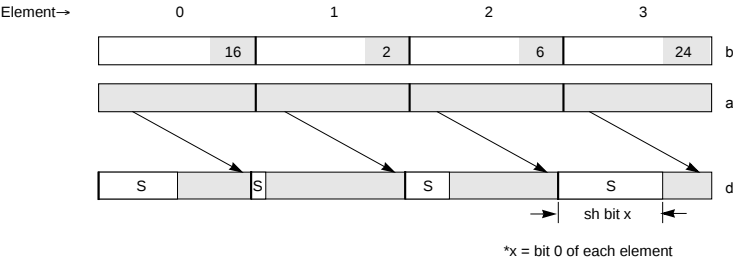
d	a	b	maps to
vector unsigned char	vector unsigned char	vector unsigned char	vsrab d,a,b
vector signed char	vector signed char	vector unsigned char	

Figure 4-121. Shift Bits Right in Sixteen Integer Elements (8-Bit)



d	a	b	maps to
vector unsigned short	vector unsigned short	vector unsigned short	vsrah d,a,b
vector signed short	vector signed short	vector unsigned short	

Figure 4-122. Shift Bits Right in Eight Integer Elements (16-bit)



d	a	b	maps to
vector unsigned int	vector unsigned int	vector unsigned int	vsraw d,a,b
vector signed int	vector signed int	vector unsigned int	

Figure 4-123. Shift Bits Right in Four Integer Elements (32-Bit)

vec_srl

Vector Shift Right Long

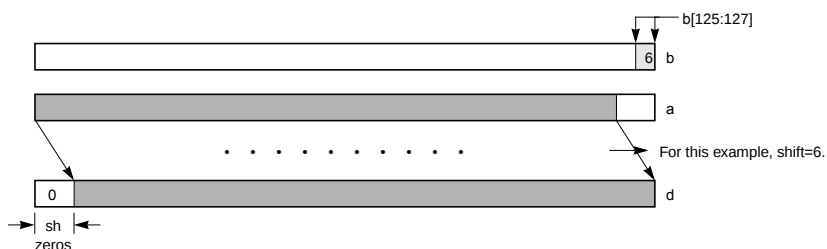
vec_srl

d = vec_srl(**a**,**b**)

```
m ← b[125:127]
if each bi[5:7] = m, where i ranges from 0 to 14
then d ← ShiftRight(a,m)
else d ← Undefined
```

The result is obtained by shifting a right by a number of bits specified by the last 3 bits of the last element of **b**. The valid combinations of argument types and the corresponding result types for **d** = vec_srl(**a**,**b**) are shown in Figure 4-124.

Note that the low-order 3 bits of all byte elements in **b** must be the same; otherwise the value placed into **d** is undefined.



d	a	b	maps to
vector unsigned char	vector unsigned char	vector unsigned char	vsr d,a,b
	vector unsigned char	vector unsigned short	
	vector unsigned char	vector unsigned int	
vector signed char	vector signed char	vector unsigned char	
	vector signed char	vector unsigned short	
	vector signed char	vector unsigned int	
vector bool char	vector bool char	vector unsigned char	
	vector bool char	vector unsigned short	
	vector bool char	vector unsigned int	
vector unsigned short	vector unsigned short	vector unsigned char	
	vector unsigned short	vector unsigned short	
	vector unsigned short	vector unsigned int	
vector signed short	vector signed short	vector unsigned char	
	vector signed short	vector unsigned short	
	vector signed short	vector unsigned int	
vector bool short	vector bool short	vector unsigned char	
	vector bool short	vector unsigned short	
	vector bool short	vector unsigned int	
vector pixel	vector pixel	vector unsigned char	
	vector pixel	vector unsigned short	
	vector pixel	vector unsigned int	
vector unsigned int	vector unsigned int	vector unsigned char	
	vector unsigned int	vector unsigned short	
	vector unsigned int	vector unsigned int	
vector signed int	vector signed int	vector unsigned char	
	vector signed int	vector unsigned short	
	vector signed int	vector unsigned int	
vector bool int	vector bool int	vector unsigned char	
	vector bool int	vector unsigned short	
	vector bool int	vector unsigned int	

Figure 4-124. Shift Bits Right in Vector (128-Bit)

vec_sro

Vector Shift Right by Octet

vec_sro

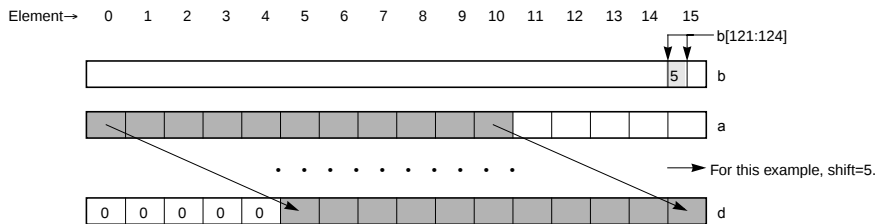
d = vec_sro(**a**,**b**)

```

m ← b[121:124]
do i=0 to 15
  j ← i - m
  if j ≥ 0
    then d{i} ← a{j}
    else d{i} ← 0
  end
end

```

The result is obtained by shifting (unsigned) **a** right by a number of bytes specified by the shifting the value of the last element of **b** by 3 bits. The valid combinations of argument types and the corresponding result types for **d** = vec_sro(**a**,**b**) are shown in Figure 4-125.



d	a	b	maps to
vector unsigned char	vector unsigned char	vector unsigned char	vsro d,a,b
	vector unsigned char	vector signed char	
vector signed char	vector signed char	vector unsigned char	
	vector signed char	vector signed char	
vector unsigned short	vector unsigned short	vector unsigned char	
	vector unsigned short	vector signed char	
vector signed short	vector signed short	vector unsigned char	
	vector signed short	vector signed char	
vector pixel	vector pixel	vector unsigned char	
	vector pixel	vector signed char	
vector unsigned int	vector unsigned int	vector unsigned char	
	vector unsigned int	vector signed char	
vector signed int	vector signed int	vector unsigned char	
	vector signed int	vector signed char	
vector float	vector float	vector unsigned char	
	vector float	vector signed char	

Figure 4-125. Right Byte Shift of Vector (128-Bit)

vec_st

Vector Store Indexed

vec_st

`vec_st(a,b,c)`

```
EA ← BoundAlign((b + c), 16)
MEM(EA,16) ← a
```

Each operation performs a 16-byte store of the value of `a` at a 16-byte aligned address. The `b` is taken to be an integer value, while `c` is a pointer. `BoundAlign(b+c,16)` is the largest value less than or equal to `b+c` that is a multiple of 16. This is not, by itself, an acceptable way to store aligned data to unaligned addresses. This store is the one that is generated for a storing dereference of a pointer to a vector type. Plain `char *` is excluded in the mapping for `c`. The valid combinations of argument types for `vec_st(a,b,c)` are shown in Table 4-18. The result type is `void`.

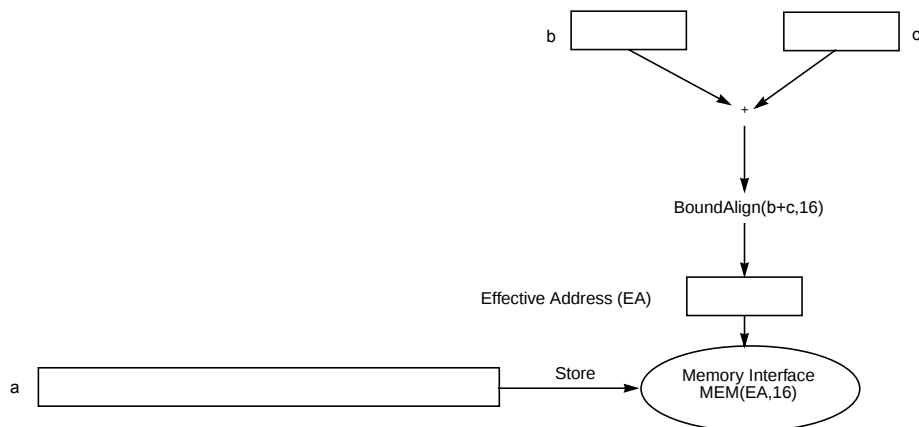


Figure 4-126. Vector Store Indexed

Table 4-18. vec_st—Vector Store Indexed Argument Types

a	b	c	Maps to
vector unsigned char	any integral type	vector unsigned char *	stvx a,b,c
vector unsigned char	any integral type	unsigned char *	
vector signed char	any integral type	vector signed char *	
vector signed char	any integral type	signed char *	
vector bool char	any integral type	vector bool char *	
vector bool char	any integral type	unsigned char *	
vector bool char	any integral type	signed char *	
vector unsigned short	any integral type	vector unsigned short *	
vector unsigned short	any integral type	unsigned short *	
vector signed short	any integral type	vector signed short *	
vector signed short	any integral type	short *	
vector bool short	any integral type	vector bool short *	
vector bool short	any integral type	unsigned short *	
vector bool short	any integral type	short *	
vector pixel	any integral type	vector pixel short *	
vector pixel	any integral type	unsigned short *	
vector pixel	any integral type	short *	
vector unsigned int	any integral type	vector unsigned int *	
vector unsigned int	any integral type	unsigned int *	
vector signed int	any integral type	vector signed int *	
vector signed int	any integral type	int *	
vector bool int	any integral type	vector bool int *	
vector bool int	any integral type	unsigned int *	
vector bool int	any integral type	int *	
vector float	any integral type	vector float *	
vector float	any integral type	float *	



vec_ste

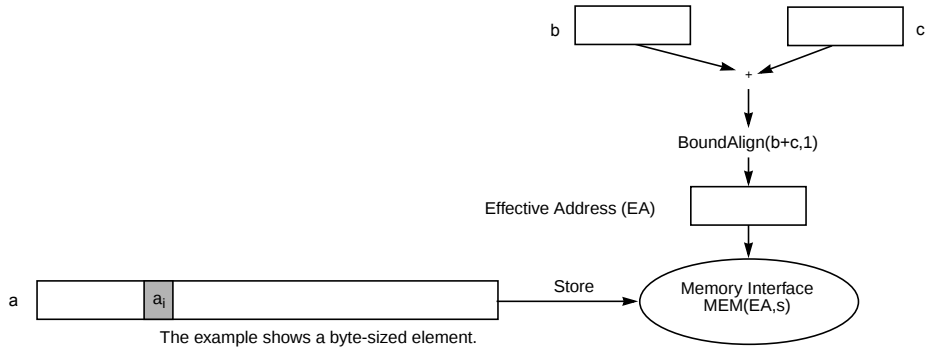
Vector Store Element Indexed

vec_ste

`vec_ste(a,b,c)`

```
s ← 16/(number of elements)
EA ← BoundAlign (b + c,s)
i ← mod(EA,16)/s
MEM(EA,s) ← ai
```

A single element of `a` is stored at the effective address. `BoundAlign(b+c,s)` is the largest value less than or equal to `b+c` that is a multiple of `s`, where `s` is 1 for `char` pointers, 2 for `short` pointers, and 4 for `int` or `float` pointers. The element stored is the one whose position in the register matches the position of the adjusted address relative to 16-byte alignment (A16). If you do not know the alignment of the sum of `b` and `c`, you will not know which element is stored. Plain `char *` is excluded in the mapping for `c`. The valid combinations of argument types for `vec_ste(a,b,c)` are shown in Figure 4-127. The result type is `void`.



a	b	c	Maps to
vector unsigned char	any integral type	unsigned char *	stvebx a,b,c
vector signed char	any integral type	signed char *	
vector bool char	any integral type	unsigned char *	
vector bool char	any integral type	signed char *	
vector unsigned short	any integral type	unsigned short *	stvehx a,b,c
vector signed short	any integral type	short *	
vector bool short	any integral type	unsigned short *	
vector bool short	any integral type	short *	
vector pixel	any integral type	unsigned short *	
vector pixel	any integral type	short *	
vector unsigned int	any integral type	unsigned int *	stvewx a,b,c
vector unsigned int	any integral type	unsigned int *	
vector signed int	any integral type	int *	
vector signed int	any integral type	int *	
vector bool int	any integral type	unsigned int *	
vector bool int	any integral type	unsigned int *	
vector bool int	any integral type	int *	
vector bool int	any integral type	int *	
vector float	any integral type	float *	

Figure 4-127. Vector Store Element

vec_stl

Vector Store Indexed LRU

vec_stl

vec_stl(a,b,c)

```
EA ← BoundAlign(b + c, 16)
MEM(EA,16) ← a
```

Each operation performs a 16-byte store of the value of a at a 16-byte aligned address. The b is taken to be an integer value, while c is a pointer. BoundAlign(b+c,16) is the largest value less than or equal to a b+c that is a multiple of 16. This is not, by itself, an acceptable way to store aligned data to unaligned addresses. The cache line stored into is marked Least Recently Used (LRU). Plain char * is excluded in the mapping for c. The valid combinations of argument types for vec_stl(a,b,c) are shown in Table 4-19. The result type is void.

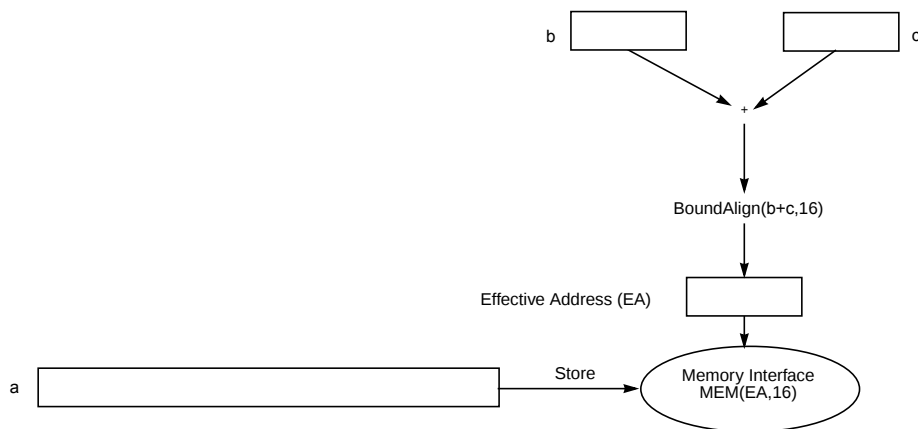


Figure 4-128. Vector Store Indexed LRU

Table 4-19vec_stl—Vector Store Index Argument Types

a	b	c	Maps to
vector unsigned char	any integral type	vector unsigned char *	stvxl a,b,c
vector unsigned char	any integral type	unsigned char *	
vector signed char	any integral type	vector signed char *	
vector signed char	any integral type	signed char *	
vector bool char	any integral type	vector bool char *	
vector bool char	any integral type	unsigned char *	
vector bool char	any integral type	signed char *	
vector unsigned short	any integral type	vector unsigned short *	
vector unsigned short	any integral type	unsigned short *	
vector signed short	any integral type	vector signed short *	
vector signed short	any integral type	short *	
vector bool short	any integral type	vector bool short *	
vector bool short	any integral type	unsigned short *	
vector bool short	any integral type	short *	
vector pixel	any integral type	vector pixel *	
vector pixel	any integral type	unsigned short *	
vector pixel	any integral type	short *	
vector unsigned int	any integral type	vector unsigned int *	
vector unsigned int	any integral type	unsigned int *	
vector signed int	any integral type	vector signed int *	
vector signed int	any integral type	int *	
vector bool int	any integral type	vector bool int *	
vector bool int	any integral type	unsigned int *	
vector bool int	any integral type	unsigned int *	
vector bool int	any integral type	int *	
vector float	any integral type	vector float *	
vector float	any integral type	float *	

vec_sub

Vector Subtract

vec_sub

d = vec_sub(**a**,**b**)

- Integer Subtract:

```
n ← number of elements
do i=0 to n-1
  di ← ai - bi
end
```

- Floating-Point Subtract:

```
do i=0 to 3
  di ← ai -fp bi
end
```

Each element of the result is the difference between the corresponding elements of **a** and **b**. The arithmetic is modular for integer types.

For `vector float` argument types, if `VSCR[NJ] = 1`, every denormalized `vector float` operand element is truncated to a 0 of the same sign before the operation is carried out, and each denormalized `vector float` result element truncates to a 0 of the same sign.

The valid combinations of argument types and the corresponding result types for **d** = vec_sub(**a**,**b**) are shown in Figure 4-129, Figure 4-130, Figure 4-131, and Figure 4-132.

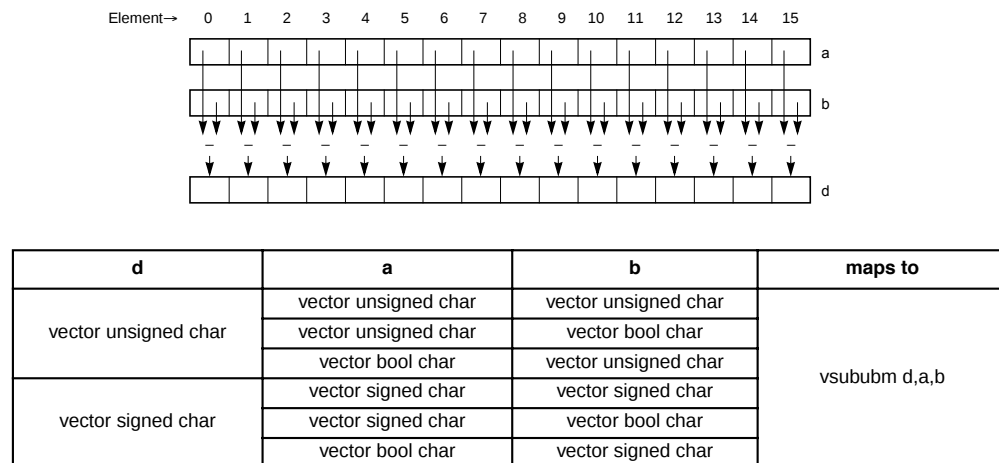


Figure 4-129. Subtract Sixteen Integer Elements (8-bit)

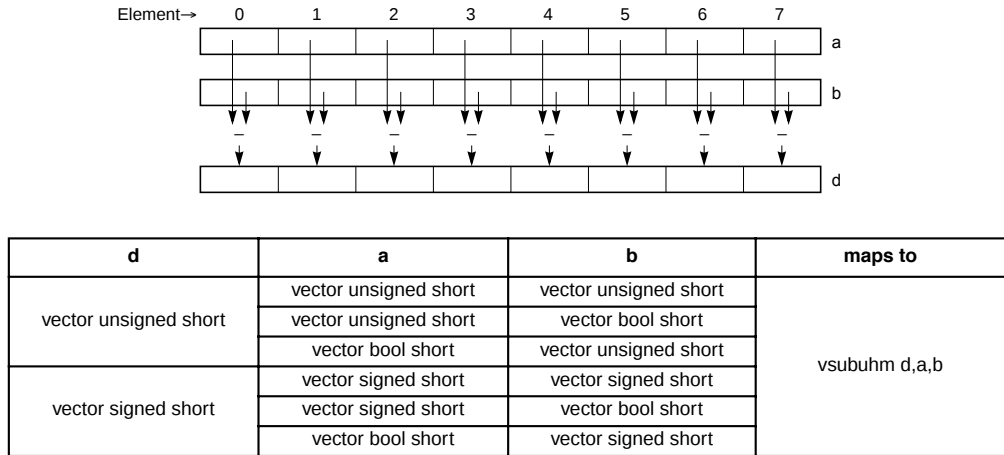


Figure 4-130. Subtract Eight Integer Elements (16-bit)

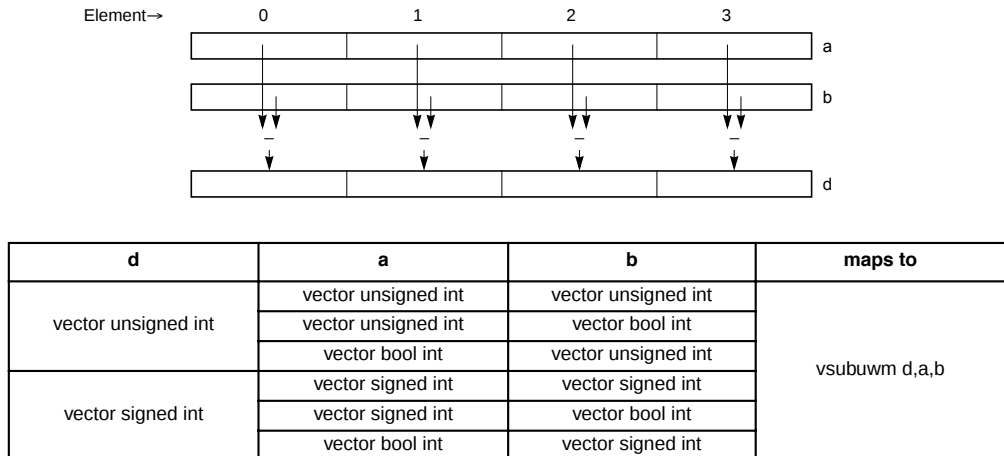


Figure 4-131. Subtract Four Integer Elements (32-bit)

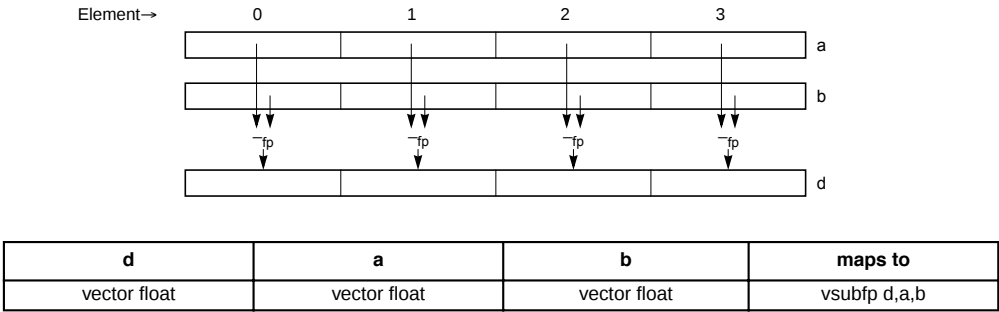


Figure 4-132. Subtract Four Floating-Point Elements (32-bit)

vec_subc

Vector Subtract Carryout

vec_subc

d = vec_subc(**a**,**b**)

```
do i=0 to 3
  di = BorrowOut(ai - bi)
end
```

Each element of **b** is subtracted from the corresponding element in **a**. The borrow from each difference is complemented and zero-extended and placed into the corresponding element of **d**. BorrowOut ($a - b$) is 0 if a borrow occurred and 1 if no borrow occurred. The valid combination of argument types and the corresponding result type for **d** = vec_subc(**a**,**b**) are shown in Figure 4-133.

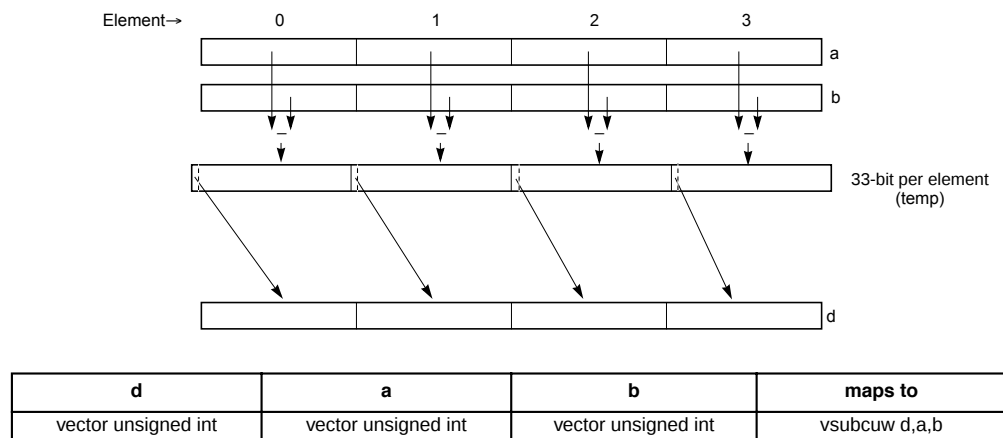


Figure 4-133. Carryout of Four Unsigned Integer Subtracts (32-bit)

vec_subs

Vector Subtract Saturated

vec_subs

```
d = vec_subs(a,b)

n ← number of elements
do i=0 to n-1
  di ← Saturate (ai - bi)
end
```

Each element of the result is the saturated difference between the corresponding elements of a and b. If saturation occurs, VSCR[SAT] is set (see Table 4-1). The valid combinations of argument types and the corresponding result types for **d = vec_subs(a,b)** are shown in Figure 4-134, Figure 4-135, and Figure 4-136.

