



man pages section 2: System Calls

Sun Microsystems, Inc.
4150 Network Circle
Santa Clara, CA 95054
U.S.A.

Part No: 817-5439-10
February 2004

Copyright 2004 Sun Microsystems, Inc. 4150 Network Circle, Santa Clara, CA 95054 U.S.A. All rights reserved.

This product or document is protected by copyright and distributed under licenses restricting its use, copying, distribution, and decompilation. No part of this product or document may be reproduced in any form by any means without prior written authorization of Sun and its licensors, if any. Third-party software, including font technology, is copyrighted and licensed from Sun suppliers.

Parts of the product may be derived from Berkeley BSD systems, licensed from the University of California. UNIX is a registered trademark in the U.S. and other countries, exclusively licensed through X/Open Company, Ltd.

Sun, Sun Microsystems, the Sun logo, docs.sun.com, AnswerBook, AnswerBook2, and Solaris are trademarks, registered trademarks, or service marks of Sun Microsystems, Inc. in the U.S. and other countries. All SPARC trademarks are used under license and are trademarks or registered trademarks of SPARC International, Inc. in the U.S. and other countries. Products bearing SPARC trademarks are based upon an architecture developed by Sun Microsystems, Inc.

The OPEN LOOK and Sun™ Graphical User Interface was developed by Sun Microsystems, Inc. for its users and licensees. Sun acknowledges the pioneering efforts of Xerox in researching and developing the concept of visual or graphical user interfaces for the computer industry. Sun holds a non-exclusive license from Xerox to the Xerox Graphical User Interface, which license also covers Sun's licensees who implement OPEN LOOK GUIs and otherwise comply with Sun's written license agreements.

Federal Acquisitions: Commercial Software—Government Users Subject to Standard License Terms and Conditions.

DOCUMENTATION IS PROVIDED "AS IS" AND ALL EXPRESS OR IMPLIED CONDITIONS, REPRESENTATIONS AND WARRANTIES, INCLUDING ANY IMPLIED WARRANTY OF MERCHANTABILITY, FITNESS FOR A PARTICULAR PURPOSE OR NON-INFRINGEMENT, ARE DISCLAIMED, EXCEPT TO THE EXTENT THAT SUCH DISCLAIMERS ARE HELD TO BE LEGALLY INVALID.

Copyright 2004 Sun Microsystems, Inc. 4150 Network Circle, Santa Clara, CA 95054 U.S.A. Tous droits réservés.

Ce produit ou document est protégé par un copyright et distribué avec des licences qui en restreignent l'utilisation, la copie, la distribution, et la décompilation. Aucune partie de ce produit ou document ne peut être reproduite sous aucune forme, par quelque moyen que ce soit, sans l'autorisation préalable et écrite de Sun et de ses bailleurs de licence, s'il y en a. Le logiciel détenu par des tiers, et qui comprend la technologie relative aux polices de caractères, est protégé par un copyright et licencié par des fournisseurs de Sun.

Des parties de ce produit pourront être dérivées du système Berkeley BSD licenciés par l'Université de Californie. UNIX est une marque déposée aux Etats-Unis et dans d'autres pays et licenciée exclusivement par X/Open Company, Ltd.

Sun, Sun Microsystems, le logo Sun, docs.sun.com, AnswerBook, AnswerBook2, et Solaris sont des marques de fabrique ou des marques déposées, ou marques de service, de Sun Microsystems, Inc. aux Etats-Unis et dans d'autres pays. Toutes les marques SPARC sont utilisées sous licence et sont des marques de fabrique ou des marques déposées de SPARC International, Inc. aux Etats-Unis et dans d'autres pays. Les produits portant les marques SPARC sont basés sur une architecture développée par Sun Microsystems, Inc.

L'interface d'utilisation graphique OPEN LOOK et Sun™ a été développée par Sun Microsystems, Inc. pour ses utilisateurs et licenciés. Sun reconnaît les efforts de pionniers de Xerox pour la recherche et le développement du concept des interfaces d'utilisation visuelle ou graphique pour l'industrie de l'informatique. Sun détient une licence non exclusive de Xerox sur l'interface d'utilisation graphique Xerox, cette licence couvrant également les licenciés de Sun qui mettent en place l'interface d'utilisation graphique OPEN LOOK et qui en outre se conforment aux licences écrites de Sun.

CETTE PUBLICATION EST FOURNIE "EN L'ETAT" ET AUCUNE GARANTIE, EXPRESSE OU IMPLICITE, N'EST ACCORDEE, Y COMPRIS DES GARANTIES CONCERNANT LA VALEUR MARCHANDE, L'APTITUDE DE LA PUBLICATION A REPOUDRE A UNE UTILISATION PARTICULIERE, OU LE FAIT QU'ELLE NE SOIT PAS CONTREFAISANTE DE PRODUIT DE TIERS. CE DENI DE GARANTIE NE S'APPLIQUERAIT PAS, DANS LA MESURE OU IL SERAIT TENU JURIDIQUEMENT NUL ET NON AVENU.



031212@7518



Contents

Preface 9

Introduction 15

Intro(2) 16

System Calls 39

access(2) 40

acct(2) 42

acl(2) 43

adjtime(2) 45

alarm(2) 46

audit(2) 48

auditon(2) 49

auditsvc(2) 53

brk(2) 55

chdir(2) 57

chmod(2) 59

chown(2) 62

chroot(2) 65

close(2) 67

creat(2) 69

dup(2) 71

exec(2) 72

exit(2) 78

fcntl(2) 81

fork(2) 89
fpathconf(2) 92
getacct(2) 95
getaudit(2) 97
getauid(2) 99
getcontext(2) 100
getdents(2) 101
getgroups(2) 102
getitimer(2) 103
getmsg(2) 107
getpid(2) 110
getrlimit(2) 111
getsid(2) 114
getuid(2) 115
ioctl(2) 116
kill(2) 118
link(2) 120
llseek(2) 122
lseek(2) 123
_lwp_cond_signal(2) 125
_lwp_cond_wait(2) 126
_lwp_create(2) 129
_lwp_exit(2) 131
_lwp_info(2) 132
_lwp_kill(2) 133
_lwp_makecontext(2) 134
_lwp_mutex_lock(2) 135
_lwp_self(2) 136
_lwp_sema_wait(2) 137
_lwp_setprivate(2) 138
_lwp_suspend(2) 139
_lwp_wait(2) 140
memcntl(2) 141
mincore(2) 145
mkdir(2) 146
mknod(2) 148
mmap(2) 151
mount(2) 157

mprotect(2) 160
msgctl(2) 161
msgget(2) 163
msgids(2) 164
msgrcv(2) 166
msgsnap(2) 168
msgsnd(2) 171
munmap(2) 173
nice(2) 174
ntp_adjtime(2) 175
ntp_gettime(2) 176
open(2) 177
pause(2) 183
pcsample(2) 184
pipe(2) 185
poll(2) 186
p_online(2) 189
priocntl(2) 191
priocntlset(2) 200
processor_bind(2) 202
processor_info(2) 204
profil(2) 205
pset_bind(2) 207
pset_create(2) 209
pset_info(2) 211
ptrace(2) 212
putmsg(2) 214
read(2) 217
readlink(2) 222
rename(2) 223
resolvepath(2) 226
rmdir(2) 227
semctl(2) 229
semget(2) 232
semids(2) 234
semop(2) 236
setpgid(2) 239
setpgrp(2) 240

setregid(2) 241
 setreuid(2) 242
 setsid(2) 243
 settaskid(2) 244
 setuid(2) 245
 shmctl(2) 247
 shmget(2) 249
 shmids(2) 251
 shmop(2) 253
 sigaction(2) 255
 sigaltstack(2) 258
 _signotifywait(2) 260
 sigpending(2) 262
 sigprocmask(2) 263
 sigsend(2) 264
 sigsuspend(2) 266
 sigwait(2) 268
 __sparc_utrap_install(2) 272
 stat(2) 277
 statvfs(2) 280
 stime(2) 282
 swapctl(2) 283
 symlink(2) 287
 sync(2) 289
 sysfs(2) 290
 sysinfo(2) 291
 time(2) 294
 times(2) 295
 uadmin(2) 297
 ulimit(2) 299
 umask(2) 300
 umount(2) 301
 uname(2) 303
 unlink(2) 304
 ustat(2) 306
 utime(2) 307
 utimes(2) 309
 vfork(2) 311

vhangup(2) 313

wait(2) 314

waitid(2) 316

waitpid(2) 318

write(2) 320

yield(2) 326

Index 327

Preface

Both novice users and those familiar with the SunOS operating system can use online man pages to obtain information about the system and its features. A man page is intended to answer concisely the question “What does it do?” The man pages in general comprise a reference manual. They are not intended to be a tutorial.

Overview

The following contains a brief description of each man page section and the information it references:

- Section 1 describes, in alphabetical order, commands available with the operating system.
- Section 1M describes, in alphabetical order, commands that are used chiefly for system maintenance and administration purposes.
- Section 2 describes all of the system calls. Most of these calls have one or more error returns. An error condition is indicated by an otherwise impossible returned value.
- Section 3 describes functions found in various libraries, other than those functions that directly invoke UNIX system primitives, which are described in Section 2.
- Section 4 outlines the formats of various files. The C structure declarations for the file formats are given where applicable.
- Section 5 contains miscellaneous documentation such as character-set tables.
- Section 6 contains available games and demos.
- Section 7 describes various special files that refer to specific hardware peripherals and device drivers. STREAMS software drivers, modules and the STREAMS-generic set of system calls are also described.

- Section 9 provides reference information needed to write device drivers in the kernel environment. It describes two device driver interface specifications: the Device Driver Interface (DDI) and the Driver/Kernel Interface (DKI).
- Section 9E describes the DDI/DKI, DDI-only, and DKI-only entry-point routines a developer can include in a device driver.
- Section 9F describes the kernel functions available for use by device drivers.
- Section 9S describes the data structures used by drivers to share information between the driver and the kernel.

Below is a generic format for man pages. The man pages of each manual section generally follow this order, but include only needed headings. For example, if there are no bugs to report, there is no BUGS section. See the `intro` pages for more information and detail about each section, and `man(1)` for more information about man pages in general.

NAME	This section gives the names of the commands or functions documented, followed by a brief description of what they do.								
SYNOPSIS	This section shows the syntax of commands or functions. When a command or file does not exist in the standard path, its full path name is shown. Options and arguments are alphabetized, with single letter arguments first, and options with arguments next, unless a different argument order is required. The following special characters are used in this section: <table> <tr> <td>[]</td><td>Brackets. The option or argument enclosed in these brackets is optional. If the brackets are omitted, the argument must be specified.</td></tr> <tr> <td>. . .</td><td>Ellipses. Several values can be provided for the previous argument, or the previous argument can be specified multiple times, for example, "filename . . .".</td></tr> <tr> <td> </td><td>Separator. Only one of the arguments separated by this character can be specified at a time.</td></tr> <tr> <td>{ }</td><td>Braces. The options and/or arguments enclosed within braces are interdependent, such that everything enclosed must be treated as a unit.</td></tr> </table>	[]	Brackets. The option or argument enclosed in these brackets is optional. If the brackets are omitted, the argument must be specified.	. . .	Ellipses. Several values can be provided for the previous argument, or the previous argument can be specified multiple times, for example, "filename . . .".		Separator. Only one of the arguments separated by this character can be specified at a time.	{ }	Braces. The options and/or arguments enclosed within braces are interdependent, such that everything enclosed must be treated as a unit.
[]	Brackets. The option or argument enclosed in these brackets is optional. If the brackets are omitted, the argument must be specified.								
. . .	Ellipses. Several values can be provided for the previous argument, or the previous argument can be specified multiple times, for example, "filename . . .".								
	Separator. Only one of the arguments separated by this character can be specified at a time.								
{ }	Braces. The options and/or arguments enclosed within braces are interdependent, such that everything enclosed must be treated as a unit.								

PROTOCOL	This section occurs only in subsection 3R to indicate the protocol description file.
DESCRIPTION	This section defines the functionality and behavior of the service. Thus it describes concisely what the command does. It does not discuss OPTIONS or cite EXAMPLES. Interactive commands, subcommands, requests, macros, and functions are described under USAGE.
IOCTL	This section appears on pages in Section 7 only. Only the device class that supplies appropriate parameters to the <code>ioctl(2)</code> system call is called <code>ioctl</code> and generates its own heading. <code>ioctl</code> calls for a specific device are listed alphabetically (on the man page for that specific device). <code>ioctl</code> calls are used for a particular class of devices all of which have an <code>io</code> ending, such as <code>mtio(7I)</code> .
OPTIONS	This section lists the command options with a concise summary of what each option does. The options are listed literally and in the order they appear in the SYNOPSIS section. Possible arguments to options are discussed under the option, and where appropriate, default values are supplied.
OPERANDS	This section lists the command operands and describes how they affect the actions of the command.
OUTPUT	This section describes the output – standard output, standard error, or output files – generated by the command.
RETURN VALUES	If the man page documents functions that return values, this section lists these values and describes the conditions under which they are returned. If a function can return only constant values, such as 0 or -1, these values are listed in tagged paragraphs. Otherwise, a single paragraph describes the return values of each function. Functions declared void do not return values, so they are not discussed in RETURN VALUES.
ERRORS	On failure, most functions place an error code in the global variable <code>errno</code> indicating why they failed. This section lists alphabetically all error codes a function can generate and describes the conditions that cause each error. When more than

	one condition can cause the same error, each condition is described in a separate paragraph under the error code.
USAGE	This section lists special rules, features, and commands that require in-depth explanations. The subsections listed here are used to explain built-in functionality: - Commands - Modifiers - Variables - Expressions - Input Grammar
EXAMPLES	This section provides examples of usage or of how to use a command or function. Wherever possible a complete example including command-line entry and machine response is shown. Whenever an example is given, the prompt is shown as <code>example%</code>, or if the user must be superuser, <code>example#</code>. Examples are followed by explanations, variable substitution rules, or returned values. Most examples illustrate concepts from the SYNOPSIS, DESCRIPTION, OPTIONS, and USAGE sections.
ENVIRONMENT VARIABLES	This section lists any environment variables that the command or function affects, followed by a brief description of the effect.
EXIT STATUS	This section lists the values the command returns to the calling program or shell and the conditions that cause these values to be returned. Usually, zero is returned for successful completion, and values other than zero for various error conditions.
FILES	This section lists all file names referred to by the man page, files of interest, and files created or required by commands. Each is followed by a descriptive summary or explanation.
ATTRIBUTES	This section lists characteristics of commands, utilities, and device drivers by defining the attribute type and its corresponding value. See <code>attributes(5)</code> for more information.
SEE ALSO	This section lists references to other man pages, in-house documentation, and outside publications.

DIAGNOSTICS	This section lists diagnostic messages with a brief explanation of the condition causing the error.
WARNINGS	This section lists warnings about special conditions which could seriously affect your working conditions. This is not a list of diagnostics.
NOTES	This section lists additional information that does not belong anywhere else on the page. It takes the form of an aside to the user, covering points of special interest. Critical information is never covered here.
BUGS	This section describes known bugs and, wherever possible, suggests workarounds.

Introduction

Intro(2)

NAME	Intro – introduction to system calls and error numbers																
SYNOPSIS	<pre>#include <errno.h></pre>																
DESCRIPTION	This section describes all of the system calls. Most of these calls return one or more error conditions. An error condition is indicated by an otherwise impossible return value. This is almost always <code>-1</code> or the null pointer; the individual descriptions specify the details. An error number is also made available in the external variable <code>errno</code>, which is not cleared on successful calls, so it should be tested only after an error has been indicated. In the case of multithreaded applications, the <code>_REENTRANT</code> flag must be defined on the command line at compilation time (<code>-D_REENTRANT</code>). When the <code>_REENTRANT</code> flag is defined, <code>errno</code> becomes a macro which enables each thread to have its own <code>errno</code>. This <code>errno</code> macro can be used on either side of the assignment, just as if it were a variable. Applications should use bound threads rather than the <code>_lwp_*()</code> functions (see <code>thr_create(3THR)</code>). Using LWPs (lightweight processes) directly is not advised because libraries are only safe to use with threads, not LWPs. Each system call description attempts to list all possible error numbers. The following is a complete list of the error numbers and their names as defined in <code><errno.h></code>. <table><tr><td>1 EPERM</td><td>Not superuser</td></tr><tr><td></td><td>Typically this error indicates an attempt to modify a file in some way forbidden except to its owner or the superuser. It is also returned for attempts by ordinary users to do things allowed only to the superuser.</td></tr><tr><td>2 ENOENT</td><td>No such file or directory</td></tr><tr><td></td><td>A file name is specified and the file should exist but doesn't, or one of the directories in a path name does not exist.</td></tr><tr><td>3 ESRCH</td><td>No such process, LWP, or thread</td></tr><tr><td></td><td>No process can be found in the system that corresponds to the specified PID, LWPID_t, or thread_t.</td></tr><tr><td>4 EINTR</td><td>Interrupted system call</td></tr><tr><td></td><td>An asynchronous signal (such as interrupt or quit), which the user has elected to catch, occurred during a system service function. If execution is resumed after processing the signal, it will appear as if the interrupted function call returned this error condition.</td></tr></table>	1 EPERM	Not superuser		Typically this error indicates an attempt to modify a file in some way forbidden except to its owner or the superuser. It is also returned for attempts by ordinary users to do things allowed only to the superuser.	2 ENOENT	No such file or directory		A file name is specified and the file should exist but doesn't, or one of the directories in a path name does not exist.	3 ESRCH	No such process, LWP, or thread		No process can be found in the system that corresponds to the specified PID, LWPID_t, or thread_t.	4 EINTR	Interrupted system call		An asynchronous signal (such as interrupt or quit), which the user has elected to catch, occurred during a system service function. If execution is resumed after processing the signal, it will appear as if the interrupted function call returned this error condition.
1 EPERM	Not superuser																
	Typically this error indicates an attempt to modify a file in some way forbidden except to its owner or the superuser. It is also returned for attempts by ordinary users to do things allowed only to the superuser.																
2 ENOENT	No such file or directory																
	A file name is specified and the file should exist but doesn't, or one of the directories in a path name does not exist.																
3 ESRCH	No such process, LWP, or thread																
	No process can be found in the system that corresponds to the specified PID, LWPID_t, or thread_t.																
4 EINTR	Interrupted system call																
	An asynchronous signal (such as interrupt or quit), which the user has elected to catch, occurred during a system service function. If execution is resumed after processing the signal, it will appear as if the interrupted function call returned this error condition.																

	In a multithreaded application, <code>EINTR</code> may be returned whenever another thread or LWP calls <code>fork(2)</code> .
5 EIO	I/O error Some physical I/O error has occurred. This error may in some cases occur on a call following the one to which it actually applies.
6 ENXIO	No such device or address I/O on a special file refers to a subdevice which does not exist, or exists beyond the limit of the device. It may also occur when, for example, a tape drive is not on-line or no disk pack is loaded on a drive.
7 E2BIG	Arg list too long An argument list longer than <code>ARG_MAX</code> bytes is presented to a member of the <code>exec</code> family of functions (see <code>exec(2)</code>). The argument list limit is the sum of the size of the argument list plus the size of the environment's exported shell variables.
8 ENOEXEC	Exec format error A request is made to execute a file which, although it has the appropriate permissions, does not start with a valid format (see <code>a.out(4)</code>).
9 EBADF	Bad file number Either a file descriptor refers to no open file, or a <code>read(2)</code> (respectively, <code>write(2)</code>) request is made to a file that is open only for writing (respectively, reading).
10 ECHILD	No child processes A <code>wait(2)</code> function was executed by a process that had no existing or unwaited-for child processes.
11 EAGAIN	No more processes, or no more LWPs For example, the <code>fork(2)</code> function failed because the system's process table is full or the user is not allowed to create any more processes, or a call failed because of insufficient memory or swap space.
12 ENOMEM	Not enough space

Intro(2)

	During execution of <code>brk()</code> or <code>sbrk()</code> (see <code>brk(2)</code>), or one of the <code>exec</code> family of functions, a program asks for more space than the system is able to supply. This is not a temporary condition; the maximum size is a system parameter. On some architectures, the error may also occur if the arrangement of text, data, and stack segments requires too many segmentation registers, or if there is not enough swap space during the <code>fork(2)</code> function. If this error occurs on a resource associated with Remote File Sharing (RFS), it indicates a memory depletion which may be temporary, dependent on system activity at the time the call was invoked.
13 EACCES	Permission denied An attempt was made to access a file in a way forbidden by the protection system.
14 EFAULT	Bad address The system encountered a hardware fault in attempting to use an argument of a routine. For example, <code>errno</code> potentially may be set to <code>EFAULT</code> any time a routine that takes a pointer argument is passed an invalid address, if the system can detect the condition. Because systems will differ in their ability to reliably detect a bad address, on some implementations passing a bad address to a routine will result in undefined behavior.
15 ENOTBLK	Block device required A non-block device or file was mentioned where a block device was required (for example, in a call to the <code>mount(2)</code> function).
16 EBUSY	Device busy An attempt was made to mount a device that was already mounted or an attempt was made to unmount a device on which there is an active file (open file, current directory, mounted-on file, active text segment). It will also occur if an attempt is made to enable accounting when it is already enabled. The device or resource is currently unavailable. <code>EBUSY</code> is also used by mutexes, semaphores, condition variables, and <code>r/w</code> locks, to indicate that a lock is held, and by the processor control function <code>P_ONLINE</code>.
17 EEXIST	File exists

	An existing file was mentioned in an inappropriate context (for example, call to the <code>link(2)</code> function).
18 EXDEV	Cross-device link
	A hard link to a file on another device was attempted.
19 ENODEV	No such device
	An attempt was made to apply an inappropriate operation to a device (for example, read a write-only device).
20 ENOTDIR	Not a directory
	A non-directory was specified where a directory is required (for example, in a path prefix or as an argument to the <code>chdir(2)</code> function).
21 EISDIR	Is a directory
	An attempt was made to write on a directory.
22 EINVAL	Invalid argument
	An invalid argument was specified (for example, unmounting a non-mounted device), mentioning an undefined signal in a call to the <code>signal(3C)</code> or <code>kill(2)</code> function.
23 ENFILE	File table overflow
	The system file table is full (that is, <code>SYS_OPEN</code> files are open, and temporarily no more files can be opened).
24 EMFILE	Too many open files
	No process may have more than <code>OPEN_MAX</code> file descriptors open at a time.
25 ENOTTY	Inappropriate <code>ioctl</code> for device
	A call was made to the <code>ioctl(2)</code> function specifying a file that is not a special character device.
26 ETXTBSY	Text file busy (obsolete)
	An attempt was made to execute a pure-procedure program that is currently open for writing. Also an attempt to open for writing or to remove a pure-procedure program that is being executed. (<i>This message is obsolete.</i>)
27 EFBIG	File too large

Intro(2)

	The size of the file exceeded the limit specified by resource <code>RLIMIT_FSIZE</code> ; the file size exceeds the maximum supported by the file system; or the file size exceeds the offset maximum of the file descriptor. See the <code>File Descriptor</code> subsection of the <code>DEFINITIONS</code> section below.
28 ENOSPC	No space left on device
	While writing an ordinary file or creating a directory entry, there is no free space left on the device. In the <code>fcntl(2)</code> function, the setting or removing of record locks on a file cannot be accomplished because there are no more record entries left on the system.
29 EPIPE	Illegal seek
	A call to the <code>lseek(2)</code> function was issued to a pipe.
30 EROFS	Read-only file system
	An attempt to modify a file or directory was made on a device mounted read-only.
31 EMLINK	Too many links
	An attempt to make more than the maximum number of links, <code>LINK_MAX</code> , to a file.
32 EPIPE	Broken pipe
	A write on a pipe for which there is no process to read the data. This condition normally generates a signal; the error is returned if the signal is ignored.
33 EDOM	Math argument out of domain of func
	The argument of a function in the math package (3M) is out of the domain of the function.
34 ERANGE	Math result not representable
	The value of a function in the math package (3M) is not representable within machine precision.
35 ENOMSG	No message of desired type
	An attempt was made to receive a message of a type that does not exist on the specified message queue (see <code>msgrcv(2)</code>).
36 EIDRM	Identifier removed

	This error is returned to processes that resume execution due to the removal of an identifier from the file system's name space (see <code>msgctl(2)</code> , <code>semctl(2)</code> , and <code>shmctl(2)</code>).
37 ECHRNG	Channel number out of range
38 EL2NSYNC	Level 2 not synchronized
39 EL3HLT	Level 3 halted
40 EL3RST	Level 3 reset
41 ELNRNG	Link number out of range
42 EUNATCH	Protocol driver not attached
43 ENOCSI	No CSI structure available
44 EL2HLT	Level 2 halted
45 EDEADLK	Deadlock condition
	A deadlock situation was detected and avoided. This error pertains to file and record locking, and also applies to mutexes, semaphores, condition variables, and r/w locks.
46 ENOLCK	No record locks available
	There are no more locks available. The system lock table is full (see <code>fcntl(2)</code>).
47 ECANCELED	Operation canceled
	The associated asynchronous operation was canceled before completion.
48 ENOTSUP	Not supported
	This version of the system does not support this feature. Future versions of the system may provide support.
49 EDQUOT	Disc quota exceeded
	A <code>write(2)</code> to an ordinary file, the creation of a directory or symbolic link, or the creation of a directory entry failed because the user's quota of disk blocks was exhausted, or the allocation of an inode for a newly created file failed because the user's quota of inodes was exhausted.
58-59	Reserved

Intro(2)

60 ENOSTR	Device not a stream A <code>putmsg(2)</code> or <code>getmsg(2)</code> call was attempted on a file descriptor that is not a STREAMS device.
61 ENODATA	No data available
62 ETIME	Timer expired The timer set for a STREAMS <code>ioctl(2)</code> call has expired. The cause of this error is device-specific and could indicate either a hardware or software failure, or perhaps a timeout value that is too short for the specific operation. The status of the <code>ioctl()</code> operation is indeterminate. This is also returned in the case of <code>_lwp_cond_timedwait(2)</code> or <code>cond_timedwait(3THR)</code> .
63 ENOSR	Out of stream resources During a STREAMS <code>open(2)</code> call, either no STREAMS queues or no STREAMS head data structures were available. This is a temporary condition; one may recover from it if other processes release resources.
64 ENONET	Machine is not on the network This error is Remote File Sharing (RFS) specific. It occurs when users try to advertise, unadvertise, mount, or unmount remote resources while the machine has not done the proper startup to connect to the network.
65 ENOPKG	Package not installed This error occurs when users attempt to use a call from a package which has not been installed.
66 EREMOTE	Object is remote This error is RFS-specific. It occurs when users try to advertise a resource which is not on the local machine, or try to mount/unmount a device (or pathname) that is on a remote machine.
67 ENOLINK	Link has been severed This error is RFS-specific. It occurs when the link (virtual circuit) connecting to a remote machine is gone.
68 EADV	Advertise error

	This error is RFS-specific. It occurs when users try to advertise a resource which has been advertised already, or try to stop RFS while there are resources still advertised, or try to force unmount a resource when it is still advertised.
69 ESRMNT	Srmount error This error is RFS-specific. It occurs when an attempt is made to stop RFS while resources are still mounted by remote machines, or when a resource is readadvertised with a client list that does not include a remote machine that currently has the resource mounted.
70 ECOMM	Communication error on send This error is RFS-specific. It occurs when the current process is waiting for a message from a remote machine, and the virtual circuit fails.
71 EPROTO	Protocol error Some protocol error occurred. This error is device-specific, but is generally not related to a hardware failure.
76 EDOTDOT	Error 76 This error is RFS-specific. A way for the server to tell the client that a process has transferred back from mount point.
77 EBADMSG	Not a data message During a <code>read(2)</code>, <code>getmsg(2)</code>, or <code>ioctl(2)</code> <code>I_RECVFD</code> call to a STREAMS device, something has come to the head of the queue that can not be processed. That something depends on the call: <code>read()</code>: control information or passed file descriptor. <code>getmsg()</code>: passed file descriptor. <code>ioctl()</code>: control or data information.
78 ENAMETOOLONG	File name too long The length of the path argument exceeds <code>PATH_MAX</code>, or the length of a path component exceeds <code>NAME_MAX</code> while <code>_POSIX_NO_TRUNC</code> is in effect; see <code>limits(4)</code>.
79 EOVERFLOW	Value too large for defined data type.

Intro(2)

80 ENOTUNIQ	Name not unique on network Given log name not unique.
81 EBADFD	File descriptor in bad state Either a file descriptor refers to no open file or a read request was made to a file that is open only for writing.
82 EREMCHG	Remote address changed
83 ELIBACC	Cannot access a needed share library Trying to <code>exec</code> an <code>a.out</code> that requires a static shared library and the static shared library does not exist or the user does not have permission to use it.
84 ELIBBAD	Accessing a corrupted shared library Trying to <code>exec</code> an <code>a.out</code> that requires a static shared library (to be linked in) and <code>exec</code> could not load the static shared library. The static shared library is probably corrupted.
85 ELIBSCN	<code>.lib</code> section in <code>a.out</code> corrupted Trying to <code>exec</code> an <code>a.out</code> that requires a static shared library (to be linked in) and there was erroneous data in the <code>.lib</code> section of the <code>a.out</code> . The <code>.lib</code> section tells <code>exec</code> what static shared libraries are needed. The <code>a.out</code> is probably corrupted.
86 ELIBMAX	Attempting to link in more shared libraries than system limit Trying to <code>exec</code> an <code>a.out</code> that requires more static shared libraries than is allowed on the current configuration of the system. See <i>System Administration Guide, Volume 3</i>
87 ELIBEXEC	Cannot <code>exec</code> a shared library directly Attempting to <code>exec</code> a shared library directly.
88 EILSEQ	Error 88 Illegal byte sequence. Handle multiple characters as a single character.
89 ENOSYS	Operation not applicable
90 ELOOP	Number of symbolic links encountered during path name traversal exceeds <code>MAXSYMLINKS</code>

91 ESTART	Restartable system call Interrupted system call should be restarted.
92 ESTRPIPE	If pipe/FIFO, don't sleep in stream head Streams pipe error (not externally visible).
93 ENOTEMPTY	Directory not empty
94 EUSERS	Too many users
95 ENOTSOCK	Socket operation on non-socket
96 EDESTADDRREQ	Destination address required A required address was omitted from an operation on a transport endpoint. Destination address required.
97 EMGSIZE	Message too long A message sent on a transport provider was larger than the internal message buffer or some other network limit.
98 EPROTOTYPE	Protocol wrong type for socket A protocol was specified that does not support the semantics of the socket type requested.
99 ENOPROTOOPT	Protocol not available A bad option or level was specified when getting or setting options for a protocol.
120 EPROTONOSUPPORT	Protocol not supported The protocol has not been configured into the system or no implementation for it exists.
121 ESOCKTNOSUPPORT	Socket type not supported The support for the socket type has not been configured into the system or no implementation for it exists.
122 EOPNOTSUPP	Operation not supported on transport endpoint For example, trying to accept a connection on a datagram transport endpoint.
123 EPFNOSUPPORT	Protocol family not supported

Intro(2)

	The protocol family has not been configured into the system or no implementation for it exists. Used for the Internet protocols.
124 EAFNOSUPPORT	Address family not supported by protocol family
	An address incompatible with the requested protocol was used.
125 EADDRINUSE	Address already in use
	User attempted to use an address already in use, and the protocol does not allow this.
126 EADDRNOTAVAIL	Cannot assign requested address
	Results from an attempt to create a transport endpoint with an address not on the current machine.
127 ENETDOWN	Network is down
	Operation encountered a dead network.
128 ENETUNREACH	Network is unreachable
	Operation was attempted to an unreachable network.
129 ENETRESET	Network dropped connection because of reset
	The host you were connected to crashed and rebooted.
130 ECONNABORTED	Software caused connection abort
	A connection abort was caused internal to your host machine.
131 ECONNRESET	Connection reset by peer
	A connection was forcibly closed by a peer. This normally results from a loss of the connection on the remote host due to a timeout or a reboot.
132 ENOBUFS	No buffer space available
	An operation on a transport endpoint or pipe was not performed because the system lacked sufficient buffer space or because a queue was full.
133 EISCONN	Transport endpoint is already connected

	A connect request was made on an already connected transport endpoint; or, a <code>sendto(3SOCKET)</code> or <code>sendmsg(3SOCKET)</code> request on a connected transport endpoint specified a destination when already connected.
134 ENOTCONN	Transport endpoint is not connected
	A request to send or receive data was disallowed because the transport endpoint is not connected and (when sending a datagram) no address was supplied.
143 ESHUTDOWN	Cannot send after transport endpoint shutdown
	A request to send data was disallowed because the transport endpoint has already been shut down.
144 ETOOMANYREFS	Too many references: cannot splice
145 ETIMEDOUT	Connection timed out
	A <code>connect(3SOCKET)</code> or <code>send(3SOCKET)</code> request failed because the connected party did not properly respond after a period of time; or a <code>write(2)</code> or <code>fsync(3C)</code> request failed because a file is on an NFS file system mounted with the <i>soft</i> option.
146 ECONNREFUSED	Connection refused
	No connection could be made because the target machine actively refused it. This usually results from trying to connect to a service that is inactive on the remote host.
147 EHOSTDOWN	Host is down
	A transport provider operation failed because the destination host was down.
148 EHOSTUNREACH	No route to host
	A transport provider operation was attempted to an unreachable host.
149 EALREADY	Operation already in progress
	An operation was attempted on a non-blocking object that already had an operation in progress.
150 EINPROGRESS	Operation now in progress

Intro(2)

	An operation that takes a long time to complete (such as a connect ()) was attempted on a non-blocking object.
	151 ESTALE Stale NFS file handle
DEFINITIONS	
Background Process Group	Any process group that is not the foreground process group of a session that has established a connection with a controlling terminal.
Controlling Process	A session leader that established a connection to a controlling terminal.
Controlling Terminal	A terminal that is associated with a session. Each session may have, at most, one controlling terminal associated with it and a controlling terminal may be associated with only one session. Certain input sequences from the controlling terminal cause signals to be sent to process groups in the session associated with the controlling terminal; see <code>termio(7I)</code> .
Directory	Directories organize files into a hierarchical system where directories are the nodes in the hierarchy. A directory is a file that catalogs the list of files, including directories (sub-directories), that are directly beneath it in the hierarchy. Entries in a directory file are called links. A link associates a file identifier with a filename. By convention, a directory contains at least two links, <code>.</code> (dot) and <code>..</code> (dot-dot). The link called dot refers to the directory itself while dot-dot refers to its parent directory. The root directory, which is the top-most node of the hierarchy, has itself as its parent directory. The pathname of the root directory is <code>/</code> and the parent directory of the root directory is <code>/</code> .
Downstream	In a stream, the direction from stream head to driver.
Driver	In a stream, the driver provides the interface between peripheral hardware and the stream. A driver can also be a pseudo-driver, such as a multiplexor or log driver (see <code>log(7D)</code>), which is not associated with a hardware device.
Effective User ID and Effective Group ID	An active process has an effective user ID and an effective group ID that are used to determine file access permissions (see below). The effective user ID and effective group ID are equal to the process's real user ID and real group ID, respectively, unless the process or one of its ancestors evolved from a file that had the set-user-ID bit or set-group-ID bit set (see <code>exec(2)</code>).
File Access Permissions	Read, write, and execute/search permissions for a file are granted to a process if one or more of the following are true: ■ The effective user ID of the process is superuser. ■ The effective user ID of the process matches the user ID of the owner of the file and the appropriate access bit of the "owner" portion (0700) of the file mode is set. ■ The effective user ID of the process does not match the user ID of the owner of the file, but either the effective group ID or one of the supplementary group IDs of the process match the group ID of the file and the appropriate access bit of the "group"

	portion (0070) of the file mode is set. ■ The effective user ID of the process does not match the user ID of the owner of the file, and neither the effective group ID nor any of the supplementary group IDs of the process match the group ID of the file, but the appropriate access bit of the “other” portion (0007) of the file mode is set. Otherwise, the corresponding permissions are denied.
File Descriptor	A file descriptor is a small integer used to perform I/O on a file. The value of a file descriptor is from 0 to (NOFILES-1). A process may have no more than NOFILES file descriptors open simultaneously. A file descriptor is returned by calls such as open(2) or pipe(2). The file descriptor is used as an argument by calls such as read(2), write(2), ioctl(2), and close(2). Each file descriptor has a corresponding offset maximum. For regular files that were opened without setting the O_LARGEFILE flag, the offset maximum is 2 Gbyte - 1 byte ($2^{31} - 1$ bytes). For regular files that were opened with the O_LARGEFILE flag set, the offset maximum is $2^{63} - 1$ bytes.
File Name	Names consisting of 1 to NAME_MAX characters may be used to name an ordinary file, special file or directory. These characters may be selected from the set of all character values excluding \0 (null) and the ASCII code for / (slash). Note that it is generally unwise to use *, ?, [, or] as part of file names because of the special meaning attached to these characters by the shell (see sh(1), csh(1), and ksh(1)). Although permitted, the use of unprintable characters in file names should be avoided. A file name is sometimes referred to as a pathname component. The interpretation of a pathname component is dependent on the values of NAME_MAX and _POSIX_NO_TRUNC associated with the path prefix of that component. If any pathname component is longer than NAME_MAX and _POSIX_NO_TRUNC is in effect for the path prefix of that component (see fpathconf(2) and limits(4)), it shall be considered an error condition in that implementation. Otherwise, the implementation shall use the first NAME_MAX bytes of the pathname component.
Foreground Process Group	Each session that has established a connection with a controlling terminal will distinguish one process group of the session as the foreground process group of the controlling terminal. This group has certain privileges when accessing its controlling terminal that are denied to background process groups.
{IOV_MAX}	Maximum number of entries in a struct iovec array.
{LIMIT}	The braces notation, {LIMIT}, is used to denote a magnitude limitation imposed by the implementation. This indicates a value which may be defined by a header file (without the braces), or the actual value may be obtained at runtime by a call to the configuration inquiry pathconf(2) with the name argument _PC_LIMIT.

Intro(2)

Masks	The file mode creation mask of the process used during any create function calls to turn off permission bits in the <i>mode</i> argument supplied. Bit positions that are set in <i>umask</i> (<i>cmask</i>) are cleared in the mode of the created file.
Message	In a stream, one or more blocks of data or information, with associated STREAMS control structures. Messages can be of several defined types, which identify the message contents. Messages are the only means of transferring data and communicating within a stream.
Message Queue	In a stream, a linked list of messages awaiting processing by a module or driver.
Message Queue Identifier	A message queue identifier (<i>msqid</i>) is a unique positive integer created by a <i>msgget</i>(2) call. Each <i>msqid</i> has a message queue and a data structure associated with it. The data structure is referred to as <i>msqid_ds</i> and contains the following members: <pre>struct ipc_perm msg_perm; struct msg *msg_first; struct msg *msg_last; ulong_t msg_cbytes; ulong_t msg_qnum; ulong_t msg_qbytes; pid_t msg_lspid; pid_t msg_lrpid; time_t msg_stime; time_t msg_rtime; time_t msg_ctime;</pre> The following are descriptions of the <i>msqid_ds</i> structure members: The <i>msg_perm</i> member is an <i>ipc_perm</i> structure that specifies the message operation permission (see below). This structure includes the following members: <pre>uid_t cuid; /* creator user id */ gid_t cgid; /* creator group id */ uid_t uid; /* user id */ gid_t gid; /* group id */ mode_t mode; /* r/w permission */ ulong_t seq; /* slot usage sequence # */ key_t key; /* key */</pre> The <i>*msg_first</i> member is a pointer to the first message on the queue. The <i>*msg_last</i> member is a pointer to the last message on the queue. The <i>msg_cbytes</i> member is the current number of bytes on the queue. The <i>msg_qnum</i> member is the number of messages currently on the queue. The <i>msg_qbytes</i> member is the maximum number of bytes allowed on the queue. The <i>msg_lspid</i> member is the process ID of the last process that performed a <i>msgsnd()</i> operation.

	The <code>msg_lrpid</code> member is the process id of the last process that performed a <code>msgrcv()</code> operation. The <code>msg_stime</code> member is the time of the last <code>msgsnd()</code> operation. The <code>msg_rtime</code> member is the time of the last <code>msgrcv()</code> operation. The <code>msg_ctime</code> member is the time of the last <code>msgctl()</code> operation that changed a member of the above structure.
Message Operation Permissions	In the <code>msgctl(2)</code>, <code>msgget(2)</code>, <code>msgrcv(2)</code>, and <code>msgsnd(2)</code> function descriptions, the permission required for an operation is given as <i>{token}</i>, where <i>token</i> is the type of permission needed, interpreted as follows: <pre> 00400 READ by user 00200 WRITE by user 00040 READ by group 00020 WRITE by group 00004 READ by others 00002 WRITE by others </pre> Read and write permissions for a <code>msgid</code> are granted to a process if one or more of the following are true: ■ The effective user ID of the process is superuser. ■ The effective user ID of the process matches <code>msg_perm.cuid</code> or <code>msg_perm.uid</code> in the data structure associated with <code>msgid</code> and the appropriate bit of the “user” portion (0600) of <code>msg_perm.mode</code> is set. ■ Any group ID in the process credentials from the set (<code>cr_gid</code>, <code>cr_groups</code>) matches <code>msg_perm.cgid</code> or <code>msg_perm.gid</code> and the appropriate bit of the “group” portion (060) of <code>msg_perm.mode</code> is set. ■ The appropriate bit of the “other” portion (006) of <code>msg_perm.mode</code> is set.” Otherwise, the corresponding permissions are denied.
Module	A module is an entity containing processing routines for input and output data. It always exists in the middle of a stream, between the stream’s head and a driver. A module is the STREAMS counterpart to the commands in a shell pipeline except that a module contains a pair of functions which allow independent bidirectional (downstream and upstream) data flow and processing.
Multiplexor	A multiplexor is a driver that allows streams associated with several user processes to be connected to a single driver, or several drivers to be connected to a single user process. STREAMS does not provide a general multiplexing driver, but does provide the facilities for constructing them and for connecting multiplexed configurations of streams.
Offset Maximum	An offset maximum is an attribute of an open file description representing the largest value that can be used as a file offset.

Intro(2)

Orphaned Process Group	A process group in which the parent of every member in the group is either itself a member of the group, or is not a member of the process group's session.
Path Name	A path name is a null-terminated character string starting with an optional slash (/), followed by zero or more directory names separated by slashes, optionally followed by a file name. If a path name begins with a slash, the path search begins at the root directory. Otherwise, the search begins from the current working directory. A slash by itself names the root directory. Unless specifically stated otherwise, the null path name is treated as if it named a non-existent file.
Process ID	Each process in the system is uniquely identified during its lifetime by a positive integer called a process ID. A process ID may not be reused by the system until the process lifetime, process group lifetime, and session lifetime ends for any process ID, process group ID, and session ID equal to that process ID. Within a process, there are threads with thread id's, called <code>thread_t</code> and <code>LWPID_t</code> . These threads are not visible to the outside process.
Parent Process ID	A new process is created by a currently active process (see <code>fork(2)</code>). The parent process ID of a process is the process ID of its creator.
Privilege	Having appropriate privilege means having the capability to override system restrictions.
Process Group	Each process in the system is a member of a process group that is identified by a process group ID. Any process that is not a process group leader may create a new process group and become its leader. Any process that is not a process group leader may join an existing process group that shares the same session as the process. A newly created process joins the process group of its parent.
Process Group Leader	A process group leader is a process whose process ID is the same as its process group ID.
Process Group ID	Each active process is a member of a process group and is identified by a positive integer called the process group ID. This ID is the process ID of the group leader. This grouping permits the signaling of related processes (see <code>kill(2)</code>).
Process Lifetime	A process lifetime begins when the process is forked and ends after it exits, when its termination has been acknowledged by its parent process. See <code>wait(2)</code> .
Process Group Lifetime	A process group lifetime begins when the process group is created by its process group leader, and ends when the lifetime of the last process in the group ends or when the last process in the group leaves the group.

Processor Set ID	The processors in a system may be divided into subsets, known as processor sets. A process bound to one of these sets will run only on processors in that set, and the processors in the set will normally run only processes that have been bound to the set. Each active processor set is identified by a positive integer. See <code>pset_create(2)</code> .
Read Queue	In a stream, the message queue in a module or driver containing messages moving upstream.
Real User ID and Real Group ID	Each user allowed on the system is identified by a positive integer (0 to <code>MAXUID</code>) called a real user ID. Each user is also a member of a group. The group is identified by a positive integer called the real group ID. An active process has a real user ID and real group ID that are set to the real user ID and real group ID, respectively, of the user responsible for the creation of the process.
Root Directory and Current Working Directory	Each process has associated with it a concept of a root directory and a current working directory for the purpose of resolving path name searches. The root directory of a process need not be the root directory of the root file system.
Saved Resource Limits	Saved resource limits is an attribute of a process that provides some flexibility in the handling of unrepresentable resource limits, as described in the <code>exec</code> family of functions and <code>setrlimit(2)</code> .
Saved User ID and Saved Group ID	The saved user ID and saved group ID are the values of the effective user ID and effective group ID just after an <code>exec</code> of a file whose set user or set group file mode bit has been set (see <code>exec(2)</code>).
Semaphore Identifier	A semaphore identifier (<code>semid</code>) is a unique positive integer created by a <code>semget(2)</code> call. Each <code>semid</code> has a set of semaphores and a data structure associated with it. The data structure is referred to as <code>semid_ds</code> and contains the following members: <pre> struct ipc_perm sem_perm; /* operation permission struct */ struct sem *sem_base; /* ptr to first semaphore in set */ ushort_t sem_nsems; /* number of sems in set */ time_t sem_otime; /* last operation time */ time_t sem_ctime; /* last change time */ /* Times measured in secs since */ /* 00:00:00 GMT, Jan. 1, 1970 */ </pre> The following are descriptions of the <code>semid_ds</code> structure members: The <code>sem_perm</code> member is an <code>ipc_perm</code> structure that specifies the semaphore operation permission (see below). This structure includes the following members: <pre> uid_t uid; /* user id */ gid_t gid; /* group id */ uid_t cuid; /* creator user id */ gid_t cgid; /* creator group id */ mode_t mode; /* r/a permission */ </pre>

Intro(2)

```
ulong_t  seq;    /* slot usage sequence number */
key_t    key;    /* key */
```

The `sem_nsems` member is equal to the number of semaphores in the set. Each semaphore in the set is referenced by a nonnegative integer referred to as a `sem_num`. `sem_num` values run sequentially from 0 to the value of `sem_nsems` minus 1.

The `sem_otime` member is the time of the last `semop(2)` operation.

The `sem_ctime` member is the time of the last `semctl(2)` operation that changed a member of the above structure.

A semaphore is a data structure called `sem` that contains the following members:

```
ushort_t  semval; /* semaphore value */
pid_t     sempid; /* pid of last operation */
ushort_t  semncnt; /* # awaiting semval > cval */
ushort_t  semzcnt; /* # awaiting semval = 0 */
```

The following are descriptions of the `sem` structure members:

The `semval` member is a non-negative integer that is the actual value of the semaphore.

The `sempid` member is equal to the process ID of the last process that performed a semaphore operation on this semaphore.

The `semncnt` member is a count of the number of processes that are currently suspended awaiting this semaphore's `semval` to become greater than its current value.

The `semzcnt` member is a count of the number of processes that are currently suspended awaiting this semaphore's `semval` to become 0.

Semaphore Operation Permissions

In the `semop(2)` and `semctl(2)` function descriptions, the permission required for an operation is given as `{token}`, where `token` is the type of permission needed interpreted as follows:

```
00400    READ by user
00200    ALTER by user
00040    READ by group
00020    ALTER by group
00004    READ by others
00002    ALTER by others
```

Read and alter permissions for a `semid` are granted to a process if one or more of the following are true:

- The effective user ID of the process is superuser.

- The effective user ID of the process matches `sem_perm.cuid` or `sem_perm.uid` in the data structure associated with `semid` and the appropriate bit of the “user” portion (0600) of `sem_perm.mode` is set.
- The effective group ID of the process matches `sem_perm.cgid` or `sem_perm.gid` and the appropriate bit of the “group” portion (060) of `sem_perm.mode` is set.
- The appropriate bit of the “other” portion (06) of `sem_perm.mode` is set.

Otherwise, the corresponding permissions are denied.

Session	A session is a group of processes identified by a common ID called a session ID, capable of establishing a connection with a controlling terminal. Any process that is not a process group leader may create a new session and process group, becoming the session leader of the session and process group leader of the process group. A newly created process joins the session of its creator.
Session ID	Each session in the system is uniquely identified during its lifetime by a positive integer called a session ID, the process ID of its session leader.
Session Leader	A session leader is a process whose session ID is the same as its process and process group ID.
Session Lifetime	A session lifetime begins when the session is created by its session leader, and ends when the lifetime of the last process that is a member of the session ends, or when the last process that is a member in the session leaves the session.
Shared Memory Identifier	A shared memory identifier (<code>shmid</code>) is a unique positive integer created by a <code>shmget(2)</code> call. Each <code>shmid</code> has a segment of memory (referred to as a shared memory segment) and a data structure associated with it. (Note that these shared memory segments must be explicitly removed by the user after the last reference to them is removed.) The data structure is referred to as <code>shmid_ds</code> and contains the following members: <pre> struct ipc_perm shm_perm; /* operation permission struct */ size_t shm_segsz; /* size of segment */ struct anon_map *shm_amp; /* ptr to region structure */ char pad[4]; /* for swap compatibility */ pid_t shm_lpid; /* pid of last operation */ pid_t shm_cpid; /* creator pid */ shmatt_t shm_nattch; /* number of current attaches */ ulong_t shm_cnattch; /* used only for shminfo */ time_t shm_atime; /* last attach time */ time_t shm_dtime; /* last detach time */ time_t shm_ctime; /* last change time */ /* Times measured in secs since */ /* 00:00:00 GMT, Jan. 1, 1970 */ </pre> The following are descriptions of the <code>shmid_ds</code> structure members: The <code>shm_perm</code> member is an <code>ipc_perm</code> structure that specifies the shared memory operation permission (see below). This structure includes the following members:

Intro(2)

```
uid_t    cuid;    /* creator user id */
gid_t    cgid;    /* creator group id */
uid_t    uid;     /* user id */
gid_t    gid;     /* group id */
mode_t    mode;   /* r/w permission */
ulong_t  seq;     /* slot usage sequence # */
key_t    key;     /* key */
```

The `shm_segsz` member specifies the size of the shared memory segment in bytes.

The `shm_cpid` member is the process ID of the process that created the shared memory identifier.

The `shm_lpid` member is the process ID of the last process that performed a `shmat()` or `shmdt()` operation (see `shmop(2)`).

The `shm_nattch` member is the number of processes that currently have this segment attached.

The `shm_atime` member is the time of the last `shmat()` operation (see `shmop(2)`).

The `shm_dtime` member is the time of the last `shmdt()` operation (see `shmop(2)`).

The `shm_ctime` member is the time of the last `shmctl(2)` operation that changed one of the members of the above structure.

Shared Memory Operation Permissions

In the `shmctl(2)`, `shmat()`, and `shmdt()` (see `shmop(2)`) function descriptions, the permission required for an operation is given as `{token}`, where `token` is the type of permission needed interpreted as follows:

```
00400    READ by user
00200    WRITE by user
00040    READ by group
00020    WRITE by group
00004    READ by others
00002    WRITE by others
```

Read and write permissions for a `shmid` are granted to a process if one or more of the following are true:

- The effective user ID of the process is superuser.
- The effective user ID of the process matches `shm_perm.cuid` or `shm_perm.uid` in the data structure associated with `shmid` and the appropriate bit of the “user” portion (0600) of `shm_perm.mode` is set.
- The effective group ID of the process matches `shm_perm.cgid` or `shm_perm.gid` and the appropriate bit of the “group” portion (060) of `shm_perm.mode` is set.
- The appropriate bit of the “other” portion (06) of `shm_perm.mode` is set.

Otherwise, the corresponding permissions are denied.

Special Processes	The process with ID 0 and the process with ID 1 are special processes referred to as <code>proc0</code> and <code>proc1</code> ; see <code>kill(2)</code> . <code>proc0</code> is the process scheduler. <code>proc1</code> is the initialization process (<i>init</i>); <code>proc1</code> is the ancestor of every other process in the system and is used to control the process structure.
STREAMS	A set of kernel mechanisms that support the development of network services and data communication drivers. It defines interface standards for character input/output within the kernel and between the kernel and user level processes. The STREAMS mechanism is composed of utility routines, kernel facilities and a set of data structures.
Stream	A stream is a full-duplex data path within the kernel between a user process and driver routines. The primary components are a stream head, a driver, and zero or more modules between the stream head and driver. A stream is analogous to a shell pipeline, except that data flow and processing are bidirectional.
Stream Head	In a stream, the stream head is the end of the stream that provides the interface between the stream and a user process. The principal functions of the stream head are processing STREAMS-related system calls and passing data and information between a user process and the stream.
Superuser	A process is recognized as a superuser process and is granted special privileges, such as immunity from file permissions, if its effective user ID is 0.
Upstream	In a stream, the direction from driver to stream head.
Write Queue	In a stream, the message queue in a module or driver containing messages moving downstream.

Intro(2)

System Calls

access(2)

NAME	access – determine accessibility of a file												
SYNOPSIS	<pre>#include <unistd.h> int access(const char *path, int amode);</pre>												
DESCRIPTION	The <code>access()</code> function checks the file named by the pathname pointed to by the <i>path</i> argument for accessibility according to the bit pattern contained in <i>amode</i>, using the real user ID in place of the effective user ID and the real group ID in place of the effective group ID. This allows a setuid process to verify that the user running it would have had permission to access this file. The value of <i>amode</i> is either the bitwise inclusive OR of the access permissions to be checked (<code>R_OK</code>, <code>W_OK</code>, <code>X_OK</code>) or the existence test, <code>F_OK</code>. These constants are defined in <code><unistd.h></code> as follows: <table><tr><td><code>R_OK</code></td><td>Test for read permission.</td></tr><tr><td><code>W_OK</code></td><td>Test for write permission.</td></tr><tr><td><code>X_OK</code></td><td>Test for execute or search permission.</td></tr><tr><td><code>F_OK</code></td><td>Check existence of file</td></tr></table> See <code>intro(2)</code> for additional information about "File Access Permission". If any access permissions are to be checked, each will be checked individually, as described in <code>intro(2)</code>. If the process has appropriate privileges, an implementation may indicate success for <code>X_OK</code> even if none of the execute file permission bits are set.	<code>R_OK</code>	Test for read permission.	<code>W_OK</code>	Test for write permission.	<code>X_OK</code>	Test for execute or search permission.	<code>F_OK</code>	Check existence of file				
<code>R_OK</code>	Test for read permission.												
<code>W_OK</code>	Test for write permission.												
<code>X_OK</code>	Test for execute or search permission.												
<code>F_OK</code>	Check existence of file												
RETURN VALUES	If the requested access is permitted, <code>access()</code> succeeds and returns 0. Otherwise, <code>-1</code> is returned and <code>errno</code> is set to indicate the error.												
ERRORS	The <code>access()</code> function will fail if: <table><tr><td><code>EACCES</code></td><td>Permission bits of the file mode do not permit the requested access, or search permission is denied on a component of the path prefix.</td></tr><tr><td><code>EFAULT</code></td><td><i>path</i> points to an illegal address.</td></tr><tr><td><code>EINTR</code></td><td>A signal was caught during the <code>access()</code> function.</td></tr><tr><td><code>ELOOP</code></td><td>Too many symbolic links were encountered in resolving <i>path</i>.</td></tr><tr><td><code>ENAMETOOLONG</code></td><td>The length of the <i>path</i> argument exceeds <code>PATH_MAX</code>, or a pathname component is longer than <code>NAME_MAX</code> while <code>_POSIX_NO_TRUNC</code> is in effect.</td></tr><tr><td><code>ENOENT</code></td><td>A component of <i>path</i> does not name an existing file or <i>path</i> is an empty string.</td></tr></table>	<code>EACCES</code>	Permission bits of the file mode do not permit the requested access, or search permission is denied on a component of the path prefix.	<code>EFAULT</code>	<i>path</i> points to an illegal address.	<code>EINTR</code>	A signal was caught during the <code>access()</code> function.	<code>ELOOP</code>	Too many symbolic links were encountered in resolving <i>path</i> .	<code>ENAMETOOLONG</code>	The length of the <i>path</i> argument exceeds <code>PATH_MAX</code> , or a pathname component is longer than <code>NAME_MAX</code> while <code>_POSIX_NO_TRUNC</code> is in effect.	<code>ENOENT</code>	A component of <i>path</i> does not name an existing file or <i>path</i> is an empty string.
<code>EACCES</code>	Permission bits of the file mode do not permit the requested access, or search permission is denied on a component of the path prefix.												
<code>EFAULT</code>	<i>path</i> points to an illegal address.												
<code>EINTR</code>	A signal was caught during the <code>access()</code> function.												
<code>ELOOP</code>	Too many symbolic links were encountered in resolving <i>path</i> .												
<code>ENAMETOOLONG</code>	The length of the <i>path</i> argument exceeds <code>PATH_MAX</code> , or a pathname component is longer than <code>NAME_MAX</code> while <code>_POSIX_NO_TRUNC</code> is in effect.												
<code>ENOENT</code>	A component of <i>path</i> does not name an existing file or <i>path</i> is an empty string.												

access(2)

ENOLINK *path* points to a remote machine and the link to that machine is no longer active.

ENOTDIR A component of the path prefix is not a directory.

EROFS Write access is requested for a file on a read-only file system.

The `access()` function may fail if:

EINVAL The value of the *amode* argument is invalid.

ENAMETOOLONG Pathname resolution of a symbolic link produced an intermediate result whose length exceeds `PATH_MAX`.

ETXTBSY Write access is requested for a pure procedure (shared text) file that is being executed.

USAGE Additional values of *amode* other than the set defined in the description may be valid, for example, if a system has extended access controls.

ATTRIBUTES See `attributes(5)` for descriptions of the following attributes:

ATTRIBUTE TYPE	ATTRIBUTE VALUE
MT-Level	Async-Signal-Safe

SEE ALSO `intro(2)`, `chmod(2)`, `stat(2)`, `attributes(5)`

acct(2)

NAME	acct – enable or disable process accounting																		
SYNOPSIS	<pre>#include <unistd.h> int acct(const char *path);</pre>																		
DESCRIPTION	The <code>acct()</code> function enables or disables the system process accounting routine. If the routine is enabled, an accounting record will be written in an accounting file for each process that terminates. The termination of a process can be caused by either an <code>exit(2)</code> call or a <code>signal(3C)</code>. The effective user ID of the process calling <code>acct()</code> must be super-user. The <i>path</i> argument points to the pathname of the accounting file, whose file format is described on the <code>acct(3HEAD)</code> manual page. The accounting routine is enabled if <i>path</i> is non-zero and no errors occur during the function. It is disabled if <i>path</i> is <code>(char *) NULL</code> and no errors occur during the function.																		
RETURN VALUES	Upon successful completion, 0 is returned. Otherwise, -1 is returned and <code>errno</code> is set to indicate the error.																		
ERRORS	The <code>acct()</code> function will fail if: <table><tr><td>EACCES</td><td>The file named by <i>path</i> is not an ordinary file.</td></tr><tr><td>EBUSY</td><td>An attempt is being made to enable accounting using the same file that is currently being used.</td></tr><tr><td>EFAULT</td><td>The <i>path</i> argument points to an illegal address.</td></tr><tr><td>ELOOP</td><td>Too many symbolic links were encountered in translating <i>path</i>.</td></tr><tr><td>ENAMETOOLONG</td><td>The length of the <i>path</i> argument exceeds <code>PATH_MAX</code>, or the length of a <i>path</i> argument exceeds <code>NAME_MAX</code> while <code>_POSIX_NO_TRUNC</code> is in effect.</td></tr><tr><td>ENOENT</td><td>One or more components of the accounting file pathname do not exist.</td></tr><tr><td>ENOTDIR</td><td>A component of the path prefix is not a directory.</td></tr><tr><td>EPERM</td><td>The effective user of the calling process is not super-user.</td></tr><tr><td>EROFS</td><td>The named file resides on a read-only file system.</td></tr></table>	EACCES	The file named by <i>path</i> is not an ordinary file.	EBUSY	An attempt is being made to enable accounting using the same file that is currently being used.	EFAULT	The <i>path</i> argument points to an illegal address.	ELOOP	Too many symbolic links were encountered in translating <i>path</i> .	ENAMETOOLONG	The length of the <i>path</i> argument exceeds <code>PATH_MAX</code> , or the length of a <i>path</i> argument exceeds <code>NAME_MAX</code> while <code>_POSIX_NO_TRUNC</code> is in effect.	ENOENT	One or more components of the accounting file pathname do not exist.	ENOTDIR	A component of the path prefix is not a directory.	EPERM	The effective user of the calling process is not super-user.	EROFS	The named file resides on a read-only file system.
EACCES	The file named by <i>path</i> is not an ordinary file.																		
EBUSY	An attempt is being made to enable accounting using the same file that is currently being used.																		
EFAULT	The <i>path</i> argument points to an illegal address.																		
ELOOP	Too many symbolic links were encountered in translating <i>path</i> .																		
ENAMETOOLONG	The length of the <i>path</i> argument exceeds <code>PATH_MAX</code> , or the length of a <i>path</i> argument exceeds <code>NAME_MAX</code> while <code>_POSIX_NO_TRUNC</code> is in effect.																		
ENOENT	One or more components of the accounting file pathname do not exist.																		
ENOTDIR	A component of the path prefix is not a directory.																		
EPERM	The effective user of the calling process is not super-user.																		
EROFS	The named file resides on a read-only file system.																		
SEE ALSO	<code>exit(2)</code> , <code>signal(3C)</code> , <code>acct(3HEAD)</code>																		

NAME	acl, facl – get or set a file's Access Control List (ACL)														
SYNOPSIS	<pre>#include <sys/acl.h> int acl(char *pathp, int cmd, int nentries, aclent_t *aclbufp); int facl(int fildes, int cmd, int nentries, aclent_t *aclbufp);</pre>														
DESCRIPTION	The <code>acl()</code> and <code>facl()</code> functions get or set the ACL of a file whose name is given by <i>pathp</i> or referenced by the open file descriptor <i>fildes</i>. The <i>nentries</i> argument specifies how many ACL entries fit into buffer <i>aclbufp</i>. The <code>acl()</code> function is used to manipulate ACL on file system objects. The following values for <i>cmd</i> are supported: <table> <tr> <td>SETACL</td><td><i>nentries</i> ACL entries, specified in buffer <i>aclbufp</i>, are stored in the file's ACL. This command can only be executed by a process that has an effective user ID equal to the owner of the file. All directories in the path name must be searchable.</td></tr> <tr> <td>GETACL</td><td>Buffer <i>aclbufp</i> is filled with the file's ACL entries. Read access to the file is not required, but all directories in the path name must be searchable.</td></tr> <tr> <td>GETACLNT</td><td>The number of entries in the file's ACL is returned. Read access to the file is not required, but all directories in the path name must be searchable.</td></tr> </table>	SETACL	<i>nentries</i> ACL entries, specified in buffer <i>aclbufp</i> , are stored in the file's ACL. This command can only be executed by a process that has an effective user ID equal to the owner of the file. All directories in the path name must be searchable.	GETACL	Buffer <i>aclbufp</i> is filled with the file's ACL entries. Read access to the file is not required, but all directories in the path name must be searchable.	GETACLNT	The number of entries in the file's ACL is returned. Read access to the file is not required, but all directories in the path name must be searchable.								
SETACL	<i>nentries</i> ACL entries, specified in buffer <i>aclbufp</i> , are stored in the file's ACL. This command can only be executed by a process that has an effective user ID equal to the owner of the file. All directories in the path name must be searchable.														
GETACL	Buffer <i>aclbufp</i> is filled with the file's ACL entries. Read access to the file is not required, but all directories in the path name must be searchable.														
GETACLNT	The number of entries in the file's ACL is returned. Read access to the file is not required, but all directories in the path name must be searchable.														
RETURN VALUES	Upon successful completion, <code>acl()</code> and <code>facl()</code> return 0 if <i>cmd</i> is SETACL. If <i>cmd</i> is GETACL or GETACLNT, the number of ACL entries is returned. Otherwise, -1 is returned and <code>errno</code> is set to indicate the error.														
ERRORS	The <code>acl()</code> function will fail if: <table> <tr> <td>EACCESS</td><td>The caller does not have access to a component of the pathname.</td></tr> <tr> <td>EFAULT</td><td>The <i>pathp</i> or <i>aclbufp</i> argument points to an illegal address.</td></tr> <tr> <td>EINVAL</td><td>The <i>cmd</i> argument is not GETACL, SETACL, or GETACLNT; the <i>cmd</i> argument is SETACL and <i>nentries</i> is less than 3; or the <i>cmd</i> argument is SETACL and the ACL specified in <i>aclbufp</i> is not valid.</td></tr> <tr> <td>EIO</td><td>A disk I/O error has occurred while storing or retrieving the ACL.</td></tr> <tr> <td>ENOENT</td><td>A component of the path does not exist.</td></tr> <tr> <td>ENOSPC</td><td>The <i>cmd</i> argument is GETACL and <i>nentries</i> is less than the number of entries in the file's ACL, or the <i>cmd</i> argument is SETACL and there is insufficient space in the file system to store the ACL.</td></tr> <tr> <td>ENOTDIR</td><td>A component of the path specified by <i>pathp</i> is not a directory, or the <i>cmd</i> argument is SETACL and an attempt is made to set a default ACL on a file type other than a directory.</td></tr> </table>	EACCESS	The caller does not have access to a component of the pathname.	EFAULT	The <i>pathp</i> or <i>aclbufp</i> argument points to an illegal address.	EINVAL	The <i>cmd</i> argument is not GETACL, SETACL, or GETACLNT; the <i>cmd</i> argument is SETACL and <i>nentries</i> is less than 3; or the <i>cmd</i> argument is SETACL and the ACL specified in <i>aclbufp</i> is not valid.	EIO	A disk I/O error has occurred while storing or retrieving the ACL.	ENOENT	A component of the path does not exist.	ENOSPC	The <i>cmd</i> argument is GETACL and <i>nentries</i> is less than the number of entries in the file's ACL, or the <i>cmd</i> argument is SETACL and there is insufficient space in the file system to store the ACL.	ENOTDIR	A component of the path specified by <i>pathp</i> is not a directory, or the <i>cmd</i> argument is SETACL and an attempt is made to set a default ACL on a file type other than a directory.
EACCESS	The caller does not have access to a component of the pathname.														
EFAULT	The <i>pathp</i> or <i>aclbufp</i> argument points to an illegal address.														
EINVAL	The <i>cmd</i> argument is not GETACL, SETACL, or GETACLNT; the <i>cmd</i> argument is SETACL and <i>nentries</i> is less than 3; or the <i>cmd</i> argument is SETACL and the ACL specified in <i>aclbufp</i> is not valid.														
EIO	A disk I/O error has occurred while storing or retrieving the ACL.														
ENOENT	A component of the path does not exist.														
ENOSPC	The <i>cmd</i> argument is GETACL and <i>nentries</i> is less than the number of entries in the file's ACL, or the <i>cmd</i> argument is SETACL and there is insufficient space in the file system to store the ACL.														
ENOTDIR	A component of the path specified by <i>pathp</i> is not a directory, or the <i>cmd</i> argument is SETACL and an attempt is made to set a default ACL on a file type other than a directory.														

acl(2)

ENOSYS	The <i>cmd</i> argument is SETACL and the file specified by <i>pathp</i> resides on a file system that does not support ACLs, or the <code>acl()</code> function is not supported by this implementation.
EPERM	The <i>cmd</i> argument is SETACL and the effective user ID of the caller does not match the owner of the file.
EROFS	The <i>cmd</i> argument is SETACL and the file specified by <i>pathp</i> resides on a file system that is mounted read-only.