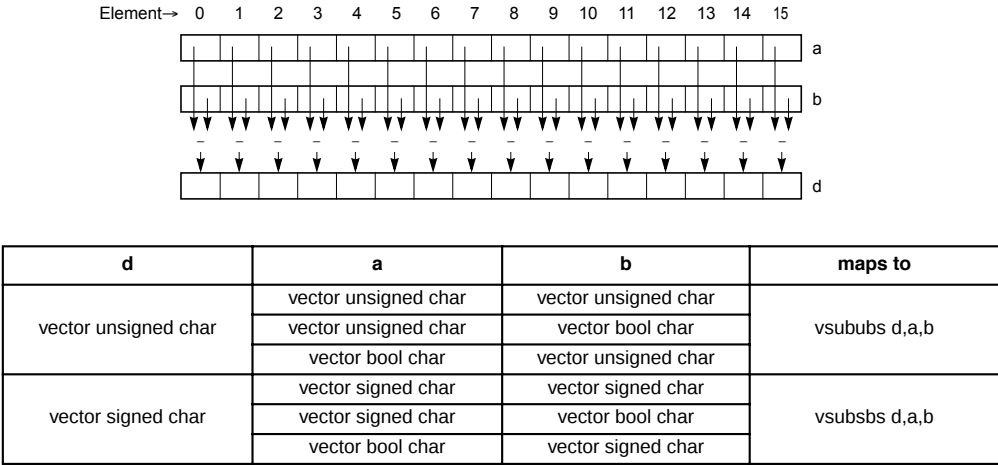


Figure 4-134. Subtract Saturating Sixteen Integer Elements (8-bit)

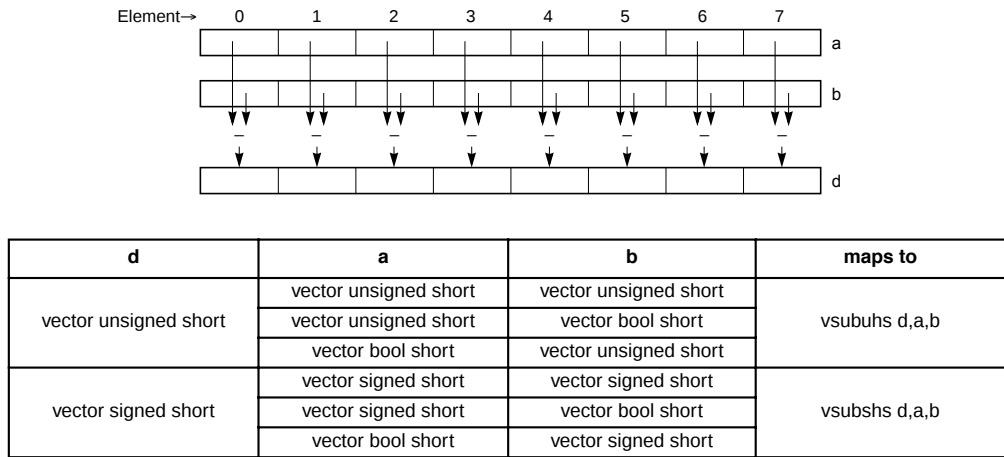


Figure 4-135. Subtract Saturating Eight Integer Elements (16-bit)

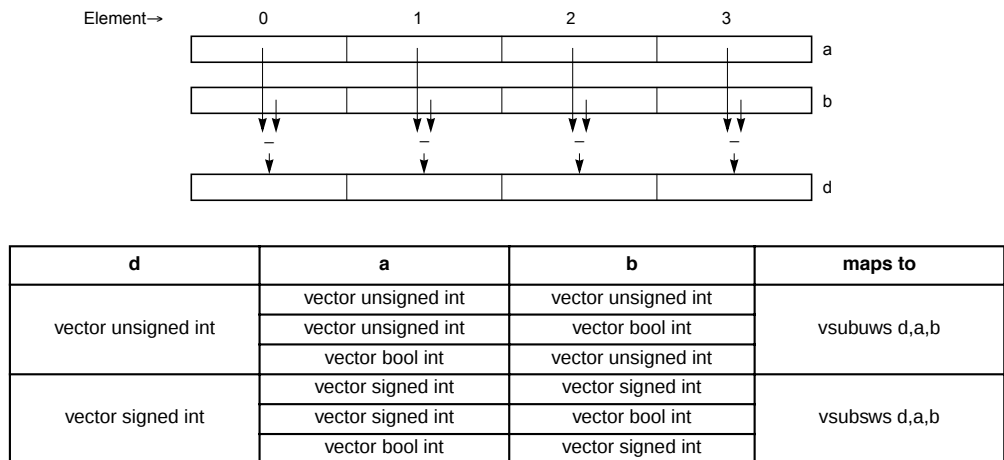


Figure 4-136. Subtract Saturating Four Integer Elements (32-bit)

vec_sum4s

Vector Sum Across Partial (1/4) Saturated

vec_sum4s

d = vec_sum4s(**a**,**b**)

- For **a** with 8-bit elements:

```
do i=0 to 3
  di ← Saturate (a4i+ a4i+1 + a4i+2 + a4i+3 + bi)
end
```

- For **a** with 16-bit elements:

```
do i=0 to 3
  di ← Saturate(a2i+ a2i+1 + bi)
end
```

Each element of the result is the 32-bit saturated sum of the corresponding element in **b** and all elements in **a** with positions overlapping those of that element. If saturation occurs, VSCR[SAT] is set (see Table 4-1). The valid combinations of argument types and the corresponding result types for **d** = vec_sum4s(**a**,**b**) are shown in Figure 4-137 and Figure 4-138.

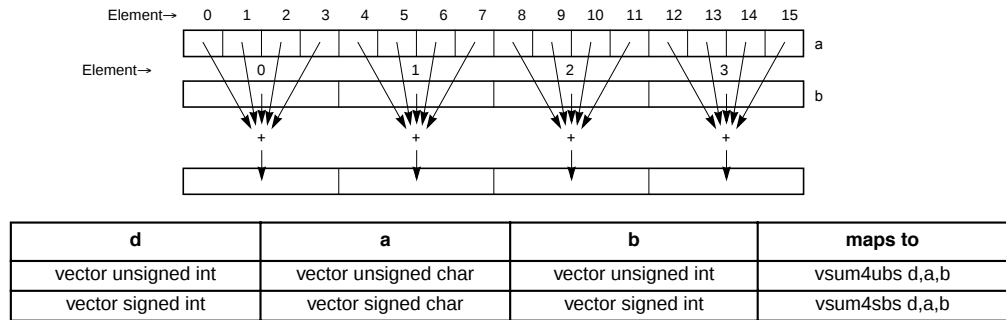


Figure 4-137. Four Sums in the Integer Elements (32-Bit)

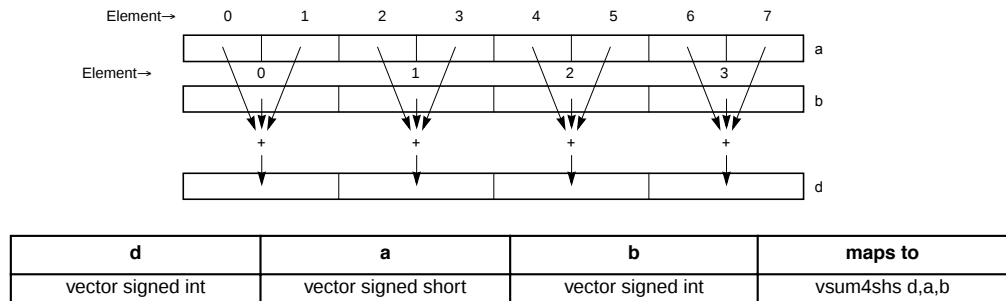


Figure 4-138. Four Sums in the Integer Elements (32-Bit)

vec_sum2s

Vector Sum Across Partial (1/2) Saturated

vec_sum2s

d = vec_sum2s(**a**,**b**)

```
do i=0 to 1
  d2i ← 0
  d2i+1 ← Saturate(a2i + a2i+1 + b2i+1)
end
```

The first and third elements of the result are 0. The second element of the result is the 32-bit saturated sum of the first two elements of a and the second element of b. The fourth element of the result is the 32-bit saturated sum of the last two elements of a and the fourth element of b. If saturation occurs, VSCR[SAT] is set (see Table 4-1). The valid combination of argument types and the corresponding result type for **d** = vec_sum2s(**a**,**b**) are shown in Figure 4-139.

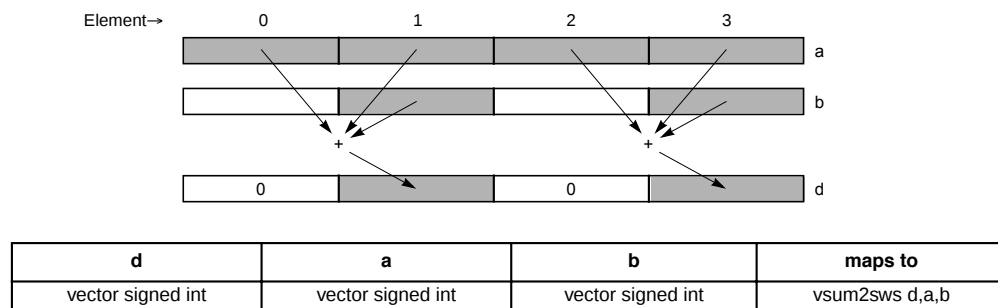


Figure 4-139. Two Saturated Sums in the Four Signed Integer Elements (32-Bit)

vec_sums

Vector Sum Saturated

```
d = vec_sums(a,b)  
do i=0 to 2  
  di ← 0  
end  
d3 ← Saturate(a0 + a1 + a2 + a3 + b3)
```

The first three elements of the result are 0. The fourth element of the result is the 32-bit saturated sum of all elements of **a** and the fourth element of **b**. If saturation occurs, VSCR[SAT] is set (see Table 4-1). The valid combination of argument types and the corresponding result type for **d** = vec_sums(**a**,**b**) are shown in Figure 4-140.

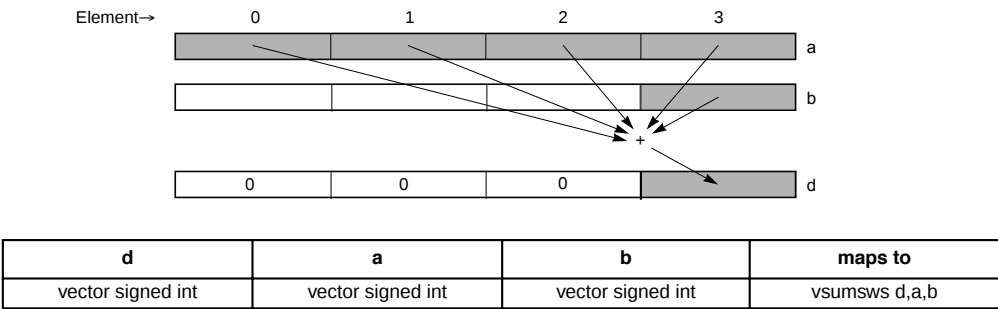


Figure 4-140. Saturated Sum of Five Signed Integer Elements (32-Bit)

vec_trunc

Vector Truncate

vec_trunc

d = vec_trunc(**a**)

```
do i=0 to 3
  di ← RndToFPITrunc(ai)
end
```

Each single-precision floating-point word element in **a** is rounded to a single-precision floating-point integer, using the Round-toward-Zero mode, and placed into the corresponding word element of **d**. Each element of the result is thus the value of the corresponding element of **a** truncated to an integral value.

The operation is independent of VSCR[NJ].

The valid argument type and corresponding result type for **d** = vec_trunc(**a**) are shown in Figure 4-141.

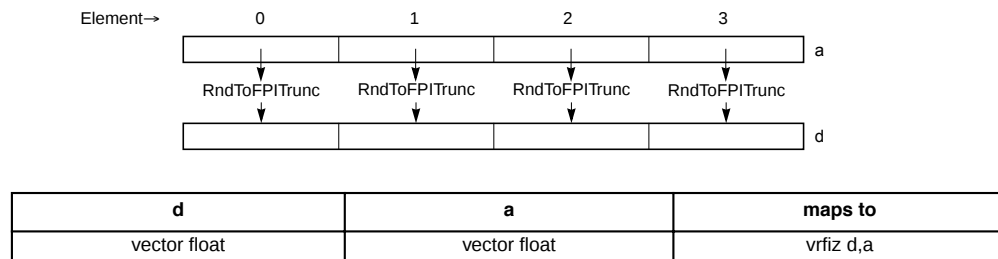


Figure 4-141. Round-to-Zero of Four Floating-Point Integer Elements (32-Bit)

vec_unpackh

Vector Unpack High Element

vec_unpackh

d = vec_unpackh(**a**)

- Integer value:

```
n ← number of elements in d
do i=0 to n-1
  di ← SignExtend(ai)
end
```

- Pixel value:

```
do i=0 to 3
  di ← SignExtend(ai[0]) || 000 || ai[1:5] || 000 || ai[6:10] || 000 || ai[11:15]
end
```

Each element of the result is the result of extending the corresponding half-width high element of **a**. The valid argument types and corresponding result types for

d = vec_unpackh(**a**) are shown in Figure 4-142, Figure 4-143, and Figure 4-144.

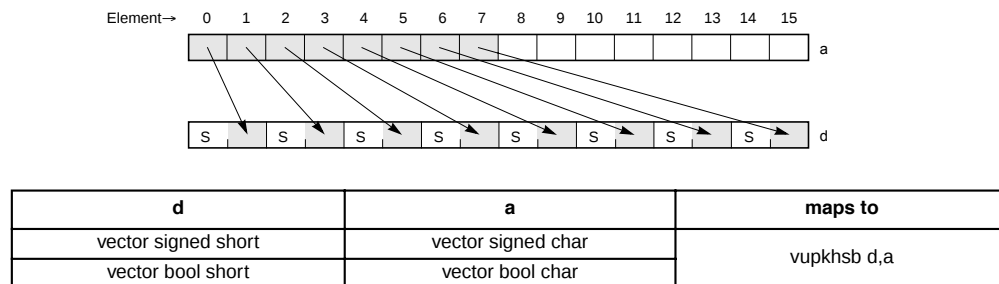


Figure 4-142. Unpack High-Order Elements (8-Bit) to Elements (16-Bit)

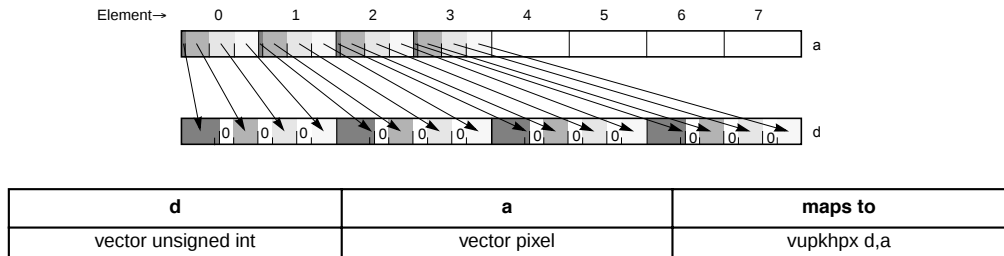


Figure 4-143. Unpack High-Order Pixel Elements (16-Bit) to Elements (32-Bit)

Programming note: Notice that the unpacking done by the vector unpack element operations for vector pixel values does not reverse the packing done by the vector pack pixel operation. Specifically, if a 16-bit pixel is unpacked to a 32-bit pixel which is then packed to a 16-bit pixel, the resulting 16-bit pixel will not, in general, be equal to the original 16-bit pixel (because, for each channel except the first, vector unpack element inserts high-order bits while vector pack pixel discards low-order bits.)

This was designed to optimize image processing where the unpacked values would be multiplied by small coefficients and accumulated in a digital filter. The usual transformation from the 16-bit pixel to a 32-bit pixel involves multiplication of the RGB channels by 255/31. This can be accomplished by replicating the 3 most significant bits in the least significant bits using the operations:

```
d = vec_unpackh(a);
d = (vector unsigned int) vec_or(vec_sl((vector unsigned char)d,
    (vector unsigned char)(3)),
    vec_sr((vector unsigned char)d,
    (vector unsigned char)(2)));
```

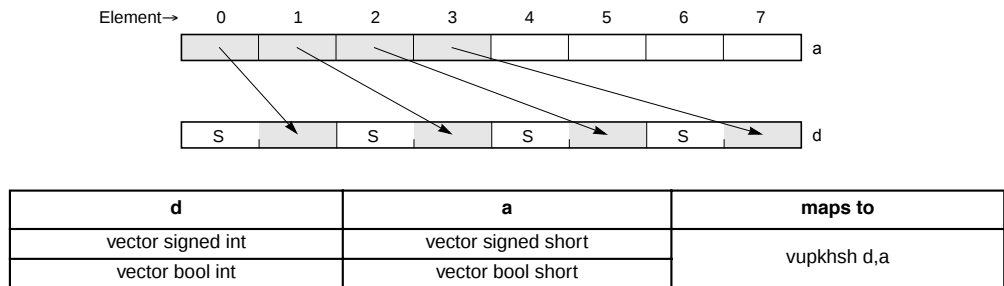


Figure 4-144. Unpack High-Order Signed Integer Elements (16-Bit) to Signed Integer Elements (32-Bit)

vec_unpackl

Vector Unpack Low Element

vec_unpackl

d = vec_unpackl(**a**)

- Integer value:

```
n ← number of elements in d
do i=0 to n-1
  di ← SignExtend(ai+n)
end
```

- Pixel value:

```
do i=0 to 3
  di ← SignExtend(ai+n[0]) || 000 || ai+n[1:5] || 000 || ai+n[6:10] || 000 || ai+n[11:15]
end
```

Each element of the result is the result of extending the corresponding half-width low element of **a**. The valid argument types and corresponding result types for

d = vec_unpackl(**a**) are shown in Figure 4-145, Figure 4-146, and Figure 4-147.

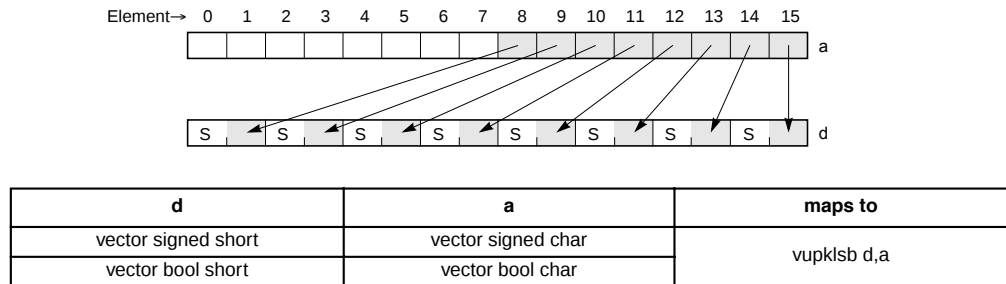


Figure 4-145. Unpack Low-Order Elements (8-Bit) to Elements (16-Bit)

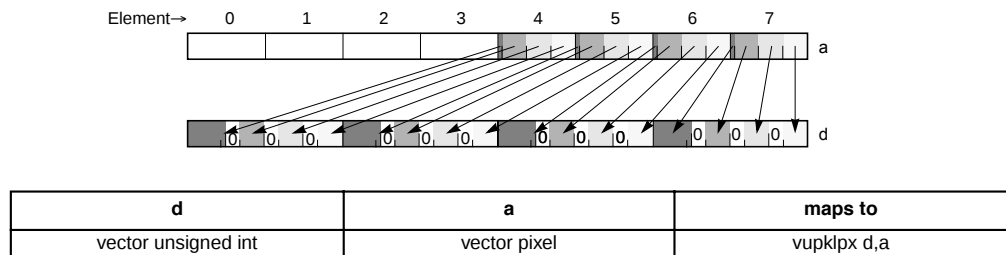


Figure 4-146. Unpack Low-Order Pixel Elements (16-Bit) to Elements (32-Bit)

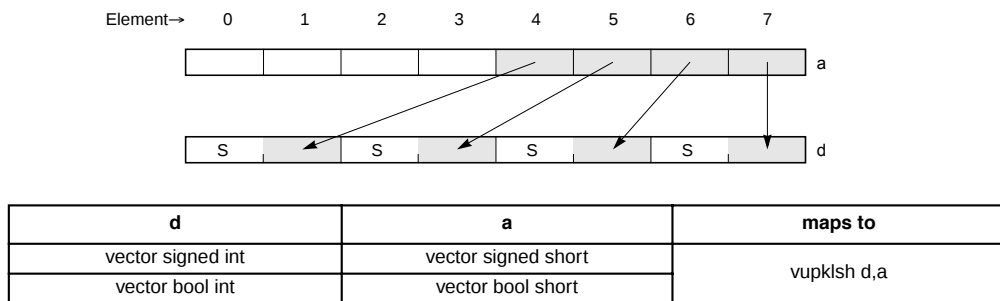


Figure 4-147. Unpack Low-Order Signed Integer Elements (16-Bit) to Signed Integer Elements (32-Bit)

Programming note: Notice that the unpacking done by the vector unpack element operations for vector pixel values does not reverse the packing done by the vector pack pixel operation. Specifically, if a 16-bit pixel is unpacked to a 32-bit pixel which is then packed to a 16-bit pixel, the resulting 16-bit pixel will not, in general, be equal to the original 16-bit pixel (because, for each channel except the first, vector unpack element inserts high-order bits while vector pack pixel discards low-order bits.)