ATTRIBUTES See `attributes(5)` for descriptions of the following attributes:

ATTRIBUTE TYPE	ATTRIBUTE VALUE
Interface Stability	Evolving

SEE ALSO `getfacl(1)`, `setfacl(1)`, `aclcheck(3SEC)`, `aclsort(3SEC)`

NAME	adjtime – correct the time to allow synchronization of the system clock								
SYNOPSIS	<pre>#include <sys/time.h> int adjtime(struct timeval *delta, struct timeval *olddelta);</pre>								
DESCRIPTION	The <code>adjtime()</code> function adjusts the system's notion of the current time as returned by <code>gettimeofday(3C)</code>, advancing or retarding it by the amount of time specified in the <code>struct timeval</code> pointed to by <i>delta</i>. The adjustment is effected by speeding up (if that amount of time is positive) or slowing down (if that amount of time is negative) the system's clock by some small percentage, generally a fraction of one percent. The time is always a monotonically increasing function. A time correction from an earlier call to <code>adjtime()</code> may not be finished when <code>adjtime()</code> is called again. If <i>delta</i> is 0, then <i>olddelta</i> returns the status of the effects of the previous <code>adjtime()</code> call with no effect on the time correction as a result of this call. If <i>olddelta</i> is not a null pointer, then the structure it points to will contain, upon successful return, the number of seconds and/or microseconds still to be corrected from the earlier call. If <i>olddelta</i> is a null pointer, the corresponding information will not be returned. This call may be used in time servers that synchronize the clocks of computers in a local area network. Such time servers would slow down the clocks of some machines and speed up the clocks of others to bring them to the average network time. Only the super-user may adjust the time of day. The adjustment value will be silently rounded to the resolution of the system clock.								
RETURN VALUES	Upon successful completion, <code>adjtime()</code> returns 0. Otherwise, it returns -1 and sets <code>errno</code> to indicate the error.								
ERRORS	The <code>adjtime()</code> function will fail if: <table> <tr> <td>EFAULT</td><td>The <i>delta</i> or <i>olddelta</i> argument points outside the process's allocated address space, or <i>olddelta</i> points to a region of the process's allocated address space that is not writable.</td></tr> <tr> <td>EINVAL</td><td>The <code>tv_usec</code> member of <i>delta</i> is not within valid range (-1000000 to 1000000).</td></tr> <tr> <td>EPERM</td><td>The effective user of the calling process is not super-user.</td></tr> </table> Additionally, the <code>adjtime()</code> function will fail for 32-bit interfaces if: <table> <tr> <td>EOVERFLOW</td><td>The size of the <code>tv_sec</code> member of the <code>timeval</code> structure pointed to by <i>olddelta</i> is too small to contain the correct number of seconds.</td></tr> </table>	EFAULT	The <i>delta</i> or <i>olddelta</i> argument points outside the process's allocated address space, or <i>olddelta</i> points to a region of the process's allocated address space that is not writable.	EINVAL	The <code>tv_usec</code> member of <i>delta</i> is not within valid range (-1000000 to 1000000).	EPERM	The effective user of the calling process is not super-user.	EOVERFLOW	The size of the <code>tv_sec</code> member of the <code>timeval</code> structure pointed to by <i>olddelta</i> is too small to contain the correct number of seconds.
EFAULT	The <i>delta</i> or <i>olddelta</i> argument points outside the process's allocated address space, or <i>olddelta</i> points to a region of the process's allocated address space that is not writable.								
EINVAL	The <code>tv_usec</code> member of <i>delta</i> is not within valid range (-1000000 to 1000000).								
EPERM	The effective user of the calling process is not super-user.								
EOVERFLOW	The size of the <code>tv_sec</code> member of the <code>timeval</code> structure pointed to by <i>olddelta</i> is too small to contain the correct number of seconds.								
SEE ALSO	<code>date(1)</code> , <code>gettimeofday(3C)</code>								

alarm(2)

NAME	alarm – set a process alarm clock
SYNOPSIS	<pre>#include <unistd.h> unsigned int alarm(unsigned int sec);</pre>
DESCRIPTION	The <code>alarm()</code> function instructs the alarm clock of the calling process to send the signal <code>SIGALRM</code> to the calling process after the number of real time seconds specified by <code>sec</code> have elapsed (see <code>signal(3C)</code>). Alarm requests are not stacked; successive calls reset the alarm clock of the calling process. If <code>sec</code> is 0, any previously made alarm request is canceled. The <code>fork(2)</code> function sets the alarm clock of a new process to 0. A process created by the <code>exec</code> family of routines inherits the time left on the old process's alarm clock. Calling <code>alarm()</code> in a multithreaded process linked with <code>-lthread</code> (Solaris threads) and not with <code>-lpthread</code> (POSIX threads) currently behaves in the following fashion: ■ if the calling thread is a bound thread, the resulting <code>SIGALRM</code> is delivered to the bound thread's LWP, i.e. to the calling thread. There is a bug currently that this signal is not maskable via <code>thr_sigsetmask(3THR)</code> on this bound thread. ■ if the calling thread is an unbound thread, the resulting <code>SIGALRM</code> is sent to the LWP on which the thread was running when it issued the call to <code>alarm()</code>. This is neither a per-process semantic, nor a per-thread semantic, since the LWP could change threads after the call to <code>alarm()</code> but before the <code>SIGALRM</code> delivery, causing some other thread to get it possibly. Hence this is basically a bug. The above documents current behavior and the bugs are not going to be fixed since the above semantics are going to be discontinued in the next release. The semantic for Solaris threads will move to the per-process semantic specified by POSIX (see <code>standards(5)</code>) at this future date. New applications should not rely on the per-thread semantic of <code>alarm()</code>, since this semantic will become obsolete. In a process linked with <code>-lpthread</code> (whether or not it is also linked with <code>-lthread</code>), the semantics of <code>alarm()</code> are per-process; the resulting <code>SIGALRM</code> is sent to the process, and not necessarily to the calling thread. This semantic will be supported in the future. This semantic is obtainable by simply linking with <code>-lpthread</code>. One can continue to use Solaris thread interfaces by linking with both <code>-lpthread</code> and <code>-lthread</code>.
RETURN VALUES	The <code>alarm()</code> function returns the amount of time previously remaining in the alarm clock of the calling process.

alarm(2)

ATTRIBUTES See attributes(5) for descriptions of the following attributes:

ATTRIBUTE TYPE	ATTRIBUTE VALUE
MT-Level	Async-Signal-Safe

SEE ALSO exec(2), fork(2), pause(2), signal(3C), thr_sigsetmask(3THR), attributes(5), standards(5)

audit(2)

NAME	audit – write a record to the audit log						
SYNOPSIS	<pre>cc [<i>flag</i> ...] <i>file</i> ... -lbsm -lsocket -lnsl -lintl [<i>library</i>...] #include <sys/param.h> #include <bsm/audit.h> int audit(caddr_t <i>record</i>, int <i>length</i>);</pre>						
DESCRIPTION	The <code>audit()</code> function is used to write a record to the system audit log. The data pointed to by <i>record</i> is written to the log after a minimal consistency check, with the <i>length</i> parameter specifying the size of the record in bytes. The data should be a well-formed audit record as described by <code>audit.log(4)</code>. The kernel validates the record header token type and length, and sets the time stamp value before writing the record to the audit log. The kernel does not do any preselection for user-level generated events. If the audit policy is set to include sequence or trailer tokens, the kernel will append them to the record.						
RETURN VALUES	Upon successful completion, 0 is returned. Otherwise, -1 is returned and <code>errno</code> is set to indicate the error.						
ERRORS	The <code>audit()</code> function will fail if: <table><tr><td>EFAULT</td><td>The <i>record</i> argument points outside the process's allocated address space.</td></tr><tr><td>EINVAL</td><td>The record header token ID is invalid or the length is either less than the header token size or greater than <code>MAXAUDITDATA</code>.</td></tr><tr><td>EPERM</td><td>The process's effective user ID is not super-user.</td></tr></table>	EFAULT	The <i>record</i> argument points outside the process's allocated address space.	EINVAL	The record header token ID is invalid or the length is either less than the header token size or greater than <code>MAXAUDITDATA</code> .	EPERM	The process's effective user ID is not super-user.
EFAULT	The <i>record</i> argument points outside the process's allocated address space.						
EINVAL	The record header token ID is invalid or the length is either less than the header token size or greater than <code>MAXAUDITDATA</code> .						
EPERM	The process's effective user ID is not super-user.						
USAGE	Only the super-user may successfully execute this call.						
SEE ALSO	<code>bsmconv(1M)</code> , <code>auditd(1M)</code> , <code>auditon(2)</code> , <code>auditsvc(2)</code> , <code>getaudit(2)</code> , <code>audit.log(4)</code>						
NOTES	The functionality described in this man page is available only if the Basic Security Module (BSM) has been enabled. See <code>bsmconv(1M)</code> for more information.						

NAME	auditon – manipulate auditing																								
SYNOPSIS	<pre>cc [<i>flag</i> ...] <i>file</i> ... -lbsm -lsocket -lnsl -lintl [<i>library</i> ...] #include <sys/param.h> #include <bsm/audit.h> int auditon(int <i>cmd</i>, caddr_t <i>data</i>, int <i>length</i>);</pre>																								
DESCRIPTION	The <code>auditon()</code> function performs various audit subsystem control operations. The <i>cmd</i> argument designates the particular audit control command. The <i>data</i> argument is a pointer to command-specific data. The <i>length</i> argument is the length in bytes of the command-specific data. The following commands are supported: <table> <tr> <td>A_GETCOND</td><td>Return the system audit on/off/disabled condition in the integer long pointed to by <i>data</i>. The following values may be returned:</td></tr> <tr> <td>AUC_AUDITING</td><td>Auditing has been turned on.</td></tr> <tr> <td>AUC_NOAUDIT</td><td>Auditing has been turned off.</td></tr> <tr> <td>AUC_DISABLED</td><td>Auditing package installed, not turned on.</td></tr> <tr> <td>A_SETCOND</td><td>Set the system's audit on/off condition to the value in the integer long pointed to by <i>data</i>. The BSM audit module must be enabled by <code>bsmconv(1M)</code> before auditing can be turned on. The following audit states may be set:</td></tr> <tr> <td>AUC_AUDITING</td><td>Turns on audit record generation.</td></tr> <tr> <td>AUC_NOAUDIT</td><td>Turns off audit record generation.</td></tr> <tr> <td>A_GETCLASS</td><td>Return the event to class mapping for the designated audit event. The <i>data</i> argument points to the <code>au_evclass_map</code> structure containing the event number. The preselection class mask is returned in the same structure.</td></tr> <tr> <td>A_SETCLASS</td><td>Set the event class preselection mask for the designated audit event. The <i>data</i> argument points to the <code>au_evclass_map</code> structure containing the event number and class mask.</td></tr> <tr> <td>A_GETKMASK</td><td>Return the kernel preselection mask in the <code>au_mask</code> structure pointed to by <i>data</i>. This is the mask used to preselect non-attributable audit events.</td></tr> <tr> <td>A_SETKMASK</td><td>Set the kernel preselection mask. The <i>data</i> argument points to the <code>au_mask</code> structure containing the class mask. This is the mask used to preselect non-attributable audit events.</td></tr> <tr> <td>A_GETPINFO</td><td>Return the audit ID, preselection mask, terminal ID and audit session ID of the specified process in the <code>auditpinfo</code> structure pointed to by <i>data</i>.</td></tr> </table>	A_GETCOND	Return the system audit on/off/disabled condition in the integer long pointed to by <i>data</i> . The following values may be returned:	AUC_AUDITING	Auditing has been turned on.	AUC_NOAUDIT	Auditing has been turned off.	AUC_DISABLED	Auditing package installed, not turned on.	A_SETCOND	Set the system's audit on/off condition to the value in the integer long pointed to by <i>data</i> . The BSM audit module must be enabled by <code>bsmconv(1M)</code> before auditing can be turned on. The following audit states may be set:	AUC_AUDITING	Turns on audit record generation.	AUC_NOAUDIT	Turns off audit record generation.	A_GETCLASS	Return the event to class mapping for the designated audit event. The <i>data</i> argument points to the <code>au_evclass_map</code> structure containing the event number. The preselection class mask is returned in the same structure.	A_SETCLASS	Set the event class preselection mask for the designated audit event. The <i>data</i> argument points to the <code>au_evclass_map</code> structure containing the event number and class mask.	A_GETKMASK	Return the kernel preselection mask in the <code>au_mask</code> structure pointed to by <i>data</i> . This is the mask used to preselect non-attributable audit events.	A_SETKMASK	Set the kernel preselection mask. The <i>data</i> argument points to the <code>au_mask</code> structure containing the class mask. This is the mask used to preselect non-attributable audit events.	A_GETPINFO	Return the audit ID, preselection mask, terminal ID and audit session ID of the specified process in the <code>auditpinfo</code> structure pointed to by <i>data</i> .
A_GETCOND	Return the system audit on/off/disabled condition in the integer long pointed to by <i>data</i> . The following values may be returned:																								
AUC_AUDITING	Auditing has been turned on.																								
AUC_NOAUDIT	Auditing has been turned off.																								
AUC_DISABLED	Auditing package installed, not turned on.																								
A_SETCOND	Set the system's audit on/off condition to the value in the integer long pointed to by <i>data</i> . The BSM audit module must be enabled by <code>bsmconv(1M)</code> before auditing can be turned on. The following audit states may be set:																								
AUC_AUDITING	Turns on audit record generation.																								
AUC_NOAUDIT	Turns off audit record generation.																								
A_GETCLASS	Return the event to class mapping for the designated audit event. The <i>data</i> argument points to the <code>au_evclass_map</code> structure containing the event number. The preselection class mask is returned in the same structure.																								
A_SETCLASS	Set the event class preselection mask for the designated audit event. The <i>data</i> argument points to the <code>au_evclass_map</code> structure containing the event number and class mask.																								
A_GETKMASK	Return the kernel preselection mask in the <code>au_mask</code> structure pointed to by <i>data</i> . This is the mask used to preselect non-attributable audit events.																								
A_SETKMASK	Set the kernel preselection mask. The <i>data</i> argument points to the <code>au_mask</code> structure containing the class mask. This is the mask used to preselect non-attributable audit events.																								
A_GETPINFO	Return the audit ID, preselection mask, terminal ID and audit session ID of the specified process in the <code>auditpinfo</code> structure pointed to by <i>data</i> .																								

auditon(2)

	Note that A_GETPINFO may fail if the terminal ID contains a network address longer than 32 bits. In this case, the A_GETPINFO_ADDR command should be used.								
A_GETPINFO_ADDR	Returns the audit ID, preselection mask, terminal ID and audit session ID of the specified process in the auditpinfo_addr structure pointed to by <i>data</i> .								
A_SETPMASK	Set the preselection mask of the specified process. The <i>data</i> argument points to the auditpinfo structure containing the process ID and the preselection mask. The other fields of the structure are ignored and should be set to NULL.								
A_SETUMASK	Set the preselection mask for all processes with the specified audit ID. The <i>data</i> argument points to the auditinfo structure containing the audit ID and the preselection mask. The other fields of the structure are ignored and should be set to NULL.								
A_SETSMASK	Set the preselection mask for all processes with the specified audit session ID. The <i>data</i> argument points to the auditinfo structure containing the audit session ID and the preselection mask. The other fields of the structure are ignored and should be set to NULL.								
A_GETQCTRL	Return the kernel audit queue control parameters. These control the high and low water marks of the number of audit records allowed in the audit queue. The high water mark is the maximum allowed number of undelivered audit records. The low water mark determines when threads blocked on the queue are wakened. Another parameter controls the size of the data buffer used by auditsvc(2) to write data to the audit trail. There is also a parameter that specifies a maximum delay before data is attempted to be written to the audit trail. The audit queue parameters are returned in the au_qctrl structure pointed to by <i>data</i> .								
A_SETQCTRL	Set the kernel audit queue control parameters as described above in the A_GETQCTRL command. The <i>data</i> argument points to the au_qctrl structure containing the audit queue control parameters. The default and maximum values 'A/B' for the audit queue control parameters are: <table> <tr> <td>high water</td><td>100/10000 (audit records)</td></tr> <tr> <td>low water</td><td>10/1024 (audit records)</td></tr> <tr> <td>output buffer size</td><td>1024/1048576 (bytes)</td></tr> <tr> <td>delay</td><td>20/20000 (hundredths second)</td></tr> </table>	high water	100/10000 (audit records)	low water	10/1024 (audit records)	output buffer size	1024/1048576 (bytes)	delay	20/20000 (hundredths second)
high water	100/10000 (audit records)								
low water	10/1024 (audit records)								
output buffer size	1024/1048576 (bytes)								
delay	20/20000 (hundredths second)								
A_GETCWD	Return the current working directory as kept by the audit subsystem. This is a path anchored on the real root, rather than on								

	the active root. The <i>data</i> argument points to a buffer into which the path is copied. The <i>length</i> argument is the length of the buffer.
A_GETCAR	Return the current active root as kept by the audit subsystem. This path may be used to anchor an absolute path for a path token generated by an application. The <i>data</i> argument points to a buffer into which the path is copied. The <i>length</i> argument is the length of the buffer.
A_GETSTAT	Return the system audit statistics in the <code>audit_stat</code> structure pointed to by <i>data</i> .
A_SETSTAT	Reset system audit statistics values. The kernel statistics value is reset if the corresponding field in the statistics structure pointed to by the <i>data</i> argument is <code>CLEAR_VAL</code> . Otherwise, the value is not changed.
A_SETFSIZE	Set the maximum size of an audit trail file. When the audit file reaches the designated size, it is closed and a new file started. If the maximum size is unset, the audit trail file generated by <code>auditsvc()</code> will grow to the size of the file system. The <i>data</i> argument points to the <code>au_fstat_t</code> structure containing the maximum audit file size in bytes. The size can not be set less than <code>0x80000</code> bytes.
A_GETFSIZE	Return the maximum audit file size and current file size in the <code>au_fstat_t</code> structure pointed to by the <i>data</i> argument.
A_GETPOLICY	Return the audit policy flags in the integer long pointed to by <i>data</i> .
A_SETPOLICY	Set the audit policy flags to the values in the integer long pointed to by <i>data</i> . The following policy flags are recognized:
AUDIT_CNT	Do not suspend processes when audit storage is full or inaccessible. The default action is to suspend processes until storage becomes available.
AUDIT_AHLT	Halt the machine when a non-attributable audit record can not be delivered. The default action is to count the number of events that could not be recorded.
AUDIT_ARGV	Include in the audit record the argument list for a member of the <i>exec</i> family of functions (see <code>exec(2)</code>). The default action is not to include this information.
AUDIT_ARGE	Include the environment variables for the <code>execv(2)</code> function in the audit record. The default action is not to include this information.

auditon(2)

	AUDIT_SEQ	Add a <i>sequence</i> token to each audit record. The default action is not to include it.
	AUDIT_TRAIL	Append a <i>trailer</i> token to each audit record. The default action is not to include it.
	AUDIT_GROUP	Include the supplementary groups list in audit records. The default action is not to include it.
	AUDIT_PATH	Include secondary paths in audit records. Examples of secondary paths are dynamically loaded shared library modules and the command shell path for executable scripts. The default action is to include only the primary path from the system call.
RETURN VALUES	Upon successful completion, <code>auditon()</code> returns 0. Otherwise, -1 is returned and <code>errno</code> is set to indicate the error.	
ERRORS	The <code>auditon()</code> function will fail if:	
	E2BIG	The <i>length</i> field for the command was too small to hold the returned value.
	EFAULT	The copy of data to/from the kernel failed.
	EINVAL	One of the arguments was illegal, or BSM has not been installed.
	EPERM	The process's effective user ID is not super-user.
USAGE	The <code>auditon()</code> function may be invoked only by processes with super-user privileges.	
SEE ALSO	<code>auditconfig(1M)</code> , <code>auditd(1M)</code> , <code>bsmconv(1M)</code> , <code>audit(2)</code> , <code>auditsvc(2)</code> , <code>exec(2)</code> , <code>audit.log(4)</code>	
NOTES	The functionality described in this man page is available only if the Basic Security Module (BSM) has been enabled. See <code>bsmconv(1M)</code> for more information.	

NAME	auditsvc – write audit log to specified file descriptor														
SYNOPSIS	<pre>cc [<i>flag</i> ...] <i>file</i>... -lbsm -lsocket -lnsl -lintl [<i>library</i> ...] #include <sys/param.h> #include <bsm/audit.h> int auditsvc(int <i>fd</i>, int <i>limit</i>);</pre>														
DESCRIPTION	The <code>auditsvc()</code> function specifies the audit log file to the kernel. The kernel writes audit records to this file until an exceptional condition occurs and then the call returns. The <i>fd</i> argument is a file descriptor that identifies the audit file. Applications should open this file for writing before calling <code>auditsvc()</code>. The <i>limit</i> argument specifies the number of free blocks that must be available in the audit file system, and causes <code>auditsvc()</code> to return when the free disk space on the audit filesystem drops below this limit. Thus, the invoking program can take action to avoid running out of disk space. The <code>auditsvc()</code> function does not return until one of the following conditions occurs: ■ The process receives a signal that is not blocked or ignored. ■ An error is encountered writing to the audit log file. ■ The minimum free space (as specified by <i>limit</i>), has been reached.														
RETURN VALUES	The <code>auditsvc()</code> function returns only on an error.														
ERRORS	The <code>auditsvc()</code> function will fail if: <table> <tr> <td>EAGAIN</td><td>The descriptor referred to a <i>stream</i>, was marked for System V-style non-blocking I/O, and no data could be written immediately.</td></tr> <tr> <td>EBADF</td><td>The <i>fd</i> argument is not a valid descriptor open for writing.</td></tr> <tr> <td>EBUSY</td><td>A second process attempted to perform this call.</td></tr> <tr> <td>EFBIG</td><td>An attempt was made to write a file that exceeds the process's file size limit or the maximum file size.</td></tr> <tr> <td>EINTR</td><td>The call is forced to terminate prematurely due to the arrival of a signal whose <code>SV_INTERRUPT</code> bit in <code>sv_flags</code> is set (see <code>sigvec(3UCB)</code>). The <code>signal(3C)</code> function sets this bit for any signal it catches.</td></tr> <tr> <td>EINVAL</td><td>Auditing is disabled (see <code>auditon(2)</code>), or the <i>fd</i> argument does not refer to a file of an appropriate type (regular files are always appropriate.)</td></tr> <tr> <td>EIO</td><td>An I/O error occurred while reading from or writing to the file system.</td></tr> </table>	EAGAIN	The descriptor referred to a <i>stream</i> , was marked for System V-style non-blocking I/O, and no data could be written immediately.	EBADF	The <i>fd</i> argument is not a valid descriptor open for writing.	EBUSY	A second process attempted to perform this call.	EFBIG	An attempt was made to write a file that exceeds the process's file size limit or the maximum file size.	EINTR	The call is forced to terminate prematurely due to the arrival of a signal whose <code>SV_INTERRUPT</code> bit in <code>sv_flags</code> is set (see <code>sigvec(3UCB)</code>). The <code>signal(3C)</code> function sets this bit for any signal it catches.	EINVAL	Auditing is disabled (see <code>auditon(2)</code>), or the <i>fd</i> argument does not refer to a file of an appropriate type (regular files are always appropriate.)	EIO	An I/O error occurred while reading from or writing to the file system.
EAGAIN	The descriptor referred to a <i>stream</i> , was marked for System V-style non-blocking I/O, and no data could be written immediately.														
EBADF	The <i>fd</i> argument is not a valid descriptor open for writing.														
EBUSY	A second process attempted to perform this call.														
EFBIG	An attempt was made to write a file that exceeds the process's file size limit or the maximum file size.														
EINTR	The call is forced to terminate prematurely due to the arrival of a signal whose <code>SV_INTERRUPT</code> bit in <code>sv_flags</code> is set (see <code>sigvec(3UCB)</code>). The <code>signal(3C)</code> function sets this bit for any signal it catches.														
EINVAL	Auditing is disabled (see <code>auditon(2)</code>), or the <i>fd</i> argument does not refer to a file of an appropriate type (regular files are always appropriate.)														
EIO	An I/O error occurred while reading from or writing to the file system.														

auditsvc(2)

	ENOSPC	The user's quota of disk blocks on the file system containing the file has been exhausted; audit filesystem space is below the specified limit; or there is no free space remaining on the file system containing the file.
	ENXIO	A hangup occurred on the <i>stream</i> being written to.
	EPERM	The process's effective user ID is not super-user.
	EWOLDBLOCK	The file was marked for 4.2 BSD-style non-blocking I/O, and no data could be written immediately.
USAGE	Only processes with an effective user ID of super-user may execute this call successfully.	
SEE ALSO	auditd(1M), bsmconv(1M), audit(2), auditon(2), sigvec(3UCB), audit.log (4)	
NOTES	The functionality described in this man page is available only if the Basic Security Module (BSM) has been enabled. See bsmconv(1M) for more information.	

NAME	brk, sbrk – change the amount of space allocated for the calling process’s data segment				
SYNOPSIS	<pre>#include <unistd.h> int brk(void *endds); void *sbrk(intptr_t incr);</pre>				
DESCRIPTION	The <code>brk()</code> and <code>sbrk()</code> functions are used to change dynamically the amount of space allocated for the calling process’s data segment (see <code>exec(2)</code>). The change is made by resetting the process’s break value and allocating the appropriate amount of space. The break value is the address of the first location beyond the end of the data segment. The amount of allocated space increases as the break value increases. Newly allocated space is set to zero. If, however, the same memory space is reallocated to the same process its contents are undefined. When a program begins execution using <code>execve()</code> the break is set at the highest location defined by the program and data storage areas. The <code>getrlimit(2)</code> function may be used to determine the maximum permissible size of the <i>data</i> segment; it is not possible to set the break beyond the <code>rlim_max</code> value returned from a call to <code>getrlimit()</code>, that is to say, “end + <code>rlim.rlim_max</code>.” See <code>end(3C)</code>. The <code>brk()</code> function sets the break value to <i>endds</i> and changes the allocated space accordingly. The <code>sbrk()</code> function adds <i>incr</i> function bytes to the break value and changes the allocated space accordingly. The <i>incr</i> function can be negative, in which case the amount of allocated space is decreased.				
RETURN VALUES	Upon successful completion, <code>brk()</code> returns 0. Otherwise, it returns -1 and sets <code>errno</code> to indicate the error. Upon successful completion, <code>sbrk()</code> returns the prior break value. Otherwise, it returns <code>(void *)-1</code> and sets <code>errno</code> to indicate the error.				
ERRORS	The <code>brk()</code> and <code>sbrk()</code> functions will fail and no additional memory will be allocated if: <table> <tr> <td>ENOMEM</td><td>The data segment size limit as set by <code>setrlimit()</code> (see <code>getrlimit(2)</code>) would be exceeded; the maximum possible size of a data segment (compiled into the system) would be exceeded; insufficient space exists in the swap area to support the expansion; or the new break value would extend into an area of the address space defined by some previously established mapping (see <code>mmap(2)</code>).</td></tr> <tr> <td>EAGAIN</td><td>Total amount of system memory available for private pages is temporarily insufficient. This may occur even though the space requested was less than the maximum data segment size (see</td></tr> </table>	ENOMEM	The data segment size limit as set by <code>setrlimit()</code> (see <code>getrlimit(2)</code>) would be exceeded; the maximum possible size of a data segment (compiled into the system) would be exceeded; insufficient space exists in the swap area to support the expansion; or the new break value would extend into an area of the address space defined by some previously established mapping (see <code>mmap(2)</code>).	EAGAIN	Total amount of system memory available for private pages is temporarily insufficient. This may occur even though the space requested was less than the maximum data segment size (see
ENOMEM	The data segment size limit as set by <code>setrlimit()</code> (see <code>getrlimit(2)</code>) would be exceeded; the maximum possible size of a data segment (compiled into the system) would be exceeded; insufficient space exists in the swap area to support the expansion; or the new break value would extend into an area of the address space defined by some previously established mapping (see <code>mmap(2)</code>).				
EAGAIN	Total amount of system memory available for private pages is temporarily insufficient. This may occur even though the space requested was less than the maximum data segment size (see				

brk(2)

ulimit(2)).

USAGE The behavior of `brk()` and `sbrk()` is unspecified if an application also uses any other memory functions (such as `malloc(3C)`, `mmap(2)`, `free(3C)`). The `brk()` and `sbrk()` functions have been used in specialized cases where no other memory allocation function provided the same capability. The use of `mmap(2)` is now preferred because it can be used portably with all other memory allocation functions and with any function that uses other allocation functions.

It is unspecified whether the pointer returned by `sbrk()` is aligned suitably for any purpose.

ATTRIBUTES See `attributes(5)` for descriptions of the following attributes:

ATTRIBUTE TYPE	ATTRIBUTE VALUE
MT-Level	MT-Safe

SEE ALSO `exec(2)`, `getrlimit(2)`, `mmap(2)`, `shmop(2)`, `ulimit(2)`, `end(3C)`, `free(3C)`, `malloc(3C)`

NOTES The value of *incr* may be adjusted by the system before setting the new break value. Upon successful completion, the implementation guarantees a minimum of *incr* bytes will be added to the data segment if *incr* is a positive value. If *incr* is a negative value, a maximum of *incr* bytes will be removed from the data segment. This adjustment may not be necessary for all machine architectures.

The value of the arguments to both `brk()` and `sbrk()` are rounded up for alignment with eight-byte boundaries.

BUGS Setting the break may fail due to a temporary lack of swap space. It is not possible to distinguish this from a failure caused by exceeding the maximum size of the data segment without consulting `getrlimit()`.

NAME	chdir, fchdir – change working directory																								
SYNOPSIS	<pre>#include <unistd.h> int chdir(const char *path) ; int fchdir(int fildes) ;</pre>																								
DESCRIPTION	The <code>chdir()</code> and <code>fchdir()</code> functions cause a directory pointed to by <i>path</i> or <i>fildes</i> to become the current working directory. The starting point for path searches for path names not beginning with / (slash). The <i>path</i> argument points to the path name of a directory. The <i>fildes</i> argument is an open file descriptor of a directory. For a directory to become the current directory, a process must have execute (search) access to the directory.																								
RETURN VALUES	Upon successful completion, 0 is returned. Otherwise, -1 is returned, the current working directory is unchanged, and <code>errno</code> is set to indicate the error.																								
ERRORS	The <code>chdir()</code> function will fail if: <table> <tr> <td>EACCES</td><td>Search permission is denied for any component of the path name.</td></tr> <tr> <td>EFAULT</td><td>The <i>path</i> argument points to an illegal address.</td></tr> <tr> <td>EINTR</td><td>A signal was caught during the execution of the <code>chdir()</code> function.</td></tr> <tr> <td>EIO</td><td>An I/O error occurred while reading from or writing to the file system.</td></tr> <tr> <td>ELOOP</td><td>Too many symbolic links were encountered in translating <i>path</i>.</td></tr> <tr> <td>ENAMETOOLONG</td><td>The length of the <i>path</i> argument exceeds <code>PATH_MAX</code>, or the length of a <i>path</i> component exceeds <code>NAME_MAX</code> while <code>_POSIX_NO_TRUNC</code> is in effect.</td></tr> <tr> <td>ENOENT</td><td>Either a component of the path prefix or the directory named by <i>path</i> does not exist or is a null pathname.</td></tr> <tr> <td>ENOLINK</td><td>The <i>path</i> argument points to a remote machine and the link to that machine is no longer active.</td></tr> <tr> <td>ENOTDIR</td><td>A component of the path name is not a directory.</td></tr> </table> The <code>fchdir()</code> function will fail if: <table> <tr> <td>EACCES</td><td>Search permission is denied for <i>fildes</i>.</td></tr> <tr> <td>EBADF</td><td>The <i>fildes</i> argument is not an open file descriptor.</td></tr> <tr> <td>EINTR</td><td>A signal was caught during the execution of the <code>fchdir()</code> function.</td></tr> </table>	EACCES	Search permission is denied for any component of the path name.	EFAULT	The <i>path</i> argument points to an illegal address.	EINTR	A signal was caught during the execution of the <code>chdir()</code> function.	EIO	An I/O error occurred while reading from or writing to the file system.	ELOOP	Too many symbolic links were encountered in translating <i>path</i> .	ENAMETOOLONG	The length of the <i>path</i> argument exceeds <code>PATH_MAX</code> , or the length of a <i>path</i> component exceeds <code>NAME_MAX</code> while <code>_POSIX_NO_TRUNC</code> is in effect.	ENOENT	Either a component of the path prefix or the directory named by <i>path</i> does not exist or is a null pathname.	ENOLINK	The <i>path</i> argument points to a remote machine and the link to that machine is no longer active.	ENOTDIR	A component of the path name is not a directory.	EACCES	Search permission is denied for <i>fildes</i> .	EBADF	The <i>fildes</i> argument is not an open file descriptor.	EINTR	A signal was caught during the execution of the <code>fchdir()</code> function.
EACCES	Search permission is denied for any component of the path name.																								
EFAULT	The <i>path</i> argument points to an illegal address.																								
EINTR	A signal was caught during the execution of the <code>chdir()</code> function.																								
EIO	An I/O error occurred while reading from or writing to the file system.																								
ELOOP	Too many symbolic links were encountered in translating <i>path</i> .																								
ENAMETOOLONG	The length of the <i>path</i> argument exceeds <code>PATH_MAX</code> , or the length of a <i>path</i> component exceeds <code>NAME_MAX</code> while <code>_POSIX_NO_TRUNC</code> is in effect.																								
ENOENT	Either a component of the path prefix or the directory named by <i>path</i> does not exist or is a null pathname.																								
ENOLINK	The <i>path</i> argument points to a remote machine and the link to that machine is no longer active.																								
ENOTDIR	A component of the path name is not a directory.																								
EACCES	Search permission is denied for <i>fildes</i> .																								
EBADF	The <i>fildes</i> argument is not an open file descriptor.																								
EINTR	A signal was caught during the execution of the <code>fchdir()</code> function.																								

chdir(2)

EIO	An I/O error occurred while reading from or writing to the file system.
ENOLINK	The <i>filides</i> argument points to a remote machine and the link to that machine is no longer active.
ENOTDIR	The open file descriptor <i>filides</i> does not refer to a directory.

ATTRIBUTES See `attributes(5)` for descriptions of the following attributes:

ATTRIBUTE TYPE	ATTRIBUTE VALUE
MT-Level	<code>chdir ()</code> is Async-Signal-Safe

SEE ALSO `chroot(2)`, `attributes(5)`

NAME	chmod, fchmod – change access permission mode of file																																													
SYNOPSIS	<pre>#include <sys/types.h> #include <sys/stat.h> int chmod(const char *path, mode_t mode); int fchmod(int fildes, mode_t mode);</pre>																																													
DESCRIPTION	The <code>chmod()</code> and <code>fchmod()</code> functions set the access permission portion of the mode of the file whose name is given by <i>path</i> or referenced by the open file descriptor <i>fildes</i> to the bit pattern contained in <i>mode</i>. Access permission bits are interpreted as follows: <table><tr><td><code>S_ISUID</code></td><td>04000</td><td>Set user ID on execution.</td></tr><tr><td><code>S_ISGID</code></td><td>020#0</td><td>Set group ID on execution if # is 7, 5, 3, or 1. Enable mandatory file/record locking if # is 6, 4, 2, or 0.</td></tr><tr><td><code>S_ISVTX</code></td><td>01000</td><td>Save text image after execution.</td></tr><tr><td><code>S_IRWXU</code></td><td>00700</td><td>Read, write, execute by owner.</td></tr><tr><td><code>S_IRUSR</code></td><td>00400</td><td>Read by owner.</td></tr><tr><td><code>S_IWUSR</code></td><td>00200</td><td>Write by owner.</td></tr><tr><td><code>S_IXUSR</code></td><td>00100</td><td>Execute (search if a directory) by owner.</td></tr><tr><td><code>S_IRWXG</code></td><td>00070</td><td>Read, write, execute by group.</td></tr><tr><td><code>S_IRGRP</code></td><td>00040</td><td>Read by group.</td></tr><tr><td><code>S_IWGRP</code></td><td>00020</td><td>Write by group.</td></tr><tr><td><code>S_IXGRP</code></td><td>00010</td><td>Execute by group.</td></tr><tr><td><code>S_IRWXO</code></td><td>00007</td><td>Read, write, execute (search) by others.</td></tr><tr><td><code>S_IROTH</code></td><td>00004</td><td>Read by others.</td></tr><tr><td><code>S_IWOTH</code></td><td>00002</td><td>Write by others.</td></tr><tr><td><code>S_IXOTH</code></td><td>00001</td><td>Execute by others.</td></tr></table> Modes are constructed by the bitwise OR operation of the access permission bits. The effective user ID of the process must match the owner of the file or the process must have the appropriate privilege to change the mode of a file. If the process is not a privileged process and the file is not a directory, mode bit 01000 (save text image on execution) is cleared. If neither the process is privileged, nor the file’s group is a member of the process’s supplementary group list, and the effective group ID of the process does not match the group ID of the file, mode bit 02000 (set group ID on execution) is cleared.	<code>S_ISUID</code>	04000	Set user ID on execution.	<code>S_ISGID</code>	020#0	Set group ID on execution if # is 7, 5, 3, or 1. Enable mandatory file/record locking if # is 6, 4, 2, or 0.	<code>S_ISVTX</code>	01000	Save text image after execution.	<code>S_IRWXU</code>	00700	Read, write, execute by owner.	<code>S_IRUSR</code>	00400	Read by owner.	<code>S_IWUSR</code>	00200	Write by owner.	<code>S_IXUSR</code>	00100	Execute (search if a directory) by owner.	<code>S_IRWXG</code>	00070	Read, write, execute by group.	<code>S_IRGRP</code>	00040	Read by group.	<code>S_IWGRP</code>	00020	Write by group.	<code>S_IXGRP</code>	00010	Execute by group.	<code>S_IRWXO</code>	00007	Read, write, execute (search) by others.	<code>S_IROTH</code>	00004	Read by others.	<code>S_IWOTH</code>	00002	Write by others.	<code>S_IXOTH</code>	00001	Execute by others.
<code>S_ISUID</code>	04000	Set user ID on execution.																																												
<code>S_ISGID</code>	020#0	Set group ID on execution if # is 7, 5, 3, or 1. Enable mandatory file/record locking if # is 6, 4, 2, or 0.																																												
<code>S_ISVTX</code>	01000	Save text image after execution.																																												
<code>S_IRWXU</code>	00700	Read, write, execute by owner.																																												
<code>S_IRUSR</code>	00400	Read by owner.																																												
<code>S_IWUSR</code>	00200	Write by owner.																																												
<code>S_IXUSR</code>	00100	Execute (search if a directory) by owner.																																												
<code>S_IRWXG</code>	00070	Read, write, execute by group.																																												
<code>S_IRGRP</code>	00040	Read by group.																																												
<code>S_IWGRP</code>	00020	Write by group.																																												
<code>S_IXGRP</code>	00010	Execute by group.																																												
<code>S_IRWXO</code>	00007	Read, write, execute (search) by others.																																												
<code>S_IROTH</code>	00004	Read by others.																																												
<code>S_IWOTH</code>	00002	Write by others.																																												
<code>S_IXOTH</code>	00001	Execute by others.																																												

chmod(2)

If a directory is writable and has `S_ISVTX` (the sticky bit) set, files within that directory can be removed or renamed only if one or more of the following is true (see `unlink(2)` and `rename(2)`):

- the user owns the file
- the user owns the directory
- the file is writable by the user
- the user is a privileged user

If a directory has the set group ID bit set, a given file created within that directory will have the same group ID as the directory, if that group ID is part of the group ID set of the process that created the file. Otherwise, the newly created file's group ID will be set to the effective group ID of the creating process.

If the mode bit 02000 (set group ID on execution) is set and the mode bit 00010 (execute or search by group) is not set, mandatory file/record locking will exist on a regular file. This may affect future calls to `open(2)`, `creat(2)`, `read(2)`, and `write(2)` on this file.

Upon successful completion, `chmod()` and `fchmod()` mark for update the `st_ctime` field of the file.

RETURN VALUES

Upon successful completion, 0 is returned. Otherwise, -1 is returned, the file mode is unchanged, and `errno` is set to indicate the error.

ERRORS

The `chmod()` function will fail if:

EACCES	Search permission is denied on a component of the path prefix of <i>path</i> .
EFAULT	The <i>path</i> argument points to an illegal address.
EINTR	A signal was caught during execution of the function.
EIO	An I/O error occurred while reading from or writing to the file system.
ELOOP	Too many symbolic links were encountered in translating <i>path</i> .
ENAMETOOLONG	The length of the <i>path</i> argument exceeds <code>PATH_MAX</code> , or the length of a <i>path</i> component exceeds <code>NAME_MAX</code> while <code>_POSIX_NO_TRUNC</code> is in effect.
ENOENT	Either a component of the path prefix or the file referred to by <i>path</i> does not exist or is a null pathname.
ENOLINK	The <i>fildev</i> argument points to a remote machine and the link to that machine is no longer active.
ENOTDIR	A component of the prefix of <i>path</i> is not a directory.
EPERM	The effective user ID does not match the owner of the file and is not super-user.

EROFS	The file referred to by <i>path</i> resides on a read-only file system.
The <code>fchmod()</code> function will fail if:	
EBADF	The <i>fildev</i> argument is not an open file descriptor
EIO	An I/O error occurred while reading from or writing to the file system.
EINTR	A signal was caught during execution of the <code>fchmod()</code> function.
ENOLINK	The <i>path</i> argument points to a remote machine and the link to that machine is no longer active.
EPERM	The effective user ID does not match the owner of the file and the process does not have appropriate privilege.
EROFS	The file referred to by <i>fildev</i> resides on a read-only file system.

ATTRIBUTES See `attributes(5)` for descriptions of the following attributes:

ATTRIBUTE TYPE	ATTRIBUTE VALUE
MT-Level	<code>chmod()</code> is Async-Signal-Safe

SEE ALSO `chmod(1)`, `chown(2)`, `creat(2)`, `fcntl(2)`, `mknod(2)`, `open(2)`, `read(2)`, `rename(2)`, `stat(2)`, `write(2)`, `mkfifo(3C)`, `attributes(5)`, `stat(3HEAD)`

System Interface Guide

NOTES If you use `chmod()` to change the file group owner permissions on a file with ACL entries, both the file group owner permissions and the ACL mask are changed to the new permissions. Be aware that the new ACL mask permissions may change the effective permissions for additional users and groups who have ACL entries on the file.

chown(2)

NAME	chown, lchown, fchown – change owner and group of a file
SYNOPSIS	<pre>#include <unistd.h> #include <sys/types.h> int chown(const char *path, uid_t owner, gid_t group); int lchown(const char *path, uid_t owner, gid_t group); int fchown(int fildes, uid_t owner, gid_t group);</pre>
DESCRIPTION	The <code>chown()</code> function sets the owner ID and group ID of the file specified by <i>path</i> or referenced by the open file descriptor <i>fildes</i> to <i>owner</i> and <i>group</i> respectively. If <i>owner</i> or <i>group</i> is specified as <code>-1</code>, <code>chown()</code> does not change the corresponding ID of the file. The <code>lchown()</code> function sets the owner ID and group ID of the named file in the same manner as <code>chown()</code>, unless the named file is a symbolic link. In this case, <code>lchown()</code> changes the ownership of the symbolic link file itself, while <code>chown()</code> changes the ownership of the file or directory to which the symbolic link refers. If <code>chown()</code>, <code>lchown()</code>, or <code>fchown()</code> is invoked by a process other than super-user, the set-user-ID and set-group-ID bits of the file mode, <code>S_ISUID</code> and <code>S_ISGID</code> respectively, are cleared (see <code>chmod(2)</code>). The operating system provides a configuration option, <code>{_POSIX_CHOWN_RESTRICTED}</code>, to restrict ownership changes for the <code>chown()</code>, <code>lchown()</code>, and <code>fchown()</code> functions. When <code>{_POSIX_CHOWN_RESTRICTED}</code> is not in effect, either the effective user ID of the process must match the owner of the file or the process must be the super-user to change the ownership of a file. When <code>{_POSIX_CHOWN_RESTRICTED}</code> is in effect (the default behavior), the <code>chown()</code>, <code>lchown()</code>, and <code>fchown()</code> functions, for users other than super-user, prevent the owner of the file from changing the owner ID of the file and restrict the change of the group of the file to the list of supplementary group IDs. To set this configuration option, include the following line in <code>/etc/system</code>: <pre>set rstchown = 1</pre> To disable this option, include the following line in <code>/etc/system</code>: <pre>set rstchown = 0</pre> See <code>system(4)</code> and <code>fpathconf(2)</code>. Upon successful completion, <code>chown()</code>, <code>fchown()</code> and <code>lchown()</code> mark for update the <code>st_ctime</code> field of the file.
RETURN VALUES	Upon successful completion, <code>0</code> is returned. Otherwise, <code>-1</code> is returned, the owner and group of the named file remain unchanged, and <code>errno</code> is set to indicate the error.
ERRORS	The <code>chown()</code> and <code>lchown()</code> functions will fail if:

EACCES	Search permission is denied on a component of the path prefix of <i>path</i> .
EFAULT	The <i>path</i> argument points to an illegal address.
EINTR	A signal was caught during the execution of the <code>chown()</code> or <code>lchown()</code> function.
EINVAL	The <i>group</i> or <i>owner</i> argument is out of range.
EIO	An I/O error occurred while reading from or writing to the file system.
ELOOP	Too many symbolic links were encountered in translating <i>path</i> .
ENAMETOOLONG	The length of the <i>path</i> argument exceeds {PATH_MAX}, or the length of a <i>path</i> component exceeds {NAME_MAX} while {_POSIX_NO_TRUNC} is in effect.
ENOLINK	The <i>path</i> argument points to a remote machine and the link to that machine is no longer active.
ENOENT	Either a component of the path prefix or the file referred to by <i>path</i> does not exist or is a null pathname.
ENOTDIR	A component of the path prefix of <i>path</i> is not a directory.
EPERM	The effective user ID does not match the owner of the file or the process is not the super-user and <code>_POSIX_CHOWN_RESTRICTED</code> indicates that such privilege is required.
EROFS	The named file resides on a read-only file system.
The <code>fchown()</code> function will fail if:	
EBADF	The <i>filides</i> argument is not an open file descriptor.
EIO	An I/O error occurred while reading from or writing to the file system.
EINTR	A signal was caught during execution of the function.
ENOLINK	The <i>filides</i> argument points to a remote machine and the link to that machine is no longer active.
EINVAL	The <i>group</i> or <i>owner</i> argument is out of range.
EPERM	The effective user ID does not match the owner of the file, or the process is not the super-user and <code>_POSIX_CHOWN_RESTRICTED</code> indicates that such privilege is required.

chown(2)

EROFS

The named file referred to by *filde*s resides on a read-only file system.

ATTRIBUTES

See `attributes(5)` for descriptions of the following attributes:

ATTRIBUTE TYPE	ATTRIBUTE VALUE
MT-Level	chown () is Async-Signal-Safe

SEE ALSO

`chgrp(1)`, `chown(1)`, `chmod(2)`, `fpathconf(2)`, `system(4)`, `attributes (5)`

NAME	chroot, fchroot – change root directory														
SYNOPSIS	<pre>#include <unistd.h> int chroot(const char *path); int fchroot(int fildes);</pre>														
DESCRIPTION	The <code>chroot()</code> and <code>fchroot()</code> functions cause a directory to become the root directory, the starting point for path searches for path names beginning with / (slash). The user's working directory is unaffected by the <code>chroot()</code> and <code>fchroot()</code> functions. The <i>path</i> argument points to a path name naming a directory. The <i>fildes</i> argument to <code>fchroot()</code> is the open file descriptor of the directory which is to become the root. The effective user ID of the process must be super-user to change the root directory. While it is always possible to change to the system root using the <code>fchroot()</code> function, it is not guaranteed to succeed in any other case, even should <i>fildes</i> be valid in all respects. The “.” entry in the root directory is interpreted to mean the root directory itself. Therefore, “.” cannot be used to access files outside the subtree rooted at the root directory. Instead, <code>fchroot()</code> can be used to reset the root to a directory that was opened before the root directory was changed.														
RETURN VALUES	Upon successful completion, 0 is returned. Otherwise, -1 is returned, the root directory remains unchanged, and <code>errno</code> is set to indicate the error.														
ERRORS	The <code>chroot()</code> function will fail if: <table> <tr> <td>EACCES</td><td>Search permission is denied for a component of the path prefix of <i>dirname</i>, or search permission is denied for the directory referred to by <i>dirname</i>.</td></tr> <tr> <td>EBADF</td><td>The descriptor is not valid.</td></tr> <tr> <td>EFAULT</td><td>The <i>path</i> argument points to an illegal address.</td></tr> <tr> <td>EINVAL</td><td>The <code>fchroot()</code> function attempted to change to a directory the is not the system root and external circumstances do not allow this.</td></tr> <tr> <td>EINTR</td><td>A signal was caught during the execution of the <code>chroot()</code> function.</td></tr> <tr> <td>EIO</td><td>An I/O error occurred while reading from or writing to the file system.</td></tr> <tr> <td>ELOOP</td><td>Too many symbolic links were encountered in translating <i>path</i>.</td></tr> </table>	EACCES	Search permission is denied for a component of the path prefix of <i>dirname</i> , or search permission is denied for the directory referred to by <i>dirname</i> .	EBADF	The descriptor is not valid.	EFAULT	The <i>path</i> argument points to an illegal address.	EINVAL	The <code>fchroot()</code> function attempted to change to a directory the is not the system root and external circumstances do not allow this.	EINTR	A signal was caught during the execution of the <code>chroot()</code> function.	EIO	An I/O error occurred while reading from or writing to the file system.	ELOOP	Too many symbolic links were encountered in translating <i>path</i> .
EACCES	Search permission is denied for a component of the path prefix of <i>dirname</i> , or search permission is denied for the directory referred to by <i>dirname</i> .														
EBADF	The descriptor is not valid.														
EFAULT	The <i>path</i> argument points to an illegal address.														
EINVAL	The <code>fchroot()</code> function attempted to change to a directory the is not the system root and external circumstances do not allow this.														
EINTR	A signal was caught during the execution of the <code>chroot()</code> function.														
EIO	An I/O error occurred while reading from or writing to the file system.														
ELOOP	Too many symbolic links were encountered in translating <i>path</i> .														

chroot(2)

ENAMETOOLONG	The length of the <i>path</i> argument exceeds <code>PATH_MAX</code> , or the length of a <i>path</i> component exceeds <code>NAME_MAX</code> while <code>_POSIX_NO_TRUNC</code> is in effect.
ENOENT	The named directory does not exist or is a null pathname.
ENOLINK	The <i>path</i> argument points to a remote machine and the link to that machine is no longer active.
ENOTDIR	Any component of the path name is not a directory.
EPERM	The effective user of the calling process is not super-user.

SEE ALSO chroot(1M), chdir(2)

WARNINGS The only use of `fchroot()` that is appropriate is to change back to the system root.

NAME	close – close a file descriptor
SYNOPSIS	<pre>#include <unistd.h> int close(int <i>fd</i>);</pre>
DESCRIPTION	The <code>close()</code> function will deallocate the file descriptor indicated by <i>fd</i>. To deallocate means to make the file descriptor available for return by subsequent calls to <code>open(2)</code> or other functions that allocate file descriptors. All outstanding record locks owned by the process on the file associated with the file descriptor will be removed (that is, unlocked). If <code>close()</code> is interrupted by a signal that is to be caught, it will return <code>-1</code> with <code>errno</code> set to <code>EINTR</code> and the state of <i>fd</i> is unspecified. When all file descriptors associated with a pipe or FIFO special file are closed, any data remaining in the pipe or FIFO will be discarded. When all file descriptors associated with an open file description have been closed the open file description will be freed. If the link count of the file is 0, when all file descriptors associated with the file are closed, the space occupied by the file will be freed and the file will no longer be accessible. If a STREAMS-based (see <code>intro(3)</code>) <i>fd</i> is closed and the calling process was previously registered to receive a <code>SIGPOL</code> signal (see <code>signal(3C)</code>) for events associated with that STREAM (see <code>I_SETSIG</code> in <code>streamio(7I)</code>), the calling process will be unregistered for events associated with the STREAM. The last <code>close()</code> for a STREAM causes the STREAM associated with <i>fd</i> to be dismantled. If <code>O_NONBLOCK</code> and <code>O_NDELAY</code> are not set and there have been no signals posted for the STREAM, and if there is data on the module's write queue, <code>close()</code> waits up to 15 seconds (for each module and driver) for any output to drain before dismantling the STREAM. The time delay can be changed via an <code>I_SETCLTIME</code> <code>ioctl(2)</code> request (see <code>streamio(7I)</code>). If the <code>O_NONBLOCK</code> or <code>O_NDELAY</code> flag is set, or if there are any pending signals, <code>close()</code> does not wait for output to drain, and dismantles the STREAM immediately. If <i>fd</i> is associated with one end of a pipe, the last <code>close()</code> causes a hangup to occur on the other end of the pipe. In addition, if the other end of the pipe has been named by <code>fattach(3C)</code>, then the last <code>close()</code> forces the named end to be detached by <code>fdetach(3C)</code>. If the named end has no open file descriptors associated with it and gets detached, the STREAM associated with that end is also dismantled. If <i>fd</i> refers to the master side of a pseudo-terminal, a <code>SIGHUP</code> signal is sent to the process group, if any, for which the slave side of the pseudo-terminal is the controlling terminal. It is unspecified whether closing the master side of the pseudo-terminal flushes all queued input and output.

close(2)

If *fildev* refers to the slave side of a STREAMS-based pseudo-terminal, a zero-length message may be sent to the master.

If *fildev* refers to a socket, `close()` causes the socket to be destroyed. If the socket is connection-mode, and the `SOCK_LINGER` option is set for the socket, and the socket has untransmitted data, then `close()` will block for up to the current linger interval until all data is transmitted.

RETURN VALUES

Upon successful completion, 0 is returned. Otherwise, -1 is returned and `errno` is set to indicate the error.

ERRORS

The `close()` function will fail if:

EBADF	The <i>fildev</i> argument is not a valid file descriptor.
EINTR	The <code>close()</code> function was interrupted by a signal.
ENOLINK	The <i>fildev</i> argument is on a remote machine and the link to that machine is no longer active.
ENOSPC	There was no free space remaining on the device containing the file.

The `close()` function may fail if:

EIO	An I/O error occurred while reading from or writing to the file system.
-----	---

USAGE

An application that used the `stdio` function `fopen(3C)` to open a file should use the corresponding `fclose(3C)` function rather than `close()`.

ATTRIBUTES

See `attributes(5)` for descriptions of the following attributes:

ATTRIBUTE TYPE	ATTRIBUTE VALUE
MT-Level	Async-Signal-Safe

SEE ALSO

`intro(3)`, `creat(2)`, `dup(2)`, `exec(2)`, `fcntl(2)`, `ioctl(2)`, `open(2)` `pipe(2)`, `fattach(3C)`, `fclose(3C)`, `fdetach(3C)`, `fopen(3C)`, `signal(3C)`, `attributes(5)`, `signal(3HEAD)`, `streamio(7I)`

NAME	creat – create a new file or rewrite an existing one						
SYNOPSIS	<pre>#include <sys/types.h> #include <sys/stat.h> #include <fcntl.h> int creat(const char *path, mode_t mode);</pre>						
DESCRIPTION	The <code>creat()</code> function creates a new ordinary file or prepares to rewrite an existing file named by the path name pointed to by <i>path</i>. If the file exists, the length is truncated to 0 and the mode and owner are unchanged. If the file does not exist the file's owner ID is set to the effective user ID of the process. The group ID of the file is set to the effective group ID of the process, or if the <code>S_ISGID</code> bit is set in the parent directory then the group ID of the file is inherited from the parent directory. The access permission bits of the file mode are set to the value of <i>mode</i> modified as follows: ■ If the group ID of the new file does not match the effective group ID or one of the supplementary group IDs, the <code>S_ISGID</code> bit is cleared. ■ All bits set in the process's file mode creation mask (see <code>umask(2)</code>) are correspondingly cleared in the file's permission mask. ■ The "save text image after execution bit" of the mode is cleared (see <code>chmod(2)</code> for the values of mode). Upon successful completion, a write-only file descriptor is returned and the file is open for writing, even if the mode does not permit writing. The file pointer is set to the beginning of the file. The file descriptor is set to remain open across <i>exec</i> functions (see <code>fcntl(2)</code>). A new file may be created with a mode that forbids writing. The call <code>creat(path, mode)</code> is equivalent to: <pre>open(path, O_WRONLY O_CREAT O_TRUNC, mode)</pre>						
RETURN VALUES	Upon successful completion, a non-negative integer representing the lowest numbered unused file descriptor is returned. Otherwise, -1 is returned, no files are created or modified, and <code>errno</code> is set to indicate the error.						
ERRORS	The <code>creat()</code> function will fail: <table> <tr> <td>EACCES</td><td>Search permission is denied on a component of the path prefix; the file does not exist and the directory in which the file is to be created does not permit writing; or the file exists and write permission is denied.</td></tr> <tr> <td>EAGAIN</td><td>The file exists, mandatory file/record locking is set, and there are outstanding record locks on the file (see <code>chmod(2)</code>).</td></tr> <tr> <td>EDQUOT</td><td>The directory where the new file entry is being placed cannot be extended because the user's quota of disk blocks on that file</td></tr> </table>	EACCES	Search permission is denied on a component of the path prefix; the file does not exist and the directory in which the file is to be created does not permit writing; or the file exists and write permission is denied.	EAGAIN	The file exists, mandatory file/record locking is set, and there are outstanding record locks on the file (see <code>chmod(2)</code>).	EDQUOT	The directory where the new file entry is being placed cannot be extended because the user's quota of disk blocks on that file
EACCES	Search permission is denied on a component of the path prefix; the file does not exist and the directory in which the file is to be created does not permit writing; or the file exists and write permission is denied.						
EAGAIN	The file exists, mandatory file/record locking is set, and there are outstanding record locks on the file (see <code>chmod(2)</code>).						
EDQUOT	The directory where the new file entry is being placed cannot be extended because the user's quota of disk blocks on that file						

creat(2)

system has been exhausted, or the user's quota of inodes on the file system where the file is being created has been exhausted.

EFAULT	The <i>path</i> argument points to an illegal address.
EINTR	A signal was caught during the execution of the <code>creat()</code> function.
EISDIR	The named file is an existing directory.
ELOOP	Too many symbolic links were encountered in translating <i>path</i> .
EMFILE	The process has too many open files (see <code>getrlimit(2)</code>).
ENFILE	The system file table is full.
ENOENT	A component of the path prefix does not exist, or the path name is null.
ENOLINK	The <i>path</i> argument points to a remote machine and the link to that machine is no longer active.
ENOSPC	The file system is out of inodes.
ENOTDIR	A component of the path prefix is not a directory.
EOVERFLOW	The file is a large file at the time of <code>creat()</code> .
EROFS	The named file resides or would reside on a read-only file system.

USAGE The `creat()` function has a transitional interface for 64-bit file offsets. See `lf64(5)`.

ATTRIBUTES See `attributes(5)` for descriptions of the following attributes:

ATTRIBUTE TYPE	ATTRIBUTE VALUE
MT-Level	Async-Signal-Safe

SEE ALSO `chmod(2)`, `close(2)`, `dup(2)`, `fcntl(2)`, `getrlimit(2)`, `lseek(2)`, `open(2)`, `read(2)`, `umask(2)`, `write(2)`, `stat(3HEAD)`, `attributes(5)`, `largefile(5)`, `lf64(5)`

NAME	dup – duplicate an open file descriptor								
SYNOPSIS	<pre>#include <unistd.h> int dup(int <i>fdes</i>);</pre>								
DESCRIPTION	The <code>dup()</code> function returns a new file descriptor having the following in common with the original open file descriptor <i>fdes</i>: ■ same open file (or pipe) ■ same file pointer (that is, both file descriptors share one file pointer) ■ same access mode (read, write or read/write). The new file descriptor is set to remain open across <i>exec</i> functions (see <code>fcntl(2)</code>). The file descriptor returned is the lowest one available. The <code>dup(<i>fdes</i>)</code> function call is equivalent to: <pre>fcntl(<i>fdes</i>, F_DUPFD, 0)</pre>								
RETURN VALUES	Upon successful completion, a non-negative integer representing the file descriptor is returned. Otherwise, <code>-1</code> is returned and <code>errno</code> is set to indicate the error.								
ERRORS	The <code>dup()</code> function will fail if: <table> <tr> <td>EBADF</td><td>The <i>fdes</i> argument is not a valid open file descriptor.</td></tr> <tr> <td>EINTR</td><td>A signal was caught during the execution of the <code>dup()</code> function.</td></tr> <tr> <td>EMFILE</td><td>The process has too many open files (see <code>getrlimit(2)</code>).</td></tr> <tr> <td>ENOLINK</td><td>The <i>fdes</i> argument is on a remote machine and the link to that machine is no longer active.</td></tr> </table>	EBADF	The <i>fdes</i> argument is not a valid open file descriptor.	EINTR	A signal was caught during the execution of the <code>dup()</code> function.	EMFILE	The process has too many open files (see <code>getrlimit(2)</code>).	ENOLINK	The <i>fdes</i> argument is on a remote machine and the link to that machine is no longer active.
EBADF	The <i>fdes</i> argument is not a valid open file descriptor.								
EINTR	A signal was caught during the execution of the <code>dup()</code> function.								
EMFILE	The process has too many open files (see <code>getrlimit(2)</code>).								
ENOLINK	The <i>fdes</i> argument is on a remote machine and the link to that machine is no longer active.								
ATTRIBUTES	See <code>attributes(5)</code> for descriptions of the following attributes: <table> <tr> <th>ATTRIBUTE TYPE</th><th>ATTRIBUTE VALUE</th></tr> <tr> <td>MT-Level</td><td>Async-Signal-Safe</td></tr> </table>	ATTRIBUTE TYPE	ATTRIBUTE VALUE	MT-Level	Async-Signal-Safe				
ATTRIBUTE TYPE	ATTRIBUTE VALUE								
MT-Level	Async-Signal-Safe								
SEE ALSO	<code>close(2)</code> , <code>creat(2)</code> , <code>exec(2)</code> , <code>fcntl(2)</code> , <code>getrlimit(2)</code> , <code>open(2)</code> , <code>pipe(2)</code> , <code>dup2(3C)</code> , <code>lockf(3C)</code> , <code>attributes(5)</code>								

exec(2)

NAME	exec, execl, execv, execl, execve, execlp, execvp – execute a file
SYNOPSIS	<pre>#include <unistd.h> int execl(const char *path, const char *arg0, ..., const char *argn, char * /*NULL*/); int execv(const char *path, char *const argv[]); int execl(const char *path, const char *arg0, ..., const char *argn, char * /*NULL*/, char *const envp[]); int execve(const char *path, char *const argv[], char *const envp[]); int execlp(const char *file, const char *arg0, ..., const char *argn, char * /*NULL*/); int execvp(const char *file, char *const argv[]);</pre>
DESCRIPTION	Each of the functions in the <code>exec</code> family replaces the current process image with a new process image. The new image is constructed from a regular, executable file called the <i>new process image file</i>. This file is either an executable object file or a file of data for an interpreter. There is no return from a successful call to one of these functions because the calling process image is overlaid by the new process image. An interpreter file begins with a line of the form <pre>#! pathname [arg]</pre> where <i>pathname</i> is the path of the interpreter, and <i>arg</i> is an optional argument. When an interpreter file is executed, the system invokes the specified interpreter. The <i>pathname</i> specified in the interpreter file is passed as <i>arg0</i> to the interpreter. If <i>arg</i> was specified in the interpreter file, it is passed as <i>arg1</i> to the interpreter. The remaining arguments to the interpreter are <i>arg0</i> through <i>argn</i> of the originally <code>exec'd</code> file. The interpreter named by <i>pathname</i> must not be an interpreter file. When a C-language program is executed as a result of this call, it is entered as a C-language function call as follows: <pre>int main (int argc, char *argv[], char *envp[]);</pre> where <i>argc</i> is the argument count, <i>argv</i> is an array of character pointers to the arguments themselves, and <i>envp</i> is an array of character pointers to the environment strings. The <i>argv</i> and <i>environ</i> arrays are each terminated by a null pointer. The null pointer terminating the <i>argv</i> array is not counted in <i>argc</i>. As indicated, <i>argc</i> is at least one and the first member of the array points to a string containing the name of the file. The arguments specified by a program with one of the <code>exec</code> functions are passed on to the new process image in the <code>main()</code> arguments. The <i>path</i> argument points to a path name that identifies the new process image file.

The *file* argument is used to construct a pathname that identifies the new process image file. If the *file* argument contains a slash character, it is used as the pathname for this file. Otherwise, the path prefix for this file is obtained by a search of the directories passed in the `PATH` environment variable (see `environ(5)`). The environment is supplied typically by the shell. If the process image file is not a valid executable object file, `execvp()` and `execvp()` use the contents of that file as standard input to the shell. In this case, the shell becomes the new process image. In a standard-conforming application (see `standards(5)`), the `exec` family of functions use `/usr/xpg4/bin/sh` (see `ksh(1)`); otherwise, they use `/usr/bin/sh` (see `sh(1)`).

The arguments represented by *arg0...* are pointers to null-terminated character strings. These strings constitute the argument list available to the new process image. The list is terminated by a null pointer. The *arg0* argument should point to a filename that is associated with the process being started by one of the `exec` functions.

The *argv* argument is an array of character pointers to null-terminated strings. The last member of this array must be a null pointer. These strings constitute the argument list available to the new process image. The value in *argv*[0] should point to a filename that is associated with the process being started by one of the `exec` functions.

The *envp* argument is an array of character pointers to null-terminated strings. These strings constitute the environment for the new process image. The *envp* array is terminated by a null pointer. For `exec1()`, `execv()`, `execvp()`, and `execvp()`, the C-language run-time start-off routine places a pointer to the environment of the calling process in the global object `extern char **environ`, and it is used to pass the environment of the calling process to the new process image.

The number of bytes available for the new process's combined argument and environment lists is `ARG_MAX`. It is implementation-dependent whether null terminators, pointers, and/or any alignment bytes are included in this total.

File descriptors open in the calling process image remain open in the new process image, except for those whose close-on-exec flag `FD_CLOEXEC` is set; (see `fcntl(2)`). For those file descriptors that remain open, all attributes of the open file description, including file locks, remain unchanged.

Directory streams open in the calling process image are closed in the new process image.

The state of conversion descriptors and message catalogue descriptors in the new process image is undefined. For the new process, the equivalent of:

```
setlocale(LC_ALL, "C")
```

is executed at startup.

Signals set to the default action (`SIG_DFL`) in the calling process image are set to the default action in the new process image (see `signal(3C)`). Signals set to be ignored (`SIG_IGN`) by the calling process image are set to be ignored by the new process

exec(2)

image. Signals set to be caught by the calling process image are set to the default action in the new process image (see `signal(3HEAD)`). After a successful call to any of the `exec` functions, alternate signal stacks are not preserved and the `SA_ONSTACK` flag is cleared for all signals.

After a successful call to any of the `exec` functions, any functions previously registered by `atexit(3C)` are no longer registered.

The saved resource limits in the new process image are set to be a copy of the process's corresponding hard and soft resource limits.

If the `ST_NOSUID` bit is set for the file system containing the new process image file, then the effective user ID, effective group ID, saved set-user-ID, and saved set-group-ID are unchanged in the new process image. If the set-user-ID mode bit of the new process image file is set (see `chmod(2)`), the effective user ID of the new process image is set to the owner ID of the new process image file. Similarly, if the set-group-ID mode bit of the new process image file is set, the effective group ID of the new process image is set to the group ID of the new process image file. The real user ID and real group ID of the new process image remain the same as those of the calling process image. The effective user ID and effective group ID of the new process image are saved (as the saved set-user-ID and the saved set-group-ID for use by `setuid(2)`).

If the effective user-ID is `root` or super-user, the set-user-ID and set-group-ID bits will be honored when the process is being controlled by `ptrace`.

Any shared memory segments attached to the calling process image will not be attached to the new process image (see `shmop(2)`). Any mappings established through `mmap()` are not preserved across an `exec`. Memory mappings created in the process are unmapped before the address space is rebuilt for the new process image. See `mmap(2)`.

Memory locks established by the calling process via calls to `mlockall(3C)` or `mlock(3C)` are removed. If locked pages in the address space of the calling process are also mapped into the address spaces the locks established by the other processes will be unaffected by the call by this process to the `exec` function. If the `exec` function fails, the effect on memory locks is unspecified.

If `_XOPEN_REALTIME` is defined and has a value other than `-1`, any named semaphores open in the calling process are closed as if by appropriate calls to `sem_close(3RT)`.

Profiling is disabled for the new process; see `profil(2)`.

Timers created by the calling process with `timer_create(3RT)` are deleted before replacing the current process image with the new process image.

For the `SCHED_FIFO` and `SCHED_RR` scheduling policies, the policy and priority settings are not changed by a call to an `exec` function.

All open message queue descriptors in the calling process are closed, as described in `mq_close(3RT)`.

Any outstanding asynchronous I/O operations may be cancelled. Those asynchronous I/O operations that are not canceled will complete as if the `exec` function had not yet occurred, but any associated signal notifications are suppressed. It is unspecified whether the `exec` function itself blocks awaiting such I/O completion. In no event, however, will the new process image created by the `exec` function be affected by the presence of outstanding asynchronous I/O operations at the time the `exec` function is called.

The new process also inherits the following attributes from the calling process:

- nice value (see `nice(2)`)
- scheduler class and priority (see `priocntl(2)`)
- process ID
- parent process ID
- process group ID
- task ID
- supplementary group IDs
- `semadj` values (see `semop(2)`)
- session membership (see `exit(2)` and `signal(3C)`)
- real user ID
- real group ID
- project ID
- trace flag (see `ptrace(2)` request 0)
- time left until an alarm clock signal (see `alarm(2)`)
- current working directory
- root directory
- file mode creation mask (see `umask(2)`)
- file size limit (see `ulimit(2)`)
- resource limits (see `getrlimit(2)`)
- `tms_utime`, `tms_stime`, `tms_cutime`, and `tms_cstime` (see `times(2)`)
- file-locks (see `fcntl(2)` and `lockf(3C)`)
- controlling terminal
- process signal mask (see `sigprocmask(2)`)
- pending signals (see `sigpending(2)`)

A call to any `exec` function from a process with more than one thread results in all threads being terminated and the new executable image being loaded and executed. No destructor functions will be called.