This was designed to optimize image processing where the unpacked values would be multiplied by small coefficients and accumulated in a digital filter. The usual transformation from the 16-bit pixel to a 32-bit pixel involves multiplication of the RGB channels by 255/31. This can be accomplished by replicating the 3 most significant bits in the least significant bits using the operations:

```
d = vec_unpackh(a);
d = (vector unsigned int) vec_or(vec_sl((vector unsigned char)d,
                                         (vector unsigned char)(3)),
                                vec_sr((vector unsigned char)d,
                                         (vector unsigned char)(2)));
```



vec_xor

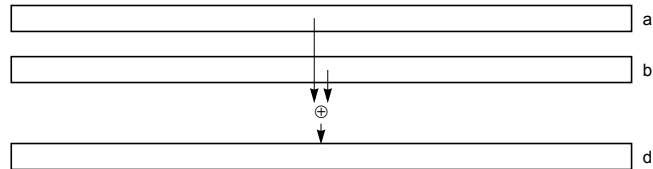
Vector Logical XOR

d = vec_xor(**a**,**b**)

d ← **a** ⊕ **b**

Each bit of the result is the logical XOR of the corresponding bits of **a** and **b**. The valid combinations of argument types and the corresponding result types for **d** = vec_xor(**a**,**b**) are shown in Figure 4-148.

vec_xor



d	a	b	maps to
vector unsigned char	vector unsigned char	vector unsigned char	vxor d,a,b
	vector unsigned char	vector bool char	
	vector bool char	vector unsigned char	
vector signed char	vector signed char	vector signed char	
	vector signed char	vector bool char	
	vector bool char	vector signed char	
vector bool char	vector bool char	vector bool char	
vector unsigned short	vector unsigned short	vector unsigned short	
	vector unsigned short	vector bool short	
	vector bool short	vector unsigned short	
vector signed short	vector signed short	vector signed short	
	vector signed short	vector bool short	
	vector bool short	vector signed short	
vector bool short	vector bool short	vector bool short	
vector unsigned int	vector unsigned int	vector unsigned int	
	vector unsigned int	vector bool int	
	vector bool int	vector unsigned int	
vector signed int	vector signed int	vector signed int	
	vector signed int	vector bool int	
	vector bool int	vector signed int	
vector bool int	vector bool int	vector bool int	
vector float	vector bool int	vector float	
	vector float	vector bool int	
	vector float	vector float	

Figure 4-148. Logical Bit-Wise XOR

4.5 AltiVec Predicates

The AltiVec predicates all begin with `vec_all_` or `vec_any_`. The AltiVec predicates are organized alphabetically by predicate name with a definition of the permitted generic AltiVec predicates. The specific operations do not exist for the predicates.

Where possible, the description is supported by reference figures indicating data modifications and including a table that lists:

- the valid set of argument types for that predicate, and
- the specific AltiVec instruction generated for that set of arguments. The AltiVec instruction is in the form `v-----. x,a,b`, where `v-----.` represents the instruction and `x,a,b` represent the operands. The `x` represents an unused vector result of the vector compare instruction used to implement the predicate. The order of operands listed after the instruction indicate the order in which they are applied for that predicate.

For example,

`vec_any_lt(vector unsigned char, vector unsigned char)`
maps to the instruction

`vcmpgtb. x,b,a`

indicating that the operands are applied in reverse order for this predicate.

vec_all_eq

All Elements Equal

vec_all_eq

d = vec_all_eq(**a**,**b**)

```
n ← number of elements
if each ai =int bi, where i ranges from 0 to n-1
then d ← 1
else d ← 0
```

The predicate vec_all_eq returns 1 if every element of a is equal to the corresponding element of b. Otherwise, it returns 0.

For vector float argument types, if VSCR[NJ] = 1, every denormalized floating-point operand element is truncated to 0 before the comparison.

The valid combinations of argument types and the corresponding result type for **d** = vec_all_eq(**a**,**b**) are shown in Figure 4-149, Figure 4-150, Figure 4-151, and Figure 4-152.

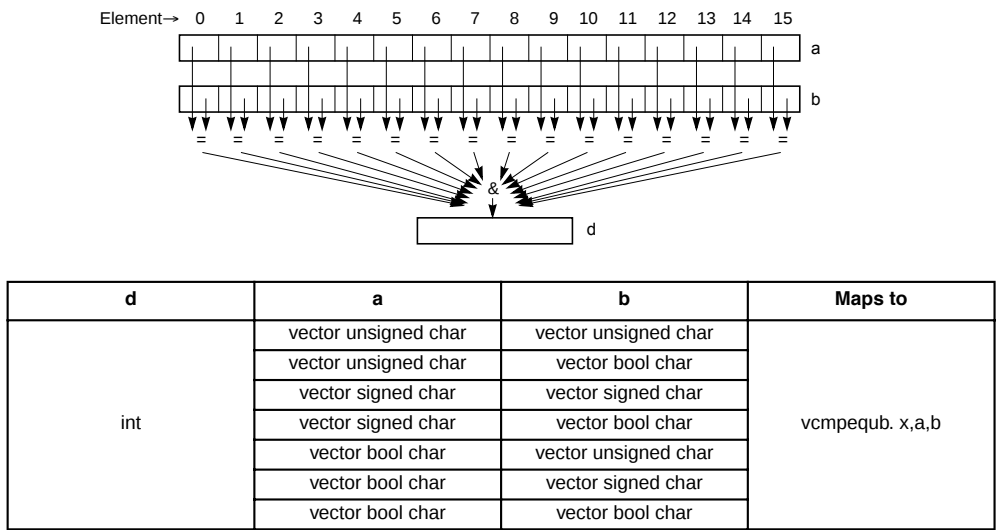
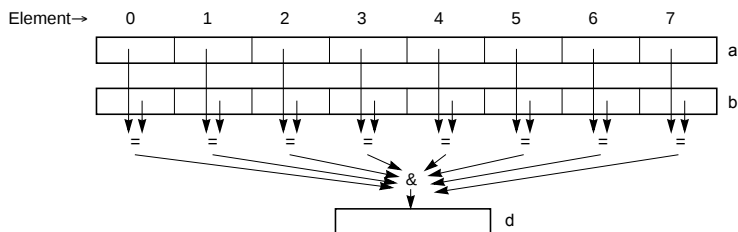
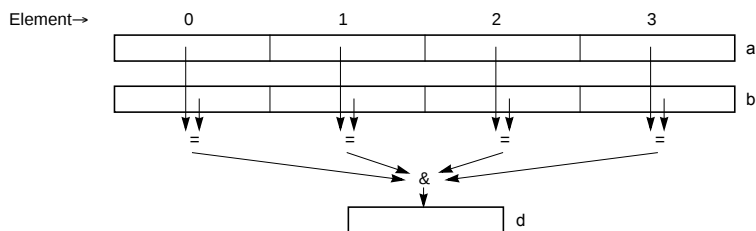


Figure 4-149. All Equal of Sixteen Integer Elements (8-bits)



d	a	b	Maps to
int	vector unsigned short	vector unsigned short	vcmpquh. x,a,b
	vector unsigned short	vector bool short	
	vector signed short	vector signed short	
	vector signed short	vector bool short	
	vector bool short	vector unsigned short	
	vector bool short	vector signed short	
	vector bool short	vector bool short	
	vector pixel	vector pixel	

Figure 4-150. All Equal of Eight Integer Elements (16-Bit)



d	a	b	Maps to
int	vector unsigned int	vector unsigned int	vcmpquw. x,a,b
	vector unsigned int	vector bool int	
	vector signed int	vector signed int	
	vector signed int	vector bool int	
	vector bool int	vector unsigned int	
	vector bool int	vector signed int	
	vector bool int	vector bool int	

Figure 4-151. All Equal of Four Integer Elements (32-Bit)

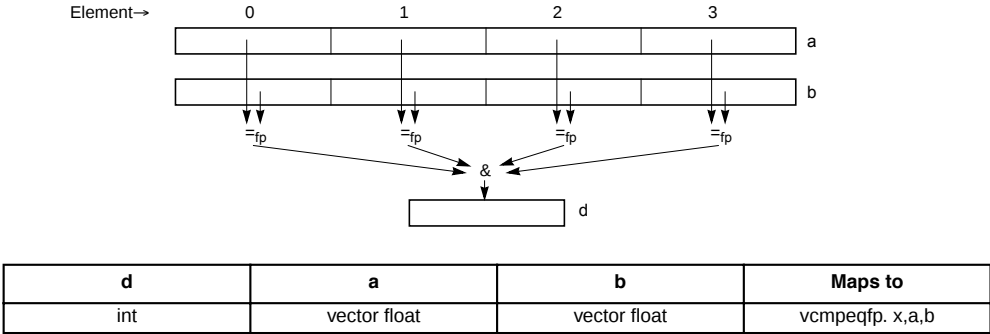


Figure 4-152. All Equal of Four Floating-Point Elements (32-Bit)

vec_all_ge

All Elements Greater Than or Equal

vec_all_ge

d = vec_all_ge(**a**,**b**)

```

n ← number of elements
if each  $a_i \geq b_i$ , where i ranges from 0 to n-1
then d ← 1
else d ← 0
    
```

The predicate vec_all_ge returns 1 if every element of a is greater than or equal to the corresponding element of b. Otherwise, it returns 0.

For vector float argument types, if VSCR[NJ] = 1, every denormalized floating-point operand element is truncated to 0 before the comparison.

The valid combinations of argument types and the corresponding result type for **d** = vec_all_ge(**a**,**b**) are shown in Figure 4-153, Figure 4-154, Figure 4-155, and Figure 4-156.

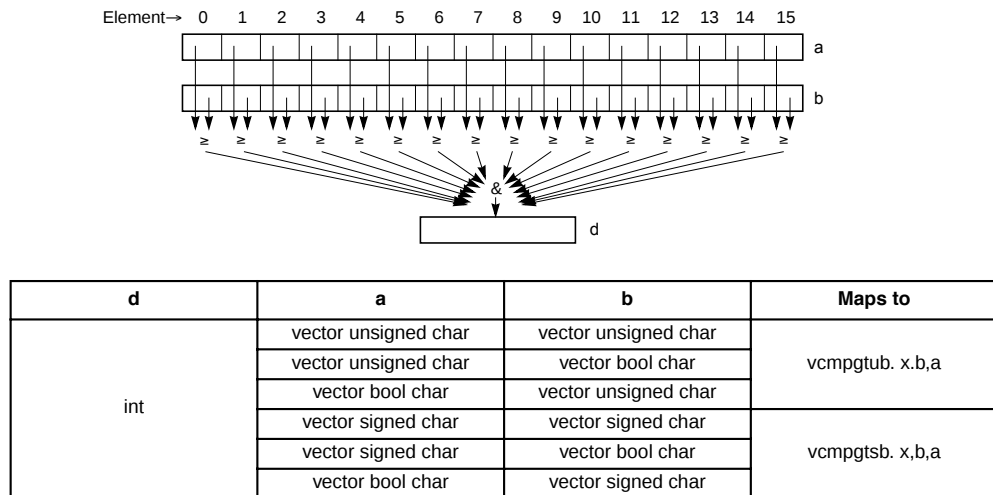
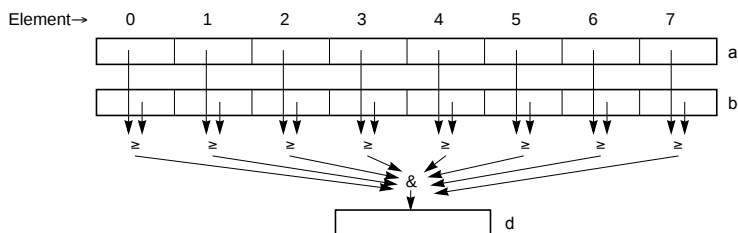
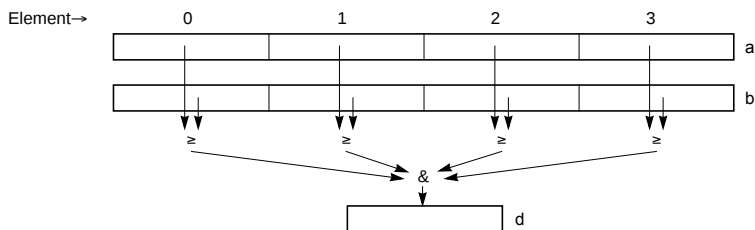


Figure 4-153. All Greater Than or Equal of Sixteen Integer Elements (8-bits)



d	a	b	Maps to
int	vector unsigned short	vector unsigned short	vcmpgtuh. x,b,a
	vector unsigned short	vector bool short	
	vector bool short	vector unsigned short	
	vector signed short	vector signed short	vcmpgtsh. x,b,a
	vector signed short	vector bool short	
	vector bool short	vector signed short	

Figure 4-154. All Greater Than or Equal of Eight Integer Elements (16-Bit)



d	a	b	Maps to
int	vector unsigned int	vector unsigned int	vcmpgtuw. x,b,a
	vector unsigned int	vector bool int	
	vector bool int	vector unsigned int	
	vector signed int	vector signed int	vcmpgtsw. x,b,a
	vector signed int	vector bool int	
	vector bool int	vector signed int	

Figure 4-155. All Greater Than or Equal of Four Integer Elements (32-Bit)

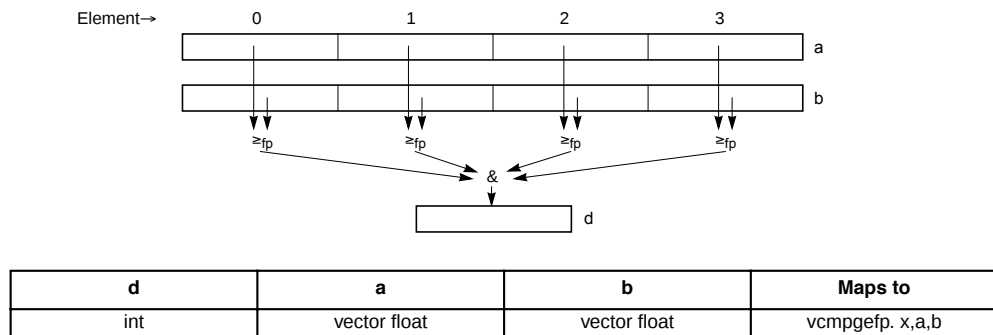


Figure 4-156. All Greater Than or Equal of Four Floating-Point Elements (32-Bit)

vec_all_gt

All Elements Greater Than

vec_all_gt

d = vec_all_gt(**a**,**b**)

```

n ← number of elements
if each  $a_i > b_i$ , where i ranges from 0 to n-1
then d ← 1
else ← 0

```

The predicate `vec_all_gt` returns 1 if every element of **a** is greater than the corresponding element of **b**. Otherwise, it returns 0.

For `vector float` argument types, if `VSCR[NJ] = 1`, every denormalized floating-point operand element is truncated to 0 before the comparison.

The valid combinations of argument types and the corresponding result type for **d** = `vec_all_gt(a,b)` are shown in Figure 4-157, Figure 4-158, Figure 4-159, and Figure 4-160.

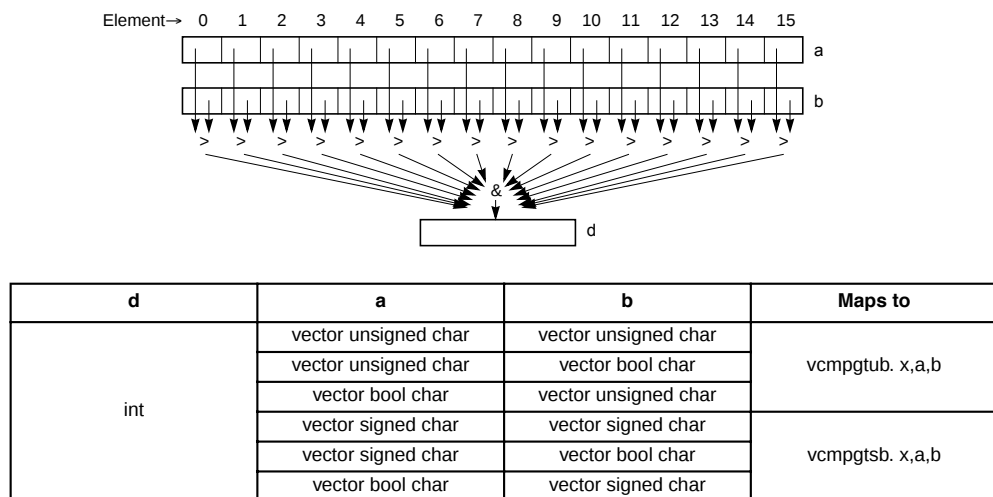
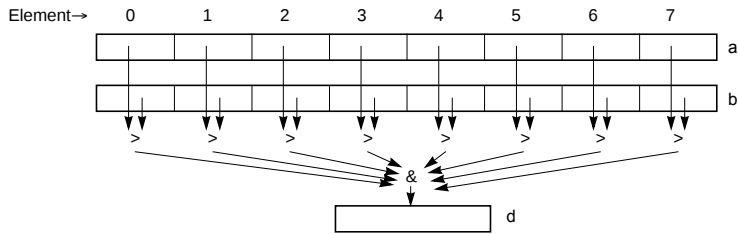
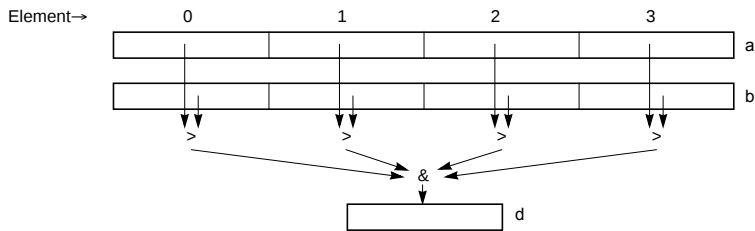


Figure 4-157. All Greater Than of Sixteen Integer Elements (8-bits)



d	a	b	Maps to
int	vector unsigned short	vector unsigned short	vcmpgtuh. x,a,b
	vector unsigned short	vector bool short	
	vector bool short	vector unsigned short	
	vector signed short	vector signed short	vcmpgtsh. x,a,b
	vector signed short	vector bool short	
	vector bool short	vector signed short	

Figure 4-158. All Greater Than of Eight Integer Elements (16-Bit)



d	a	b	Maps to
int	vector unsigned int	vector unsigned int	vcmpgtuw. x,a,b
	vector unsigned int	vector bool int	
	vector bool int	vector unsigned int	
	vector signed int	vector signed int	vcmpgtsw. x,a,b
	vector signed int	vector bool int	
	vector bool int	vector signed int	

Figure 4-159. All Greater Than of Four Integer Elements (32-Bit)

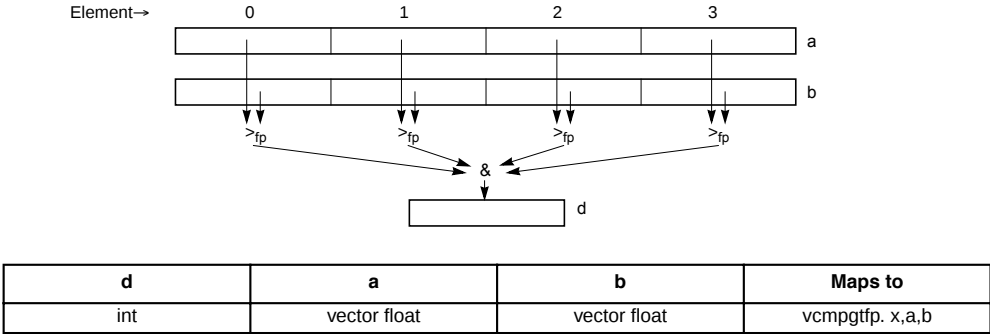


Figure 4-160. All Greater Than of Four Floating-Point Elements (32-Bit)

vec_all_in

All Elements in Bounds

vec_all_in

d = vec_all_in(**a**,**b**)

```

if each  $a_i \leq b_i$  and  $a_i \geq -b_i$ , where  $i$  ranges from 0 to 3
then d  $\leftarrow$  1
else  $\leftarrow$  0
    
```

The predicate `vec_all_in` returns 1 if every element of **a** is less than or equal to the corresponding element of **b** (high bound) and greater than or equal to the negative (NEG) of the corresponding element of **b** (low bound). Otherwise, it returns 0.

If `VSCR[NJ] = 1`, every denormalized floating-point operand element is truncated to 0 before the comparison.

The valid argument types and the corresponding result type for **d** = `vec_all_in(a,b)` are shown in Figure 4-161.

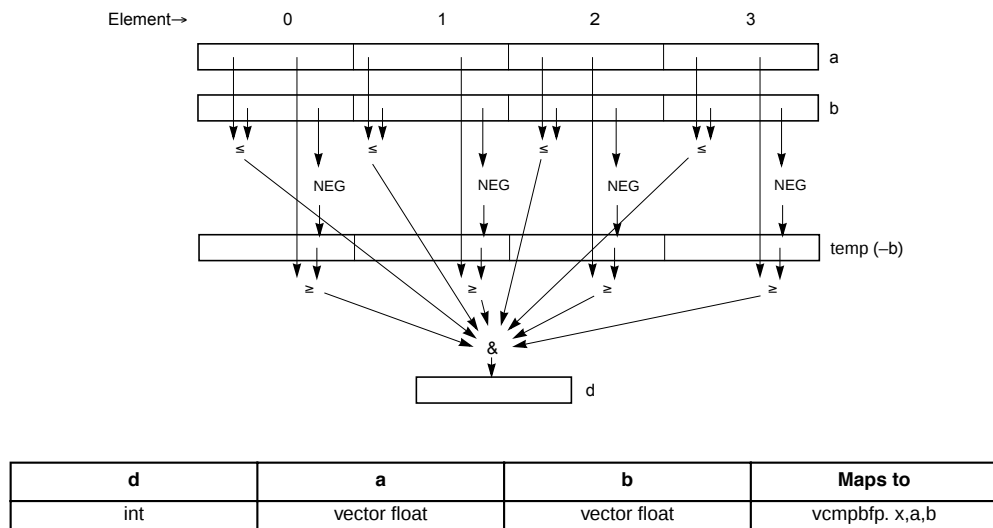


Figure 4-161. All in Bounds of Four Floating-Point Elements (32-Bit)

vec_all_le

All Elements Less Than or Equal

vec_all_le

d = vec_all_le(**a**,**b**)

```

n ← number of elements
if each  $a_i \leq b_i$ , where i ranges from 0 to n-1
then d ← 0
else d ← 1
    
```

The predicate vec_all_le returns 1 if every element of a is less than or equal to the corresponding element of b. Otherwise, it returns 0.

For vector float argument types, if VSCR[NJ] = 1, every denormalized floating-point operand element is truncated to 0 before the comparison.

The valid combinations of argument types and the corresponding result type for **d** = vec_all_le(**a**,**b**) are shown in Figure 4-162, Figure 4-163, Figure 4-164, and Figure 4-165.