Upon successful completion, each of the functions in the `exec` family marks for update the `st_atime` field of the file. If an `exec` function failed but was able to locate the *process image file*, whether the `st_atime` field is marked for update is unspecified. Should the function succeed, the process image file is considered to have been opened with `open(2)`. The corresponding `close(2)` is considered to occur at a time after this open, but before process termination or successful completion of a subsequent call to

exec(2)

	one of the <code>exec</code> functions. The <code>argv[]</code> and <code>envp[]</code> arrays of pointers and the strings to which those arrays point will not be modified by a call to one of the <code>exec</code> functions, except as a consequence of replacing the process image.	
	The saved resource limits in the new process image are set to be a copy of the process's corresponding hard and soft limits.	
RETURN VALUES	If a function in the <code>exec</code> family returns to the calling process image, an error has occurred; the return value is <code>-1</code> and <code>errno</code> is set to indicate the error.	
ERRORS	The <code>exec</code> functions will fail if:	
	E2BIG	The number of bytes in the new process's argument list is greater than the system-imposed limit of <code>ARG_MAX</code> bytes. The argument list limit is sum of the size of the argument list plus the size of the environment's exported shell variables.
	EACCES	Search permission is denied for a directory listed in the new process file's path prefix; the new process file is not an ordinary file; or the new process file mode denies execute permission.
	EAGAIN	Total amount of system memory available when reading using raw I/O is temporarily insufficient.
	EFAULT	An argument points to an illegal address.
	EINTR	A signal was caught during the execution of one of the functions in the <code>exec</code> family.
	ELOOP	Too many symbolic links were encountered in translating <i>path</i> or <i>file</i> .
	ENAMETOOLONG	The length of the <i>file</i> or <i>path</i> argument exceeds <code>PATH_MAX</code> , or the length of a <i>file</i> or <i>path</i> component exceeds <code>{NAME_MAX}</code> while <code>{_POSIX_NO_TRUNC}</code> is in effect.
	ENOENT	One or more components of the new process path name of the file do not exist or is a null pathname.
	ENOLINK	The <i>path</i> argument points to a remote machine and the link to that machine is no longer active.
	ENOTDIR	A component of the new process path of the file prefix is not a directory.
	The <code>exec</code> functions, except for <code>execvp()</code> and <code>execvp()</code> , will fail if:	
	ENOEXEC	The new process image file has the appropriate access permission but is not in the proper format.

The exec functions may fail if:

ENAMETOOLONG	Pathname resolution of a symbolic link produced an intermediate result whose length exceeds <code>PATH_MAX</code> .
ENOMEM	The new process image requires more memory than is allowed by the hardware or system-imposed by memory management constraints. (see <code>brk(2)</code>).
ETXTBSY	The new process image file is a pure procedure (shared text) file that is currently open for writing by some process.

USAGE As the state of conversion descriptors and message catalogue descriptors in the new process image is undefined, portable applications should not rely on their use and should close them prior to calling one of the `exec` functions.

Applications that require other than the default POSIX locale should call `setlocale(3C)` with the appropriate parameters to establish the locale of the new process.

The `environ` array should not be accessed directly by the application.

ATTRIBUTES See `attributes(5)` for descriptions of the following attributes:

ATTRIBUTE TYPE	ATTRIBUTE VALUE
MT-Level	<code>execle()</code> and <code>execve()</code> are Async-Signal-Safe

SEE ALSO `ksh(1)`, `ps(1)`, `sh(1)`, `alarm(2)`, `brk(2)`, `chmod(2)`, `exit(2)`, `fcntl(2)`, `fork(2)`, `getrlimit(2)`, `mmap(2)`, `nice(2)`, `prctl(2)`, `profil(2)`, `ptrace(2)`, `semop(2)`, `shmop(2)`, `sigpending(2)`, `sigprocmask(2)`, `times(2)`, `umask(2)`, `lockf(3C)`, `setlocale(3C)`, `signal(3C)`, `system(3C)`, `timer_create(3RT)`, `a.out(4)`, `attributes(5)`, `environ(5)`, `standards(5)`

WARNINGS If a program is setuid to a user ID other than the super-user, and the program is executed when the real user ID is super-user, then the program has some of the powers of a super-user as well.

exit(2)

NAME	exit, _exit – terminate process
SYNOPSIS	<pre>#include <stdlib.h> void exit(int <i>status</i>) ; #include <unistd.h> void _exit(int <i>status</i>) ;</pre>
DESCRIPTION	The <code>exit()</code> function first calls all functions registered by <code>atexit(3C)</code>, in the reverse order of their registration. Each function is called as many times as it was registered. If a function registered by a call to <code>atexit(3C)</code> fails to return, the remaining registered functions are not called and the rest of the <code>exit()</code> processing is not completed. If <code>exit()</code> is called more than once, the effects are undefined. The <code>exit()</code> function then flushes all output streams, closes all open streams, and removes all files created by <code>tmpfile(3C)</code>. The <code>_exit()</code> and <code>exit()</code> functions terminate the calling process with the following consequences: ■ All of the file descriptors, directory streams, conversion descriptors and message catalogue descriptors open in the calling process are closed. ■ If the parent process of the calling process is executing a <code>wait(2)</code>, <code>wait3(3C)</code>, <code>waitid(2)</code> or <code>waitpid(2)</code>, and has neither set its <code>SA_NOCLDWAIT</code> flag nor set <code>SIGCHLD</code> to <code>SIG_IGN</code>, it is notified of the calling process's termination and the low-order eight bits (that is, bits 0377) of <i>status</i> are made available to it. If the parent is not waiting, the child's status will be made available to it when the parent subsequently executes <code>wait(2)</code>, <code>wait3(3C)</code>, <code>waitid(2)</code> or <code>waitpid(2)</code>. ■ If the parent process of the calling process is not executing a <code>wait(2)</code>, <code>wait3(3C)</code>, <code>waitid(2)</code> or <code>waitpid(2)</code>, and has not set its <code>SA_NOCLDWAIT</code> flag, or set <code>SIGCHLD</code> to <code>SIG_IGN</code>, the calling process is transformed into a <i>zombie process</i>. A <i>zombie process</i> is an inactive process and it will be deleted at some later time when its parent process executes <code>wait(2)</code>, <code>wait3(3C)</code>, <code>waitid(2)</code> or <code>waitpid(2)</code>. A zombie process only occupies a slot in the process table; it has no other space allocated either in user or kernel space. The process table slot that it occupies is partially overlaid with time accounting information (see <code><sys/proc.h></code>) to be used by the <code>times(2)</code> function. ■ Termination of a process does not directly terminate its children. The sending of a <code>SIGHUP</code> signal as described below indirectly terminates children in some circumstances. ■ A <code>SIGCHLD</code> will be sent to the parent process. ■ The parent process ID of all of the calling process's existing child processes and zombie processes is set to 1. That is, these processes are inherited by the initialization process (see <code>intro(3)</code>). ■ Each mapped memory object is unmapped.

exit(2)

- Each attached shared-memory segment is detached and the value of `shm_nattch` (see `shmget(2)`) in the data structure associated with its shared memory ID is decremented by 1.
- For each semaphore for which the calling process has set a `semadj` value (see `semop(2)`), that value is added to the `semval` of the specified semaphore.
- If the process is a controlling process, the `SIGHUP` signal will be sent to each process in the foreground process group of the controlling terminal belonging to the calling process.
- If the process is a controlling process, the controlling terminal associated with the session is disassociated from the session, allowing it to be acquired by a new controlling process.
- If the exit of the process causes a process group to become orphaned, and if any member of the newly-orphaned process group is stopped, then a `SIGHUP` signal followed by a `SIGCONT` signal will be sent to each process in the newly-orphaned process group.
- If the parent process has set its `SA_NOCLDWAIT` flag, or set `SIGCHLD` to `SIG_IGN`, the status will be discarded, and the lifetime of the calling process will end immediately.
- If the process has process, text or data locks, an `UNLOCK` is performed (see `plock(3C)` and `memcntl(2)`).
- All open named semaphores in the process are closed as if by appropriate calls to `sem_close(3RT)`. All open message queues in the process are closed as if by appropriate calls to `mq_close(3RT)`. Any outstanding asynchronous I/O operations may be cancelled.
- An accounting record is written on the accounting file if the system's accounting routine is enabled (see `acct(2)`).
- An extended accounting record is written to the extended process accounting file if the system's extended process accounting facility is enabled (see `acctadm(1M)`).
- If the current process is the last process within its task and if the system's extended task accounting facility is enabled (see `acctadm(1M)`), an extended accounting record is written to the extended task accounting file.

RETURN VALUES These functions do not return.

ERRORS No errors are defined.

USAGE Normally applications should use `exit()` rather than `_exit()`.

ATTRIBUTES See `attributes(5)` for descriptions of the following attributes:

ATTRIBUTE TYPE	ATTRIBUTE VALUE
MT-Level	<code>_exit()</code> is Async-Signal Safe

exit(2)

SEE ALSO | acctadm(1M), acct(2), close(2), memcntl(2), semop(2), shmget(2),
sigaction(2), times(2), wait(2), waitid(2), waitpid(2), intro(3), atexit(3C),
fclose(3C), mq_close(3RT), plock(3C), tmpfile(3C), wait3(3C), attributes(5),
signal(3HEAD)

NAME	fcntl – file control												
SYNOPSIS	<pre>#include <sys/types.h> #include <unistd.h> #include <fcntl.h> int fcntl(int <i>fildev</i>, int <i>cmd</i>, /* <i>arg</i> */ ...);</pre>												
DESCRIPTION	The <code>fcntl()</code> function provides for control over open files. The <i>fildev</i> argument is an open file descriptor. The <code>fcntl()</code> function may take a third argument, <i>arg</i>, whose data type, value and use depend upon the value of <i>cmd</i>. The <i>cmd</i> argument specifies the operation to be performed by <code>fcntl()</code>. The available values for <i>cmd</i> are defined in the header <code><fcntl.h></code>, which include: <table> <tr> <td>F_DUPFD</td><td>Return a new file descriptor which is the lowest numbered available (that is, not already open) file descriptor greater than or equal to the third argument, <i>arg</i>, taken as an integer of type <code>int</code>. The new file descriptor refers to the same open file description as the original file descriptor, and shares any locks. The <code>FD_CLOEXEC</code> flag associated with the new file descriptor is cleared to keep the file open across calls to one of the <code>exec(2)</code> functions.</td></tr> <tr> <td>F_DUP2FD</td><td>Similar to <code>F_DUPFD</code>, but always returns <i>arg</i>. <code>F_DUP2FD</code> closes <i>arg</i> if it is open and not equal to <i>fildev</i>. <code>F_DUP2FD</code> is equivalent to <code>dup2(fildev, arg)</code>.</td></tr> <tr> <td>F_GETFD</td><td>Get the file descriptor flags defined in <code><fcntl.h></code> that are associated with the file descriptor <i>fildev</i>. File descriptor flags are associated with a single file descriptor and do not affect other file descriptors that refer to the same file.</td></tr> <tr> <td>F_SETFD</td><td>Set the file descriptor flags defined in <code><fcntl.h></code>, that are associated with <i>fildev</i>, to the third argument, <i>arg</i>, taken as type <code>int</code>. If the <code>FD_CLOEXEC</code> flag in the third argument is 0, the file will remain open across the <code>exec()</code> functions; otherwise the file will be closed upon successful execution of one of the <code>exec()</code> functions.</td></tr> <tr> <td>F_GETFL</td><td>Get the file status flags and file access modes, defined in <code><fcntl.h></code>, for the file description associated with <i>fildev</i>. The file access modes can be extracted from the return value using the mask <code>O_ACCMODE</code>, which is defined in <code><fcntl.h></code>. File status flags and file access modes are associated with the file description and do not affect other file descriptors that refer to the same file with different open file descriptions.</td></tr> <tr> <td>F_SETFL</td><td>Set the file status flags, defined in <code><fcntl.h></code>, for the file description associated with <i>fildev</i> from the corresponding bits in the third argument, <i>arg</i>, taken as type <code>int</code>. Bits corresponding to</td></tr> </table>	F_DUPFD	Return a new file descriptor which is the lowest numbered available (that is, not already open) file descriptor greater than or equal to the third argument, <i>arg</i> , taken as an integer of type <code>int</code> . The new file descriptor refers to the same open file description as the original file descriptor, and shares any locks. The <code>FD_CLOEXEC</code> flag associated with the new file descriptor is cleared to keep the file open across calls to one of the <code>exec(2)</code> functions.	F_DUP2FD	Similar to <code>F_DUPFD</code> , but always returns <i>arg</i> . <code>F_DUP2FD</code> closes <i>arg</i> if it is open and not equal to <i>fildev</i> . <code>F_DUP2FD</code> is equivalent to <code>dup2(fildev, arg)</code> .	F_GETFD	Get the file descriptor flags defined in <code><fcntl.h></code> that are associated with the file descriptor <i>fildev</i> . File descriptor flags are associated with a single file descriptor and do not affect other file descriptors that refer to the same file.	F_SETFD	Set the file descriptor flags defined in <code><fcntl.h></code> , that are associated with <i>fildev</i> , to the third argument, <i>arg</i> , taken as type <code>int</code> . If the <code>FD_CLOEXEC</code> flag in the third argument is 0, the file will remain open across the <code>exec()</code> functions; otherwise the file will be closed upon successful execution of one of the <code>exec()</code> functions.	F_GETFL	Get the file status flags and file access modes, defined in <code><fcntl.h></code> , for the file description associated with <i>fildev</i> . The file access modes can be extracted from the return value using the mask <code>O_ACCMODE</code> , which is defined in <code><fcntl.h></code> . File status flags and file access modes are associated with the file description and do not affect other file descriptors that refer to the same file with different open file descriptions.	F_SETFL	Set the file status flags, defined in <code><fcntl.h></code> , for the file description associated with <i>fildev</i> from the corresponding bits in the third argument, <i>arg</i> , taken as type <code>int</code> . Bits corresponding to
F_DUPFD	Return a new file descriptor which is the lowest numbered available (that is, not already open) file descriptor greater than or equal to the third argument, <i>arg</i> , taken as an integer of type <code>int</code> . The new file descriptor refers to the same open file description as the original file descriptor, and shares any locks. The <code>FD_CLOEXEC</code> flag associated with the new file descriptor is cleared to keep the file open across calls to one of the <code>exec(2)</code> functions.												
F_DUP2FD	Similar to <code>F_DUPFD</code> , but always returns <i>arg</i> . <code>F_DUP2FD</code> closes <i>arg</i> if it is open and not equal to <i>fildev</i> . <code>F_DUP2FD</code> is equivalent to <code>dup2(fildev, arg)</code> .												
F_GETFD	Get the file descriptor flags defined in <code><fcntl.h></code> that are associated with the file descriptor <i>fildev</i> . File descriptor flags are associated with a single file descriptor and do not affect other file descriptors that refer to the same file.												
F_SETFD	Set the file descriptor flags defined in <code><fcntl.h></code> , that are associated with <i>fildev</i> , to the third argument, <i>arg</i> , taken as type <code>int</code> . If the <code>FD_CLOEXEC</code> flag in the third argument is 0, the file will remain open across the <code>exec()</code> functions; otherwise the file will be closed upon successful execution of one of the <code>exec()</code> functions.												
F_GETFL	Get the file status flags and file access modes, defined in <code><fcntl.h></code> , for the file description associated with <i>fildev</i> . The file access modes can be extracted from the return value using the mask <code>O_ACCMODE</code> , which is defined in <code><fcntl.h></code> . File status flags and file access modes are associated with the file description and do not affect other file descriptors that refer to the same file with different open file descriptions.												
F_SETFL	Set the file status flags, defined in <code><fcntl.h></code> , for the file description associated with <i>fildev</i> from the corresponding bits in the third argument, <i>arg</i> , taken as type <code>int</code> . Bits corresponding to												

fcntl(2)

	the file access mode and the <i>oflag</i> values that are set in <i>arg</i> are ignored. If any bits in <i>arg</i> other than those mentioned here are changed by the application, the result is unspecified.
F_GETOWN	If <i>fildev</i> refers to a socket, get the process or process group ID specified to receive SIGURG signals when out-of-band data is available. Positive values indicate a process ID; negative values, other than -1, indicate a process group ID. If <i>fildev</i> does not refer to a socket, the results are unspecified.
F_SETOWN	If <i>fildev</i> refers to a socket, set the process or process group ID specified to receive SIGURG signals when out-of-band data is available, using the value of the third argument, <i>arg</i> , taken as type int. Positive values indicate a process ID; negative values, other than -1, indicate a process group ID. If <i>fildev</i> does not refer to a socket, the results are unspecified.
F_FREESP	Free storage space associated with a section of the ordinary file <i>fildev</i> . The section is specified by a variable of data type struct flock pointed to by <i>arg</i> . The data type struct flock is defined in the <fcntl.h> header (see fcntl(3HEAD)) and is described below. Note that all file systems might not support all possible variations of F_FREESP arguments. In particular, many file systems allow space to be freed only at the end of a file.
The following commands are available for advisory record locking. Record locking is supported for regular files, and may be supported for other files.	
F_GETLK	Get the first lock which blocks the lock description pointed to by the third argument, <i>arg</i> , taken as a pointer to type struct flock, defined in <fcntl.h>. The information retrieved overwrites the information passed to fcntl() in the structure flock. If no lock is found that would prevent this lock from being created, then the structure will be left unchanged except for the lock type which will be set to F_UNLCK.
F_GETLK64	Equivalent to F_GETLK, but takes a struct flock64 argument rather than a struct flock argument.
F_SETLK	Set or clear a file segment lock according to the lock description pointed to by the third argument, <i>arg</i> , taken as a pointer to type struct flock, defined in <fcntl.h>. F_SETLK is used to establish shared (or read) locks (F_RDLCK) or exclusive (or write) locks (F_WRLCK), as well as to remove either type of lock (F_UNLCK). F_RDLCK, F_WRLCK and F_UNLCK are defined in <fcntl.h>. If a shared or exclusive lock cannot be set, fcntl() will return immediately with a return value of -1.
F_SETLK64	Equivalent to F_SETLK, but takes a struct flock64 argument rather than a struct flock argument.

<code>F_SETLKW</code>	This command is the same as <code>F_SETLK</code> except that if a shared or exclusive lock is blocked by other locks, the process will wait until the request can be satisfied. If a signal that is to be caught is received while <code>fcntl()</code> is waiting for a region, <code>fcntl()</code> will be interrupted. Upon return from the process' signal handler, <code>fcntl()</code> will return <code>-1</code> with <code>errno</code> set to <code>EINTR</code> , and the lock operation will not be done.
<code>F_SETLKW64</code>	Equivalent to <code>F_SETLKW</code> , but takes a <code>struct flock64</code> argument rather than a <code>struct flock</code> argument.

When a shared lock is set on a segment of a file, other processes will be able to set shared locks on that segment or a portion of it. A shared lock prevents any other process from setting an exclusive lock on any portion of the protected area. A request for a shared lock will fail if the file descriptor was not opened with read access.

An exclusive lock will prevent any other process from setting a shared lock or an exclusive lock on any portion of the protected area. A request for an exclusive lock will fail if the file descriptor was not opened with write access.

The `flock` structure contains at least the following elements:

```

short  l_type;      /* lock operation type */
short  l_whence;    /* lock base indicator */
off_t  l_start;     /* starting offset from base */
off_t  l_len;       /* lock length; l_len == 0 means
                    until end of file */
long   l_sysid;     /* system ID running process holding lock */
pid_t  l_pid;       /* process ID of process holding lock */

```

The value of `l_whence` is `SEEK_SET`, `SEEK_CUR`, or `SEEK_END`, to indicate that the relative offset `l_start` bytes will be measured from the start of the file, current position or end of the file, respectively. The value of `l_len` is the number of consecutive bytes to be locked. The value of `l_len` may be negative (where the definition of `off_t` permits negative values of `l_len`). After a successful `F_GETLK` or `F_GETLK64` request, that is, one in which a lock was found, the value of `l_whence` will be `SEEK_SET`.

The `l_pid` and `l_sysid` fields are used only with `F_GETLK` or `F_GETLK64` to return the process ID of the process holding a blocking lock and to indicate which system is running that process.

If `l_len` is positive, the area affected starts at `l_start` and ends at `l_start + l_len - 1`. If `l_len` is negative, the area affected starts at `l_start + l_len` and ends at `l_start - 1`. Locks may start and extend beyond the current end of a file, but must not be negative relative to the beginning of the file. A lock will be set to extend to the largest possible value of the file offset for that file by setting `l_len` to 0. If such a lock also has `l_start` set to 0 and `l_whence` is set to `SEEK_SET`, the whole file will be locked.

fcntl(2)

If a process has an existing lock in which `l_len` is 0 and which includes the last byte of the requested segment, and an unlock (`F_UNLCK`) request is made in which `l_len` is non-zero and the offset of the last byte of the requested segment is the maximum value for an object of type `off_t`, then the `F_UNLCK` request will be treated as a request to unlock from the start of the requested segment with an `l_len` equal to 0. Otherwise, the request will attempt to unlock only the requested segment.

There will be at most one type of lock set for each byte in the file. Before a successful return from an `F_SETLK`, `F_SETLK64`, `F_SETLKW`, or `F_SETLKW64` request when the calling process has previously existing locks on bytes in the region specified by the request, the previous lock type for each byte in the specified region will be replaced by the new lock type. As specified above under the descriptions of shared locks and exclusive locks, an `F_SETLK`, `F_SETLK64`, `F_SETLKW`, or `F_SETLKW64` request will (respectively) fail or block when another process has existing locks on bytes in the specified region and the type of any of those locks conflicts with the type specified in the request.

All locks associated with a file for a given process are removed when a file descriptor for that file is closed by that process or the process holding that file descriptor terminates. Locks are not inherited by a child process created using `fork(2)`.

A potential for deadlock occurs if a process controlling a locked region is put to sleep by attempting to lock another process' locked region. If the system detects that sleeping until a locked region is unlocked would cause a deadlock, `fcntl()` will fail with an `EDEADLK` error.

The following values for *cmd* are used for file share reservations. A share reservation is placed on an entire file to allow cooperating processes to control access to the file.

`F_SHARE` Sets a share reservation on a file with the specified access mode and designates which types of access to deny.

`F_UNSHARE` Remove an existing share reservation.

File share reservations are an advisory form of access control among cooperating processes, on both local and remote machines. They are most often used by DOS or Windows emulators and DOS based NFS clients. However, native UNIX versions of DOS or Windows applications may also choose to use this form of access control.

A share reservation is described by an `fshare` structure defined in `<sys/fcntl.h>`, which is included in `<fcntl.h>` as follows:

```
typedef struct fshare {
    short    f_access;
    short    f_deny;
    long     f_id;
} fshare_t;
```

A share reservation specifies the type of access, `f_access`, to be requested on the open file descriptor. If access is granted, it further specifies what type of access to deny other processes, `f_deny`. A single process on the same file may hold multiple non-conflicting reservations by specifying an identifier, `f_id`, unique to the process, with each request.

An `F_UNSHARE` request releases the reservation with the specified `f_id`. The `f_access` and `f_deny` fields are ignored.

Valid `f_access` values are:

<code>F_RDACC</code>	Set a file share reservation for read-only access.
<code>F_WRACC</code>	Set a file share reservation for write-only access.
<code>F_RWACC</code>	Set a file share reservation for read and write access.

Valid `f_deny` values are:

<code>F_COMPAT</code>	Set a file share reservation to compatibility mode.
<code>F_RDDNY</code>	Set a file share reservation to deny read access to other processes.
<code>F_WRDNY</code>	Set a file share reservation to deny write access to other processes.
<code>F_RWDNY</code>	Set a file share reservation to deny read and write access to other processes.
<code>F_NODNY</code>	Do not deny read or write access to any other process.

RETURN VALUES

Upon successful completion, the value returned depends on `cmd` as follows:

<code>F_DUPFD</code>	A new file descriptor.
<code>F_GETFD</code>	Value of flags defined in <code><fcntl.h></code> . The return value will not be negative.
<code>F_SETFD</code>	Value other than <code>-1</code> .
<code>F_GETFL</code>	Value of file status flags and access modes. The return value will not be negative.
<code>F_SETFL</code>	Value other than <code>-1</code> .
<code>F_GETOWN</code>	Value of the socket owner process or process group; this will not be <code>-1</code> .
<code>F_SETOWN</code>	Value other than <code>-1</code> .
<code>F_FREESP</code>	Value of <code>0</code> .
<code>F_GETLK</code>	Value other than <code>-1</code> .
<code>F_GETLK64</code>	Value other than <code>-1</code> .
<code>F_SETLK</code>	Value other than <code>-1</code> .

fcntl(2)

F_SETLK64	Value other than -1.
F_SETLKW	Value other than -1.
F_SETLKW64	Value other than -1.
F_SHARE	Value other than -1.
F_UNSHARE	Value other than -1.

Otherwise, -1 is returned and `errno` is set to indicate the error.

ERRORS

The `fcntl()` function will fail if:

EAGAIN	The <i>cmd</i> argument is <code>F_SETLK</code> or <code>F_SETLK64</code>, the type of lock (<i>l_type</i>) is a shared (<code>F_RDLCK</code>) or exclusive (<code>F_WRLCK</code>) lock, and the segment of a file to be locked is already exclusive-locked by another process; or the type is an exclusive lock and some portion of the segment of a file to be locked is already shared-locked or exclusive-locked by another process. The <i>cmd</i> argument is <code>F_FREESP</code>, the file exists, mandatory file/record locking is set, and there are outstanding record locks on the file; or the <i>cmd</i> argument is <code>F_SETLK</code>, <code>F_SETLK64</code>, <code>F_SETLKW</code>, or <code>F_SETLKW64</code>, mandatory file/record locking is set, and the file is currently being mapped to virtual memory using <code>mmap(2)</code>. The <i>cmd</i> argument is <code>F_SHARE</code> and <code>f_access</code> conflicts with an existing <code>f_deny</code> share reservation.
EBADF	The <i>fildev</i> argument is not a valid open file descriptor; or the <i>cmd</i> argument is <code>F_SETLK</code>, <code>F_SETLK64</code>, <code>F_SETLKW</code>, or <code>F_SETLKW64</code>, the type of lock, <i>l_type</i>, is a shared lock (<code>F_RDLCK</code>), and <i>fildev</i> is not a valid file descriptor open for reading; or the type of lock <i>l_type</i> is an exclusive lock (<code>F_WRLCK</code>) and <i>fildev</i> is not a valid file descriptor open for writing. The <i>cmd</i> argument is <code>F_FREESP</code> and <i>fildev</i> is not a valid file descriptor open for writing. The <i>cmd</i> argument is <code>F_DUP2FD</code>, and <i>arg</i> is negative or is not less than the current resource limit for <code>RLIMIT_NOFILE</code>. The <i>cmd</i> argument is <code>F_SHARE</code>, the <code>f_access</code> share reservation is for write access, and <i>fildev</i> is not a valid file descriptor open for writing. The <i>cmd</i> argument is <code>F_SHARE</code>, the <code>f_access</code> share reservation is for read access, and <i>fildev</i> is not a valid file descriptor open for reading.

EFAULT	The <i>cmd</i> argument is <code>F_GETLK</code>, <code>F_GETLK64</code>, <code>F_SETLK</code>, <code>F_SETLK64</code>, <code>F_SETLKW</code>, <code>F_SETLKW64</code>, or <code>F_FREESP</code> and the <i>arg</i> argument points to an illegal address. The <i>cmd</i> argument is <code>F_SHARE</code> or <code>F_UNSHARE</code> and <i>arg</i> points to an illegal address.
EINTR	The <i>cmd</i> argument is <code>F_SETLKW</code> or <code>F_SETLKW64</code> and the function was interrupted by a signal.
EINVAL	The <i>cmd</i> argument is invalid; or the <i>cmd</i> argument is <code>F_DUPFD</code> and <i>arg</i> is negative or greater than or equal to <code>OPEN_MAX</code>; or the <i>cmd</i> argument is <code>F_GETLK</code>, <code>F_GETLK64</code>, <code>F_SETLK</code>, <code>F_SETLK64</code>, <code>F_SETLKW</code>, or <code>F_SETLKW64</code> and the data pointed to by <i>arg</i> is not valid; or <i>fildev</i> refers to a file that does not support locking. The <i>cmd</i> argument is <code>F_UNSHARE</code> and a reservation with this <i>f_id</i> for this process does not exist.
EIO	An I/O error occurred while reading from or writing to the file system.
EMFILE	The <i>cmd</i> argument is <code>F_DUPFD</code> and either <code>OPEN_MAX</code> file descriptors are currently open in the calling process, or no file descriptors greater than or equal to <i>arg</i> are available.
ENOLCK	The <i>cmd</i> argument is <code>F_SETLK</code> , <code>F_SETLK64</code> , <code>F_SETLKW</code> , or <code>F_SETLKW64</code> and satisfying the lock or unlock request would result in the number of locked regions in the system exceeding a system-imposed limit.
ENOLINK	Either the <i>fildev</i> argument is on a remote machine and the link to that machine is no longer active; or the <i>cmd</i> argument is <code>F_FREESP</code> , the file is on a remote machine, and the link to that machine is no longer active.
EOVERFLOW	One of the values to be returned cannot be represented correctly. The <i>cmd</i> argument is <code>F_GETLK</code>, <code>F_SETLK</code>, or <code>F_SETLKW</code> and the smallest or, if <i>l_len</i> is non-zero, the largest, offset of any byte in the requested segment cannot be represented correctly in an object of type <code>off_t</code>. The <i>cmd</i> argument is <code>F_GETLK64</code>, <code>F_SETLK64</code>, or <code>F_SETLKW64</code> and the smallest or, if <i>l_len</i> is non-zero, the largest, offset of any byte in the requested segment cannot be represented correctly in an object of type <code>off64_t</code>.

The `fcntl()` function may fail if:

fcntl(2)

EAGAIN	The <i>cmd</i> argument is F_SETLK, F_SETLK64, F_SETLKW, or F_SETLKW64, and the file is currently being mapped to virtual memory using mmap(2).
EDEADLK	The <i>cmd</i> argument is F_SETLKW or F_SETLKW64, the lock is blocked by some lock from another process and putting the calling process to sleep, waiting for that lock to become free would cause a deadlock. The <i>cmd</i> argument is F_FREESP, mandatory record locking is enabled, O_NDELAY and O_NONBLOCK are clear and a deadlock condition was detected.

ATTRIBUTES See attributes(5) for descriptions of the following attributes:

ATTRIBUTE TYPE	ATTRIBUTE VALUE
MT-Level	Async-Signal Safe

SEE ALSO lockd(1M), chmod(2), close(2), creat(2), dup(2), exec(2), fork(2), mmap(2), open(2), pipe(2), read(2), sigaction(2), write(2), dup2(3C), attributes(5), fcntl(3HEAD)

System Interface Guide

NOTES In the past, the variable `errno` was set to `EACCES` rather than `EAGAIN` when a section of a file is already locked by another process. Therefore, portable application programs should expect and test for either value.

Advisory locks allow cooperating processes to perform consistent operations on files, but do not guarantee exclusive access. Files can be accessed without advisory locks, but inconsistencies may result. The network share locking protocol does not support the `f_deny` value of `F_COMPAT`. For network file systems, if `f_access` is `F_RDACC`, `f_deny` is mapped to `F_RDDNY`. Otherwise, it is mapped to `F_RWDNY`.

To prevent possible file corruption, the system may reject `mmap()` requests for advisory locked files, or it may reject advisory locking requests for mapped files. Applications that require a file be both locked and mapped should lock the entire file (`l_start` and `l_len` both set to 0). If a file is mapped, the system may reject an unlock request, resulting in a lock that does not cover the entire file.

If the file server crashes and has to be rebooted, the lock manager (see `lockd(1M)`) attempts to recover all locks that were associated with that server. If a lock cannot be reclaimed, the process that held the lock is issued a `SIGLOST` signal.

NAME	fork, fork1 – create a new process
SYNOPSIS	<pre>#include <sys/types.h> #include <unistd.h> pid_t fork(void); pid_t fork1(void);</pre>
DESCRIPTION	The <code>fork()</code> and <code>fork1()</code> functions create a new process. The new process (child process) is an exact copy of the calling process (parent process). The child process inherits the following attributes from the parent process: ■ real user ID, real group ID, effective user ID, effective group ID ■ environment ■ open file descriptors ■ close-on-exec flags (see <code>exec(2)</code>) ■ signal handling settings (that is, <code>SIG_DFL</code>, <code>SIG_IGN</code>, <code>SIG_HOLD</code>, function address) ■ supplementary group IDs ■ set-user-ID mode bit ■ set-group-ID mode bit ■ profiling on/off status ■ nice value (see <code>nice(2)</code>) ■ scheduler class (see <code>prctl(2)</code>) ■ all attached shared memory segments (see <code>shmop(2)</code>) ■ process group ID -- memory mappings (see <code>mmap(2)</code>) ■ session ID (see <code>exit(2)</code>) ■ current working directory ■ root directory ■ file mode creation mask (see <code>umask(2)</code>) ■ resource limits (see <code>getrlimit(2)</code>) ■ controlling terminal ■ saved user ID and group ID ■ task ID and project ID Scheduling priority and any per-process scheduling parameters that are specific to a given scheduling class may or may not be inherited according to the policy of that particular class (see <code>prctl(2)</code>). The child process differs from the parent process in the following ways: ■ The child process has a unique process ID which does not match any active process group ID.

fork(2)

- The child process has a different parent process ID (that is, the process ID of the parent process).
- The child process has its own copy of the parent's file descriptors and directory streams. Each of the child's file descriptors shares a common file pointer with the corresponding file descriptor of the parent.
- Each shared memory segment remains attached and the value of `shm_nattach` is incremented by 1.
- All `semadj` values are cleared (see `semop(2)`).
- Process locks, text locks, data locks, and other memory locks are not inherited by the child (see `plock(3C)` and `memcntl(2)`).
- The child process's `tms` structure is cleared: `tms_utime`, `stime`, `cutime`, and `cstime` are set to 0 (see `times(2)`).
- The child processes resource utilizations are set to 0; see `getrlimit(2)`. The `it_value` and `it_interval` values for the `ITIMER_REAL` timer are reset to 0; see `getitimer(2)`.
- The set of signals pending for the child process is initialized to the empty set.
- Timers created by `timer_create(3RT)` are not inherited by the child process.
- No asynchronous input or asynchronous output operations are inherited by the child.

Record locks set by the parent process are not inherited by the child process (see `fcntl(2)`).

Solaris Threads In applications that use the Solaris threads API rather than the POSIX threads API (applications linked with `-lthread` but not `-lpthread`), `fork()` duplicates in the child process all threads (see `thr_create(3THR)`) and LWPs in the parent process. The `fork1()` function duplicates only the calling thread (LWP) in the child process.

POSIX Threads In applications that use the POSIX threads API rather than the Solaris threads API (applications linked with `-lpthread`, whether or not linked with `-lthread`), a call to `fork()` is like a call to `fork1()`, which replicates only the calling thread. There is no call that forks a child with all threads and LWPs duplicated in the child.

Note that if a program is linked with both libraries (`-lthread` and `-lpthread`), the POSIX semantic of `fork()` prevails.

fork() Safety If a Solaris threads application calls `fork1()` or a POSIX threads application calls `fork()`, and the child does more than simply call `exec()`, there is a possibility of deadlock occurring in the child. The application should use `pthread_atfork(3THR)` to ensure safety with respect to this deadlock. A Solaris threads application must explicitly link with `-lpthread` to access `pthread_atfork()`. Should there be any outstanding mutexes throughout the process, the application should call `pthread_atfork()` to wait for and acquire those mutexes prior to calling `fork()` or `fork1()`. See "MT-Level of Libraries" on the `attributes(5)` manual page.