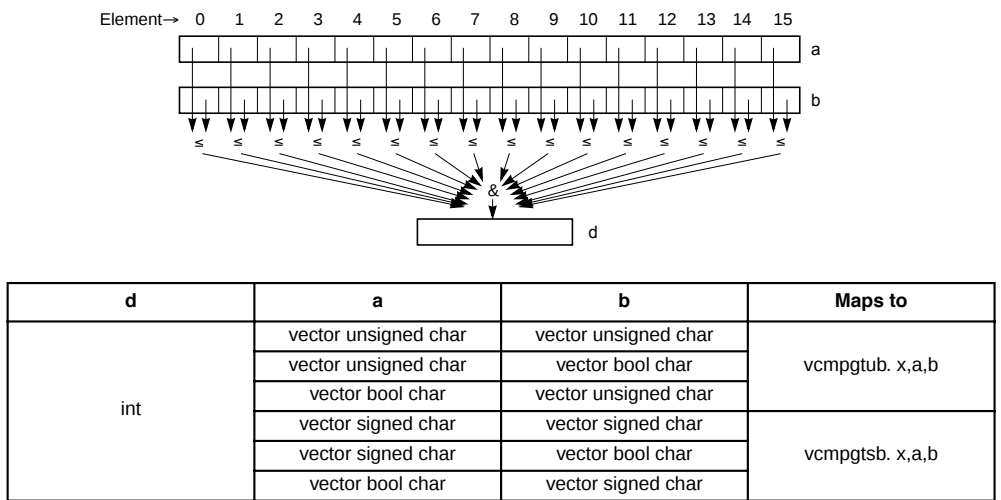


Figure 4-162. All Less Than or Equal of Sixteen Integer Elements (8-bits)

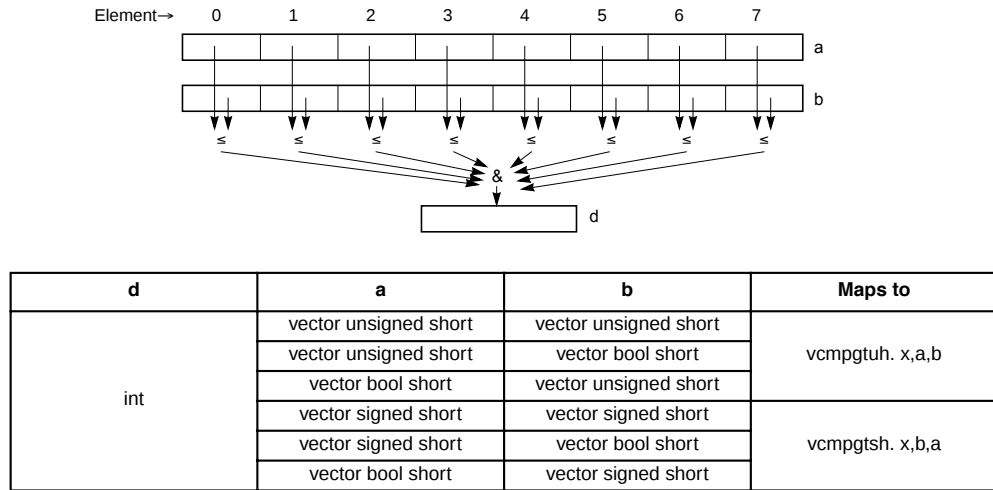


Figure 4-163. All Less Than or Equal of Eight Integer Elements (16-Bit)

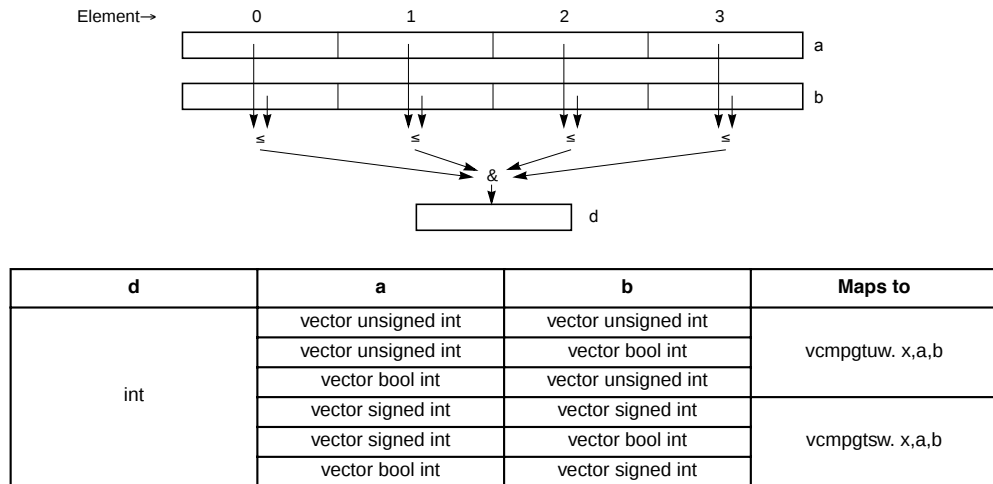


Figure 4-164. All Less Than or Equal of Four Integer Elements (32-Bit)

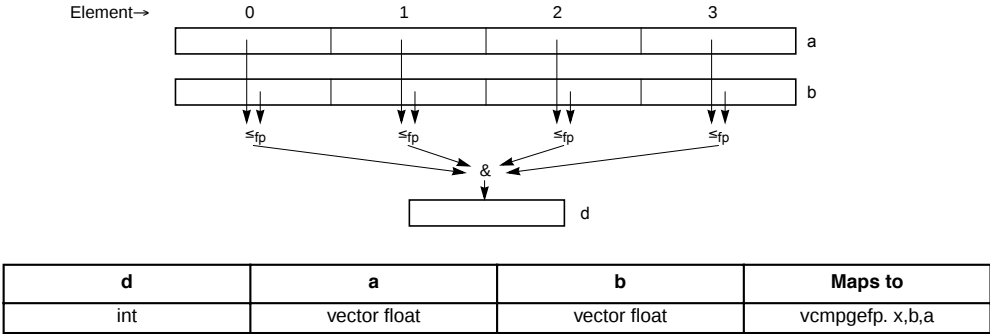


Figure 4-165. All Less Than or Equal of Four Floating-Point Elements (32-Bit)

vec_all_lt

All Elements Less Than

vec_all_lt

d = vec_all_lt(**a**,**b**)

```

n ← number of elements
if each ai < bi, where i ranges from 0 to n-1
then d ← 1
else d ← 0
    
```

The predicate `vec_all_lt` returns 1 if every element of `a` is less than the corresponding element of `b`. Otherwise, it returns 0.

For `vector float` argument types, if `VSCR[NJ] = 1`, every denormalized floating-point operand element is truncated to 0 before the comparison.

The valid combinations of argument types and the corresponding result type for `d = vec_all_lt(a,b)` are shown in Figure 4-166, Figure 4-167, Figure 4-168, and Figure 4-169.

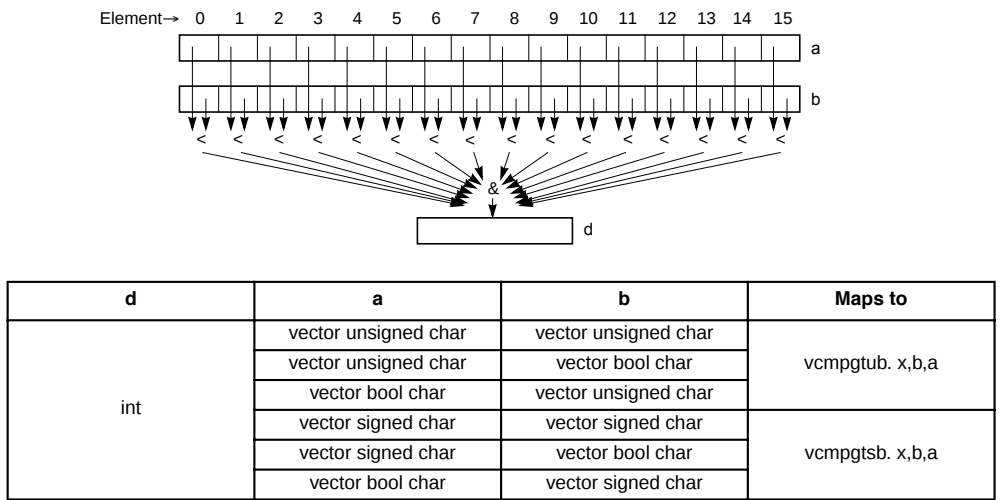
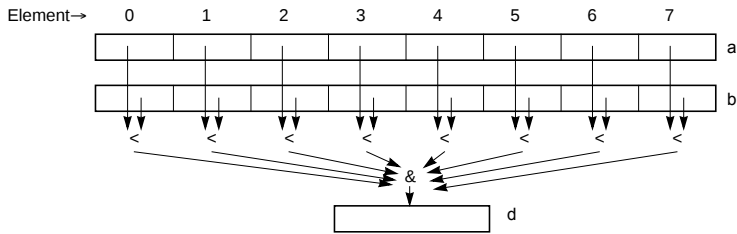
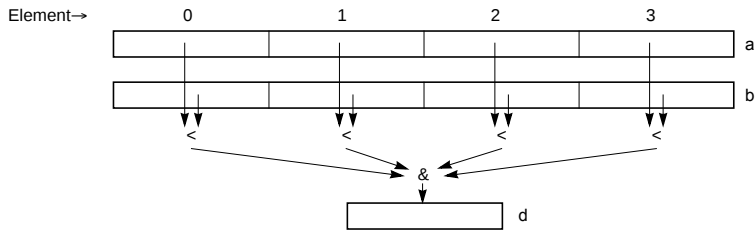


Figure 4-166. All Less Than of Sixteen Integer Elements (8-bits)



d	a	b	Maps to
int	vector unsigned short	vector unsigned short	vcmpgtuh. x,b,a
	vector unsigned short	vector bool short	
	vector bool short	vector unsigned short	
	vector signed short	vector signed short	vcmpgtsh. x,b,a
	vector signed short	vector bool short	
	vector bool short	vector signed short	

Figure 4-167. All Less Than of Eight Integer Elements (16-Bit)



d	a	b	Maps to
int	vector unsigned int	vector unsigned int	vcmpgtuw. x,b,a
	vector unsigned int	vector bool int	
	vector bool int	vector unsigned int	
	vector signed int	vector signed int	vcmpgtsw. x,b,a
	vector signed int	vector bool int	
	vector bool int	vector signed int	

Figure 4-168. All Less Than of Four Integer Elements (32-Bit)

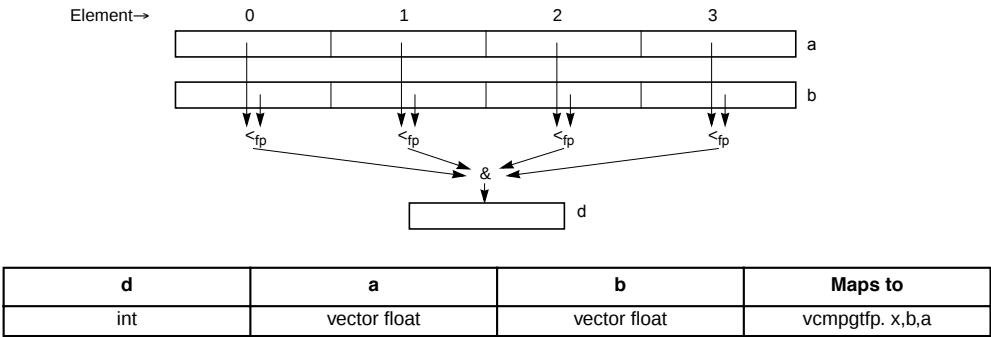


Figure 4-169. All Less Than of Four Floating-Point Elements (32-Bit)

vec_all_nan

All Elements Not a Number

vec_all_nan

```
d = vec_all_nan(a)  
  
if each ISNaN(ai) = 1, where i ranges from 0 to 3  
then d ← 1  
else d ← 0
```

The predicate `vec_all_nan` returns 1 if every element of `a` is Not a Number (NaN). Otherwise, it returns 0.

The operation is independent of VSCR[NJ].

The valid argument type and corresponding result type for `d = vec_all_nan(a)` are shown in Figure 4-170.

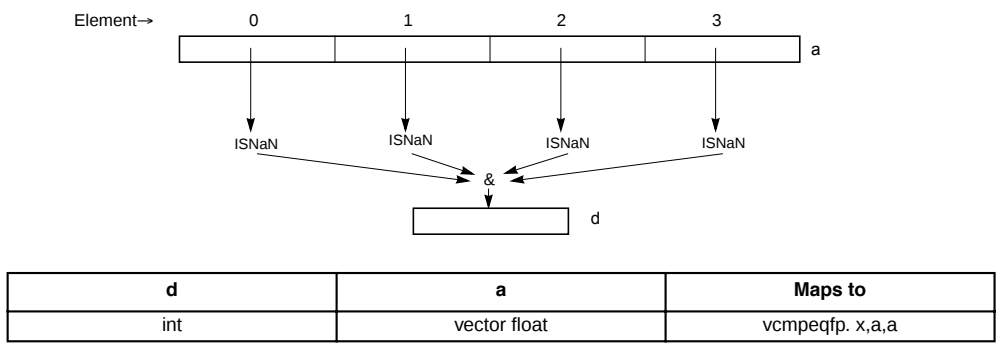


Figure 4-170. All NaN of Four Floating-Point Elements (32-Bit)

vec_all_ne

All Elements Not Equal

vec_all_ne

d = vec_all_ne(**a**,**b**)

```

n ← number of elements
if each  $a_i \neq b_i$ , where i ranges from 0 to n-1
then d ← 1
else d ← 0
    
```

The predicate `vec_all_ne` returns 1 if every element of `a` is not equal to ($\neq$) the corresponding element of `b`. Otherwise, it returns 0.

For `vector float` argument types, if `VSCR[NJ] = 1`, every denormalized floating-point operand element is truncated to 0 before the comparison.

The valid combinations of argument types and the corresponding result type for **d** = `vec_all_ne(a,b)` are shown in Figure 4-171, Figure 4-172, Figure 4-173, and Figure 4-174.

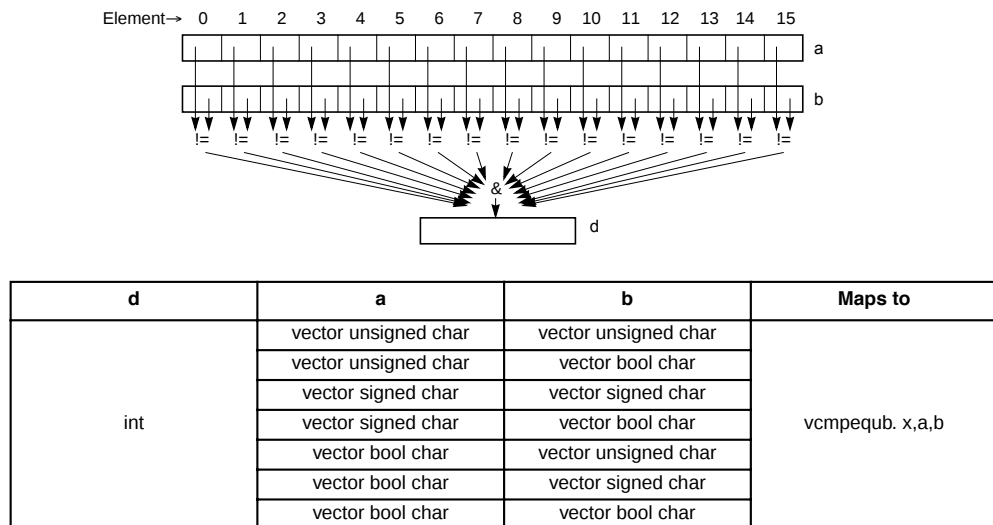


Figure 4-171. All Not Equal of Sixteen Integer Elements (8-bits)

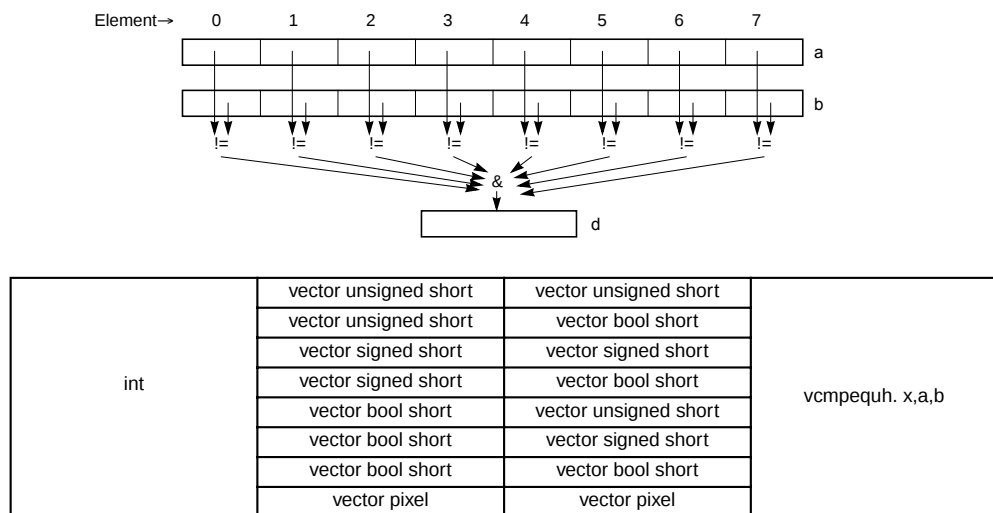


Figure 4-172. All Not Equal of Eight Integer Elements (16-Bit)

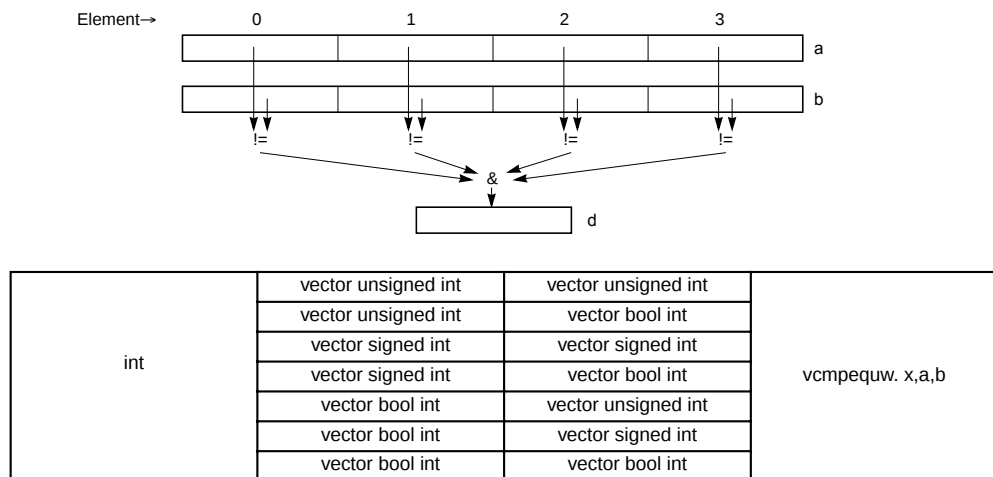


Figure 4-173. All Not Equal of Four Integer Elements (32-Bit)

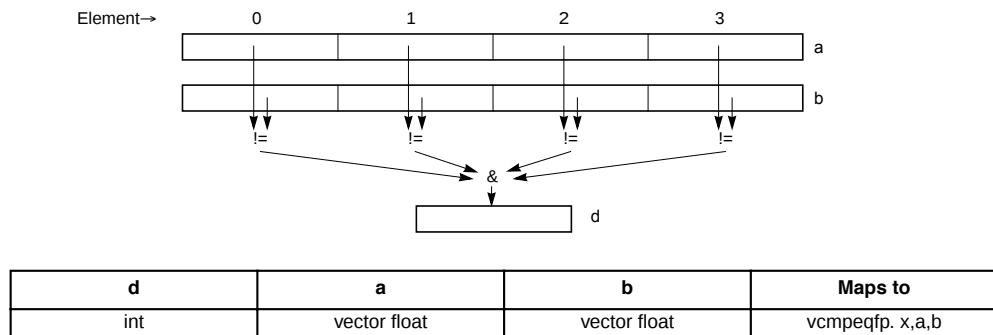


Figure 4-174. All Not Equal of Four Floating-Point Elements (32-Bit)

vec_all_nge

All Elements Not Greater Than or Equal

vec_all_nge

```
d = vec_all_nge(a,b)

if each NGE(ai, bi) = 1, where i ranges from 0 to 3
then d ← 1
else d ← 0
```

The predicate `vec_all_nge` returns 1 if every element of `a` is not greater than or equal to (NGE) the corresponding element of `b`. Otherwise, it returns 0. Not greater than or equal can mean either less than or that one of the elements is NaN.

If `VSCR[NJ] = 1`, every denormalized floating-point operand element is truncated to 0 before the comparison.

The valid argument types and the corresponding result type for `d = vec_all_nge(a,b)` are shown in Figure 4-175.

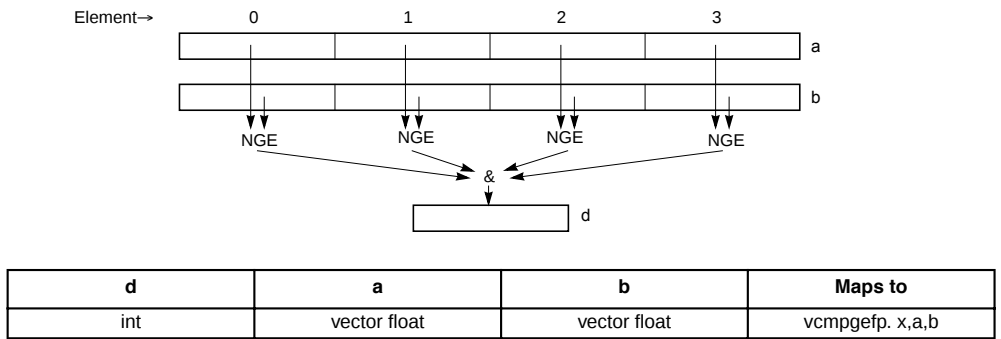


Figure 4-175. All Not Greater Than or Equal of Four Floating-Point Elements (32-Bit)

vec_all_ngt

All Elements Not Greater Than

vec_all_ngt

d = vec_all_ngt(**a**,**b**)

```

if each NGT( $a_i$ ,  $b_i$ ) = 1, where i ranges from 0 to 3
then d ← 1
else d ← 0
    
```

The predicate `vec_all_ngt` returns 1 if every element of **a** is not greater than (NGT) the corresponding element of **b**. Otherwise, it returns 0. Not greater than can either mean less than or equal to or that one of the elements is NaN.

If `VSCR[NJ] = 1`, every denormalized floating-point operand element is truncated to 0 before the comparison.

The valid argument types and the corresponding result type for **d** = `vec_all_ngt(a,b)` is shown in Figure 4-176.

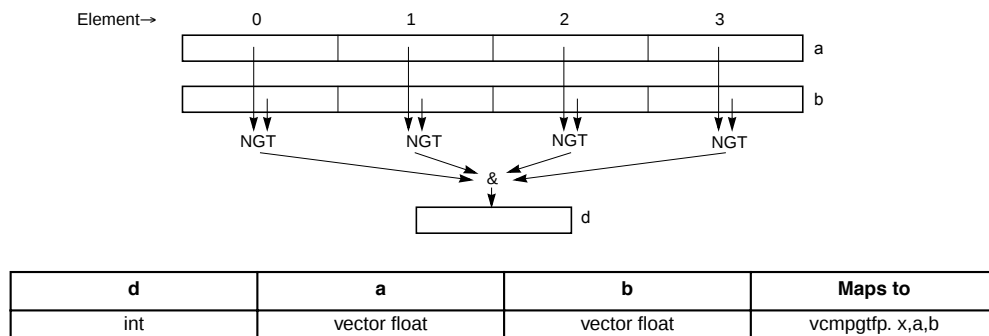


Figure 4-176. All Not Greater Than of Four Floating-Point Elements (32-Bit)

vec_all_nle

vec_all_nle

All Elements Not Less Than or Equal

```

d = vec_all_nle(a,b)

    if each NLE(ai, bi) = 1, where i ranges from 0 to 3
    then d ← 1
    else d ← 0
  
```

The predicate `vec_all_nle` returns 1 if every element of `a` is not less than or equal to (NLE) the corresponding element of `b`. Otherwise, it returns 0. Not less than or equal to can either mean greater than or that one of the elements is NaN.

If `VSCR[NJ] = 1`, every denormalized floating-point operand element is truncated to 0 before the comparison.

The valid argument types and the corresponding result type for `d = vec_all_nle(a,b)` are shown in Figure 4-177.

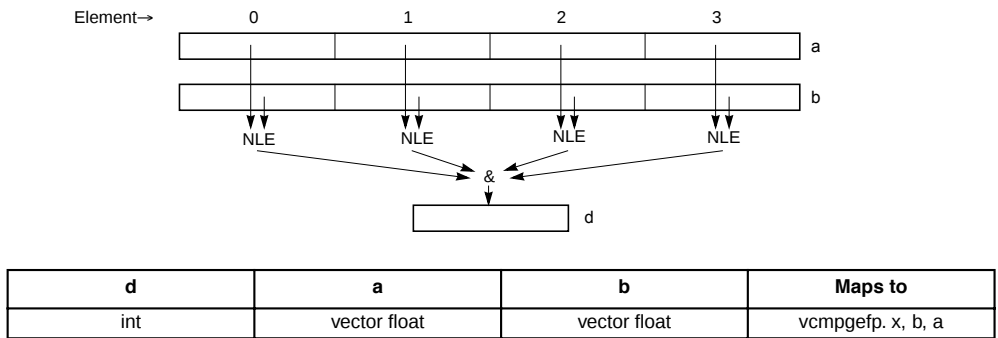


Figure 4-177. All Not Less Than or Equal of Four Floating-Point Elements (32-Bit)

vec_all_nlt

All Elements Not Less Than

vec_all_nlt

d = vec_all_nlt(**a**,**b**)

```

if each  $NLT(a_i, b_i)$ , where  $i$  ranges from 0 to 3
then d  $\leftarrow$  1
else d  $\leftarrow$  0
    
```

The predicate `vec_all_nlt` returns 1 if every element of **a** is not less than (NLT) the corresponding element of **b**. Otherwise, it returns 0. Not less than can either mean greater than or equal to or that one of the elements is NaN.

If `VSCR[NJ] = 1`, every denormalized floating-point operand element is truncated to 0 before the comparison.

The valid argument types and the corresponding result type for **d** = `vec_all_nlt(a,b)` are shown in Figure 4-178.

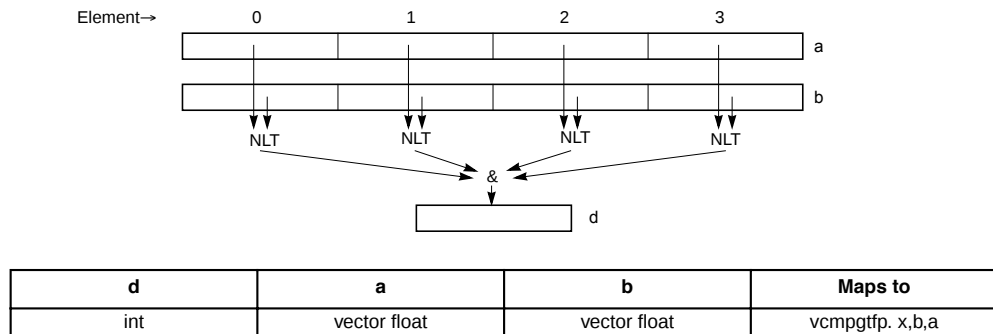


Figure 4-178. All Not Less Than of Four Floating-Point Elements (32-Bit)

vec_all_numeric

vec_all_numeric

All Elements Numeric

```
d = vec_all_numeric(a)  
  
if each ISNUM(ai) = 1, where i ranges from 0 to 3  
then d ← 1  
else d ← 0
```

The predicate `vec_all_numeric` returns 1 if every element of `a` is numeric. Otherwise, it returns 0.

The operation is independent of `VSCR[NJ]`.

The valid argument types and the corresponding result type for `d = vec_all_numeric( )` are shown in Figure 4-179.

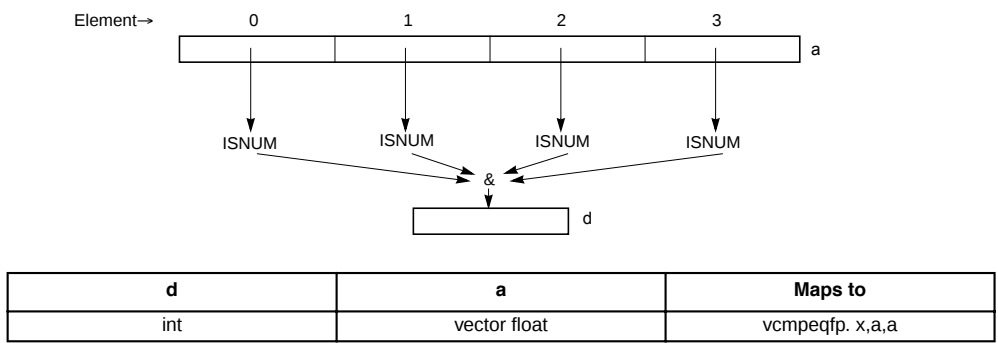


Figure 4-179. All Numeric of Four Floating-Point Elements (32-Bit)

vec_any_eq

Any Element Equal

vec_any_eq

d = vec_any_eq(**a**,**b**)

```

n ← number of elements
if any ai =int bi, where i ranges from 0 to n-1
then d ← 1
else d ← 0
    
```

The predicate `vec_any_eq` returns 1 if any element of `a` is equal to the corresponding element of `b`. Otherwise, it returns 0.

For `vector float` argument types, if `VSCR[NJ] = 1`, every denormalized floating-point operand element is truncated to 0 before the comparison.

The valid combinations of argument types and the corresponding result type for `d = vec_any_eq(a,b)` are shown in Figure 4-180, Figure 4-181, Figure 4-182, and Figure 4-183.

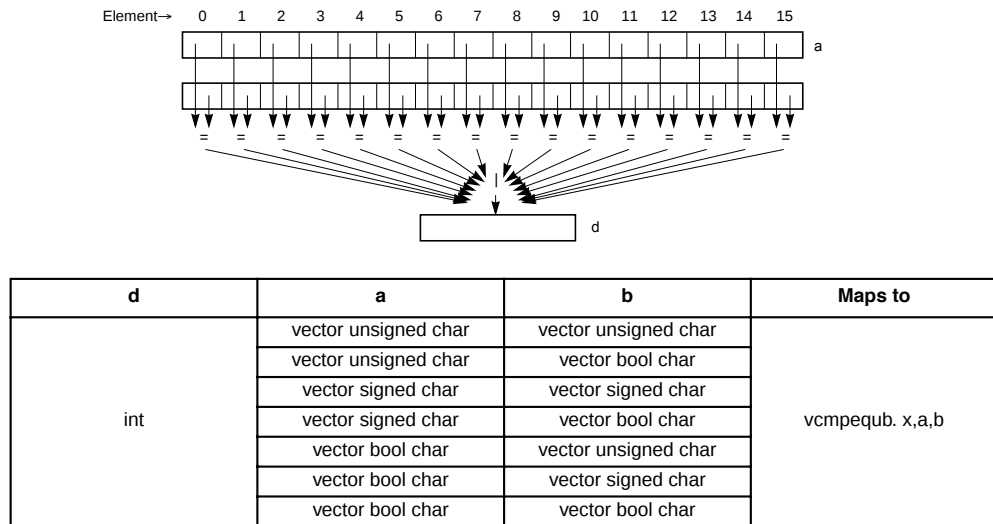
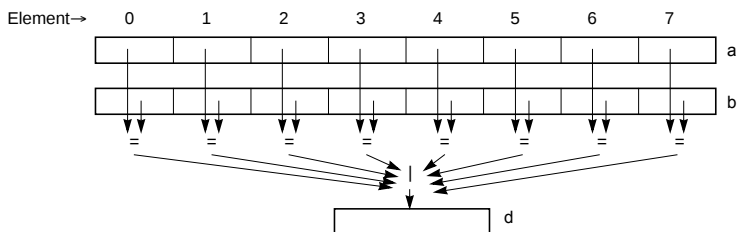
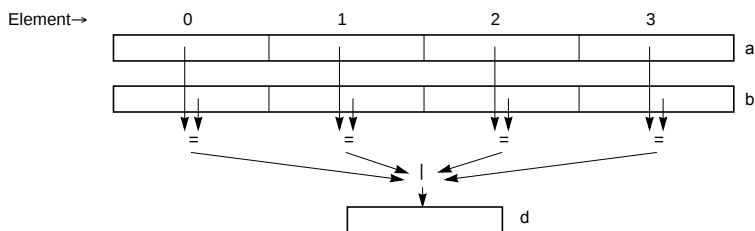


Figure 4-180. Any Equal of Sixteen Integer Elements (8-bits)



d	a	b	Maps to
int	vector unsigned short	vector unsigned short	vcmpquh. x,a,b
	vector unsigned short	vector bool short	
	vector signed short	vector signed short	
	vector signed short	vector bool short	
	vector bool short	vector unsigned short	
	vector bool short	vector signed short	
	vector bool short	vector bool short	
	vector pixel	vector pixel	

Figure 4-181. Any Equal of Eight Integer Elements (16-Bit)



d	a	b	Maps to
int	vector unsigned int	vector unsigned int	vcmpquw. x,a,b
	vector unsigned int	vector bool int	
	vector signed int	vector signed int	
	vector signed int	vector bool int	
	vector bool int	vector unsigned int	
	vector bool int	vector signed int	
	vector bool int	vector bool int	

Figure 4-182. Any Equal of Four Integer Elements (32-Bit)

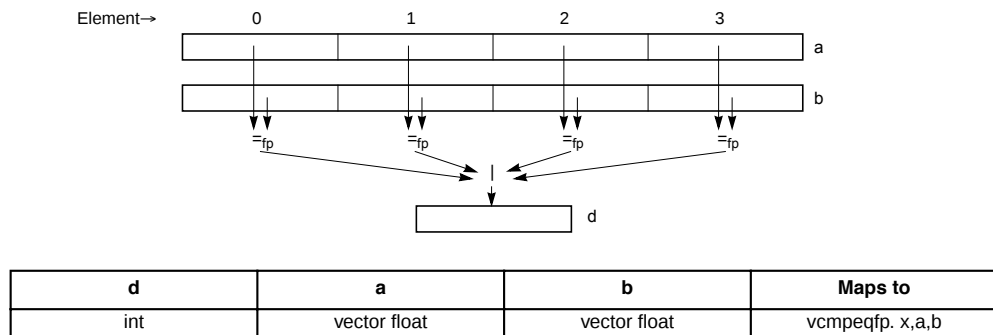


Figure 4-183. Any Equal of Four Floating-Point Elements (32-Bit)

vec_any_ge

Any Element Greater Than or Equal

vec_any_ge

d = vec_any_ge(**a**,**b**)

```
n ← number of elements
if any  $a_i \geq b_i$ , where i ranges from 0 to n-1
then d ← 1
else d ← 0
```

The predicate `vec_any_ge` returns 1 if any element of `a` is greater than or equal to the corresponding element of `b`. Otherwise, it returns 0.

For `vector float` argument types, if `VSCR[NJ] = 1`, every denormalized floating-point operand element is truncated to 0 before the comparison.

The valid combinations of argument types and the corresponding result type for `d = vec_any_ge(a,b)` are shown in Figure 4-184, Figure 4-185, Figure 4-186, and Figure 4-187.

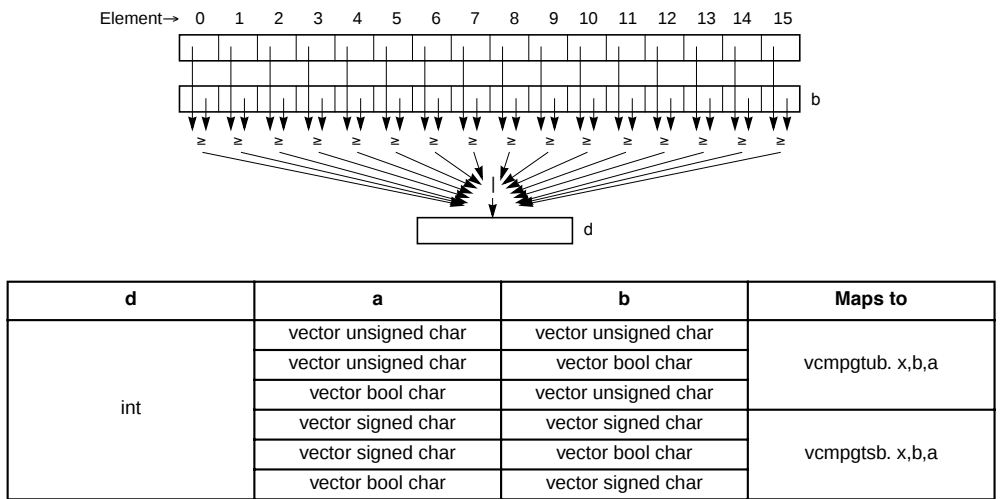
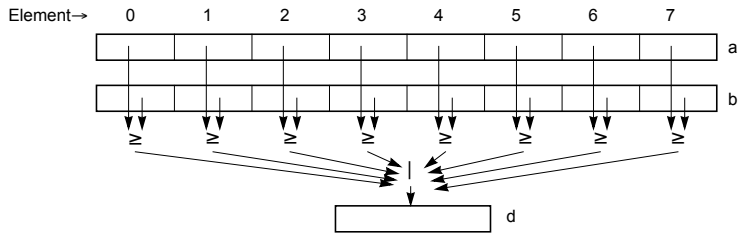
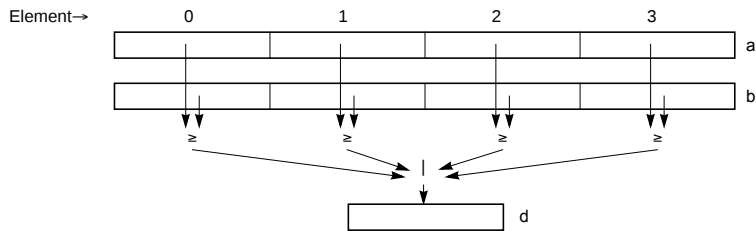


Figure 4-184. Any Greater Than or Equal of Sixteen Integer Elements (8-bits)



d	a	b	Maps to
int	vector unsigned short	vector unsigned short	vcmpgtuh. x,b,a
	vector unsigned short	vector bool short	
	vector bool short	vector unsigned short	
	vector signed short	vector signed short	vcmpgtsh. x,b,a
	vector signed short	vector bool short	
	vector bool short	vector signed short	

Figure 4-185. Any Greater Than or Equal of Eight Integer Elements (16-Bit)



d	a	b	Maps to
int	vector unsigned int	vector unsigned int	vcmpgtuw. x,b,a
	vector unsigned int	vector bool int	
	vector bool int	vector unsigned int	
	vector signed int	vector signed int	vcmpgtsw. x,b,a
	vector signed int	vector bool int	
	vector bool int	vector signed int	

Figure 4-186. Any Greater Than or Equal of Four Integer Elements (32-Bit)

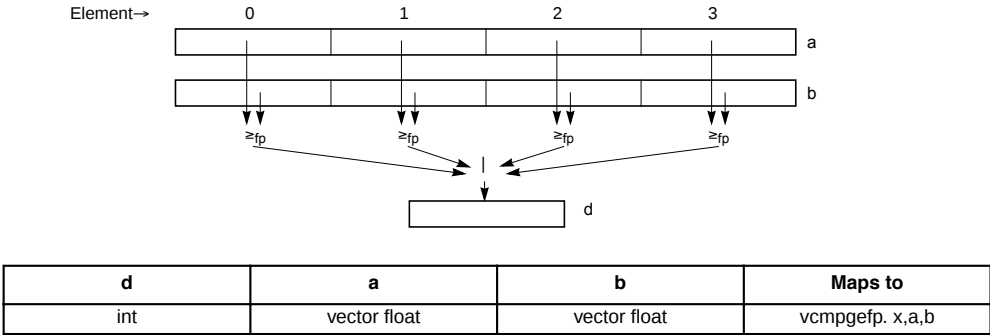


Figure 4-187. Any Greater Than or Equal of Four Floating-Point Elements (32-Bit)

vec_any_gt

Any Element Greater Than

vec_any_gt

d = vec_any_gt(**a**,**b**)

```

n ← number of elements
if any  $a_i > b_i$ , where i ranges from 0 to n-1
then d ← 1
else d ← 0
    
```

The predicate `vec_any_gt` returns 1 if any element of `a` is greater than the corresponding element of `b`. Otherwise, it returns 0.

For `vector float` argument types, if `VSCR[NJ] = 1`, every denormalized floating-point operand element is truncated to 0 before the comparison.

The valid combinations of argument types and the corresponding result type for `d = vec_any_gt(a,b)` are shown in Figure 4-188, Figure 4-189, Figure 4-190, and Figure 4-191.

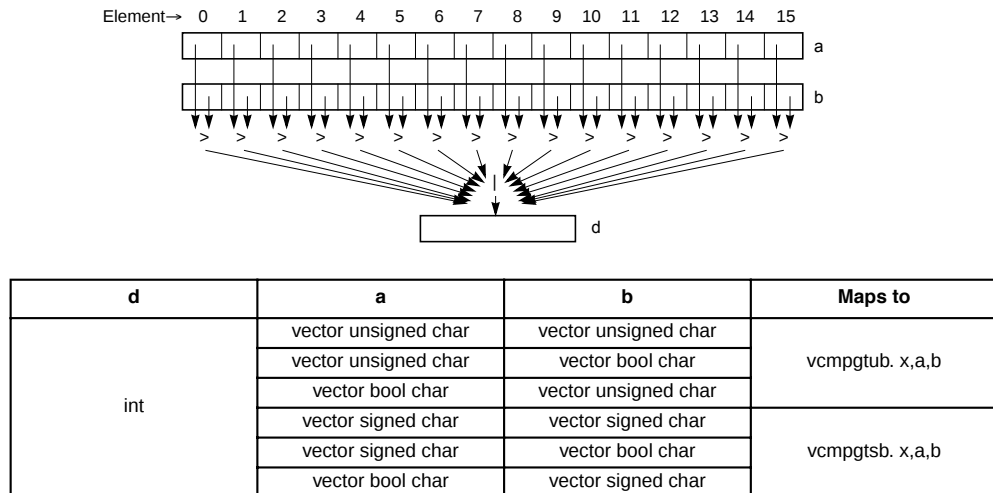
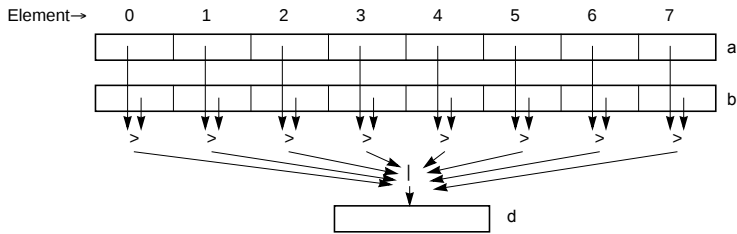
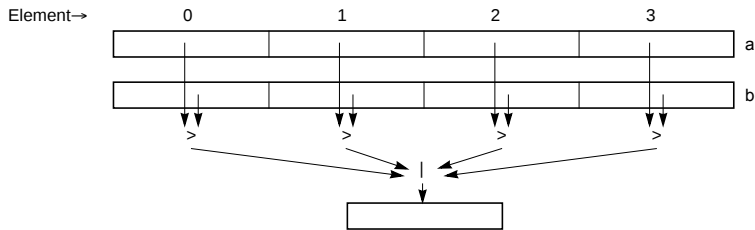


Figure 4-188. Any Greater Than of Sixteen Integer Elements (8-bits)



d	a	b	Maps to
int	vector unsigned short	vector unsigned short	vcmpgtuh. x,a,b
	vector unsigned short	vector bool short	
	vector bool short	vector unsigned short	
	vector signed short	vector signed short	vcmpgtsh. x,a,b
	vector signed short	vector bool short	
	vector bool short	vector signed short	

Figure 4-189. Any Greater Than of Eight Integer Elements (16-Bit)



d	a	b	Maps to
int	vector unsigned int	vector unsigned int	vcmpgtuw. x,a,b
	vector unsigned int	vector bool int	
	vector bool int	vector unsigned int	
	vector signed int	vector signed int	vcmpgtsw. x,a,b
	vector signed int	vector bool int	
	vector bool int	vector signed int	

Figure 4-190. Any Greater Than of Four Integer Elements (32-Bit)

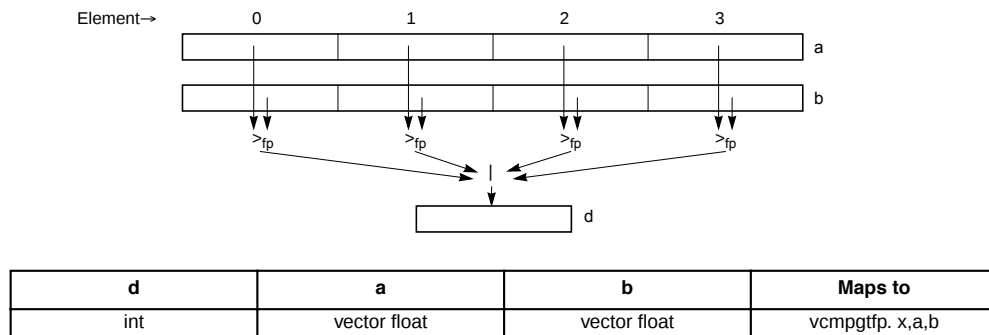


Figure 4-191. Any Greater Than of Four Floating-Point Elements (32-Bit)

vec_any_le

Any Element Less Than or Equal

vec_any_le

d = vec_any_le(**a**,**b**)

```

n ← number of elements
if any  $a_i \leq b_i$ , where i ranges from 0 to n-1
then d ← 1
else d ← 0

```

The predicate `vec_any_le` returns 1 if any element of **a** is less than or equal to the corresponding element of **b**. Otherwise, it returns 0.

For `vector float` argument types, if `VSCR[NJ] = 1`, every denormalized floating-point operand element is truncated to 0 before the comparison.

The valid combinations of argument types and the corresponding result type for **d** = `vec_any_le(a,b)` are shown in Figure 4-192, Figure 4-193, Figure 4-194, and Figure 4-195.

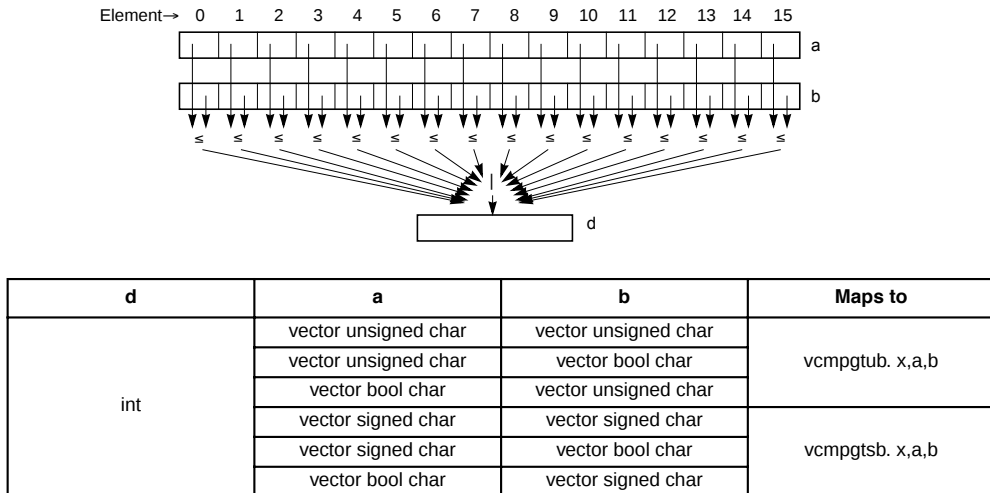
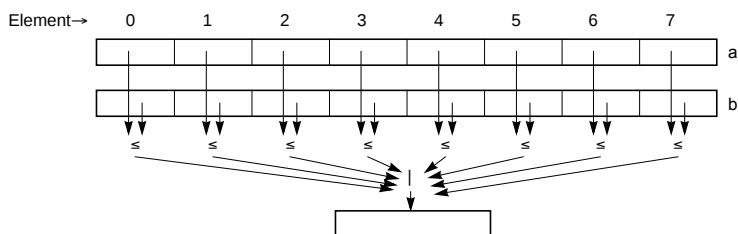
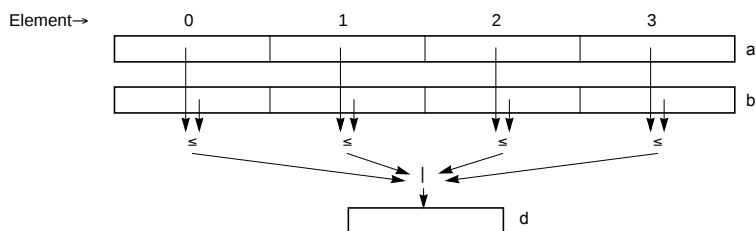


Figure 4-192. Any Less Than or Equal of Sixteen Integer Elements (8-bits)



d	a	b	Maps to
int	vector unsigned short	vector unsigned short	vcmpgtuh. x,a,b
	vector unsigned short	vector bool short	
	vector bool short	vector unsigned short	
	vector signed short	vector signed short	vcmpgtsh. x,a,b
	vector signed short	vector bool short	
	vector bool short	vector signed short	

Figure 4-193. Any Less Than or Equal of Eight Integer Elements (16-Bit)



d	a	b	Maps to
int	vector unsigned int	vector unsigned int	vcmpgtuw. x,a,b
	vector unsigned int	vector bool int	
	vector bool int	vector unsigned int	
	vector signed int	vector signed int	vcmpgtsw. x,a,b
	vector signed int	vector bool int	
	vector bool int	vector signed int	

Figure 4-194. Any Less Than or Equal of Four Integer Elements (32-Bit)

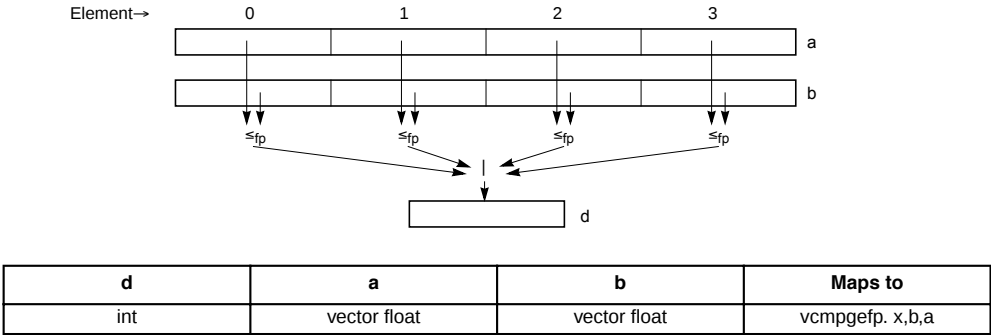


Figure 4-195. Any Less Than or Equal of Four Floating-Point Elements (32-Bit)

vec_any_lt

Any Element Less Than

vec_any_lt

d = vec_any_lt(**a**,**b**)

```

n ← number of elements
if any  $a_i < b_i$ , where i ranges from 0 to n-1
then d ← 1
else d ← 0
    
```

The predicate `vec_any_lt` returns 1 if any element of **a** is less than the corresponding element of **b**. Otherwise, it returns 0.