RETURN VALUES	Upon successful completion, <code>fork()</code> and <code>fork1()</code> return 0 to the child process and return the process ID of the child process to the parent process. Otherwise, <code>(pid_t)-1</code> is returned to the parent process, no child process is created, and <code>errno</code> is set to indicate the error.				
ERRORS	The <code>fork()</code> function will fail if: EAGAIN The system-imposed limit on the total number of processes under execution by a single user has been exceeded; or the total amount of system memory available is temporarily insufficient to duplicate this process. ENOMEM There is not enough swap space.				
ATTRIBUTES	See <code>attributes(5)</code> for descriptions of the following attributes: <table border="1"> <thead> <tr> <th>ATTRIBUTE TYPE</th><th>ATTRIBUTE VALUE</th></tr> </thead> <tbody> <tr> <td>MT-Level</td><td><code>fork()</code> is Async-Signal-Safe</td></tr> </tbody> </table>	ATTRIBUTE TYPE	ATTRIBUTE VALUE	MT-Level	<code>fork()</code> is Async-Signal-Safe
ATTRIBUTE TYPE	ATTRIBUTE VALUE				
MT-Level	<code>fork()</code> is Async-Signal-Safe				
SEE ALSO	<code>alarm(2)</code> , <code>exec(2)</code> , <code>exit(2)</code> , <code>fcntl(2)</code> , <code>getitimer(2)</code> , <code>getrlimit(2)</code> , <code>memcntl(2)</code> , <code>mmap(2)</code> , <code>nice(2)</code> , <code>prctl(2)</code> , <code>ptrace(2)</code> , <code>semop(2)</code> , <code>shmop(2)</code> , <code>times(2)</code> , <code>umask(2)</code> , <code>wait(2)</code> , <code>exit(3C)</code> , <code>plock(3C)</code> , <code>pthread_atfork(3THR)</code> , <code>signal(3C)</code> , <code>system(3C)</code> , <code>thr_create(3THR)</code> , <code>timer_create(3RT)</code> , <code>attributes(5)</code> , <code>standards(5)</code>				
NOTES	An applications should call <code>_exit()</code> rather than <code>exit(3C)</code> if it cannot <code>execve()</code>, since <code>exit()</code> will flush and close standard I/O channels and thereby corrupt the parent process's standard I/O data structures. Using <code>exit(3C)</code> will flush buffered data twice. See <code>exit(2)</code>. The thread (or LWP) in the child that calls <code>fork1()</code> must not depend on any resources held by threads (or LWPs) that no longer exist in the child. In particular, locks held by these threads (or LWPs) will not be released. In a multithreaded process, <code>fork()</code> or <code>fork1()</code> can cause blocking system calls to be interrupted and return with an <code>EINTR</code> error. The <code>fork()</code> and <code>fork1()</code> functions suspend all threads in the process before proceeding. Threads that are executing in the kernel and are in an uninterruptible wait cannot be suspended immediately and therefore cause a delay before <code>fork()</code> and <code>fork1()</code> can complete. During this delay, since all other threads will have already been suspended, the process will appear "hung."				

fpathconf(2)

<

2. If *path* or *fildes* does not refer to a terminal file, it is unspecified whether an implementation supports an association of the variable name with the specified file.
3. If *path* or *fildes* refers to a directory, the value returned applies to filenames within the directory.
4. If *path* or *fildes* does not refer to a directory, it is unspecified whether an implementation supports an association of the variable name with the specified file.
5. If *path* or *fildes* refers to a directory, the value returned is the maximum length of a relative pathname when the specified directory is the working directory.
6. If *path* refers to a FIFO, or *fildes* refers to a pipe or FIFO, the value returned applies to the referenced object. If *path* or *fildes* refers to a directory, the value returned applies to any FIFO that exists or can be created within the directory. If *path* or *fildes* refers to any other type of file, it is unspecified whether an implementation supports an association of the variable name with the specified file.
7. If *path* or *fildes* refers to a directory, the value returned applies to any files, other than directories, that exist or can be created within the directory.
8. If *path* or *fildes* refers to a directory, it is unspecified whether an implementation supports an association of the variable name with the specified file.

RETURN VALUES

If *name* is an invalid value, both `pathconf()` and `fpathconf()` return `-1` and `errno` is set to indicate the error.

If the variable corresponding to *name* has no limit for the *path* or file descriptor, both `pathconf()` and `fpathconf()` return `-1` without changing `errno`. If the implementation needs to use *path* to determine the value of *name* and the implementation does not support the association of *name* with the file specified by *path*, or if the process did not have appropriate privileges to query the *appropriate privileges* file specified by *path*, or *path* does not exist, `pathconf()` returns `-1` and `errno` is set to indicate the error.

If the implementation needs to use *fildes* to determine the value of *name* and the implementation does not support the association of *name* with the file specified by *fildes*, or if *fildes* is an invalid file descriptor, `fpathconf()` will return `-1` and `errno` is set to indicate the error.

Otherwise `pathconf()` or `fpathconf()` returns the current variable value for the file or directory without changing `errno`. The value returned will not be more restrictive than the corresponding value available to the application when it was compiled with the implementation's `<limits.h>` or `<unistd.h>`.

ERRORS

The `pathconf()` function will fail if:

- | | |
|--------|---|
| EINVAL | The value of <i>name</i> is not valid. |
| ELOOP | Too many symbolic links were encountered in resolving <i>path</i> . |

fpathconf(2)

The `pathconf()` function may fail if:

EACCES	Search permission is denied for a component of the path prefix.
EINVAL	The implementation does not support an association of the variable <i>name</i> with the specified file.
ENAMETOOLONG	The length of the <i>path</i> argument exceeds <code>PATH_MAX</code> or a pathname component is longer than <code>NAME_MAX</code> .
ENAMETOOLONG	Pathname resolution of a symbolic link produced an intermediate result whose length exceeds <code>PATH_MAX</code> .
ENOENT	A component of <i>path</i> does not name an existing file or <i>path</i> is an empty string.
ENOTDIR	A component of the path prefix is not a directory.

The `fpathconf()` function will fail if:

EINVAL	The value of <i>name</i> is not valid.
--------	--

The `fpathconf()` function may fail if:

EBADF	The <i>filides</i> argument is not a valid file descriptor.
EINVAL	The implementation does not support an association of the variable <i>name</i> with the specified file.

ATTRIBUTES See `attributes(5)` for descriptions of the following attributes:

ATTRIBUTE TYPE	ATTRIBUTE VALUE
MT-Level	<code>pathconf()</code> is Async-Signal-Safe

SEE ALSO `sysconf(3C)`, `limits(4)`, `attributes(5)`, `standards(5)`

NAME	getacct, putacct, wracct – get, put, or write extended accounting data				
SYNOPSIS	<pre>#include <sys/exacct.h> size_t getacct(idtype_t idtype, id_t id, void *buf, size_t bufsize); int putacct(idtype_t idtype, id_t id, void *buf, size_t bufsize, int flags); int wracct(idtype_t idtype, id_t id, int flags);</pre>				
DESCRIPTION	These functions provide access to the extended accounting facility. The <code>getacct()</code> function returns extended accounting buffers from the kernel for currently executing tasks and processes. The resulting data buffer can be examined using the functions of the extended accounting library, <code>libexacct(3LIB)</code>. The <code>putacct()</code> function provides privileged processes the ability to tag accounting records with additional data specific to that process. For instance, a queuing facility may wish to record which queue a given task or process was submitted to prior to running. The <i>flags</i> field determines whether the contents of <i>buf</i> should be treated as raw data or as an embedded <code>exacct</code> structure, as described in <code><sys/exacct.h></code>. The use of an inappropriate flag or the inclusion of corrupt <code>exacct</code> data will likely corrupt the enclosing <code>exacct</code> file. The <code>wracct()</code> function requests the kernel to write, given its internal state of resource usage, the appropriate data for the specified task or process. The <i>flags</i> field determines whether a partial (<code>EW_PARTIAL</code>) or interval record (<code>EW_INTERVAL</code>) is written. These functions require root privilege, as they allow inquiry or reporting relevant to system tasks and processes other than the invoking process. The <code>putacct()</code> and <code>wracct()</code> functions also cause the kernel to write records to the system's extended accounting files.				
RETURN VALUES	The <code>getacct()</code> function returns the number of bytes required to represent the extended accounting record for the requested system task or process. If <i>bufsize</i> exceeds the returned size, <i>buf</i> will contain a valid accounting record buffer. If <i>bufsize</i> is less than the return value, <i>buf</i> will contain the first <i>bufsize</i> bytes of the record. In the event of failure, <code>-1</code> is returned and <code>errno</code> is set to indicate the error. The <code>putacct()</code> and <code>wracct()</code> functions return <code>0</code> if the record was successfully written. Otherwise, <code>-1</code> is returned and <code>errno</code> is set to indicate the error.				
ERRORS	The <code>getacct()</code>, <code>putacct()</code>, and <code>wracct()</code> functions will fail if: <table> <tr> <td><code>EINVAL</code></td><td>The <i>idtype</i> argument was not <code>P_TASKID</code> or <code>P_PID</code>.</td></tr> <tr> <td><code>ENOSPC</code></td><td>The file system containing the extended accounting file is full. The <code>wracct()</code> or <code>putacct()</code> function will fail if the record size would exceed the amount of space remaining on the file system.</td></tr> </table>	<code>EINVAL</code>	The <i>idtype</i> argument was not <code>P_TASKID</code> or <code>P_PID</code> .	<code>ENOSPC</code>	The file system containing the extended accounting file is full. The <code>wracct()</code> or <code>putacct()</code> function will fail if the record size would exceed the amount of space remaining on the file system.
<code>EINVAL</code>	The <i>idtype</i> argument was not <code>P_TASKID</code> or <code>P_PID</code> .				
<code>ENOSPC</code>	The file system containing the extended accounting file is full. The <code>wracct()</code> or <code>putacct()</code> function will fail if the record size would exceed the amount of space remaining on the file system.				

getacct(2)

ENOTACTIVE	The extended accounting facility for the requested <code>idtype_t</code> is not active. Either <code>putacct()</code> attempted to write a task record when the task accounting file was unset, or <code>getacct()</code> attempted to retrieve accounting data for a process when extended process accounting was inactive.
EPERM	The invoking process lacks sufficient permission to perform the request operation.
ERSCH	The <i>id</i> argument does not refer to a presently active system task ID or process ID.

ATTRIBUTES See `attributes(5)` for descriptions of the following attributes:

ATTRIBUTE TYPE	ATTRIBUTE VALUE
MT-Level	Async-Signal-Safe

SEE ALSO `libexacct(3LIB)`, `attributes(5)`

NAME	getaudit, setaudit, getaudit_addr, setaudit_addr – get and set process audit information
SYNOPSIS	<pre>cc [<i>flag</i> ...] <i>file</i> ... -lbsm -lsocket -lnsl -lintl [<i>library</i> ...] #include <sys/param.h> #include <bsm/audit.h> int getaudit(struct auditinfo *<i>info</i>); int setaudit(struct auditinfo *<i>info</i>); int getaudit_addr(struct auditinfo_addr *<i>info</i>, int <i>length</i>); int setaudit_addr(struct auditinfo_addr *<i>info</i>, int <i>length</i>);</pre>
DESCRIPTION	The <code>getaudit()</code> function gets the audit ID, the preselection mask, the terminal ID and the audit session ID for the current process. Note that <code>getaudit()</code> may fail and return an E2BIG <code>errno</code> if the address field in the terminal ID is larger than 32 bits. In this case, <code>getaudit_addr()</code> should be used. The <code>setaudit()</code> function sets the audit ID, the preselection mask, the terminal ID and the audit session ID for the current process. The <code>getaudit_addr()</code> function returns a variable length <code>auditinfo_addr</code> structure that contains the audit ID, the preselection mask, the terminal ID, and the audit session ID for the current process. The terminal ID contains a size field that indicates the size of the network address. The <code>setaudit_addr()</code> function sets the audit ID, the preselection mask, the terminal ID, and the audit session ID for the current process. The values are taken from the variable length structure <code>auditinfo_addr</code>. The terminal ID contains a size field that indicates the size of the network address. The <code>auditinfo</code> structure is used to pass the process audit information and contains the following members: <pre>au_id_t ai_auid; /* audit user ID */ au_mask_t ai_mask; /* preselection mask */ au_tid_t ai_termid; /* terminal ID */ au_asid_t ai_asid; /* audit session ID */</pre> The <code>auditinfo_addr</code> structure is used to pass the process audit information and contains the following members: <pre>au_id_t ai_auid; / audit user ID / au_mask_t ai_mask; / preselection mask / au_tid_addr_t ai_termid; / terminal ID / au_asid_t ai_asid; / audit session ID /</pre>
RETURN VALUES	Upon successful completion, <code>getaudit()</code> and <code>setaudit()</code> return 0. Otherwise, -1 is returned and <code>errno</code> is set to indicate the error.
ERRORS	The <code>getaudit()</code> and <code>setaudit()</code> functions will fail if:

getaudit(2)

	EFAULT	The <i>info</i> parameter points outside the process's allocated address space.
	EPERM	The process's effective user ID is not super-user.
USAGE	Only processes with the effective user ID of the super-user may successfully execute these calls.	
SEE ALSO	bsmconv(1M), audit(2)	
NOTES	The functionality described in this man page is available only if the Basic Security Module (BSM) has been enabled. See bsmconv(1M) for more information.	

NAME	getaudit, setaudit – get and set user audit identity				
SYNOPSIS	<pre>cc [<i>flag</i> ...] <i>file</i> ... -lbsm -lsocket -lnsl -lintl [<i>library</i> ...] #include <sys/param.h> #include <bsm/audit.h> int getaudit(au_id_t *<i>au</i>); int setaudit(au_id_t *<i>au</i>);</pre>				
DESCRIPTION	The <code>getaudit()</code> function returns the audit user ID for the current process. This value is initially set at login time and inherited by all child processes. This value does not change when the real/effective user IDs change, so it can be used to identify the logged-in user even when running a <code>setuid</code> program. The audit user ID governs audit decisions for a process. The <code>setaudit()</code> function sets the audit user ID for the current process.				
RETURN VALUES	Upon successful completion, the <code>getaudit()</code> function returns the audit user ID of the current process on success. Otherwise, it returns <code>-1</code> and sets <code>errno</code> to indicate the error. Upon successful completion the <code>setaudit()</code> function returns <code>0</code>. Otherwise, <code>-1</code> is returned and <code>errno</code> is set to indicate the error.				
ERRORS	The <code>getaudit()</code> and <code>setaudit()</code> functions will fail if: <table> <tr> <td>EFAULT</td><td>The <i>au</i> argument points to an invalid address.</td></tr> <tr> <td>EPERM</td><td>The process's effective user ID is not super-user.</td></tr> </table>	EFAULT	The <i>au</i> argument points to an invalid address.	EPERM	The process's effective user ID is not super-user.
EFAULT	The <i>au</i> argument points to an invalid address.				
EPERM	The process's effective user ID is not super-user.				
USAGE	Only the super-user may successfully execute these calls.				
SEE ALSO	<code>bsmconv(1M)</code> , <code>audit(2)</code> , <code>getaudit(2)</code>				
NOTES	The functionality described in this man page is available only if the Basic Security Module (BSM) has been enabled. See <code>bsmconv(1M)</code> for more information. These system calls have been superseded by <code>getaudit()</code> and <code>setaudit()</code>.				

getcontext(2)

NAME	getcontext, setcontext – get and set current user context
SYNOPSIS	<pre>#include <ucontext.h> int getcontext (ucontext_t *ucp) ; int setcontext (const ucontext_t *ucp) ;</pre>
DESCRIPTION	The <code>getcontext()</code> function initializes the structure pointed to by <code>ucp</code> to the current user context of the calling process. The <code>ucontext_t</code> type that <code>ucp</code> points to defines the user context and includes the contents of the calling process' machine registers, the signal mask, and the current execution stack. The <code>setcontext()</code> function restores the user context pointed to by <code>ucp</code>. A successful call to <code>setcontext()</code> does not return; program execution resumes at the point specified by the <code>ucp</code> argument passed to <code>setcontext()</code>. The <code>ucp</code> argument should be created either by a prior call to <code>getcontext()</code>, or by being passed as an argument to a signal handler. If the <code>ucp</code> argument was created with <code>getcontext()</code>, program execution continues as if the corresponding call of <code>getcontext()</code> had just returned. If the <code>ucp</code> argument was created with <code>makecontext(3C)</code>, program execution continues with the function passed to <code>makecontext(3C)</code>. When that function returns, the process continues as if after a call to <code>setcontext()</code> with the <code>ucp</code> argument that was input to <code>makecontext(3C)</code>. If the <code>ucp</code> argument was passed to a signal handler, program execution continues with the program instruction following the instruction interrupted by the signal. If the <code>uc_link</code> member of the <code>ucontext_t</code> structure pointed to by the <code>ucp</code> argument is equal to 0, then this context is the main context, and the process will exit when this context returns. The effects of passing a <code>ucp</code> argument obtained from any other source are unspecified.
RETURN VALUES	On successful completion, <code>setcontext()</code> does not return and <code>getcontext()</code> returns 0. Otherwise, -1 is returned.
ERRORS	No errors are defined.
USAGE	When a signal handler is executed, the current user context is saved and a new context is created. If the process leaves the signal handler via <code>longjmp(3UCB)</code>, then it is unspecified whether the context at the time of the corresponding <code>setjmp(3UCB)</code> call is restored and thus whether future calls to <code>getcontext()</code> will provide an accurate representation of the current context, since the context restored by <code>longjmp(3UCB)</code> may not contain all the information that <code>setcontext()</code> requires. Signal handlers should use <code>siglongjmp(3C)</code> or <code>setcontext()</code> instead. Portable applications should not modify or access the <code>uc_mcontext</code> member of <code>ucontext_t</code>. A portable application cannot assume that context includes any process-wide static data, possibly including <code>errno</code>. Users manipulating contexts should take care to handle these explicitly when required.
SEE ALSO	<code>sigaction(2)</code> , <code>sigaltstack(2)</code> , <code>sigprocmask(2)</code> , <code>bsd_signal(3C)</code> , <code>makecontext(3C)</code> , <code>setjmp(3UCB)</code> , <code>sigsetjmp(3C)</code> , <code>ucontext(3HEAD)</code>

NAME	getdents – read directory entries and put in a file system independent format																
SYNOPSIS	<pre>#include <sys/dirent.h> int getdents(int <i>fildev</i>, struct dirent *<i>buf</i>, size_t <i>nbyte</i>);</pre>																
DESCRIPTION	The <code>getdents()</code> function attempts to read <i>nbyte</i> bytes from the directory associated with the file descriptor <i>fildev</i> and to format them as file system independent directory entries in the buffer pointed to by <i>buf</i>. Since the file system independent directory entries are of variable lengths, in most cases the actual number of bytes returned will be less than <i>nbyte</i>. The file system independent directory entry is specified by the <code>dirent</code> structure. See <code>dirent(3HEAD)</code>. On devices capable of seeking, <code>getdents()</code> starts at a position in the file given by the file pointer associated with <i>fildev</i>. Upon return from <code>getdents()</code>, the file pointer is incremented to point to the next directory entry.																
RETURN VALUES	Upon successful completion, a non-negative integer is returned indicating the number of bytes actually read. A return value of 0 indicates the end of the directory has been reached. Otherwise, -1 is returned and <code>errno</code> is set to indicate the error.																
ERRORS	The <code>getdents()</code> function will fail if: <table> <tr> <td>EBADF</td><td>The <i>fildev</i> argument is not a valid file descriptor open for reading.</td></tr> <tr> <td>EFAULT</td><td>The <i>buf</i> argument points to an illegal address.</td></tr> <tr> <td>EINVAL</td><td>The <i>nbyte</i> argument is not large enough for one directory entry.</td></tr> <tr> <td>EIO</td><td>An I/O error occurred while accessing the file system.</td></tr> <tr> <td>ENOENT</td><td>The current file pointer for the directory is not located at a valid entry.</td></tr> <tr> <td>ENOLINK</td><td>The <i>fildev</i> argument points to a remote machine and the link to that machine is no longer active.</td></tr> <tr> <td>ENOTDIR</td><td>The <i>fildev</i> argument is not a directory.</td></tr> <tr> <td>EOVERFLOW</td><td>The value of the <code>dirent</code> structure member <code>d_ino</code> or <code>d_off</code> cannot be represented in an <code>ino_t</code> or <code>off_t</code>.</td></tr> </table>	EBADF	The <i>fildev</i> argument is not a valid file descriptor open for reading.	EFAULT	The <i>buf</i> argument points to an illegal address.	EINVAL	The <i>nbyte</i> argument is not large enough for one directory entry.	EIO	An I/O error occurred while accessing the file system.	ENOENT	The current file pointer for the directory is not located at a valid entry.	ENOLINK	The <i>fildev</i> argument points to a remote machine and the link to that machine is no longer active.	ENOTDIR	The <i>fildev</i> argument is not a directory.	EOVERFLOW	The value of the <code>dirent</code> structure member <code>d_ino</code> or <code>d_off</code> cannot be represented in an <code>ino_t</code> or <code>off_t</code> .
EBADF	The <i>fildev</i> argument is not a valid file descriptor open for reading.																
EFAULT	The <i>buf</i> argument points to an illegal address.																
EINVAL	The <i>nbyte</i> argument is not large enough for one directory entry.																
EIO	An I/O error occurred while accessing the file system.																
ENOENT	The current file pointer for the directory is not located at a valid entry.																
ENOLINK	The <i>fildev</i> argument points to a remote machine and the link to that machine is no longer active.																
ENOTDIR	The <i>fildev</i> argument is not a directory.																
EOVERFLOW	The value of the <code>dirent</code> structure member <code>d_ino</code> or <code>d_off</code> cannot be represented in an <code>ino_t</code> or <code>off_t</code> .																
USAGE	The <code>getdents()</code> function was developed to implement the <code>readdir(3C)</code> function and should not be used for other purposes. The <code>getdents()</code> function has a transitional interface for 64-bit file offsets. See <code>lf64(5)</code>.																
SEE ALSO	<code>readdir(3C)</code> , <code>dirent(3HEAD)</code> , <code>lf64(5)</code>																

getgroups(2)

NAME	getgroups, setgroups – get or set supplementary group access list IDs								
SYNOPSIS	<pre>#include <unistd.h> int getgroups(int <i>gidsetsize</i>, gid_t *<i>grouplist</i>); int setgroups(int <i>ngroups</i>, const gid_t *<i>grouplist</i>);</pre>								
DESCRIPTION	The <code>getgroups()</code> function gets the current supplemental group access list of the calling process and stores the result in the array of group IDs specified by <i>grouplist</i>. This array has <i>gidsetsize</i> entries and must be large enough to contain the entire list. This list cannot be larger than <code>NGROUPS_MAX</code>. If <i>gidsetsize</i> equals 0, <code>getgroups()</code> will return the number of groups to which the calling process belongs without modifying the array pointed to by <i>grouplist</i>. The <code>setgroups()</code> function sets the supplementary group access list of the calling process from the array of group IDs specified by <i>grouplist</i>. The number of entries is specified by <i>ngroups</i> and can not be greater than <code>NGROUPS_MAX</code>.								
RETURN VALUES	Upon successful completion, <code>getgroups()</code> returns the number of supplementary group IDs set for the calling process and <code>setgroups()</code> returns 0. Otherwise, -1 is returned and <code>errno</code> is set to indicate the error.								
ERRORS	The <code>getgroups()</code> and <code>setgroups()</code> functions will fail if: <table><tr><td>EFAULT</td><td>A referenced part of the array pointed to by <i>grouplist</i> is an illegal address.</td></tr></table> The <code>getgroups()</code> function will fail if: <table><tr><td>EINVAL</td><td>The value of <i>gidsetsize</i> is non-zero and less than the number of supplementary group IDs set for the calling process.</td></tr></table> The <code>setgroups()</code> function will fail if: <table><tr><td>EINVAL</td><td>The value of <i>ngroups</i> is greater than <code>NGROUPS_MAX</code>.</td></tr><tr><td>EPERM</td><td>The effective user of the calling process is not super-user.</td></tr></table>	EFAULT	A referenced part of the array pointed to by <i>grouplist</i> is an illegal address.	EINVAL	The value of <i>gidsetsize</i> is non-zero and less than the number of supplementary group IDs set for the calling process.	EINVAL	The value of <i>ngroups</i> is greater than <code>NGROUPS_MAX</code> .	EPERM	The effective user of the calling process is not super-user.
EFAULT	A referenced part of the array pointed to by <i>grouplist</i> is an illegal address.								
EINVAL	The value of <i>gidsetsize</i> is non-zero and less than the number of supplementary group IDs set for the calling process.								
EINVAL	The value of <i>ngroups</i> is greater than <code>NGROUPS_MAX</code> .								
EPERM	The effective user of the calling process is not super-user.								
USAGE	This function may be invoked only by the super-user.								
ATTRIBUTES	See <code>attributes(5)</code> for descriptions of the following attributes: <table><thead><tr><th>ATTRIBUTE TYPE</th><th>ATTRIBUTE VALUE</th></tr></thead><tbody><tr><td>MT-Level</td><td>Async-Signal-Safe</td></tr></tbody></table>	ATTRIBUTE TYPE	ATTRIBUTE VALUE	MT-Level	Async-Signal-Safe				
ATTRIBUTE TYPE	ATTRIBUTE VALUE								
MT-Level	Async-Signal-Safe								
SEE ALSO	<code>groups(1)</code> , <code>chown(2)</code> , <code>getuid(2)</code> , <code>setuid(2)</code> , <code>getgrnam(3C)</code> , <code>initgroups(3C)</code> , <code>attributes(5)</code>								

NAME	getitimer, setitimer – get or set value of interval timer
SYNOPSIS	<pre>#include <sys/time.h> int getitimer(int <i>which</i>, struct itimerval <i>*value</i>); int setitimer(int <i>which</i>, const struct itimerval <i>*value</i>, struct itimerval <i>*ovalue</i>);</pre>
DESCRIPTION	The system provides each process with four interval timers, defined in <code>sys/time.h</code>. The <code>getitimer()</code> function stores the current value of the timer specified by <i>which</i> into the structure pointed to by <i>value</i>. The <code>setitimer()</code> function call sets the value of the timer specified by <i>which</i> to the value specified in the structure pointed to by <i>value</i>, and if <i>ovalue</i> is not NULL, stores the previous value of the timer in the structure pointed to by <i>ovalue</i>. A timer value is defined by the <code>itimerval</code> structure (see <code>gettimeofday(3C)</code>) for the definition of <code>timeval</code>), which includes the following members: <pre>struct timeval it_interval; /* timer interval */ struct timeval it_value; /* current value */</pre> The <code>it_value</code> member indicates the time to the next timer expiration. The <code>it_interval</code> member specifies a value to be used in reloading <code>it_value</code> when the timer expires. Setting <code>it_value</code> to 0 disables a timer, regardless of the value of <code>it_interval</code>. Setting <code>it_interval</code> to 0 disables a timer after its next expiration (assuming <code>it_value</code> is non-zero). Time values smaller than the resolution of the system clock are rounded up to the resolution of the system clock, except for <code>ITIMER_REALPROF</code>, whose values are rounded up to the resolution of the profiling clock. The four timers are as follows: ITIMER_REAL Decrements in real time. A <code>SIGALRM</code> signal is delivered when this timer expires. In the current and previous releases, when <code>setitimer(ITIMER_REAL, ...)</code> is called in a multithreaded process linked with <code>-lthread</code> (Solaris threads) or <code>-lpthread</code> (POSIX threads; see <code>standards(5)</code>), the resulting <code>SIGALRM</code> is sent to the bound thread that called <code>setitimer()</code>; <code>setitimer()</code> has a per-thread semantic when called from a bound thread. This semantic will become obsolete in a future release. The semantic will move to a per-process semantic, with the resulting <code>SIGALRM</code> being sent to the process. The <code>SIGALRM</code> so generated is not maskable on this bound thread by any signal masking function, <code>pthread_sigmask(3THR)</code>, <code>thr_sigsetmask(3THR)</code>, or <code>sigprocmask(2)</code>. This is a bug that will not be fixed, since the per-thread semantic will be discontinued in the next release. Also, calling this routine from an unbound thread is not guaranteed to work as in the case of bound threads. The resulting <code>SIGALRM</code> may be sent to some other thread (see <code>alarm(2)</code>). This is a bug and will not be fixed since the per-thread semantic is going to be discontinued.

getitimer(2)

Calling `setitimer(ITIMER_REAL, . . .)` from a process linked with `-lpthread` (POSIX threads) has the same behavior as Solaris threads described above, where a Solaris bound thread is the same as a POSIX thread in system scheduling scope and a Solaris unbound thread is the same as a POSIX thread in local scheduling scope.

Hence, for multithreaded (Solaris or POSIX) programs in the current and previous releases, the only reliable way to use the `ITIMER_REAL` flag is to call it from a bound thread which does not mask `SIGALRM` and to expect the `SIGALRM` to be delivered to this bound thread.

The current working of this flag is not being improved since some applications might depend on the current (slightly broken) semantic. When this semantic is discontinued in the future, it will be replaced with a per-process semantic, i.e. using this flag from any thread, bound or unbound, will result in the `SIGALRM` being sent to the process.

New MT applications should not use this flag, and should use `alarm(2)` instead.

ITIMER_VIRTUAL

Decrements in process virtual time. It runs only when the process is executing. A `SIGVTALRM` signal is delivered when it expires. (For multithreaded programs see the `WARNINGS` section below).

ITIMER_PROF

Decrements both in process virtual time and when the system is running on behalf of the process. It is designed to be used by interpreters in statistically profiling the execution of interpreted programs. Each time the `ITIMER_PROF` timer expires, the `SIGPROF` signal is delivered. Because this signal may interrupt in-progress functions, programs using this timer must be prepared to restart interrupted functions. (For multithreaded programs see the `WARNINGS` section below).

ITIMER_REALPROF

Decrements in real time. It is designed to be used for real-time profiling of multithreaded programs. Each time the `ITIMER_REALPROF` timer expires, one counter in a set of counters maintained by the system for each lightweight process (lwp) is incremented. The counter corresponds to the state of the lwp at the time of the timer tick. All lwps executing in user mode when the timer expires are interrupted into system mode. When each lwp resumes execution in user mode, if any of the elements in its set of counters are non-zero, the `SIGPROF` signal is delivered to the lwp. The `SIGPROF` signal is delivered before any other signal except `SIGKILL`. This signal does not interrupt any in-progress function. A `siginfo` structure, defined in `<sys/siginfo.h>`, is associated with the delivery of the `SIGPROF` signal, and includes the following members:

```
si_tstamp;      /* high resolution timestamp */
si_syscall;     /* current syscall */
si_nsysarg;     /* number of syscall arguments */
si_sysarg[ ];   /* actual syscall arguments */
si_fault;       /* last fault type */
si_faddr;       /* last fault address */
si_mstate[ ];   /* ticks in each microstate */
```


	The enumeration of microstates (indices into <code>si_mstate</code>) is defined in <code><sys/msacct.h></code> . (For multithreaded programs see the WARNINGS section below).				
RETURN VALUES	Upon successful completion, 0 is returned. Otherwise, -1 is returned and <code>errno</code> is set to indicate the error.				
ERRORS	The <code>getitimer()</code> and <code>setitimer()</code> functions will fail if: <code>EINVAL</code> The specified number of seconds is greater than 100,000,000, the number of microseconds is greater than or equal to 1,000,000, or the <i>which</i> argument is unrecognized. The <code>setitimer()</code> function will fail if: <code>EACCES</code> Either an unbound Solaris thread or a POSIX thread in local scheduling scope with a flag other than <code>ITIMER_REAL</code> called <code>setitimer()</code>.				
ATTRIBUTES	See <code>attributes(5)</code> for descriptions of the following attributes:				
	<table> <tr> <th>ATTRIBUTE TYPE</th><th>ATTRIBUTE VALUE</th></tr> <tr> <td>MT-Level</td><td>MT-Safe</td></tr> </table>	ATTRIBUTE TYPE	ATTRIBUTE VALUE	MT-Level	MT-Safe
ATTRIBUTE TYPE	ATTRIBUTE VALUE				
MT-Level	MT-Safe				
SEE ALSO	<code>alarm(2)</code> , <code>sigprocmask(2)</code> , <code>gettimeofday(3C)</code> , <code>pthread_attr_setscope(3THR)</code> , <code>pthread_sigmask(3THR)</code> , <code>sleep(3C)</code> , <code>sysconf(3C)</code> , <code>attributes(5)</code> , <code>standards(5)</code>				
WARNINGS	All flags to <code>setitimer()</code> other than <code>ITIMER_REAL</code> behave as documented only with “bound” threads. Their ability to mask the signal works only with bound threads. If the call is made using one of these flags from an unbound thread, the system call returns -1 and sets <code>errno</code> to <code>EACCES</code>. These behaviors are the same for bound or unbound POSIX threads. A POSIX thread with system-wide scope, created by the call <pre>pthread_attr_setscope(&attr, PTHREAD_SCOPE_SYSTEM);</pre> is equivalent to a Solaris bound thread. A POSIX thread with local process scope, created by the call <pre>pthread_attr_setscope(&attr, PTHREAD_SCOPE_PROCESS);</pre> is equivalent to a Solaris unbound thread.				
NOTES	The microseconds field should not be equal to or greater than one second. The <code>setitimer()</code> function is independent of the <code>alarm()</code> function.				

getitimer(2)

Do not use `setitimer(ITIMER_REAL)` with the `sleep()` routine. A `sleep(3C)` call wipes out knowledge of the user signal handler for `SIGALRM`.

The `ITIMER_PROF` and `ITIMER_REALPROF` timers deliver the same signal and have different semantics. They cannot be used together.

The granularity of the resolution of alarm time is platform-dependent.

NAME	getmsg, getpmsg – get next message off a stream
SYNOPSIS	<pre>#include <stropts.h> int getmsg(int <i>fildev</i>, struct strbuf *<i>ctlptr</i>, struct strbuf *<i>dataptr</i>, int *<i>flagsp</i>); int getpmsg(int <i>fildev</i>, struct strbuf *<i>ctlptr</i>, struct strbuf *<i>dataptr</i>, int *<i>bandp</i>, int *<i>flagsp</i>);</pre>
DESCRIPTION	The <code>getmsg()</code> function retrieves the contents of a message (see <code>intro(2)</code>) located at the stream head read queue from a STREAMS file, and places the contents into user specified buffer(s). The message must contain either a data part, a control part, or both. The data and control parts of the message are placed into separate buffers, as described below. The semantics of each part is defined by the STREAMS module that generated the message. The <code>getpmsg()</code> function behaved like <code>getmsg()</code>, but provides finer control over the priority of the messages received. Except where noted, all information pertaining to <code>getmsg()</code> also pertains to <code>getpmsg()</code>. The <i>fildev</i> argument specifies a file descriptor referencing an open stream. The <i>ctlptr</i> and <i>dataptr</i> arguments each point to a <code>strbuf</code> structure, which contains the following members: <pre>int maxlen; /* maximum buffer length */ int len; /* length of data */ char *buf; /* ptr to buffer */</pre> The <code>buf</code> member points to a buffer into which the data or control information is to be placed, and the <code>maxlen</code> member indicates the maximum number of bytes this buffer can hold. On return, the <code>len</code> member contains the number of bytes of data or control information actually received; 0 if there is a zero-length control or data part; or -1 if no data or control information is present in the message. The <i>flagsp</i> argument should point to an integer that indicates the type of message the user is able to receive, as described below. The <i>ctlptr</i> argument holds the control part from the message and the <i>dataptr</i> argument holds the data part from the message. If <i>ctlptr</i> (or <i>dataptr</i>) is NULL or the <code>maxlen</code> member is -1, the control (or data) part of the message is not processed and is left on the stream head read queue. If <i>ctlptr</i> (or <i>dataptr</i>) is not NULL and there is no corresponding control (or data) part of the messages on the stream head read queue, <code>len</code> is set to -1. If the <code>maxlen</code> member is set to 0 and there is a zero-length control (or data) part, that zero-length part is removed from the read queue and <code>len</code> is set to 0. If the <code>maxlen</code> member is set to 0 and there are more than zero bytes of control (or data) information, that information is left on the read queue and <code>len</code> is set to 0. If the <code>maxlen</code> member in <i>ctlptr</i> or <i>dataptr</i> is less than, respectively, the control or data part of the message, <code>maxlen</code> bytes are retrieved. In this case, the remainder of the message is left on the stream head read queue and a non-zero return value is provided, as described below under RETURN VALUES.

getmsg(2)

By default, `getmsg()` processes the first available message on the stream head read queue. A user may, however, choose to retrieve only high priority messages by setting the integer pointed to by *flagsp* to `RS_HIPRI`. In this case, `getmsg()` processes the next message only if it is a high priority message.