For `vector float` argument types, if `VSCR[NJ] = 1`, every denormalized floating-point operand element is truncated to 0 before the comparison.

The valid combinations of argument types and the corresponding result type for **d** = `vec_any_lt(a,b)` are shown in Figure 4-196, Figure 4-197, Figure 4-198, and Figure 4-199.

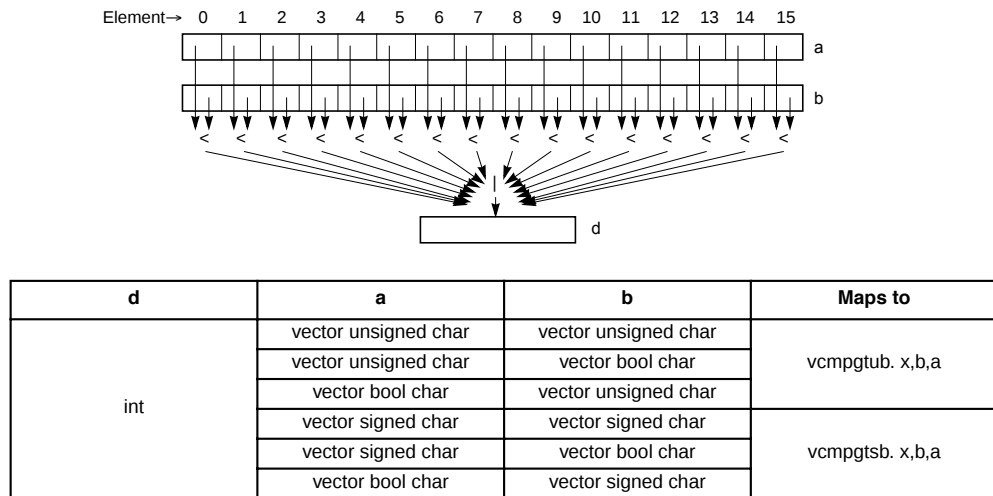


Figure 4-196. Any Less Than of Sixteen Integer Elements (8-bits)

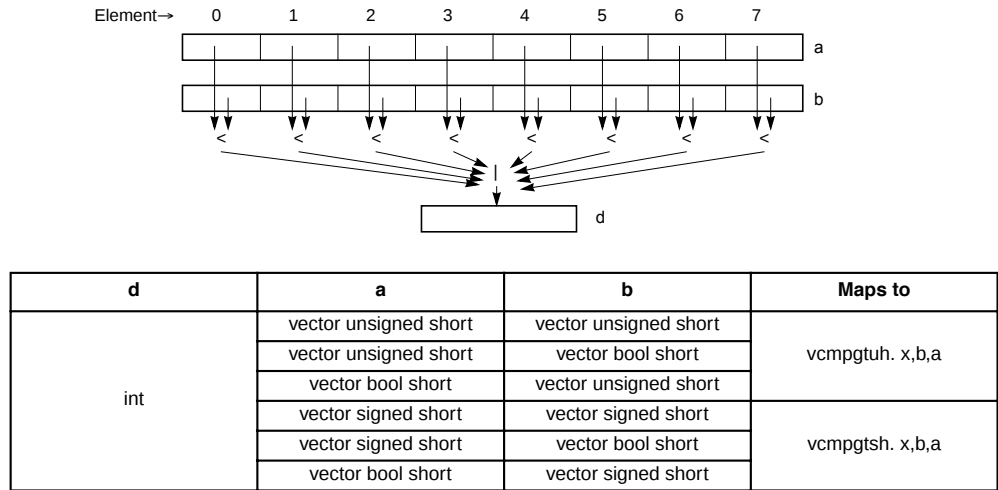


Figure 4-197. Any Less Than of Eight Integer Elements (16-Bit)

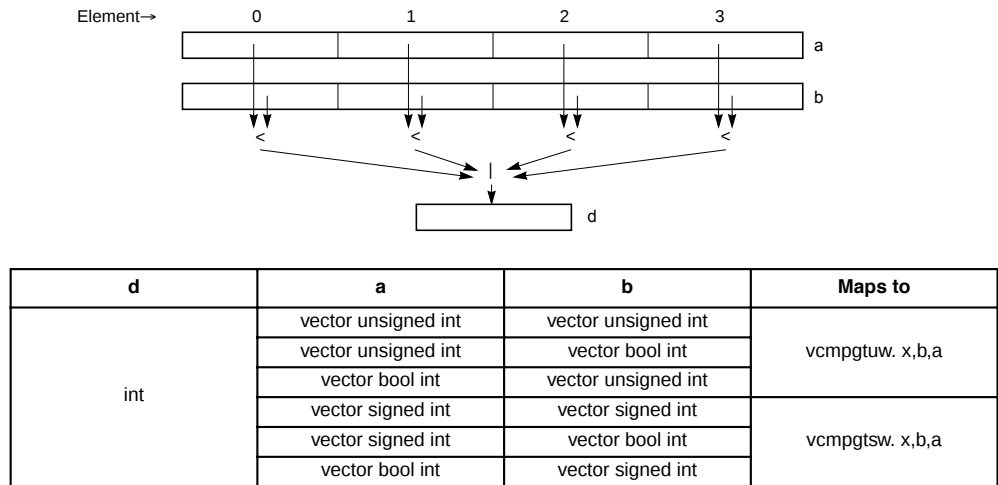


Figure 4-198. Any Less Than of Four Integer Elements (32-Bit)

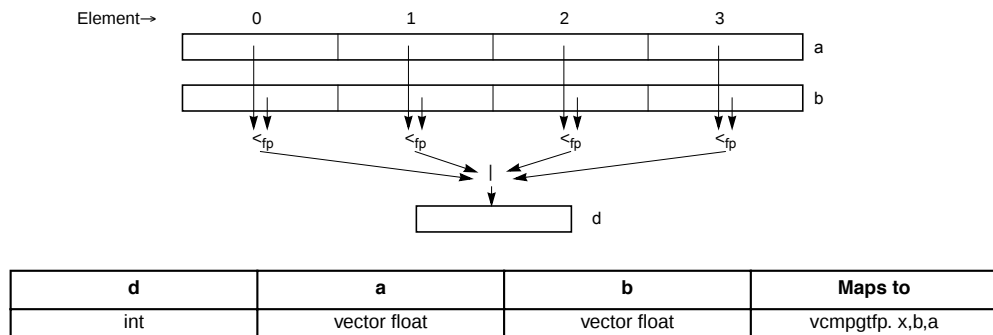


Figure 4-199. Any Less Than of Four Floating-Point Elements (32-Bit)

vec_any_nan

Any Element Not a Number

vec_any_nan

```
d = vec_any_nan(a)  
  
if any ISNaN(ai) = 1, where i ranges from 0 to 3  
then d ← 1  
else d ← 0
```

The predicate `vec_any_nan` returns 1 if any element of `a` is Not a Number (NaN). Otherwise, it returns 0.

The operation is independent of VSCR[NJ].

The valid argument type and corresponding result type for `d = vec_any_nan(a)` are shown in Figure 4-200.

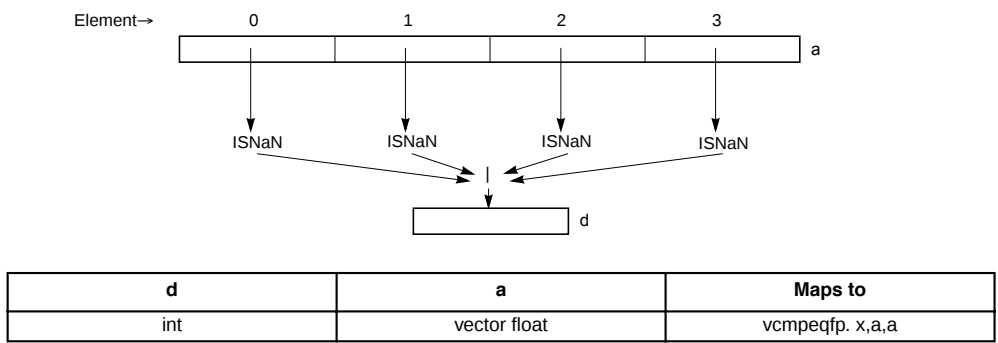


Figure 4-200. Any NaN of Four Floating-Point Elements (32-Bit)

vec_any_ne

Any Element Not Equal

vec_any_ne

d = vec_any_ne(**a**,**b**)

```

n ← number of elements
if any ai != bi, where i ranges from 0 to n-1
then d ← 1
else d ← 0
    
```

The predicate `vec_any_ne` returns 1 if any element of `a` is not equal to (`!=`) the corresponding element of `b`. Otherwise, it returns 0.

For vector float argument types, if `VSCR[NJ] = 1`, every denormalized floating-point operand element is truncated to 0 before the comparison.

The valid combinations of argument types and the corresponding result types for `d = vec_any_ne(a,b)` are shown in Figure 4-201, Figure 4-202, Figure 4-203, and Figure 4-204.

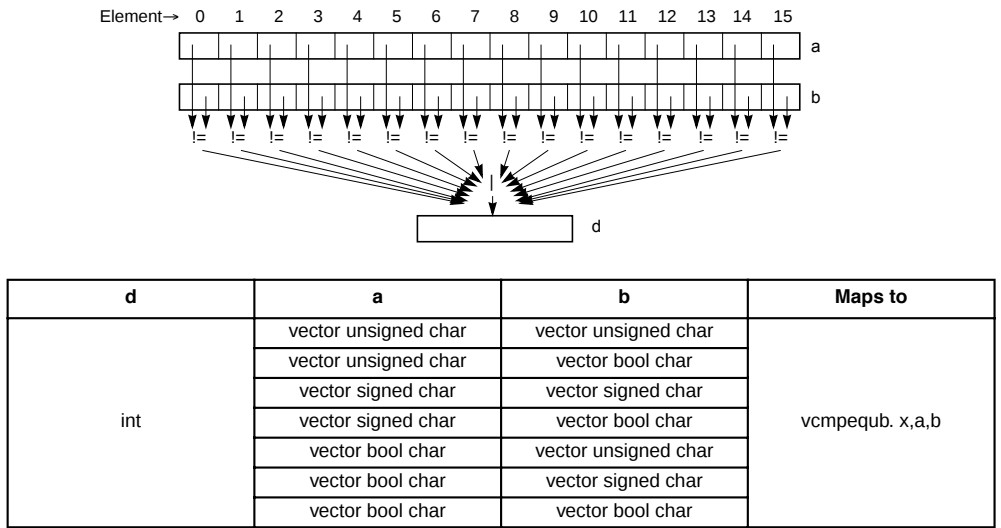
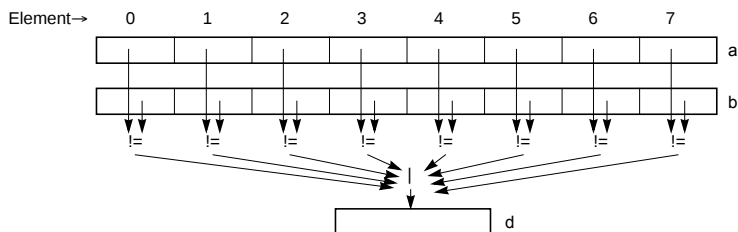
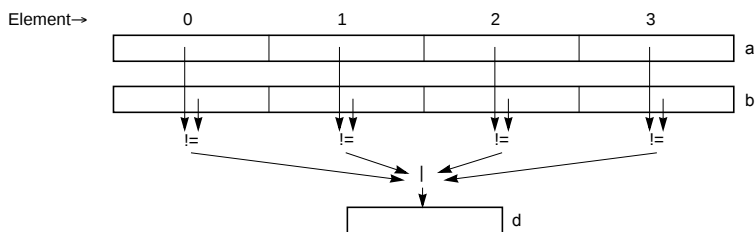


Figure 4-201. Any Not Equal of Sixteen Integer Elements (8-bits)



d	a	b	Maps to
int	vector unsigned short	vector unsigned short	vcmpequh. x,a,b
	vector unsigned short	vector bool short	
	vector signed short	vector signed short	
	vector signed short	vector bool short	
	vector bool short	vector unsigned short	
	vector bool short	vector signed short	
	vector bool short	vector bool short	
	vector pixel	vector pixel	

Figure 4-202. Any Not Equal of Eight Integer Elements (16-Bit)



d	a	b	Maps to
int	vector unsigned int	vector unsigned int	vcmpequw. x,a,b
	vector unsigned int	vector bool int	
	vector signed int	vector signed int	
	vector signed int	vector bool int	
	vector bool int	vector unsigned int	
	vector bool int	vector signed int	
	vector bool int	vector bool int	

Figure 4-203. Any Not Equal of Four Integer Elements (32-Bit)

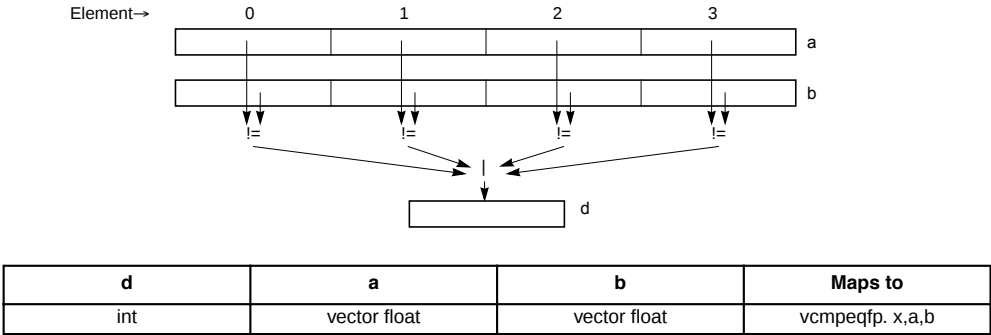


Figure 4-204. Any Not Equal of Four Floating-Point Elements (32-Bit)

vec_any_nge

Any Element Not Greater Than or Equal

vec_any_nge

d = vec_any_nge(**a**,**b**)

```
if any NGE(ai, bi) = 1, where i ranges from 0 to 3
then d ← 1
else d ← 0
```

The predicate `vec_any_nge` returns 1 if any element of `a` is not greater than or equal to (NGE) the corresponding element of `b`. Otherwise, it returns 0. Not greater than or equal can either mean less than or that one of the elements is NaN.

If `VSCR[NJ] = 1`, every denormalized floating-point operand element is truncated to 0 before the comparison.

The valid combination of argument types and the corresponding result type for `d = vec_any_nge(a,b)` are shown in Figure 4-205.

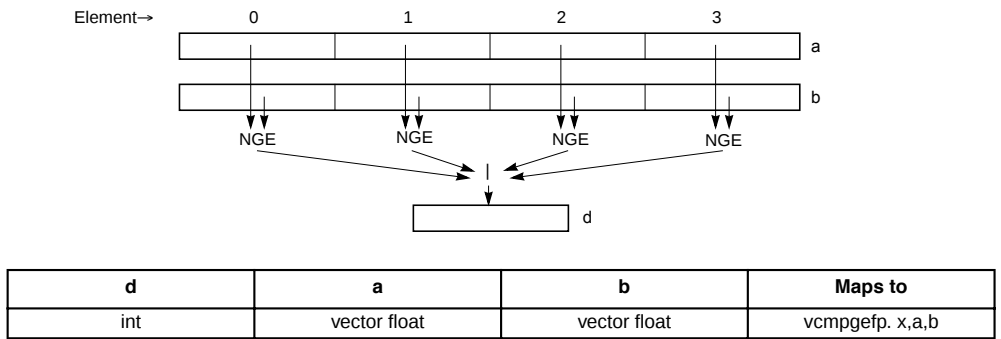


Figure 4-205. Any Not Greater Than or Equal of Four Floating-Point Elements (32-Bit)

vec_any_ngt

Any Element Not Greater Than

vec_any_ngt

d = vec_any_ngt(**a**,**b**)

```

if any NGT( $a_i$ ,  $b_i$ ) = 1, where i ranges from 0 to 3
then d  $\leftarrow$  1
else d  $\leftarrow$  0
    
```

The predicate `vec_any_ngt` returns 1 if any element of `a` is not greater than (NGT) the corresponding element of `b`. Otherwise, it returns 0. Not greater than can either mean less than or equal to or that one of the elements is NaN.

If `VSCR[NJ] = 1`, every denormalized floating-point operand element is truncated to 0 before the comparison.

The valid combination of argument types and the corresponding result type for

d = vec_any_ngt(**a**,**b**) are shown in Figure 4-206.

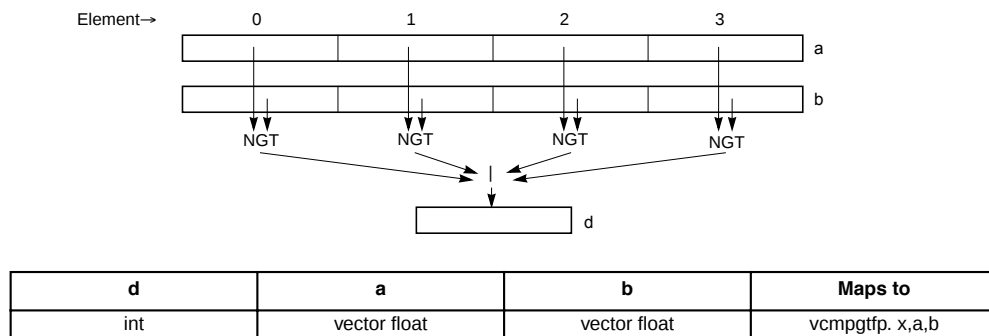


Figure 4-206. Any Not Greater Than of Four Floating-Point Elements (32-Bit)

vec_any_nle

Any Element Not Less Than or Equal

vec_any_nle

d = vec_any_nle(**a**,**b**)

```
if any NLE(ai, bi) = 1, where i ranges from 0 to 3
then d ← 1
else d ← 0
```

The predicate `vec_any_nle` returns 1 if any element of `a` is not less than or equal to (NLE) the corresponding element of `b`. Otherwise, it returns 0. Not less than or equal to can either mean greater than or that one of the elements is NaN.

If `VSCR[NJ] = 1`, every denormalized floating-point operand element is truncated to 0 before the comparison.

The valid combination of argument types and the corresponding result type for `d = vec_any_nle(a,b)` are shown in Figure 4-207.

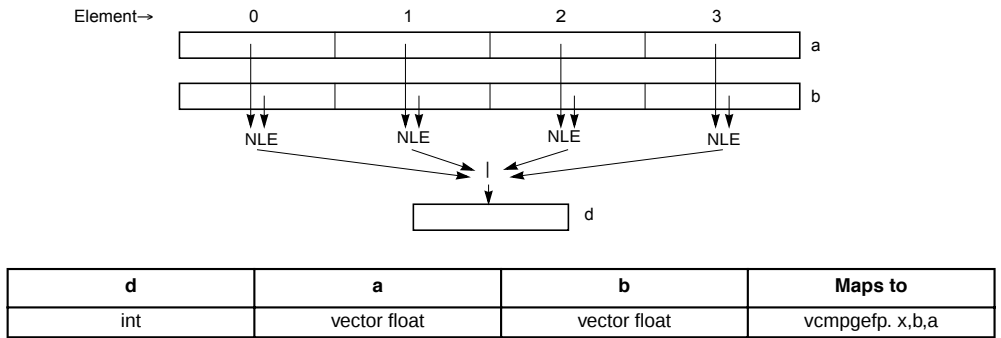


Figure 4-207. Any Not Less Than or Equal of Four Floating-Point Elements (32-Bit)

vec_any_nlt

Any Element Not Less Than

vec_any_nlt

d = vec_any_nlt(**a**,**b**)

```
if any  $NLT(a_i, b_i) = 1$ , where  $i$  ranges from 0 to 3
then d  $\leftarrow$  1
else d  $\leftarrow$  0
```

The predicate `vec_any_nlt` returns 1 if any element of `a` is not less than (NLT) the corresponding element of `b`. Otherwise, it returns 0. Not less than can either mean greater than or equal to or that one of the elements is NaN.

If `VSCR[NJ] = 1`, every denormalized floating-point operand element is truncated to 0 before the comparison.

The valid combination of argument types and the corresponding result type for

d = `vec_any_nlt(a,b)` are shown in Figure 4-208.

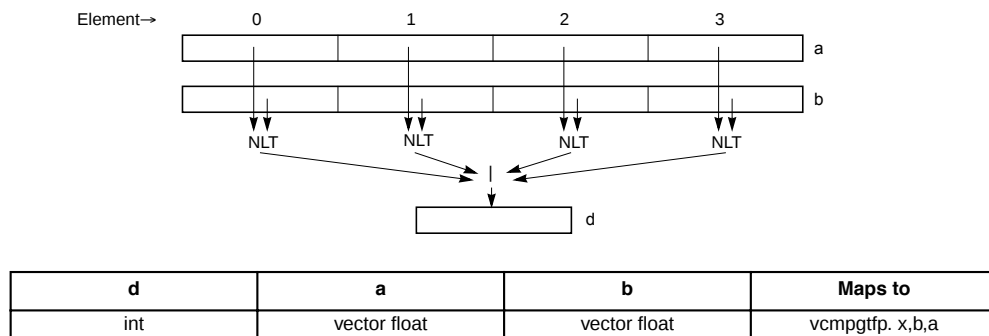


Figure 4-208. Any Not Less Than of Four Floating-Point Elements (32-Bit)

vec_any_numeric

Any Element Numeric

vec_any_numeric

```
d = vec_any_numeric(a)

if any ISNUM(ai) = 1, where i ranges from 0 to 3
then d ← 1
else d ← 0
```

The predicate `vec_any_numeric` returns 1 if any element of `a` is numeric. Otherwise, it returns 0.