If the integer pointed to by *flagsp* is 0, `getmsg()` retrieves any message available on the stream head read queue. In this case, on return, the integer pointed to by *flagsp* will be set to `RS_HIPRI` if a high priority message was retrieved, or to 0 otherwise.

For `getpmsg()`, the *flagsp* argument points to a bitmask with the following mutually-exclusive flags defined: `MSG_HIPRI`, `MSG_BAND`, and `MSG_ANY`. Like `getmsg()`, `getpmsg()` processes the first available message on the stream head read queue. A user may choose to retrieve only high-priority messages by setting the integer pointed to by *flagsp* to `MSG_HIPRI` and the integer pointed to by *bandp* to 0. In this case, `getpmsg()` will only process the next message if it is a high-priority message. In a similar manner, a user may choose to retrieve a message from a particular priority band by setting the integer pointed to by *flagsp* to `MSG_BAND` and the integer pointed to by *bandp* to the priority band of interest. In this case, `getpmsg()` will only process the next message if it is in a priority band equal to, or greater than, the integer pointed to by *bandp*, or if it is a high-priority message. If a user just wants to get the first message off the queue, the integer pointed to by *flagsp* should be set to `MSG_ANY` and the integer pointed to by *bandp* should be set to 0. On return, if the message retrieved was a high-priority message, the integer pointed to by *flagsp* will be set to `MSG_HIPRI` and the integer pointed to by *bandp* will be set to 0. Otherwise, the integer pointed to by *flagsp* will be set to `MSG_BAND` and the integer pointed to by *bandp* will be set to the priority band of the message.

If `O_NDELAY` and `O_NONBLOCK` are clear, `getmsg()` blocks until a message of the type specified by *flagsp* is available on the stream head read queue. If `O_NDELAY` or `O_NONBLOCK` has been set and a message of the specified type is not present on the read queue, `getmsg()` fails and sets `errno` to `EAGAIN`.

If a hangup occurs on the stream from which messages are to be retrieved, `getmsg()` continues to operate normally, as described above, until the stream head read queue is empty. Thereafter, it returns 0 in the `len` member of *ctlptr* and *dataptr*.

RETURN VALUES

Upon successful completion, a non-negative value is returned. A return value of 0 indicates that a full message was read successfully. A return value of `MORECTL` indicates that more control information is waiting for retrieval. A return value of `MOREDATA` indicates that more data are waiting for retrieval. A return value of `MORECTL | MOREDATA` indicates that both types of information remain. Subsequent `getmsg()` calls retrieve the remainder of the message. However, if a message of higher priority has been received by the stream head read queue, the next call to `getmsg()` will retrieve that higher priority message before retrieving the remainder of the previously received partial message.

ERRORS

The `getmsg()` and `getpmsg()` functions will fail if:

<code>EAGAIN</code>	The <code>O_NDELAY</code> or <code>O_NONBLOCK</code> flag is set and no messages are available.
---------------------	---

getmsg(2)

EBADF	The <i>fildes</i> argument is not a valid file descriptor open for reading.
EBADMSG	Queued message to be read is not valid for <code>getmsg</code> .
EFAULT	The <i>ctlptr</i> , <i>dataptr</i> , <i>bandp</i> , or <i>flagsp</i> argument points to an illegal address.
EINTR	A signal was caught during the execution of the <code>getmsg</code> function.
EINVAL	An illegal value was specified in <i>flagsp</i> , or the stream referenced by <i>fildes</i> is linked under a multiplexor.
ENOSTR	A stream is not associated with <i>fildes</i> .

The `getmsg()` function can also fail if a STREAMS error message had been received at the stream head before the call to `getmsg()`. The error returned is the value contained in the STREAMS error message.

SEE ALSO `intro(2)`, `poll(2)`, `putmsg(2)`, `read(2)`, `write(2)`

STREAMS Programming Guide

getpid(2)

NAME	getpid, getpgrp, getppid, getpgid – get process, process group, and parent process IDs						
SYNOPSIS	<pre>#include <unistd.h> pid_t getpid(void); pid_t getpgrp(void); pid_t getppid(void); pid_t getpgid(pid_t <i>pid</i>);</pre>						
DESCRIPTION	The <code>getpid()</code> function returns the process ID of the calling process. The <code>getpgrp()</code> function returns the process group ID of the calling process. The <code>getppid()</code> function returns the parent process ID of the calling process. The <code>getpgid()</code> function returns the process group ID of the process whose process ID is equal to <i>pid</i>, or the process group ID of the calling process, if <i>pid</i> is equal to 0.						
RETURN VALUES	Upon successful completion, these functions return the process group ID. Otherwise, <code>getpgid()</code> returns <code>(pid_t)-1</code> and sets <code>errno</code> to indicate the error.						
ERRORS	The <code>getpgid()</code> function will fail if: <table> <tr> <td>EPERM</td><td>The process whose process ID is equal to <i>pid</i> is not in the same session as the calling process, and the implementation does not allow access to the process group ID of that process from the calling process.</td></tr> <tr> <td>ESRCH</td><td>There is no process with a process ID equal to <i>pid</i>.</td></tr> </table> The <code>getpgid()</code> function may fail if: <table> <tr> <td>EINVAL</td><td>The value of the <i>pid</i> argument is invalid.</td></tr> </table>	EPERM	The process whose process ID is equal to <i>pid</i> is not in the same session as the calling process, and the implementation does not allow access to the process group ID of that process from the calling process.	ESRCH	There is no process with a process ID equal to <i>pid</i> .	EINVAL	The value of the <i>pid</i> argument is invalid.
EPERM	The process whose process ID is equal to <i>pid</i> is not in the same session as the calling process, and the implementation does not allow access to the process group ID of that process from the calling process.						
ESRCH	There is no process with a process ID equal to <i>pid</i> .						
EINVAL	The value of the <i>pid</i> argument is invalid.						
ATTRIBUTES	See <code>attributes(5)</code> for descriptions of the following attributes: <table> <tr> <th>ATTRIBUTE TYPE</th><th>ATTRIBUTE VALUE</th></tr> <tr> <td>MT-Level</td><td>Async-Signal-Safe</td></tr> </table>	ATTRIBUTE TYPE	ATTRIBUTE VALUE	MT-Level	Async-Signal-Safe		
ATTRIBUTE TYPE	ATTRIBUTE VALUE						
MT-Level	Async-Signal-Safe						
SEE ALSO	<code>intro(3)</code> , <code>exec(2)</code> , <code>fork(2)</code> , <code>getsid(2)</code> , <code>setpgid(2)</code> , <code>setpgrp(2)</code> , <code>setsid(2)</code> , <code>signal(3C)</code> , <code>attributes(5)</code>						

NAME	getrlimit, setrlimit – control maximum system resource consumption										
SYNOPSIS	<pre>#include <sys/resource.h> int getrlimit(int <i>resource</i>, struct rlimit <i>*rlp</i>); int setrlimit(int <i>resource</i>, const struct rlimit <i>*rlp</i>);</pre>										
DESCRIPTION	Limits on the consumption of a variety of system resources by a process and each process it creates may be obtained with the <code>getrlimit()</code> and set with <code>setrlimit()</code> functions. Each call to either <code>getrlimit()</code> or <code>setrlimit()</code> identifies a specific resource to be operated upon as well as a resource limit. A resource limit is a pair of values: one specifying the current (soft) limit, the other a maximum (hard) limit. Soft limits may be changed by a process to any value that is less than or equal to the hard limit. A process may (irreversibly) lower its hard limit to any value that is greater than or equal to the soft limit. Only a process with an effective user ID of super-user can raise a hard limit. Both hard and soft limits can be changed in a single call to <code>setrlimit()</code> subject to the constraints described above. Limits may have an “infinite” value of <code>RLIM_INFINITY</code>. The <i>rlp</i> argument is a pointer to <code>struct rlimit</code> that includes the following members: <pre>rlim_t rlim_cur; /* current (soft) limit */ rlim_t rlim_max; /* hard limit */</pre> The type <code>rlim_t</code> is an arithmetic data type to which objects of type <code>int</code>, <code>size_t</code>, and <code>off_t</code> can be cast without loss of information. The possible resources, their descriptions, and the actions taken when the current limit is exceeded are summarized as follows: <table> <tr> <td><code>RLIMIT_CORE</code></td><td>The maximum size of a core file in bytes that may be created by a process. A limit of 0 will prevent the creation of a core file. The writing of a core file will terminate at this size.</td></tr> <tr> <td><code>RLIMIT_CPU</code></td><td>The maximum amount of CPU time in seconds used by a process. This is a soft limit only. The <code>SIGXCPU</code> signal is sent to the process. If the process is holding or ignoring <code>SIGXCPU</code>, the behavior is scheduling class defined.</td></tr> <tr> <td><code>RLIMIT_DATA</code></td><td>The maximum size of a process’s heap in bytes. The <code>brk(2)</code> function will fail with <code>errno</code> set to <code>ENOMEM</code>.</td></tr> <tr> <td><code>RLIMIT_FSIZE</code></td><td>The maximum size of a file in bytes that may be created by a process. A limit of 0 will prevent the creation of a file. The <code>SIGXFSZ</code> signal is sent to the process. If the process is holding or ignoring <code>SIGXFSZ</code>, continued attempts to increase the size of a file beyond the limit will fail with <code>errno</code> set to <code>EFBIG</code>.</td></tr> <tr> <td><code>RLIMIT_NOFILE</code></td><td>One more than the maximum value that the system may assign to a newly created descriptor. This limit constrains the number of file descriptors that a process may create.</td></tr> </table>	<code>RLIMIT_CORE</code>	The maximum size of a core file in bytes that may be created by a process. A limit of 0 will prevent the creation of a core file. The writing of a core file will terminate at this size.	<code>RLIMIT_CPU</code>	The maximum amount of CPU time in seconds used by a process. This is a soft limit only. The <code>SIGXCPU</code> signal is sent to the process. If the process is holding or ignoring <code>SIGXCPU</code> , the behavior is scheduling class defined.	<code>RLIMIT_DATA</code>	The maximum size of a process’s heap in bytes. The <code>brk(2)</code> function will fail with <code>errno</code> set to <code>ENOMEM</code> .	<code>RLIMIT_FSIZE</code>	The maximum size of a file in bytes that may be created by a process. A limit of 0 will prevent the creation of a file. The <code>SIGXFSZ</code> signal is sent to the process. If the process is holding or ignoring <code>SIGXFSZ</code> , continued attempts to increase the size of a file beyond the limit will fail with <code>errno</code> set to <code>EFBIG</code> .	<code>RLIMIT_NOFILE</code>	One more than the maximum value that the system may assign to a newly created descriptor. This limit constrains the number of file descriptors that a process may create.
<code>RLIMIT_CORE</code>	The maximum size of a core file in bytes that may be created by a process. A limit of 0 will prevent the creation of a core file. The writing of a core file will terminate at this size.										
<code>RLIMIT_CPU</code>	The maximum amount of CPU time in seconds used by a process. This is a soft limit only. The <code>SIGXCPU</code> signal is sent to the process. If the process is holding or ignoring <code>SIGXCPU</code> , the behavior is scheduling class defined.										
<code>RLIMIT_DATA</code>	The maximum size of a process’s heap in bytes. The <code>brk(2)</code> function will fail with <code>errno</code> set to <code>ENOMEM</code> .										
<code>RLIMIT_FSIZE</code>	The maximum size of a file in bytes that may be created by a process. A limit of 0 will prevent the creation of a file. The <code>SIGXFSZ</code> signal is sent to the process. If the process is holding or ignoring <code>SIGXFSZ</code> , continued attempts to increase the size of a file beyond the limit will fail with <code>errno</code> set to <code>EFBIG</code> .										
<code>RLIMIT_NOFILE</code>	One more than the maximum value that the system may assign to a newly created descriptor. This limit constrains the number of file descriptors that a process may create.										

getrlimit(2)

RLIMIT_STACK	The maximum size of a process's stack in bytes. The system will not automatically grow the stack beyond this limit. Within a process, <code>setrlimit()</code> will increase the limit on the size of your stack, but will not move current memory segments to allow for that growth. To guarantee that the process stack can grow to the limit, the limit must be altered prior to the execution of the process in which the new stack size is to be used. Within a multithreaded process, <code>setrlimit()</code> has no impact on the stack size limit for the calling thread if the calling thread is not the main thread. A call to <code>setrlimit()</code> for <code>RLIMIT_STACK</code> impacts only the main thread's stack, and should be made only from the main thread, if at all. The <code>SIGSEGV</code> signal is sent to the process. If the process is holding or ignoring <code>SIGSEGV</code>, or is catching <code>SIGSEGV</code> and has not made arrangements to use an alternate stack (see <code>sigaltstack(2)</code>), the disposition of <code>SIGSEGV</code> will be set to <code>SIG_DFL</code> before it is sent.
RLIMIT_VMEM	The maximum size of a process's mapped address space in bytes. If this limit is exceeded, the <code>brk(2)</code> and <code>mmap(2)</code> functions will fail with <code>errno</code> set to <code>ENOMEM</code>. In addition, the automatic stack growth will fail with the effects outlined above.
RLIMIT_AS	This is the maximum size of a process's total available memory, in bytes. If this limit is exceeded, the <code>brk(2)</code>, <code>malloc(3C)</code>, <code>mmap(2)</code> and <code>sbrk(2)</code> functions will fail with <code>errno</code> set to <code>ENOMEM</code>. In addition, the automatic stack growth will fail with the effects outlined above.

Because limit information is stored in the per-process information, the shell builtin `ulimit` command must directly execute this system call if it is to affect all future processes created by the shell.

The value of the current limit of the following resources affect these implementation defined parameters:

Limit	Implementation Defined Constant
RLIMIT_FSIZE	FCHR_MAX
RLIMIT_NOFILE	OPEN_MAX

When using the `getrlimit()` function, if a resource limit can be represented correctly in an object of type `rlim_t`, then its representation is returned; otherwise, if the value of the resource limit is equal to that of the corresponding saved hard limit, the value returned is `RLIM_SAVED_MAX`; otherwise the value returned is `RLIM_SAVED_CUR`.

When using the `setrlimit()` function, if the requested new limit is `RLIM_INFINITY`, the new limit will be "no limit"; otherwise if the requested new limit is `RLIM_SAVED_MAX`, the new limit will be the corresponding saved hard limit; otherwise, if the requested new limit is `RLIM_SAVED_CUR`, the new limit will be the corresponding saved soft limit; otherwise, the new limit will be the requested value. In addition, if the corresponding saved limit can be represented correctly in an object of type `rlim_t`, then it will be overwritten with the new limit.

The result of setting a limit to `RLIM_SAVED_MAX` or `RLIM_SAVED_CUR` is unspecified unless a previous call to `getrlimit()` returned that value as the soft or hard limit for the corresponding resource limit.

A limit whose value is greater than `RLIM_INFINITY` is permitted.

The `exec` family of functions also cause resource limits to be saved. See `exec(2)`.

RETURN VALUES Upon successful completion, `getrlimit()` and `setrlimit()` return 0. Otherwise, these functions return -1 and set `errno` to indicate the error.

ERRORS The `getrlimit()` and `setrlimit()` functions will fail if:

<code>EFAULT</code>	The <i>rlp</i> argument points to an illegal address.
<code>EINVAL</code>	An invalid <i>resource</i> was specified; or in a <code>setrlimit()</code> call, the new <code>rlim_cur</code> exceeds the new <code>rlim_max</code> .
<code>EPERM</code>	The limit specified to <code>setrlimit()</code> would have raised the maximum limit value, and the effective user of the calling process is not super-user.

The `setrlimit()` function may fail if:

<code>EINVAL</code>	The limit specified cannot be lowered because current usage is already higher than the limit.
---------------------	---

USAGE The `getrlimit()` and `setrlimit()` functions have transitional interfaces for 64-bit file offsets. See `lf64(5)`.

SEE ALSO `brk(2)`, `exec(2)`, `fork(2)`, `open(2)`, `sigaltstack(2)`, `ulimit(2)`, `getdtablesize(3C)`, `malloc(3C)`, `signal(3C)`, `sysconf(3C)`, `lf64(5)`, `signal(3HEAD)`

getsid(2)

NAME	getsid – get process group ID of session leader				
SYNOPSIS	<pre>#include <unistd.h> pid_t getsid(pid_t pid) ;</pre>				
DESCRIPTION	The <code>getsid()</code> function obtains the process group ID of the process that is the session leader of the process specified by <i>pid</i> . If <i>pid</i> is <code>(pid_t) 0</code> , it specifies the calling process.				
RETURN VALUES	Upon successful completion, <code>getsid()</code> returns the process group ID of the session leader of the specified process. Otherwise, it returns <code>(pid_t)-1</code> and sets <code>errno</code> to indicate the error.				
ERRORS	The <code>getsid()</code> function will fail if: <table><tr><td>EPERM</td><td>The process specified by <i>pid</i> is not in the same session as the calling process, and the implementation does not allow access to the process group ID of the session leader of that process from the calling process.</td></tr><tr><td>ESRCH</td><td>There is no process with a process ID equal to <i>pid</i>.</td></tr></table>	EPERM	The process specified by <i>pid</i> is not in the same session as the calling process, and the implementation does not allow access to the process group ID of the session leader of that process from the calling process.	ESRCH	There is no process with a process ID equal to <i>pid</i> .
EPERM	The process specified by <i>pid</i> is not in the same session as the calling process, and the implementation does not allow access to the process group ID of the session leader of that process from the calling process.				
ESRCH	There is no process with a process ID equal to <i>pid</i> .				
SEE ALSO	<code>exec(2)</code> , <code>fork(2)</code> , <code>getpid(2)</code> , <code>getpgid(2)</code> , <code>setpgid(2)</code> , <code>setsid(2)</code>				

NAME	getuid, geteuid, getgid, getegid – get real user, effective user, real group, and effective group IDs				
SYNOPSIS	<pre>#include <sys/types.h> #include <unistd.h> uid_t getuid(void); uid_t geteuid(void); gid_t getgid(void); gid_t getegid(void);</pre>				
DESCRIPTION	The <code>getuid()</code> function returns the real user ID of the calling process. The real user ID identifies the person who is logged in. The <code>geteuid()</code> function returns the effective user ID of the calling process. The effective user ID gives the process various permissions during execution of “set-user-ID” mode processes which use <code>getuid()</code> to determine the real user ID of the process that invoked them. The <code>getgid()</code> function returns the real group ID of the calling process. The <code>getegid()</code> function returns the effective group ID of the calling process.				
ATTRIBUTES	See <code>attributes(5)</code> for descriptions of the following attributes: <table border="1"> <thead> <tr> <th>ATTRIBUTE TYPE</th><th>ATTRIBUTE VALUE</th></tr> </thead> <tbody> <tr> <td>MT-Level</td><td>Async-Signal-Safe</td></tr> </tbody> </table>	ATTRIBUTE TYPE	ATTRIBUTE VALUE	MT-Level	Async-Signal-Safe
ATTRIBUTE TYPE	ATTRIBUTE VALUE				
MT-Level	Async-Signal-Safe				
SEE ALSO	<code>intro(2)</code> , <code>setuid(2)</code> , <code>attributes(5)</code>				

ioctl(2)

NAME	ioctl – control device												
SYNOPSIS	<pre>#include <unistd.h> #include <stropts.h> int ioctl(int <i>fildev</i>, int <i>request</i>, /* <i>arg</i> */ ...);</pre>												
DESCRIPTION	The <code>ioctl()</code> function performs a variety of control functions on devices and STREAMS. For non-STREAMS files, the functions performed by this call are device-specific control functions. The <i>request</i> argument and an optional third argument with varying type are passed to the file designated by <i>fildev</i> and are interpreted by the device driver. For STREAMS files, specific functions are performed by the <code>ioctl()</code> function as described in <code>streamio(7I)</code>. The <i>fildev</i> argument is an open file descriptor that refers to a device. The <i>request</i> argument selects the control function to be performed and depends on the device being addressed. The <i>arg</i> argument represents a third argument that has additional information that is needed by this specific device to perform the requested function. The data type of <i>arg</i> depends upon the particular control request, but it is either an <code>int</code> or a pointer to a device-specific data structure. In addition to device-specific and STREAMS functions, generic functions are provided by more than one device driver (for example, the general terminal interface.) See <code>termio(7I)</code>.												
RETURN VALUES	Upon successful completion, the value returned depends upon the device control function, but must be a non-negative integer. Otherwise, <code>-1</code> is returned and <code>errno</code> is set to indicate the error.												
ERRORS	The <code>ioctl()</code> function will fail for any type of file if: <table> <tr> <td>EBADF</td><td>The <i>fildev</i> argument is not a valid open file descriptor.</td></tr> <tr> <td>EINTR</td><td>A signal was caught during the execution of the <code>ioctl()</code> function.</td></tr> <tr> <td>EINVAL</td><td>The STREAM or multiplexer referenced by <i>fildev</i> is linked (directly or indirectly) downstream from a multiplexer.</td></tr> </table> The <code>ioctl()</code> function will also fail if the device driver detects an error. In this case, the error is passed through <code>ioctl()</code> without change to the caller. A particular driver might not have all of the following error cases. Under the following conditions, requests to device drivers may fail and set <code>errno</code> to indicate the error <table> <tr> <td>EFAULT</td><td>The <i>request</i> argument requires a data transfer to or from a buffer pointed to by <i>arg</i>, but <i>arg</i> points to an illegal address.</td></tr> <tr> <td>EINVAL</td><td>The <i>request</i> or <i>arg</i> argument is not valid for this device.</td></tr> <tr> <td>EIO</td><td>Some physical I/O error has occurred.</td></tr> </table>	EBADF	The <i>fildev</i> argument is not a valid open file descriptor.	EINTR	A signal was caught during the execution of the <code>ioctl()</code> function.	EINVAL	The STREAM or multiplexer referenced by <i>fildev</i> is linked (directly or indirectly) downstream from a multiplexer.	EFAULT	The <i>request</i> argument requires a data transfer to or from a buffer pointed to by <i>arg</i> , but <i>arg</i> points to an illegal address.	EINVAL	The <i>request</i> or <i>arg</i> argument is not valid for this device.	EIO	Some physical I/O error has occurred.
EBADF	The <i>fildev</i> argument is not a valid open file descriptor.												
EINTR	A signal was caught during the execution of the <code>ioctl()</code> function.												
EINVAL	The STREAM or multiplexer referenced by <i>fildev</i> is linked (directly or indirectly) downstream from a multiplexer.												
EFAULT	The <i>request</i> argument requires a data transfer to or from a buffer pointed to by <i>arg</i> , but <i>arg</i> points to an illegal address.												
EINVAL	The <i>request</i> or <i>arg</i> argument is not valid for this device.												
EIO	Some physical I/O error has occurred.												

ENOLINK	The <i>fildev</i> argument is on a remote machine and the link to that machine is no longer active.
ENOTTY	The <i>fildev</i> argument is not associated with a STREAMS device that accepts control functions.
ENXIO	The <i>request</i> and <i>arg</i> arguments are valid for this device driver, but the service requested can not be performed on this particular subdevice.
ENODEV	The <i>fildev</i> argument refers to a valid STREAMS device, but the corresponding device driver does not support the <code>ioctl()</code> function.

STREAMS errors are described in `streamio(7I)`.

SEE ALSO `streamio(7I)`, `termio(7I)`

kill(2)

NAME	kill – send a signal to a process or a group of processes						
SYNOPSIS	<pre>#include <sys/types.h> #include <signal.h> int kill(pid_t pid, int sig);</pre>						
DESCRIPTION	The <code>kill()</code> function sends a signal to a process or a group of processes. The process or group of processes to which the signal is to be sent is specified by <i>pid</i>. The signal that is to be sent is specified by <i>sig</i> and is either one from the list given in <code>signal(3HEAD)</code>, or 0. If <i>sig</i> is 0 (the null signal), error checking is performed but no signal is actually sent. This can be used to check the validity of <i>pid</i>. The real or effective user ID of the sending process must match the real or saved (from one of functions in the <i>exec</i> family, see <code>exec(2)</code>) user ID of the receiving process unless the effective user ID of the sending process is super-user, (see <code>intro(3)</code>), or <i>sig</i> is <code>SIGCONT</code> and the sending process has the same session ID as the receiving process. If <i>pid</i> is greater than 0, <i>sig</i> will be sent to the process whose process ID is equal to <i>pid</i>. If <i>pid</i> is negative but not <code>(pid_t)-1</code>, <i>sig</i> will be sent to all processes whose process group ID is equal to the absolute value of <i>pid</i> and for which the process has permission to send a signal. If <i>pid</i> is 0, <i>sig</i> will be sent to all processes excluding special processes (see <code>intro(3)</code>) whose process group ID is equal to the process group ID of the sender. If <i>pid</i> is <code>(pid_t)-1</code> and the effective user ID of the sender is not super-user, <i>sig</i> will be sent to all processes excluding special processes whose real user ID is equal to the effective user ID of the sender. If <i>pid</i> is <code>(pid_t)-1</code> and the effective user ID of the sender is super-user, <i>sig</i> will be sent to all processes excluding special processes.						
RETURN VALUES	Upon successful completion, 0 is returned. Otherwise, -1 is returned, no signal is sent, and <code>errno</code> is set to indicate the error.						
ERRORS	The <code>kill()</code> function will fail if: <table> <tr> <td><code>EINVAL</code></td><td>The <i>sig</i> argument is not a valid signal number.</td></tr> <tr> <td><code>EPERM</code></td><td>The <i>sig</i> argument is <code>SIGKILL</code> and the <i>pid</i> argument is <code>(pid_t)1</code> (that is, the calling process does not have permission to send the signal to any of the processes specified by <i>pid</i>); or the effective user of the calling process does not match the real or saved user and is not super-user, and the calling process is not sending <code>SIGCONT</code> to a process that shares the same session ID.</td></tr> <tr> <td><code>ESRCH</code></td><td>No process or process group can be found corresponding to that specified by <i>pid</i>.</td></tr> </table>	<code>EINVAL</code>	The <i>sig</i> argument is not a valid signal number.	<code>EPERM</code>	The <i>sig</i> argument is <code>SIGKILL</code> and the <i>pid</i> argument is <code>(pid_t)1</code> (that is, the calling process does not have permission to send the signal to any of the processes specified by <i>pid</i>); or the effective user of the calling process does not match the real or saved user and is not super-user, and the calling process is not sending <code>SIGCONT</code> to a process that shares the same session ID.	<code>ESRCH</code>	No process or process group can be found corresponding to that specified by <i>pid</i> .
<code>EINVAL</code>	The <i>sig</i> argument is not a valid signal number.						
<code>EPERM</code>	The <i>sig</i> argument is <code>SIGKILL</code> and the <i>pid</i> argument is <code>(pid_t)1</code> (that is, the calling process does not have permission to send the signal to any of the processes specified by <i>pid</i>); or the effective user of the calling process does not match the real or saved user and is not super-user, and the calling process is not sending <code>SIGCONT</code> to a process that shares the same session ID.						
<code>ESRCH</code>	No process or process group can be found corresponding to that specified by <i>pid</i> .						
USAGE	The <code>sigsend(2)</code> function provides a more versatile way to send signals to processes.						

kill(2)

ATTRIBUTES See attributes(5) for descriptions of the following attributes:

ATTRIBUTE TYPE	ATTRIBUTE VALUE
MT-Level	Async-Signal-Safe

SEE ALSO kill(1), intro(3), exec(2), getpid(2), getsid(2), setpgrp(2), sigaction(2), sigsend(2), signal(3C), attributes(5), signal(3HEAD)

link(2)

NAME	link – link to a file																
SYNOPSIS	<pre>#include <unistd.h> int link(const char *existing, const char *new);</pre>																
DESCRIPTION	The <code>link()</code> function creates a new link (directory entry) for the existing file and increments its link count by one. The <i>existing</i> argument points to a path name naming an existing file. The <i>new</i> argument points to a pathname naming the new directory entry to be created. To create hard links, both files must be on the same file system. Both the old and the new link share equal access and rights to the underlying object. The super-user may make multiple links to a directory. Unless the caller is the super-user, the file named by <i>existing</i> must not be a directory. Upon successful completion, <code>link()</code> marks for update the <code>st_ctime</code> field of the file. Also, the <code>st_ctime</code> and <code>st_mtime</code> fields of the directory that contains the new entry are marked for update.																
RETURN VALUES	Upon successful completion, 0 is returned. Otherwise, -1 is returned, no link is created, and <code>errno</code> is set to indicate the error.																
ERRORS	The <code>link()</code> function will fail if: <table> <tr> <td>EACCES</td><td>A component of either path prefix denies search permission, or the requested link requires writing in a directory with a mode that denies write permission.</td></tr> <tr> <td>EDQUOT</td><td>The directory where the entry for the new link is being placed cannot be extended because the user's quota of disk blocks on that file system has been exhausted.</td></tr> <tr> <td>EEXIST</td><td>The link named by <i>new</i> exists.</td></tr> <tr> <td>EFAULT</td><td>The <i>existing</i> or <i>new</i> argument points to an illegal address.</td></tr> <tr> <td>EINTR</td><td>A signal was caught during the execution of the <code>link()</code> function.</td></tr> <tr> <td>ELOOP</td><td>Too many symbolic links were encountered in translating <i>path</i>.</td></tr> <tr> <td>EMLINK</td><td>The maximum number of links to a file would be exceeded.</td></tr> <tr> <td>ENAMETOOLONG</td><td>The length of the <i>existing</i> or <i>new</i> argument exceeds <code>PATH_MAX</code>, or the length of a <i>existing</i> or <i>new</i> component exceeds <code>NAME_MAX</code> while <code>_POSIX_NO_TRUNC</code> is in effect.</td></tr> </table>	EACCES	A component of either path prefix denies search permission, or the requested link requires writing in a directory with a mode that denies write permission.	EDQUOT	The directory where the entry for the new link is being placed cannot be extended because the user's quota of disk blocks on that file system has been exhausted.	EEXIST	The link named by <i>new</i> exists.	EFAULT	The <i>existing</i> or <i>new</i> argument points to an illegal address.	EINTR	A signal was caught during the execution of the <code>link()</code> function.	ELOOP	Too many symbolic links were encountered in translating <i>path</i> .	EMLINK	The maximum number of links to a file would be exceeded.	ENAMETOOLONG	The length of the <i>existing</i> or <i>new</i> argument exceeds <code>PATH_MAX</code> , or the length of a <i>existing</i> or <i>new</i> component exceeds <code>NAME_MAX</code> while <code>_POSIX_NO_TRUNC</code> is in effect.
EACCES	A component of either path prefix denies search permission, or the requested link requires writing in a directory with a mode that denies write permission.																
EDQUOT	The directory where the entry for the new link is being placed cannot be extended because the user's quota of disk blocks on that file system has been exhausted.																
EEXIST	The link named by <i>new</i> exists.																
EFAULT	The <i>existing</i> or <i>new</i> argument points to an illegal address.																
EINTR	A signal was caught during the execution of the <code>link()</code> function.																
ELOOP	Too many symbolic links were encountered in translating <i>path</i> .																
EMLINK	The maximum number of links to a file would be exceeded.																
ENAMETOOLONG	The length of the <i>existing</i> or <i>new</i> argument exceeds <code>PATH_MAX</code> , or the length of a <i>existing</i> or <i>new</i> component exceeds <code>NAME_MAX</code> while <code>_POSIX_NO_TRUNC</code> is in effect.																

link(2)

ENOENT	The <i>existing</i> or <i>new</i> argument is a null pathname; a component of either path prefix does not exist; or the file named by <i>existing</i> does not exist.
ENOLINK	The <i>existing</i> or <i>new</i> argument points to a remote machine and the link to that machine is no longer active.
ENOSPC	The directory that would contain the link cannot be extended.
ENOTDIR	A component of either path prefix is not a directory.
EPERM	The file named by <i>existing</i> is a directory and the effective user of the calling process is not super-user.
EROFS	The requested link requires writing in a directory on a read-only file system.
EXDEV	The link named by <i>new</i> and the file named by <i>existing</i> are on different logical devices (file systems).