The operation is independent of `VSCR[NJ]`.

The valid argument type and the corresponding result type for `d = vec_any_numeric(a)` are shown in Figure 4-209.

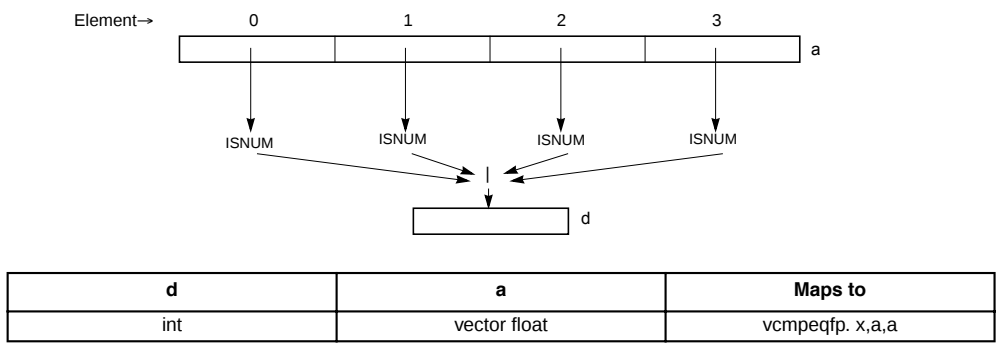


Figure 4-209. Any Numeric of Four Floating-Point Elements (32-Bit)

vec_any_out

Any Element Out of Bounds

vec_any_out

d = vec_any_out(**a**,**b**)

```

if any  $NLE(a_i, b_i) = 1$  or any  $NGE(a_i, -b_i) = 1$ , where  $i$  ranges from 0 to 3
then d  $\leftarrow$  1
else d  $\leftarrow$  0
    
```

The predicate `vec_any_out` returns 1 if any element of `a` is greater than the corresponding element of `b` (high bound) or is less than the negative (NEG) of the corresponding element of `b` (low bound). Otherwise, it returns 0.

If `VSCR[NJ] = 1`, every denormalized floating-point operand element is truncated to 0 before the comparison.

The valid combination of argument types and the corresponding result type for **d** = vec_any_out(**a**,**b**) are shown in Figure 4-210.

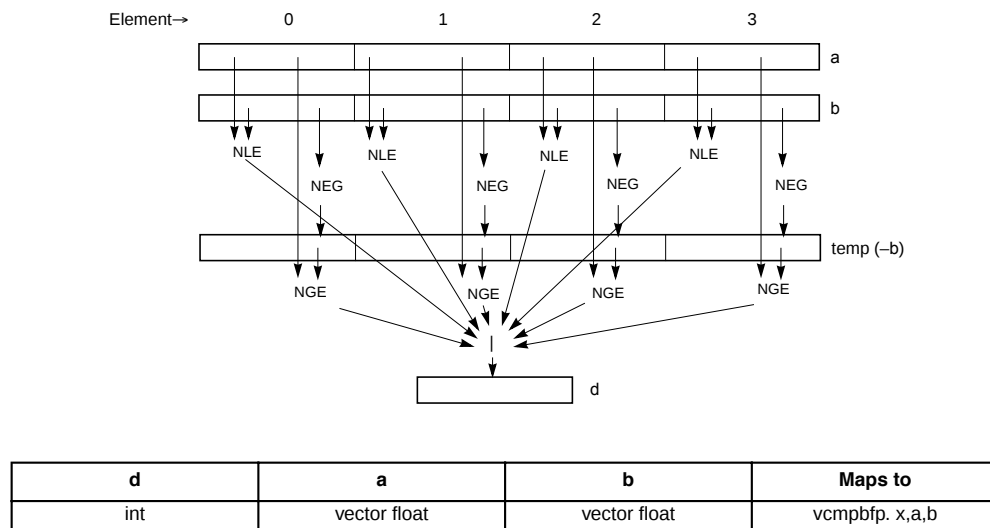


Figure 4-210. Any Out of Bounds of Four Floating-Point Elements (32-Bit)



Appendix A

AltiVec Instruction Set/Operation/Predicate Cross-Reference

This appendix cross-references the instruction set for the AltiVec™ technology, the AltiVec vector operations, and the AltiVec predicates. Table A-1 lists the instructions and the alternate vector operation form cross-referenced to the vector operations and predicates.

Table A-1. Instructions to Operations/Predicates Cross-Reference

AltiVec Instruction	Specific Operation	Generic Operation/Predicate
dss	vec_dss	vec_dss
dssall	vec_dssall	vec_dssall
dst	vec_dst	vec_dst
dstst	vec_dstst	vec_dstst
dststt	vec_dststt	vec_dststt
dstt	vec_dstt	vec_dstt
lvebx	vec_lvebx	vec_lde
lvehx	vec_lvehx	vec_lde
lviewx	vec_lviewx	vec_lde
lvsl	vec_lvsl	vec_lvsl
lvslr	vec_lvslr	vec_lvslr
lvx	vec_lvsl	vec_ld
lvxl	vec_lvxl	vec_lvxl
mfvscr	vec_mfvscr	vec_mfvscr
mtvscr	vec_mtvscr	vec_mtvscr
stvebx	vec_stvebx	vec_ste
stvehx	vec_stvehx	vec_ste
stviewx	vec_stviewx	vec_ste

Table A-1. Instructions to Operations/Predicates Cross-Reference (Continued)

Altivec Instruction	Specific Operation	Generic Operation/Predicate
stvx	vec_stvx	vec_st
stvxl	vec_stvxl	vec_stl
vaddcuw	vec_vaddcuw	vec_addc
vaddfp	vec_vaddfp	vec_add
vaddsbs	vec_vaddsbs	vec_adds
vaddshs	vec_vaddshs	vec_adds
vaddsws	vec_vaddsws	vec_adds
vaddubm	vec_vaddubm	vec_add
vaddubs	vec_vaddubs	vec_adds
vadduhm	vec_vadduhm	vec_add
vadduhs	vec_vadduhs	vec_adds
vadduwm	vec_vadduwm	vec_add
vadduws	vec_vadduws	vec_adds
vand	vec_vand	vec_and
vandc	vec_vandc	vec_andc
vavgsb	vec_vavgsb	vec_avg
vavgsh	vec_vavgsh	vec_avg
vavgsw	vec_vavgsw	vec_avg
vavgub	vec_vavgub	vec_avg
vavguh	vec_vavguh	vec_avg
vavguw	vec_vavguw	vec_avg
vcfsx	vec_vcfsx	vec_ctf
vcfux	vec_vcfux	vec_ctf
vcmpbfp	vec_vcmpbfp	vec_cmpb
vcmpbfp.	—	vec_all_in, vec_any_out
vcmpqfx	vec_vcmpqfx	vec_cmpeq
vcmpqfp.	—	vec_all_eq, vec_all_nan, vec_all_ne, vec_all_numeric, vec_any_eq, vec_any_nan, vec_any_ne, vec_any_numeric
vcmpqubx	vec_vcmpqubx	vec_cmpeq
vcmpqub.	—	vec_all_eq, vec_all_ne, vec_any_eq, vec_any_ne
vcmpquhx	vec_vcmpquhx	vec_cmpeq

Table A-1. Instructions to Operations/Predicates Cross-Reference (Continued)

AltiVec Instruction	Specific Operation	Generic Operation/Predicate
vcmpequh.	—	vec_all_eq, vec_all_ne, vec_any_eq, vec_any_ne
vcmpequwx	vec_vcmpequwx	vec_cmpeq
vcmpequw.	—	vec_all_eq, vec_all_ne, vec_any_eq, vec_any_ne
vcmpgefp	vec_vcmpgefp	vec_cmpge, vec_cmple
vcmpgefp.	—	vec_all_ge, vec_all_le, vec_all_nge, vec_all_nle, vec_any_ge, vec_any_le, vec_any_nge, vec_any_nle
vcmpgtfp	vec_vcmpgtfp	vec_cmpgt, vec_cmplt
vcmpgtfp.	—	vec_all_gt, vec_all_lt, vec_all_ngt, vec_all_nlt, vec_any_gt, vec_any_lt, vec_any_ngt, vec_any_nlt
vcmpgtsbx	vec_vcmpgtsbx	vec_cmpgt, vec_cmplt
vcmpgtsb.	—	vec_all_ge, vec_all_gt, vec_all_le, vec_all_lt, vec_any_ge, vec_any_gt, vec_any_le, vec_any_lt
vcmpgtshx	vec_vcmpgtshx	vec_cmpgt, vec_cmplt
vcmpgtsh.	—	vec_all_ge, vec_all_gt, vec_all_le, vec_all_lt, vec_any_ge, vec_any_gt, vec_any_le, vec_any_lt
vcmpgtswx	vec_vcmpgtswx	vec_cmpgt, vec_cmplt
vcmpgtsw.	—	vec_all_ge, vec_all_gt, vec_all_le, vec_all_lt, vec_any_ge, vec_any_gt, vec_any_le, vec_any_lt
vcmpgtubx	vec_vcmpgtubx	vec_cmpgt, vec_cmplt
vcmpgtub.	—	vec_all_ge, vec_all_gt, vec_all_le, vec_all_lt, vec_any_ge, vec_any_gt, vec_any_le, vec_any_lt
vcmpgtuhx	vec_vcmpgtuhx	vec_cmpgt, vec_cmplt
vcmpgtuh.	—	vec_all_ge, vec_all_gt, vec_all_le, vec_all_lt, vec_any_ge, vec_any_gt, vec_any_le, vec_any_lt
vcmpgtuwx	vec_vcmpgtuwx	vec_cmpgt, vec_cmplt
vcmpgtuw.	—	vec_all_ge, vec_all_gt, vec_all_le, vec_all_lt, vec_any_ge, vec_any_gt, vec_any_le, vec_any_lt
vctxs	vec_vctxs	vec_cts
vctuxs	vec_vctuxs	vec_ctu
vexptefp	vec_vexptefp	vec_expte

Table A-1. Instructions to Operations/Predicates Cross-Reference (Continued)

AltiVec Instruction	Specific Operation	Generic Operation/Predicate
vlogefp	vec_vlogefp	vec_loge
vmaddfp	vec_vmaddfp	vec_madd
vmaxfp	vec_vmaxfp	vec_max
vmaxsb	vec_vmaxsb	vec_max
vmaxsh	vec_vmaxsh	vec_max
vmaxsw	vec_vmaxsw	vec_max
vmaxub	vec_vmaxub	vec_max
vmaxuh	vec_vmaxuh	vec_max
vmaxuw	vec_vmaxuw	vec_max
vmhaddshs	vec_vmhaddshs	vec_madds
vmhraddshs	vec_vmhraddshs	vec_mradds
vminfp	vec_vminfp	vec_min
vminsb	vec_vminsb	vec_min
vminsh	vec_vminsh	vec_min
vminsw	vec_vminsw	vec_min
vminub	vec_vminub	vec_min
vminuh	vec_vminuh	vec_min
vminuw	vec_vminuw	vec_min
vmladduhm	vec_vmladduhm	vec_mladd
vmrghb	vec_vmrghb	vec_mergeh
vmrghh	vec_vmrghh	vec_mergeh
vmrghw	vec_vmrghw	vec_mergeh
vmrglb	vec_vmrglb	vec_mergel
vmrglh	vec_vmrglh	vec_mergel
vmrglw	vec_vmrglw	vec_mergel
vmsummbm	vec_vmsummbm	vec_msum
vmsumshm	vec_vmsumshm	vec_msum
vmsumshs	vec_vmsumshs	vec_msums
vmsumubm	vec_vmsumubm	vec_msum
vmsumuhm	vec_vmsumuhm	vec_msum
vsumuhs	vec_vsumuhs	vec_msums
vmulesb	vec_vmulesb	vec_mule

**Table A-1. Instructions to Operations/Predicates Cross-Reference (Continued)**

AltiVec Instruction	Specific Operation	Generic Operation/Predicate
vmulesh	vec_vmulesh	vec_mule
vmuleub	vec_vmuleub	vec_mule
vmuleuh	vec_vmuleuh	vec_mule
vmulosb	vec_vmulosb	vec_mulo
vmulosh	vec_vmulosh	vec_mulo
vmuloub	vec_vmuloub	vec_mulo
vmulouh	vec_vmulouh	vec_mulo
vnmsubfp	vec_vnmsubfp	vec_nmsub
vnor	vec_vnor	vec_nor
vor	vec_vor	vec_or
vperm	vec_vperm	vec_perm
vpkpx	vec_vpkpx	vec_packpx
vpkshss	vpkshss	vec_packs
vpkshus	vec_vpkshus	vec_packsu
vpkswss	vec_vpkswss	vec_packs
vpkswus	vec_vpkswus	vec_packsu
vpkuhum	vec_vpkuhum	vec_pack
vpkuhus	vec_vpkuhus	vec_packs, vec_packsu
vpkuwum	vec_vpkuwum	vec_pack
vpkuwus	vec_vpkuwus	vec_packs, vec_packsu
vrefp	vec_vrefp	vec_re
vrfim	vec_vrfim	vec_floor
vrfin	vec_vrfin	vec_round
vrrip	vec_vrrip	vec_ceil
vrfiz	vec_vrfiz	vec_trunc
vrlb	vec_vrlb	vec_rl
vrlh	vec_vrlh	vec_rl
vrlw	vec_vrlw	vec_rl
vrsqrtefp	vec_vrsqrtefp	vec_rsqrte
vsel	vec_vsel	vec_sel
vsl	vec_vsl	vec_sll
vslb	vec_vslb	vec_sl

**Table A-1. Instructions to Operations/Predicates Cross-Reference (Continued)**

Altivec Instruction	Specific Operation	Generic Operation/Predicate
vldoi	vec_vldoi	vec_sld
vslh	vec_vslh	vec_sl
vslo	vec_vslo	vec_slo
vslw	vec_vslw	vec_sl
vspltb	vec_vspltb	vec_splat
vsplth	vec_vsplth	vec_splat
vspltisb	vec_vspltisb	vec_splat_s8, vec_splat_u8
vspltish	vec_vspltish	vec_splat_s16, vec_splat_u16
vspltisw	vec_vspltisw	vec_splat_s32, vec_splat_u32
vspltw	vec_vspltw	vec_splat
vsr	vec_vsr	vec_srl
vsrab	vec_vsrab	vec_sra
vsrah	vec_vsrab	vec_sra
vsraw	vec_vsrab	vec_sra
vsrb	vec_vsr	vec_sr
vsrh	vec_vsr	vec_sr
vsro	vec_vsr	vec_sr
vsrw	vec_vsr	vec_sr
vsubcuw	vec_vsubcuw	vec_subc
vsubfp	vec_vsubfp	vec_sub
vsubsbbs	vec_vsubsbbs	vec_subs
vsubshs	vec_vsubshs	vec_subs
vsubsws	vec_vsubsws	vec_subs
vsububm	vec_vsububm	vec_sub
vsububs	vec_vsububs	vec_subs
vsubuhm	vec_vsubuhm	vec_sub
vsubuhs	vec_vsubuhs	vec_subs
vsubuwm	vec_vsubuwm	vec_sub
vsubuws	vec_vsubuws	vec_subs
vsumsws	vec_vsumsws	vec_sums
vsum2sws	vec_vsum2sws	vec_sum2s
vsum4sbs	vec_vsum4sbs	vec_sum4s

Table A-1. Instructions to Operations/Predicates Cross-Reference (Continued)

AltiVec Instruction	Specific Operation	Generic Operation/Predicate
vsum4shs	vec_vsum4shs	vec_sum4s
vsum4ubs	vec_vsum4ubs	vec_sum4s
vupkhp	vec_vupkhp	vec_unpackh
vupkhsb	vec_vupkhsb	vec_unpackh
vupkhsh	vec_vupkhsh	vec_unpackh
vupklp	vec_vupklp	vec_unpackl
vupklsb	vec_vupklsb	vec_unpackl
vupklsh	vec_vupklsh	vec_unpackl
vxor	vec_vxor	vec_xor

Table A-2 lists the vector operations cross-referenced to the AltiVec instructions.

Table A-2. Operations to Instructions Cross-Reference

Specific Operation	AltiVec Instruction(s)
vec_abs	vspltisb, vsbubm, vmaxsb
	vspltisb, vsbuhm, vmaxsh
	vspltisb, vsbuwm, vmaxsw
	vspltisw, vsllw, vandc
vec_abss	vspltisb, vsbubs, vmaxsb
	vspltisb, vsbushs, vmaxsh
	vspltisb, vsbushs, vmaxsw
vec_add	vaddfp
	vaddubm
	vadduhm
	vadduwm
vec_addc	vaddcuw
vec_adds	vaddsbs
	vaddshs
	vaddsws
	vaddubs
	vadduhs
	vadduws
vec_and	vand

Table A-2. Operations to Instructions Cross-Reference (Continued)

Specific Operation	AltiVec Instruction(s)
vec_andc	vandc
vec_avg	vavgsb
	vavgsh
	vavgsw
	vavgub
	vavguh
	vavguw
vec_ceil	vrrip
vec_cmpb	vcmpbfp
vec_cmpeq	vcmpqfp
	vcmpqub
	vcmpquh
	vcmpquw
vec_cmpge	vcmpgef
vec_cmpgt	vcmpgtfp
	vcmpgtb
	vcmpgtsh
	vcmpgtsw
	vcmpgtub
	vcmpgtuh
	vcmpgtuw
vec_cmple	vcmpgef
vec_cmplt	vcmpgtfp
	vcmpgtb
	vcmpgtsh
	vcmpgtsw
	vcmpgtub
	vcmpgtuh
	vcmpgtuw
vec_ctf	vcfs
	vcfux
vec_cts	vctxs

**Table A-2. Operations to Instructions Cross-Reference (Continued)**

Specific Operation	AltiVec Instruction(s)
vec_ctu	vctuxs
vec_dss	dss
vec_dssall	dssall
vec_dst	dst
vec_dstst	dstst
vec_dststt	dststt
vec_dstt	dstt
vec_expte	vexptefp
vec_floor	vrfim
vec_ld	lvx
vec_lde	lvebx
	lvehx
	lvewx
vec_ldl	lvxl
vec_loge	vlogefp
vec_lvsl	lvsl
vec_lvsr	lvsr
vec_madd	vmaddfp
vec_madds	vmhaddshs
vec_max	vmaxfp
	vmaxsb
	vmaxsh
	vmaxsw
	vmaxub
	vmaxuh
	vmaxuw
vec_mergeh	vmrghw
	vmrghb
	vmrghh
vec_mergel	vmrglw
	vmrglb
	vmrglh

Table A-2. Operations to Instructions Cross-Reference (Continued)

Specific Operation	AltiVec Instruction(s)
vec_mfvscr	mfvscr
vec_min	vminfp
	vminsb
	vminsh
	vminsw
	vminub
	vminuh
	vminuw
vec_mladd	vmladduhm
vec_mradds	vmhraddshs
vec_msum	vmsummbm
	vmsumshm
	vmsumubm
	vmsumuhm
vec_msums	vmsumshs
vec_msums	vmsumuhs
vec_mtvscr	mtvscr
vec_mule	vmulesb
	vmulesh
	vmuleub
	vmuleuh
vec_mulo	vmulosb
	vmulosh
	vmuloub
	vmulouh
vec_nmsub	vnmsubfp
vec_nor	vnor
vec_or	vor
vec_pack	vpkuhum
	vpkuwum
vec_packpx	vpkpx

**Table A-2. Operations to Instructions Cross-Reference (Continued)**

Specific Operation	AltiVec Instruction(s)
vec_packs	vpkshss
	vpkswss
	vpkuhus
	vpkuwus
vec_packsu	vpkuhus
	vpkuwus
	vpkshus
	vpkswus
vec_perm	vperm
vec_re	vrefp
vec_rl	vrlb
	vrlh
	vrlw
vec_round	vrfn
vec_rsqrte	vrsqrtefp
vec_sel	vsel
vec_sl	vslb
	vslh
	vslw
vec_sld	vsldoi
vec_sll	vsl
vec_slo	vslo
vec_splat	vspltb
	vsplth
	vspltw
vec_splat_s16	vspltish
vec_splat_s32	vspltisw
vec_splat_s8	vspltisb
vec_splat_u16	vspltish
vec_splat_u32	vspltisw
vec_splat_u8	vspltisb

**Table A-2. Operations to Instructions Cross-Reference (Continued)**

Specific Operation	AltiVec Instruction(s)
vec_sr	vsrb
	vsrh
	vsrw
vec_sra	vsrab
	vsrah
	vsraw
vec_srl	vsr
vec_sro	vsro
vec_st	stvx
vec_ste	stvebx
	stvehx
	stviewx
vec_stl	stvxl
vec_sub	vsubfp
	vsububm
	vsubuhm
	vsubuwm
vec_subc	vsubcuw
vec_subs	vsubsbs
	vsubshs
	vsubsws
	vsububs
	vsubuhs
	vsubuws
vec_sum2s	vsum2sws
vec_sum4s	vsum4sbs
	vsum4shs
	vsum4ubs
vec_sums	vsumsws
vec_trunc	vrfiz



Table A-2. Operations to Instructions Cross-Reference (Continued)

Specific Operation	AltiVec Instruction(s)
vec_unpackh	vupkhpX
	vupkhsb
	vupkhsh
vec_unpackl	vupklpX
	vupklsb
	vupklsh
vec_xor	vxor



Table A-3 lists the predicates cross-referenced to the AltiVec instructions.

Table A-3. Predicate to Instruction Cross-Reference

Predicate	AltiVec Instruction
vec_all_eq	vcmppeqfp.
	vcmppequb.
	vcmppequh.
	vcmppequw.
vec_all_ge	vcmpgtsb.
	vcmpgtsh.
	vcmpgtsw.
	vcmpgtub.
	vcmpgtuh.
	vcmpgtuw.
	vcmpgefp.
vec_all_gt	vcmpgtsb.
	vcmpgtsh.
	vcmpgtsw.
	vcmpgtub.
	vcmpgtuh.
	vcmpgtuw.
	vcmpgtfp.
vec_all_in	vcmpbfp.
vec_all_le	vcmpgtsb.
	vcmpgtsh.
	vcmpgtsw.
	vcmpgtub.
	vcmpgtuh.
	vcmpgtuw.
	vcmpgefp.

**Table A-3. Predicate to Instruction Cross-Reference (Continued)**

Predicate	Altivec Instruction
vec_all_lt	vcmpgtsb.
	vcmpgtsh.
	vcmpgtsw.
	vcmpgtub.
	vcmpgtuh.
	vcmpgtuw.
	vcmpgtfp.
vec_all_nan	vcmpeqfp.
vec_all_ne	vcmpeqfp.
	vcmpequb.
	vcmpequh.
	vcmpequw.
vec_all_nge	vcmpgefp.
vec_all_ngt	vcmpgtfp.
vec_all_nle	vcmpgefp.
vec_all_nlt	vcmpgtfp.
vec_all_numeric	vcmpeqfp.
vec_any_eq	vcmpeqfp.
	vcmpequb.
	vcmpequh.
	vcmpequw.
vec_any_ge	vcmpgtsb.
	vcmpgtsh.
	vcmpgtsw.
	vcmpgtub.
	vcmpgtuh.
	vcmpgtuw.
	vcmpgefp.

Table A-3. Predicate to Instruction Cross-Reference (Continued)

Predicate	Altivec Instruction
vec_any_gt	vcmpgtsb.
	vcmpgtsh.
	vcmpgtsw.
	vcmpgtub.
	vcmpgtuh.
	vcmpgtuw.
	vcmpgtfp.
vec_any_le	vcmpgtsb.
	vcmpgtsh.
	vcmpgtsw.
	vcmpgtub.
	vcmpgtuh.
	vcmpgtuw.
	vcmpgefp.
vec_any_lt	vcmpgtsb.
	vcmpgtsh.
	vcmpgtsw.
	vcmpgtub.
	vcmpgtuh.
	vcmpgtuw.
	vcmpgtfp.
vec_any_nan	vcmpeqfp.
vec_any_ne	vcmpeqfp.
	vcmpequb.
	vcmpequh.
	vcmpequw.
vec_any_nge	vcmpgefp.
vec_any_ngt	vcmpgtfp.
vec_any_nle	vcmpgefp.
vec_any_nlt	vcmpgtfp.
vec_any_numeric	vcmpeqfp.
vec_any_out	vcmpbfp.

Glossary of Terms and Abbreviations

The glossary contains an alphabetical list of terms, phrases, and abbreviations used in this book. Some of the terms and definitions included in the glossary are reprinted from *IEEE Std. 754-1985, IEEE Standard for Binary Floating-Point Arithmetic*, copyright ©1985 by the Institute of Electrical and Electronics Engineers, Inc. with the permission of the IEEE.

Note that some terms are defined in the context of how they are used in this book.

A **Architecture.** A detailed specification of requirements for a processor or computer system. It does not specify details of how the processor or computer system must be implemented; instead it provides a template for a family of compatible *implementations*.

B **Biased exponent.** An *exponent* whose range of values is shifted by a constant (bias). Typically a bias is provided to allow a range of positive values to express a range that includes both positive and negative values.

Big-endian. A byte-ordering method in memory where the address *n* of a word corresponds to the *most-significant byte*. In an addressed memory word, the bytes are ordered (left to right) 0, 1, 2, 3, with 0 being the most-significant byte. *See* Little-endian.

C **Cache.** High-speed memory component containing recently-accessed data and/or instructions (subset of main memory).

Cast. A cast expression consists of a left parenthesis, a type name, a right parenthesis, and an operand expression. The cast causes the operand value to be converted to the type name within the parentheses.

D **Denormalized number.** A nonzero floating-point number whose *exponent* has a reserved value, usually the format's minimum, and whose explicit or implicit leading significand bit is zero.

-
- E** **Effective address (EA).** The 32- or 64-bit address specified for a load, store, or an instruction fetch. This address is then submitted to the MMU for translation to either a *physical memory* address or an I/O address.
- Exponent.** In the binary representation of a floating-point number, the exponent is the component that normally signifies the integer power to which the value two is raised in determining the value of the represented number. *See also* Biased exponent.
-
- F** **Floating-point register (FPR).** Any of the 32 registers in the floating-point register file. These registers provide the source operands and destination results for floating-point instructions. Load instructions move data from memory to FPRs and store instructions move data from FPRs to memory. The FPRs are 64 bits wide and store floating-point values in double-precision format.
- Fraction.** In the binary representation of a floating-point number, the field of the *significand* that lies to the right of its implied binary point.
-
- G** **General-purpose register (GPR).** Any of the 32 registers in the general-purpose register file. These registers provide the source operands and destination results for all integer data manipulation instructions. Integer load instructions move data from memory to GPRs and store instructions move data from GPRs to memory.
-
- I** **IEEE 754.** A standard written by the Institute of Electrical and Electronics Engineers that defines operations and representations of binary floating-point arithmetic.
- Inexact.** Loss of accuracy in an arithmetic operation when the rounded result differs from the infinitely precise value with unbounded range.
-
- L** **Least-significant bit (lsb).** The bit of least value in an address, register, data element, or instruction encoding.
- Little-endian.** A byte-ordering method in memory where the address *n* of a word corresponds to the *least-significant byte*. In an addressed memory word, the bytes are ordered (left to right) 3, 2, 1, 0, with 3 being the *most-significant byte*. *See* Big-endian.
-
- M** **Mnemonic.** The abbreviated name of an instruction used for coding.
-

Modulo. A value v which lies outside the range of numbers representable by an n -bit wide destination type is replaced by the low-order n bits of the two's complement representation of v .

Most-significant bit (msb). The highest-order bit in an address, registers, data element, or instruction encoding.

N

NaN. An abbreviation for 'Not a Number'; a symbolic entity encoded in floating-point format. There are two types of NaNs—signaling NaNs (SNaNs) and quiet NaNs (QNaNs).

Normalization. A process by which a floating-point value is manipulated such that it can be represented in the format for the appropriate precision (single- or double-precision). For a floating-point value to be representable in the single- or double-precision format, the leading implied bit must be a 1.

O

Overflow. An error condition that occurs during arithmetic operations when the result cannot be stored accurately in the destination register(s). For example, if two 32-bit numbers are multiplied, the result may not be representable in 32 bits.

Q

Quad word. A group of 16 contiguous locations starting at an address divisible by 16.

Quiet NaN. A type of *NaN* that can propagate through most arithmetic operations without signaling exceptions. A quiet NaN is used to represent the results of certain invalid operations, such as invalid arithmetic operations on infinities or on NaNs, when invalid. *See* Signaling NaN.

R

Record bit. Bit 31 (or the Rc bit) in the instruction encoding. When it is set, updates the condition register (CR) to reflect the result of the operation. Its presence is denoted by a “.” following the mnemonic.

Reserved field. In a register, a reserved field is one that is not assigned a function. A reserved field may be a single bit. The handling of reserved bits is *implementation-dependent*. Software is permitted to write any value to such a bit. A subsequent reading of the bit returns 0 if the value last written to the bit was 0 and returns an undefined value (0 or 1) otherwise.

RISC (reduced instruction set computing). An *architecture* characterized by fixed-length instructions with nonoverlapping functionality and by a separate set of load and store instructions that perform memory accesses.

S

Saturate. A value v which lies outside the range of numbers representable by a destination type is replaced by the representable number closest to v .

Signaling NaN. A type of *NaN* that generates an invalid operation program exception when it is specified as arithmetic operands. *See* Quiet NaN.

Significand. The component of a binary floating-point number that consists of an explicit or implicit leading bit to the left of its implied binary point and a fraction field to the right.

Splat. A splat instruction will take one element and replicate (splat) that value into a vector register.

Sticky bit. A bit that when *set* must be cleared explicitly.

Supervisor mode. The privileged operation state of a processor. In supervisor mode, software, typically the operating system, can access all control registers and can access the supervisor memory space, among other privileged operations.

T

Tiny. A floating-point value that is too small to be represented for a particular precision format, including *denormalized* numbers; they do not include ± 0 .

U

Underflow. An error condition that occurs during arithmetic operations when the result cannot be represented accurately in the destination register. For example, underflow can happen if two floating-point fractions are multiplied and the result requires a smaller *exponent* and/or mantissa than the single-precision format can provide. In other words, the result is too small to be represented accurately.

User mode. The unprivileged operating state of a processor used typically by application software. In user mode, software can only access certain control registers and can access only user memory space. No privileged operations can be performed. Also referred to as problem state.

V

Vector Literal. A vector literal is a constant expression with a value that is taken as a vector type. See Section 2.5.1, “Vector Literals” for details.

Vector Register (VR). Any of the 32 registers in the vector register file. Each vector register is 128 bits wide. These registers can provide the source operands and destination results for AltiVec instructions.

W

Word. A 32-bit data element.



INDEX

Symbols

#pragma altivec_codegen 2-10
 #pragma altivec_model 2-10
 #pragma altivec_vrsave 2-10
 __pixel 2-2, 2-3
 __va_arg 3-9
 __vector 2-2, 2-3

A

ABI 1-1, 1-2, 3-1
 ABS 4-4, 4-8, 4-10
 AIX ABI 3-1, 3-2, 3-10
 stack frame 3-5
 aligning data from an unaligned address 4-54, 4-55
 alignment
 aggregates and unions containing vector types 2-3
 non-vector types 2-3
 vector types 2-3
 AltiVec registers 3-1
 Apple Macintosh ABI 3-1, 3-2, 3-10
 stack frame 3-5

B

bool 2-2, 2-3
 BorrowOut 4-4, 4-121
 BoundAlign 4-4, 4-48, 4-50, 4-51, 4-112, 4-114,
 4-116
 byte ordering 4-3

C

cache touches
 all 4-37
 loads 4-38
 stores 4-40
 tag a 4-36
 transient loads 4-44
 transient stores 4-42
 calloc 3-10
 CarryOut 4-4, 4-15
 casts 2-5
 Ceil 4-4, 4-23
 condition register CR6 2-9
 cross-reference
 AltiVec Instructions to Operations/Predicates A-1
 AltiVec Operations to Instructions A-7
 AltiVec Predicates to Instructions A-14

D

data stream 4-36, 4-37, 4-38, 4-40, 4-42, 4-44
 DataStreamPrefetchControl 4-36, 4-37, 4-38, 4-40, 4-42, 4-44
 debugging information 3-11
 DWARF 3-12

E

EABI 3-1, 3-2, 3-3, 3-9
 Effective Address 4-48, 4-50, 4-51, 4-54, 4-55, 4-112,
 4-114, 4-116

F

Floor 4-5, 4-47
 FP2⁸Est 4-5, 4-46
 FPLog₂Est 4-5, 4-53
 FPRecipEst 4-5, 4-85
 fprintf 3-12
 fscanf 3-12

G

generic AltiVec operation 2-8

H

high-level language interface 1-1, 2-1
 high-order byte numbering 4-3

I

ISNaN 4-5, 4-150, 4-174
 ISNUM 4-5, 4-158, 4-182

L

longjmp 3-11

M

malloc 3-10
 MAX 4-5, 4-58
 MEM 4-5, 4-48, 4-50, 4-51, 4-112, 4-114, 4-116
 MIN 4-5, 4-66
 mod 4-50

N

NaN 4-5, 4-24, 4-58, 4-66, 4-85, 4-150, 4-154, 4-155,

INDEX

4-156, 4-157, 4-174, 4-178, 4-179, 4-180, 4-181
 NEG 4-5, 4-183
 NGE 4-5, 4-154, 4-178, 4-183
 NGT 4-5, 4-155, 4-179
 NJ bit 4-2, 4-8, 4-12, 4-23, 4-24, 4-25, 4-27, 4-28, 4-30, 4-31, 4-33, 4-34, 4-35, 4-46, 4-47, 4-53, 4-56, 4-58, 4-66, 4-77, 4-85, 4-88, 4-89, 4-118, 4-127, 4-134, 4-137, 4-140, 4-143, 4-144, 4-147, 4-150, 4-151, 4-154, 4-155, 4-156, 4-157, 4-158, 4-159, 4-162, 4-165, 4-168, 4-171, 4-174, 4-175, 4-178, 4-179, 4-180, 4-181, 4-182, 4-183
 NLE 4-5, 4-156, 4-180, 4-183
 NLT 4-5, 4-157, 4-181
 non-Java mode. See NJ bit
 notation and conventions 4-4

O

operation description format 4-7
 operator new 3-10

P

parameter passing 3-9, 3-10
 pixel 2-2, 2-3, 4-81, 4-128, 4-130
 pointer arithmetic 2-4
 pointer dereferencing 2-4
 precedence rules 4-6
 predicate 2-8, 4-133
 printf 3-12
 pseudocode 4-4

Q

QNaN 4-5, 4-58, 4-66, 4-85

R

realloc 3-10
 RecipSQRTTest 4-5, 4-89
 register usage conventions 3-1
 RndToFPINear 4-5, 4-88
 RndToFPITrunc 4-5, 4-127
 RndToFPNearest 4-5, 4-56, 4-77
 ROTL 4-5, 4-86
 Round to Nearest 4-88
 Round toward +Infinity 4-23
 Round toward Zero 4-127
 Round towards -Infinity 4-47

S

SAT bit 4-1, 4-2, 4-10, 4-16, 4-34, 4-35, 4-57, 4-70, 4-73, 4-82, 4-83, 4-122, 4-124, 4-125, 4-126
 Saturate 4-5, 4-10, 4-16, 4-34, 4-35, 4-57, 4-70, 4-73, 4-82, 4-83, 4-122, 4-124, 4-125, 4-126

saturation. See SAT bit
 save and restore functions 3-7
 scanf 3-12
 setjmp 3-11
 ShiftLeft 4-5, 4-91, 4-94
 ShiftRight 4-5, 4-105, 4-109
 ShiftRightA 4-5, 4-107
 SignExtend 4-5, 4-99, 4-100, 4-101, 4-102, 4-103, 4-104, 4-128, 4-130
 SIToFP 4-5, 4-33
 sizeof 2-4
 specific AltiVec operation 2-8
 sprintf 3-12
 sscanf 3-12
 stack frame 1-2, 3-2, 3-5
 SVR4 ABI 3-1, 3-2, 3-3, 3-9

T

type casting 2-5
 types 2-5

U

UIToUImod 4-6, 4-80
 Undefined 4-6, 4-50, 4-94, 4-109
 user-level cache operations
 vec_dss 4-36
 vec_dssall 4-37
 vec_dst 4-38
 vec_dstst 4-40
 vec_dststt 4-42
 vec_dstt 4-44

V

va_arg 3-10
 Varargs 3-9
 vec_abs 4-8
 vec_abss 4-10
 vec_add 2-8, 2-9, 4-12
 vec_addc 4-15
 vec_adds 4-16
 vec_addubm 2-8
 vec_all_eq 2-8, 4-134
 vec_all_ge 4-137
 vec_all_gt 2-9, 4-140
 vec_all_in 4-143
 vec_all_le 4-144
 vec_all_lt 2-9, 4-147
 vec_all_nan 2-9, 4-150
 vec_all_ne 4-151
 vec_all_nge 4-154
 vec_all_ngt 4-155
 vec_all_nle 4-156

INDEX

vec_all_nlt 4-157
 vec_all_numeric 4-158
 vec_alloc 3-10
 vec_and 4-18
 vec_andc 4-19
 vec_any_eq 4-159
 vec_any_ge 4-162
 vec_any_gt 4-165
 vec_any_le 4-168
 vec_any_lt 4-171
 vec_any_nan 4-174
 vec_any_ne 4-175
 vec_any_nge 4-178
 vec_any_ngt 4-179
 vec_any_nle 4-180
 vec_any_nlt 4-181
 vec_any_numeric 4-182
 vec_any_out 4-183
 vec_avg 4-21
 vec_calloc 3-10
 vec_ceil 4-23
 vec_cmpb 4-24
 vec_cmpeq 4-25
 vec_cmpge 4-27
 vec_cmpgt 4-28
 vec_cmples 4-30
 vec_cmplt 4-31
 vec_ctf 4-33
 vec_cts 4-34
 vec_ctu 4-35
 vec_data 2-2
 vec_dss 4-36
 vec_dssall 4-37
 vec_dst 4-38
 vec_dstst 4-40
 vec_dststt 4-42
 vec_dstt 4-44
 vec_expte 4-46
 vec_floor 4-47
 vec_free 3-10
 vec_ld 2-4, 4-48
 vec_lde 4-50
 vec_ldl 2-4, 4-51
 vec_log2 4-53
 vec_lvs1 2-3, 4-54
 vec_lvsr 2-3, 4-55
 vec_madd 4-56
 vec_madds 4-57
 vec_malloc 3-10
 vec_max 4-8, 4-10, 4-58
 vec_mergeh 4-61
 vec_mergel 4-63
 vec_mfvscr 4-2, 4-65
 vec_min 4-8, 4-10, 4-66
 vec_mladd 4-69
 vec_mradds 4-70
 vec_msum 4-71
 vec_msums 4-73
 vec_mtvscr 4-74
 vec_mule 4-75
 vec_mulo 4-76
 vec_nmsub 4-77
 vec_nor 4-78
 vec_or 4-79, 4-129, 4-131
 vec_pack 4-80
 vec_packpx 4-81
 vec_packs 4-82
 vec_packsu 4-83
 vec_perm 2-3, 4-84
 vec_re 4-85
 vec_realloc 3-10
 vec_rl 4-81, 4-86
 vec_round 4-88
 vec_rsqte 4-89
 vec_sel 4-90
 vec_sl 4-91, 4-129, 4-131
 vec_sld 4-93
 vec_sll 4-94
 vec_slo 4-96
 vec_splat 4-97
 vec_splat_s16 4-100
 vec_splat_s32 4-101
 vec_splat_s8 4-99
 vec_splat_u16 4-103
 vec_splat_u32 4-104
 vec_splat_u8 4-102
 vec_sr 4-105, 4-129, 4-131
 vec_sra 4-107
 vec_srl 4-109
 vec_sro 4-111
 vec_st 2-4, 4-112
 vec_ste 4-114
 vec_step 2-8
 vec_stl 2-4, 4-116
 vec_sub 4-8, 4-118
 vec_subc 4-121
 vec_subs 4-10, 4-122
 vec_sum2s 4-125
 vec_sum4s 4-124
 vec_sums 4-126
 vec_trunc 4-127
 vec_unpackh 4-128, 4-129, 4-131
 vec_unpackl 4-130
 vec_vaddubh 2-9
 vec_vaddubm 2-9
 vec_vaddubs 2-9
 vec_vadduhm 2-9
 vec_xor 4-132
 vector 2-2, 2-3
 vector bool char 2-1, 2-5

INDEX

- vector bool int 2-2, 2-5
- vector bool long 2-2
- vector bool long int 2-2
- vector bool short 2-1, 2-5
- vector bool short int 2-1
- vector cast 2-7
- vector data types 3-1
- vector float 2-2, 2-5, 2-7
- vector literal 2-7
- vector operations, arithmetic
 - vec_abs 4-8
 - vec_abss 4-10
 - vec_add 4-12
 - vec_addc 4-15
 - vec_adds 4-16
 - vec_avg 4-21
 - vec_max 4-58
 - vec_min 4-66
 - vec_mule 4-75
 - vec_mulo 4-76
 - vec_sub 4-118
 - vec_subc 4-121
 - vec_subs 4-122
- vector operations, compare
 - vec_cmpb 4-24
 - vec_cmpeq 4-25
 - vec_cmpge 4-27
 - vec_cmpgt 4-28
 - vec_cmple 4-30
 - vec_cmplt 4-31
- vector operations, function estimate
 - vec_expte 4-46
 - vec_logc 4-53
 - vec_re 4-85
 - vec_rsqrte 4-89
- vector operations, load/store
 - vec_ld 4-48
 - vec_lde 4-50
 - vec_ldl 4-51
 - vec_st 4-112
 - vec_ste 4-114
 - vec_stl 4-116
- vector operations, logical
 - vec_and 4-18
 - vec_andc 4-19
 - vec_nor 4-78
 - vec_or 4-79
 - vec_sel 4-90
 - vec_xor 4-132
- vector operations, merge
 - vec_mergeh 4-61
 - vec_mergel 4-63
- vector operations, miscellaneous
 - vec_alloc 3-10
 - vec_calloc 3-10
 - vec_free 3-10
 - vec_malloc 3-10
 - vec_mfvscr 4-65
 - vec_mtvscr 4-74
 - vec_realloc 3-10
 - vec_step 2-8
 - vector cast 2-7
 - vector literals 2-7
- vector operations, mixed arithmetic
 - vec_madd 4-56
 - vec_madds 4-57
 - vec_mladd 4-69
 - vec_mradds 4-70
 - vec_msum 4-71
 - vec_msums 4-73
 - vec_nmsub 4-77
 - vec_sum2s 4-125
 - vec_sum4s 4-124
 - vec_sums 4-126
- vector operations, pack and unpack
 - vec_pack 4-80
 - vec_packpx 4-81
 - vec_packs 4-82
 - vec_packsu 4-83
 - vec_unpackh 4-128
 - vec_unpackl 4-130
- vector operations, permute
 - vec_perm 4-84
- vector operations, rounding and conversion
 - vec_ceil 4-23
 - vec_ctf 4-33
 - vec_cts 4-34
 - vec_ctu 4-35
 - vec_floor 4-47
 - vec_round 4-88
 - vec_trunc 4-127
- vector operations, shift
 - vec_sld 4-93
 - vec_sll 4-94
 - vec_slo 4-96
 - vec_srl 4-109
 - vec_sro 4-111
- vector operations, shift and rotate
 - vec_rl 4-86
 - vec_sl 4-91
 - vec_sr 4-105
 - vec_sra 4-107
- vector operations, splat
 - vec_splat 4-97
 - vec_splat_32 4-101
 - vec_splat_s16 4-100
 - vec_splat_s8 4-99
 - vec_splat_u16 4-103

INDEX

- vec_splat_u32 4-104
- vec_splat_u8 4-102
- vector operations, supporting alignment
 - vec_lvsr 4-54
 - vec_lvsl 4-55
- vector pixel 2-2, 2-5
- vector predicates
 - vec_all_eq 4-134
 - vec_all_ge 4-137
 - vec_all_gt 4-140
 - vec_all_in 4-143
 - vec_all_le 4-144
 - vec_all_lt 4-147
 - vec_all_nan 4-150
 - vec_all_ne 4-151
 - vec_all_nge 4-154
 - vec_all_ngt 4-155
 - vec_all_nle 4-156
 - vec_all_nlt 4-157
 - vec_all_numeric 4-158
 - vec_any_eq 4-159
 - vec_any_ge 4-162
 - vec_any_gt 4-165
 - vec_any_le 4-168
 - vec_any_lt 4-171
 - vec_any_nan 4-174
 - vec_any_ne 4-175
 - vec_any_nge 4-178
 - vec_any_ngt 4-179
 - vec_any_nle 4-180
 - vec_any_nlt 4-181
 - vec_any_numeric 4-182
 - vec_any_out 4-183
- vector register 1-2
- vector register saving and restoring functions 3-7
- vector signed char 2-1, 2-5, 2-7
- vector signed int 2-2, 2-5, 2-7
- vector signed long 2-2
- vector signed long int 2-2
- vector signed short 2-1, 2-5, 2-7
- vector signed short int 2-1
- vector unsigned char 2-1, 2-5, 2-7
- vector unsigned int 2-1, 2-5, 2-7
- vector unsigned long 2-1
- vector unsigned long int 2-1
- vector unsigned short 2-1, 2-5, 2-7
- vector unsigned short int 2-1
- vfprintf 3-12
- vprintf 3-12
- VRSAVE 3-2, 3-4, 3-6, 3-11
- VSCR 4-1, 4-65, 4-74
- vsprintf 3-12

W

website xv, xviii, 1-1

X

xcoff stabstrings 3-12



INDEX



Overview 1

High-Level Language Interface 2

Application Binary Interface 3

AltiVec Operations and Predicates 4

AltiVec Instruction Set/Operations/Predicates Cross-Reference A

Glossary of Terms and Abbreviations GLO

Index IND



1

Overview

2

High-Level Language Interface

3

Application Binary Interface

4

AltiVec Operations and Predicates

A

AltiVec Instruction Set/Operations/Predicates Cross-Reference

GLO

Glossary of Terms and Abbreviations

IND

Index



Attention!

This book is a companion to the *PowerPC Microprocessor Family: The Programming Environments*, referred to as *The Programming Environments Manual*. Note that the companion *Programming Environments Manual* exists in two versions. See the Preface for a description of the following two versions:

- *PowerPC Microprocessor Family: The Programming Environments*, Rev 1
Order #: MPCFPE/AD
- *PowerPC Microprocessor Family: The Programming Environments for 32-Bit Microprocessors*, Rev 1
Order #: MPCFPE32B/AD