ATTRIBUTES See attributes(5) for descriptions of the following attributes:

ATTRIBUTE TYPE	ATTRIBUTE VALUE
MT-Level	Async-Signal-Safe

SEE ALSO symlink(2), unlink(2), attributes(5)

llseek(2)

NAME	llseek – move extended read/write file pointer						
SYNOPSIS	<pre>#include <sys/types.h> #include <unistd.h> offset_t llseek(int <i>fildes</i>, offset_t <i>offset</i>, int <i>whence</i>);</pre>						
DESCRIPTION	The <code>llseek()</code> function sets the 64-bit extended file pointer associated with the open file descriptor specified by <i>fildes</i> as follows: ■ If <i>whence</i> is <code>SEEK_SET</code>, the pointer is set to <i>offset</i> bytes. ■ If <i>whence</i> is <code>SEEK_CUR</code>, the pointer is set to its current location plus <i>offset</i>. ■ If <i>whence</i> is <code>SEEK_END</code>, the pointer is set to the size of the file plus <i>offset</i>. Although each file has a 64-bit file pointer associated with it, some existing file system types (such as <code>tmpfs</code>) do not support the full range of 64-bit offsets. In particular, on such file systems, non-device files remain limited to offsets of less than two gigabytes. Device drivers may support offsets of up to 1024 gigabytes for device special files. Some devices are incapable of seeking. The value of the file pointer associated with such a device is undefined.						
RETURN VALUES	Upon successful completion, <code>llseek()</code> returns the resulting pointer location as measured in bytes from the beginning of the file. Remote file descriptors are the only ones that allow negative file pointers. Otherwise, <code>-1</code> is returned, the file pointer remains unchanged, and <code>errno</code> is set to indicate the error.						
ERRORS	The <code>llseek()</code> function will fail if: <table><tr><td>EBADF</td><td>The <i>fildes</i> argument is not an open file descriptor.</td></tr><tr><td>EINVAL</td><td>The <i>whence</i> argument is not <code>SEEK_SET</code>, <code>SEEK_CUR</code>, or <code>SEEK_END</code>; the <i>offset</i> argument is not a valid offset for this file system type; or the <i>fildes</i> argument is not a remote file descriptor and the resulting file pointer would be negative.</td></tr><tr><td>ESPIPE</td><td>The <i>fildes</i> argument is associated with a pipe or FIFO.</td></tr></table>	EBADF	The <i>fildes</i> argument is not an open file descriptor.	EINVAL	The <i>whence</i> argument is not <code>SEEK_SET</code> , <code>SEEK_CUR</code> , or <code>SEEK_END</code> ; the <i>offset</i> argument is not a valid offset for this file system type; or the <i>fildes</i> argument is not a remote file descriptor and the resulting file pointer would be negative.	ESPIPE	The <i>fildes</i> argument is associated with a pipe or FIFO.
EBADF	The <i>fildes</i> argument is not an open file descriptor.						
EINVAL	The <i>whence</i> argument is not <code>SEEK_SET</code> , <code>SEEK_CUR</code> , or <code>SEEK_END</code> ; the <i>offset</i> argument is not a valid offset for this file system type; or the <i>fildes</i> argument is not a remote file descriptor and the resulting file pointer would be negative.						
ESPIPE	The <i>fildes</i> argument is associated with a pipe or FIFO.						
SEE ALSO	<code>creat(2)</code> , <code>dup(2)</code> , <code>fcntl(2)</code> , <code>lseek(2)</code> , <code>open(2)</code>						

NAME	lseek – move read/write file pointer								
SYNOPSIS	<pre>#include <sys/types.h> #include <unistd.h> off_t lseek(int <i>fildev</i>, off_t <i>offset</i>, int <i>whence</i>);</pre>								
DESCRIPTION	The <code>lseek()</code> function sets the file pointer associated with the open file descriptor specified by <i>fildev</i> as follows: ■ If <i>whence</i> is <code>SEEK_SET</code>, the pointer is set to <i>offset</i> bytes. ■ If <i>whence</i> is <code>SEEK_CUR</code>, the pointer is set to its current location plus <i>offset</i>. ■ If <i>whence</i> is <code>SEEK_END</code>, the pointer is set to the size of the file plus <i>offset</i>. The symbolic constants <code>SEEK_SET</code>, <code>SEEK_CUR</code>, and <code>SEEK_END</code> are defined in the header <code><unistd.h></code>. Some devices are incapable of seeking. The value of the file pointer associated with such a device is undefined. The <code>lseek()</code> function allows the file pointer to be set beyond the existing data in the file. If data are later written at this point, subsequent reads in the gap between the previous end of data and the newly written data will return bytes of value 0 until data are written into the gap. If <i>fildev</i> is a remote file descriptor and <i>offset</i> is negative, <code>lseek()</code> returns the file pointer even if it is negative. The <code>lseek()</code> function will not, by itself, extend the size of a file.								
RETURN VALUES	Upon successful completion, the resulting offset, as measured in bytes from the beginning of the file, is returned. Otherwise, <code>(off_t)-1</code> is returned, the file offset remains unchanged, and <code>errno</code> is set to indicate the error.								
ERRORS	The <code>lseek()</code> function will fail if: <table> <tr> <td><code>EBADF</code></td><td>The <i>fildev</i> argument is not an open file descriptor.</td></tr> <tr> <td><code>EINVAL</code></td><td>The <i>whence</i> argument is not <code>SEEK_SET</code>, <code>SEEK_CUR</code>, or <code>SEEK_END</code>; or the <i>fildev</i> argument is not a remote file descriptor and the resulting file pointer would be negative.</td></tr> <tr> <td><code>EOVERFLOW</code></td><td>The resulting file offset would be a value which cannot be represented correctly in an object of type <code>off_t</code> for regular files.</td></tr> <tr> <td><code>ESPIPE</code></td><td>The <i>fildev</i> argument is associated with a pipe, a FIFO, or a socket.</td></tr> </table>	<code>EBADF</code>	The <i>fildev</i> argument is not an open file descriptor.	<code>EINVAL</code>	The <i>whence</i> argument is not <code>SEEK_SET</code> , <code>SEEK_CUR</code> , or <code>SEEK_END</code> ; or the <i>fildev</i> argument is not a remote file descriptor and the resulting file pointer would be negative.	<code>EOVERFLOW</code>	The resulting file offset would be a value which cannot be represented correctly in an object of type <code>off_t</code> for regular files.	<code>ESPIPE</code>	The <i>fildev</i> argument is associated with a pipe, a FIFO, or a socket.
<code>EBADF</code>	The <i>fildev</i> argument is not an open file descriptor.								
<code>EINVAL</code>	The <i>whence</i> argument is not <code>SEEK_SET</code> , <code>SEEK_CUR</code> , or <code>SEEK_END</code> ; or the <i>fildev</i> argument is not a remote file descriptor and the resulting file pointer would be negative.								
<code>EOVERFLOW</code>	The resulting file offset would be a value which cannot be represented correctly in an object of type <code>off_t</code> for regular files.								
<code>ESPIPE</code>	The <i>fildev</i> argument is associated with a pipe, a FIFO, or a socket.								
USAGE	The <code>lseek()</code> function has a transitional interface for 64-bit file offsets. See <code>1f64(5)</code>. In multithreaded applications, using <code>lseek()</code> in conjunction with a <code>read(2)</code> or <code>write(2)</code> call on a file descriptor shared by more than one thread is not an atomic operation. To ensure atomicity, use <code>pread()</code> or <code>pwrite()</code>.								

lseek(2)

ATTRIBUTES See `attributes(5)` for descriptions of the following attributes:

ATTRIBUTE TYPE	ATTRIBUTE VALUE
MT-Level	Async-Signal-Safe

SEE ALSO `creat(2)`, `dup(2)`, `fcntl(2)`, `open(2)`, `read(2)`, `write(2)`, `attributes(5)`, `lf64(5)`

NAME	<code>_lwp_cond_signal</code> , <code>_lwp_cond_broadcast</code> – signal a condition variable
SYNOPSIS	<pre>#include <sys/lwp.h> int _lwp_cond_signal(lwp_cond_t *cvp); int _lwp_cond_broadcast(lwp_cond_t *cvp);</pre>
DESCRIPTION	The <code>_lwp_cond_signal()</code> function unblocks one LWP that is blocked on the LWP condition variable pointed to by <i>cvp</i>. The <code>_lwp_cond_broadcast()</code> function unblocks all LWPs that are blocked on the LWP condition variable pointed to by <i>cvp</i>. If no LWPs are blocked on the LWP condition variable, then <code>_lwp_cond_signal()</code> and <code>_lwp_cond_broadcast()</code> have no effect. Both functions should be called under the protection of the same LWP mutex lock that is used with the LWP condition variable being signaled. Otherwise, the condition variable may be signalled between the test of the associated condition and blocking in <code>_lwp_cond_wait()</code>. This can cause an infinite wait.
RETURN VALUES	Upon successful completion, 0 is returned. A non-zero value indicates an error.
ERRORS	The <code>_lwp_cond_signal()</code> and <code>_lwp_cond_broadcast()</code> functions will fail if: EINVAL The <i>cvp</i> argument points to an invalid LWP condition variable. EFAULT The <i>cvp</i> argument points to an invalid address.
SEE ALSO	<code>_lwp_cond_wait(2)</code> , <code>_lwp_mutex_lock(2)</code>

`_lwp_cond_wait(2)`

NAME	<code>_lwp_cond_wait</code> , <code>_lwp_cond_timedwait</code> , <code>_lwp_cond_reltimedwait</code> – wait on a condition variable				
SYNOPSIS	<pre>#include <sys/lwp.h> int _lwp_cond_wait(lwp_cond_t *cwp, lwp_mutex_t *mp); int _lwp_cond_timedwait(lwp_cond_t *cwp, lwp_mutex_t *mp, timestruc_t *abstime); int _lwp_cond_reltimedwait(lwp_cond_t *cwp, lwp_mutex_t *mp, timestruc_t *reltime);</pre>				
DESCRIPTION	These functions are used to wait for the occurrence of a condition represented by an LWP condition variable. LWP condition variables must be initialized to 0 before use. The <code>_lwp_cond_wait()</code> function atomically releases the LWP mutex pointed to by <i>mp</i> and causes the calling LWP to block on the LWP condition variable pointed to by <i>cwp</i>. The blocked LWP may be awakened by <code>_lwp_cond_signal(2)</code>, <code>_lwp_cond_broadcast(2)</code>, or when interrupted by delivery of a signal. Any change in value of a condition associated with the condition variable cannot be inferred by the return of <code>_lwp_cond_wait()</code> and any such condition must be re-evaluated. The <code>_lwp_cond_timedwait()</code> function is similar to <code>_lwp_cond_wait()</code>, except that the calling LWP will not block past the time of day specified by <i>abstime</i>. If the time of day becomes greater than <i>abstime</i>, <code>_lwp_cond_timedwait()</code> returns with the error code <code>ETIME</code>. The <code>_lwp_cond_reltimedwait()</code> function is similar to <code>_lwp_cond_wait()</code>, except that the calling LWP will not block past the relative time specified by <i>reltime</i>. If the time of day becomes greater than the starting time of day plus <i>reltime</i>, <code>_lwp_cond_reltimedwait()</code> returns with the error code <code>ETIME</code>. The <code>_lwp_cond_wait()</code>, <code>_lwp_cond_timedwait()</code>, and <code>_lwp_cond_reltimedwait()</code> functions always return with the mutex locked and owned by the calling lightweight process.				
RETURN VALUES	Upon successful completion, 0 is returned. A non-zero value indicates an error.				
ERRORS	If any of the following conditions are detected, <code>_lwp_cond_wait()</code>, <code>_lwp_cond_timedwait()</code>, and <code>_lwp_cond_reltimedwait()</code> fail and return the corresponding value: <table><tr><td><code>EINVAL</code></td><td>The <i>cwp</i> argument points to an invalid LWP condition variable or the <i>mp</i> argument points to an invalid LWP mutex.</td></tr><tr><td><code>EFAULT</code></td><td>The <i>mp</i>, <i>cwp</i>, or <i>abstime</i> argument points to an illegal address.</td></tr></table> If any of the following conditions occur, <code>_lwp_cond_wait()</code>, <code>_lwp_cond_timedwait()</code>, and <code>_lwp_cond_reltimedwait()</code> fail and return the corresponding value:	<code>EINVAL</code>	The <i>cwp</i> argument points to an invalid LWP condition variable or the <i>mp</i> argument points to an invalid LWP mutex.	<code>EFAULT</code>	The <i>mp</i> , <i>cwp</i> , or <i>abstime</i> argument points to an illegal address.
<code>EINVAL</code>	The <i>cwp</i> argument points to an invalid LWP condition variable or the <i>mp</i> argument points to an invalid LWP mutex.				
<code>EFAULT</code>	The <i>mp</i> , <i>cwp</i> , or <i>abstime</i> argument points to an illegal address.				

`_lwp_cond_wait(2)`

EINTR The call was interrupted by a signal or `fork(2)`.

If any of the following conditions occur, `_lwp_cond_timedwait()` and `_lwp_cond_reltimedwait()` fail and return the corresponding value:

ETIME The time specified *inabstime* or *reltime* has passed.

EXAMPLES

EXAMPLE 1 Use the `_lwp_cond_wait()` function in a loop testing some condition.

The `_lwp_cond_wait()` function is normally used in a loop testing some condition, as follows:

```
lwp_mutex_t m;
lwp_cond_t cv;
int cond;
(void) _lwp_mutex_lock(&m);
while (cond == FALSE) {
    (void) _lwp_cond_wait(&cv, &m);
}
(void) _lwp_mutex_unlock(&m);
```

EXAMPLE 2 Use the `_lwp_cond_timedwait()` function in a loop testing some condition.

The `_lwp_cond_timedwait()` function is also normally used in a loop testing some condition. It uses an absolute timeout value as follows:

```
timestruct_t to;
lwp_mutex_t m;
lwp_cond_t cv;
int cond, err;
(void) _lwp_mutex_lock(&m);
to.tv_sec = time(NULL) + TIMEOUT;
to.tv_nsec = 0;
while (cond == FALSE) {
    err = _lwp_cond_timedwait(&cv, &m, &to);
    if (err == ETIME) {
        /* timeout, do something */
        break;
    }
    SENDwhom}
}
(void) _lwp_mutex_unlock(&m);
```

This example sets a bound on the total wait time even though the `_lwp_cond_timedwait()` may return several times due to the condition being signalled or the wait being interrupted.

EXAMPLE 3 Use the `_lwp_cond_reltimedwait()` function in a loop testing some condition.

The `_lwp_cond_reltimedwait()` function is also normally used in a loop testing some condition. It uses a relative timeout value as follows:

```
timestruct_t to;
lwp_mutex_t m;
```

`_lwp_cond_wait(2)`

EXAMPLE 3 Use the `_lwp_cond_reltimedwait()` function in a loop testing some condition. *(Continued)*

```
lwp_cond_t cv;
int cond, err;
(void) _lwp_mutex_lock(&m);
while (cond == FALSE) {
    to.tv_sec = TIMEOUT;
    to.tv_nsec = 0;
    err = _lwp_cond_reltimedwait(&cv, &m, &to);
    if (err == ETIME) {
        /* timeout, do something */
        break;
    }
}
(void) _lwp_mutex_unlock(&m);
```

SEE ALSO `_lwp_cond_broadcast(2)`, `_lwp_cond_signal(2)`, `_lwp_kill(2)`,
`_lwp_mutex_lock(2)`, `fork(2)`, `kill(2)`

NAME	<code>_lwp_create</code> – create a new light-weight process						
SYNOPSIS	<pre>#include <sys/lwp.h> int _lwp_create(ucontext_t *contextp, unsigned long flags, lwpid_t *new_lwp) ;</pre>						
DESCRIPTION	The <code>_lwp_create()</code> function adds a lightweight process (LWP) to the current process. The <code>contextp</code> argument specifies the initial signal mask, stack, and machine context (including the program counter and stack pointer) for the new LWP. The new LWP inherits the scheduling class and priority of the caller. If <code>_lwp_create()</code> is successful and <code>new_lwp</code> is not null, the ID of the new LWP is stored in the location pointed to by <code>new_lwp</code>. The <code>flags</code> argument specifies additional attributes for the new LWP. The value in <code>flags</code> is constructed by the bitwise inclusive OR operation of the following values: <table><tr><td><code>LWP_DETACHED</code></td><td>The LWP is created detached.</td></tr><tr><td><code>LWP_SUSPENDED</code></td><td>The LWP is created suspended.</td></tr><tr><td><code>__LWP_ASLWP</code></td><td>The LWP created is the ASLWP (Asynchronous Signals LWP) (see <code>signal(3HEAD)</code>). The ASLWP should always be created with all signals blocked. If <code>__LWP_ASLWP</code> is specified, then the LWP created is the special, designated LWP that handles signals sent to a multithreaded process (ASLWP). There can be only one ASLWP in a multithreaded process, so the creation of another ASLWP will return <code>EINVAL</code>. It should never exit by way of <code>_lwp_exit(2)</code> or <code>exit(2)</code>. This is a reserved flag and should not be used by any user program. It is documented here for the sake of completion and not for use by an application.</td></tr></table> If <code>LWP_DETACHED</code> is specified, then the LWP is created in the <i>detached</i> state. Otherwise the LWP is created in the undetached state. The ID (and system resources) associated with a detached LWP can be automatically reclaimed when the LWP exits. The ID of an undetached LWP cannot be reclaimed until it exits and another LWP has reported its termination by way of <code>_lwp_wait(2)</code>. This allows the waiting LWP to determine that the waited for LWP has terminated and to reclaim any process resources that it was using. If <code>LWP_SUSPENDED</code> is specified, then the LWP is created in a suspended state. This allows the creator to change the LWP's inherited attributes before it starts to execute. The suspended LWP can only be resumed by way of <code>_lwp_continue(2)</code>. If <code>LWP_SUSPENDED</code> is not specified the LWP can begin to run immediately after it has been created.	<code>LWP_DETACHED</code>	The LWP is created detached.	<code>LWP_SUSPENDED</code>	The LWP is created suspended.	<code>__LWP_ASLWP</code>	The LWP created is the ASLWP (Asynchronous Signals LWP) (see <code>signal(3HEAD)</code>). The ASLWP should always be created with all signals blocked. If <code>__LWP_ASLWP</code> is specified, then the LWP created is the special, designated LWP that handles signals sent to a multithreaded process (ASLWP). There can be only one ASLWP in a multithreaded process, so the creation of another ASLWP will return <code>EINVAL</code> . It should never exit by way of <code>_lwp_exit(2)</code> or <code>exit(2)</code> . This is a reserved flag and should not be used by any user program. It is documented here for the sake of completion and not for use by an application.
<code>LWP_DETACHED</code>	The LWP is created detached.						
<code>LWP_SUSPENDED</code>	The LWP is created suspended.						
<code>__LWP_ASLWP</code>	The LWP created is the ASLWP (Asynchronous Signals LWP) (see <code>signal(3HEAD)</code>). The ASLWP should always be created with all signals blocked. If <code>__LWP_ASLWP</code> is specified, then the LWP created is the special, designated LWP that handles signals sent to a multithreaded process (ASLWP). There can be only one ASLWP in a multithreaded process, so the creation of another ASLWP will return <code>EINVAL</code> . It should never exit by way of <code>_lwp_exit(2)</code> or <code>exit(2)</code> . This is a reserved flag and should not be used by any user program. It is documented here for the sake of completion and not for use by an application.						
RETURN VALUES	Upon successful completion, 0 is returned. A non-zero value indicates an error.						
ERRORS	If any of the following conditions are detected, <code>_lwp_create()</code> fails and returns the corresponding value:						

`_lwp_create(2)`

EFAULT	Either the <i>context</i> parameter or the <i>new_lwp</i> parameter point to invalid addresses.
EAGAIN	A system limit is exceeded, (for example, too many LWP were created for this real user ID).
EINVAL	The <code>__LWP_ASLWP</code> flag was used to create more than one ASLWP in the process. There can be only one ASLWP within a process.

EXAMPLES

EXAMPLE 1 How a stack is allocated to a new LWP.

This example shows how a stack is allocated to a new LWP. The `_lwp_makecontext` () function is used to set up the *context* parameter so that the new LWP begins executing a function.

```
contextp = (ucontext_t *)malloc(sizeof(ucontext_t));
stackbase = malloc(stacksize);
sigprocmask(SIGSETMASK, NULL, &contextp->uc_sigmask);
_lwp_makecontext(contextp, func, arg, private, stackbase, stacksize);
error = _lwp_create(contextp, NULL, &new_lwp);
```

SEE ALSO

`_lwp_cond_timedwait(2)`, `_lwp_continue(2)`, `_lwp_exit(2)`,
`_lwp_makecontext(2)`, `_lwp_wait(2)`, `alarm(2)`, `exit(2)`, `poll(2)`, `sleep(3C)`,
`thr_create(3THR)`, `signal(3HEAD)`, `ucontext(3HEAD)`

NOTES

Applications should use bound threads rather than the `_lwp_*` () functions (see `thr_create(3THR)`). Using LWPs directly is not advised because libraries are only safe to use with threads, not LWPs.

In Solaris releases 2.5 through 7, the signal `SIGALRM` is defined to be per-process. This does not affect the behavior of single-threaded or multithreaded applications. If the application was using LWPs directly, and was relying on `alarm(2)` or `sleep(3C)`, then the application's behavior might be impacted. The calling LWP will not necessarily be the recipient of the `SIGALRM` signal when `SIGALRM` is sent to the process. You might have to use a substitute like `poll(2)`, or `_lwp_cond_timedwait(2)` to simulate the old per-LWP semantic of `SIGALRM`.

NAME	<code>_lwp_exit</code> – terminate the calling LWP
SYNOPSIS	<pre>#include <sys/lwp.h> void _lwp_exit(void);</pre>
DESCRIPTION	The <code>_lwp_exit()</code> function causes the calling LWP to terminate. If it is the last LWP in the process, then the process exits with a status of 0 (see <code>exit(2)</code>). If the LWP was created undetached, it is transformed into a “zombie LWP” that retains at least the LWP’s ID until it is waited for (see <code>_lwp_wait(2)</code>). Otherwise, its ID and system resources may be reclaimed immediately.
SEE ALSO	<code>_lwp_create(2)</code> , <code>_lwp_wait(2)</code> , <code>exit(2)</code>

`_lwp_info(2)`

NAME	<code>_lwp_info</code> – return the time-accounting information of a single LWP				
SYNOPSIS	<pre>#include <sys/time.h> #include <sys/lwp.h> int _lwp_info(struct lwpinfo *<i>buffer</i>);</pre>				
DESCRIPTION	The <code>_lwp_info()</code> function fills the <code>lwpinfo</code> structure pointed to by <i>buffer</i> with time-accounting information pertaining to the calling LWP. This call may be extended in the future to return other information to the <code>lwpinfo</code> structure as needed. The <code>lwpinfo</code> structure in <code><sys/lwp.h></code> includes the following members: <pre>timestruct_t lwp_utime; timestruct_t lwp_stime;</pre> The <code>lwp_utime</code> member is the CPU time used while executing instructions in the user space of the calling LWP. The <code>lwp_stime</code> member is the CPU time used by the system on behalf of the calling LWP.				
RETURN VALUES	Upon successful completion, <code>_lwp_info()</code> returns 0 and fills in the <code>lwpinfo</code> structure pointed to by <i>buffer</i> .				
ERRORS	If the following condition is detected, <code>_lwp_info()</code> returns the corresponding value: <table><tr><td>EFAULT</td><td>The <i>buffer</i> argument points to an illegal address.</td></tr></table> Additionally, the <code>_lwp_info()</code> function will fail for 32-bit interfaces if: <table><tr><td>E_OVERFLOW</td><td>The size of the <code>tv_sec</code> member of the <code>timestruct_t</code> type pointed to by <code>lwp_utime</code> and <code>lwp_stime</code> is too small to contain the correct number of seconds.</td></tr></table>	EFAULT	The <i>buffer</i> argument points to an illegal address.	E_OVERFLOW	The size of the <code>tv_sec</code> member of the <code>timestruct_t</code> type pointed to by <code>lwp_utime</code> and <code>lwp_stime</code> is too small to contain the correct number of seconds.
EFAULT	The <i>buffer</i> argument points to an illegal address.				
E_OVERFLOW	The size of the <code>tv_sec</code> member of the <code>timestruct_t</code> type pointed to by <code>lwp_utime</code> and <code>lwp_stime</code> is too small to contain the correct number of seconds.				
SEE ALSO	<code>times(2)</code>				

NAME	<code>_lwp_kill</code> – send a signal to a LWP				
SYNOPSIS	<pre>#include <sys/lwp.h> #include <signal.h> int _lwp_kill(lwpid_t <i>target_lwp</i>, int <i>sig</i>);</pre>				
DESCRIPTION	The <code>_lwp_kill()</code> function sends a signal to the LWP specified by <i>target_lwp</i>. The signal that is to be sent is specified by <i>sig</i> and must be one from the list given in <code>signal(3HEAD)</code>. If <i>sig</i> is 0 (the null signal), error checking is performed but no signal is actually sent. This can be used to check the validity of <i>target_lwp</i>. The <i>target_lwp</i> must be an LWP within the same process as the calling LWP.				
RETURN VALUES	Upon successful completion, 0 is returned. A non-zero value indicates an error.				
ERRORS	If any of the following conditions occur, <code>_lwp_kill()</code> fails and returns the corresponding value: <table><tr><td>EINVAL</td><td>The <i>sig</i> argument is not a valid signal number.</td></tr><tr><td>ESRCH</td><td>The <i>target_lwp</i> argument cannot be found in the current process.</td></tr></table>	EINVAL	The <i>sig</i> argument is not a valid signal number.	ESRCH	The <i>target_lwp</i> argument cannot be found in the current process.
EINVAL	The <i>sig</i> argument is not a valid signal number.				
ESRCH	The <i>target_lwp</i> argument cannot be found in the current process.				
SEE ALSO	<code>kill(2)</code> , <code>sigaction(2)</code> , <code>sigprocmask(2)</code> , <code>signal(3HEAD)</code>				

`_lwp_makecontext(2)`

NAME	<code>_lwp_makecontext</code> – initialize an LWP context
SYNOPSIS	<pre>#include <sys/types.h> #include <sys/lwp.h> #include <ucontext.h> void _lwp_makecontext(ucontext_t *ucp, void (*start_routine)(void *), void *arg, void *private, caddr_t stack_base, size_t stack_size);</pre>
DESCRIPTION	The <code>_lwp_makecontext()</code> function initializes the user context structure pointed to by <i>ucp</i>. The user context is defined by <code>ucontext(3HEAD)</code>. The resulting user context can be used by <code>_lwp_create(2)</code> for specifying the initial state of the new LWP. The user context is set up to start executing the function <i>start_routine</i> with a single argument, <i>arg</i>, and to call <code>_lwp_exit(2)</code> if <i>start_routine</i> returns. The new LWP will use the storage starting at <i>stack_base</i> and continuing for <i>stack_size</i> bytes as an execution stack. The initial value in LWP-private memory will be set to <i>private</i> (see <code>_lwp_setprivate(2)</code>). The signal mask in the user context is not initialized.
SEE ALSO	<code>_lwp_create(2)</code> , <code>_lwp_exit(2)</code> , <code>_lwp_setprivate(2)</code> , <code>ucontext(3HEAD)</code>

NAME	<code>_lwp_mutex_lock</code> , <code>_lwp_mutex_unlock</code> , <code>_lwp_mutex_trylock</code> – mutual exclusion						
SYNOPSIS	<pre>#include <sys/lwp.h> int _lwp_mutex_lock(lwp_mutex_t *mp); int _lwp_mutex_trylock(lwp_mutex_t *mp); int _lwp_mutex_unlock(lwp_mutex_t *mp);</pre>						
DESCRIPTION	These functions serialize the execution of lightweight processes. They are useful for ensuring that only one lightweight process can execute a critical section of code at any one time (mutual exclusion). LWP mutexes must be initialized to 0 before use. The <code>_lwp_mutex_lock()</code> function locks the LWP mutex pointed to by <i>mp</i>. If the mutex is already locked, the calling LWP blocks until the mutex becomes available. When <code>_lwp_mutex_lock()</code> returns, the mutex is locked and the calling LWP is the "owner". The <code>_lwp_mutex_trylock()</code> function attempts to lock the mutex. If the mutex is already locked it returns with an error. If the mutex is unlocked, it is locked and <code>_lwp_mutex_trylock()</code> returns. The <code>_lwp_mutex_unlock()</code> function unlocks a locked mutex. The mutex must be locked and the calling LWP must be the one that last locked the mutex (the owner). If any other LWPs are waiting for the mutex to become available, one of them is unblocked.						
RETURN VALUES	Upon successful completion, 0 is returned. A non-zero value indicates an error.						
ERRORS	If any of the following conditions are detected, <code>_lwp_mutex_lock()</code>, <code>_lwp_mutex_trylock()</code>, and <code>_lwp_mutex_unlock()</code> fail and return the corresponding value: <table><tr><td>EINVAL</td><td>The <i>mp</i> argument points to an invalid LWP mutex.</td></tr><tr><td>EFAULT</td><td>The <i>mp</i> argument points to an illegal address.</td></tr></table> If any of the following conditions occur, <code>_lwp_mutex_trylock()</code> fails and returns the corresponding value: <table><tr><td>EBUSY</td><td>The <i>mp</i> argument points to a locked mutex.</td></tr></table>	EINVAL	The <i>mp</i> argument points to an invalid LWP mutex.	EFAULT	The <i>mp</i> argument points to an illegal address.	EBUSY	The <i>mp</i> argument points to a locked mutex.
EINVAL	The <i>mp</i> argument points to an invalid LWP mutex.						
EFAULT	The <i>mp</i> argument points to an illegal address.						
EBUSY	The <i>mp</i> argument points to a locked mutex.						
SEE ALSO	<code>intro(2)</code> , <code>_lwp_cond_wait(2)</code>						

`_lwp_self(2)`

NAME	<code>_lwp_self</code> – get LWP identifier
SYNOPSIS	<pre>#include <sys/lwp.h> lwpid_t _lwp_self(void);</pre>
DESCRIPTION	The <code>_lwp_self()</code> function returns the ID of the calling LWP.
SEE ALSO	<code>_lwp_create(2)</code>

NAME	<code>_lwp_sema_wait</code> , <code>_lwp_sema_trywait</code> , <code>_lwp_sema_init</code> , <code>_lwp_sema_post</code> – semaphore operations										
SYNOPSIS	<pre>#include <sys/lwp.h> int _lwp_sema_wait(lwp_sema_t *sema); int _lwp_sema_trywait(lwp_sema_t *sema); int _lwp_sema_init(lwp_sema_t *sema, int count); int _lwp_sema_post(lwp_sema_t *sema);</pre>										
DESCRIPTION	Conceptually, a semaphore is an non-negative integer count that is atomically incremented and decremented. Typically this represents the number of resources available. The <code>_lwp_sema_init()</code> function initializes the count, <code>_lwp_sema_post()</code> atomically increments the count, and <code>_lwp_sema_wait()</code> waits for the count to become greater than 0 and then atomically decrements it. LWP semaphores must be initialized before use. The <code>_lwp_sema_init()</code> function initializes the count associated with the LWP semaphore pointed to by <i>sema</i> to <i>count</i>. The <code>_lwp_sema_wait()</code> function blocks the calling LWP until the semaphore count becomes greater than 0 and then atomically decrements it. The <code>_lwp_sema_trywait()</code> function atomically decrements the count if it is greater than zero. Otherwise it returns an error. The <code>_lwp_sema_post()</code> function atomically increments the semaphore count. If there are any LWPs blocked on the semaphore, one is unblocked.										
RETURN VALUES	Upon successful completion, 0 is returned. A non-zero value indicates an error.										
ERRORS	The <code>_lwp_sema_init()</code>, <code>_lwp_sema_trywait()</code>, <code>_lwp_sema_wait()</code>, and <code>_lwp_sema_post()</code> functions will fail if: <table><tr><td>EINVAL</td><td>The <i>sema</i> argument points to an invalid semaphore.</td></tr><tr><td>EFAULT</td><td>The <i>sema</i> argument points to an illegal address.</td></tr></table> The <code>_lwp_sema_wait()</code> function will fail if: <table><tr><td>EINTR</td><td>The function execution was interrupted by a signal or <code>fork(2)</code>.</td></tr></table> The <code>_lwp_sema_trywait()</code> function will fail if: <table><tr><td>EBUSY</td><td>The function was called on a semaphore with a zero count.</td></tr></table> The <code>_lwp_sema_post()</code> function will fail if: <table><tr><td>EOVERFLOW</td><td>The value of the <i>sema</i> argument exceeds <code>SEM_VALUE_MAX</code>.</td></tr></table>	EINVAL	The <i>sema</i> argument points to an invalid semaphore.	EFAULT	The <i>sema</i> argument points to an illegal address.	EINTR	The function execution was interrupted by a signal or <code>fork(2)</code> .	EBUSY	The function was called on a semaphore with a zero count.	EOVERFLOW	The value of the <i>sema</i> argument exceeds <code>SEM_VALUE_MAX</code> .
EINVAL	The <i>sema</i> argument points to an invalid semaphore.										
EFAULT	The <i>sema</i> argument points to an illegal address.										
EINTR	The function execution was interrupted by a signal or <code>fork(2)</code> .										
EBUSY	The function was called on a semaphore with a zero count.										
EOVERFLOW	The value of the <i>sema</i> argument exceeds <code>SEM_VALUE_MAX</code> .										
SEE ALSO	<code>fork(2)</code>										

`_lwp_setprivate(2)`

NAME	<code>_lwp_setprivate</code> , <code>_lwp_getprivate</code> – set or get LWP specific storage
SYNOPSIS	<pre>#include <sys/lwp.h> void _lwp_setprivate(void *<i>buffer</i>) ; void *_lwp_getprivate(void) ;</pre>
DESCRIPTION	The <code>_lwp_setprivate()</code> function stores the value specified by <i>buffer</i> in LWP-private memory that is unique to the calling LWP. This is typically used by thread library implementations to maintain a pointer to information about the thread currently running on the calling LWP. The <code>_lwp_getprivate()</code> function returns the value stored in LWP-private memory.
SEE ALSO	<code>_lwp_makecontext(2)</code>

NAME	_lwp_suspend, _lwp_continue – continue or suspend LWP execution				
SYNOPSIS	<pre>#include <sys/lwp.h> int _lwp_suspend(lwpid_t target_lwp); int _lwp_continue(lwpid_t target_lwp);</pre>				
DESCRIPTION	The <code>_lwp_suspend()</code> function immediately suspends the execution of the LWP specified by <i>target_lwp</i>. On successful return from <code>_lwp_suspend()</code>, <i>target_lwp</i> is no longer executing. Once a thread is suspended, subsequent calls to <code>_lwp_suspend()</code> have no affect. The <code>_lwp_continue()</code> function resumes the execution of a suspended LWP. Once a suspended LWP is continued, subsequent calls to <code>_lwp_continue()</code> have no effect. A suspended LWP will not be awakened by a signal. The signal stays pending until the execution of the LWP is resumed by <code>_lwp_continue()</code>.				
RETURN VALUES	Upon successful completion, 0 is returned. A non-zero value indicates an error.				
ERRORS	If the following condition occurs, <code>_lwp_suspend()</code> and <code>_lwp_continue()</code> fail and return the corresponding value: <table><tr><td>ESRCH</td><td>The <i>target_lwpid</i> argument cannot be found in the current process.</td></tr></table> If the following condition is detected, <code>_lwp_suspend()</code> fails and returns the corresponding value: <table><tr><td>EDEADLK</td><td>Suspending <i>target_lwpid</i> will cause all LWPs in the process to be suspended.</td></tr></table>	ESRCH	The <i>target_lwpid</i> argument cannot be found in the current process.	EDEADLK	Suspending <i>target_lwpid</i> will cause all LWPs in the process to be suspended.
ESRCH	The <i>target_lwpid</i> argument cannot be found in the current process.				
EDEADLK	Suspending <i>target_lwpid</i> will cause all LWPs in the process to be suspended.				
SEE ALSO	<code>_lwp_create(2)</code>				

`_lwp_wait(2)`

NAME	<code>_lwp_wait</code> – wait for a LWP to terminate						
SYNOPSIS	<pre>#include <sys/lwp.h> int _lwp_wait(lwpid_t wait_for, lwpid_t *departed_lwp);</pre>						
DESCRIPTION	The <code>_lwp_wait()</code> function blocks the current LWP until the LWP specified by <i>wait_for</i> terminates. If the specified LWP terminated prior to the call to <code>_lwp_wait()</code>, then <code>_lwp_wait()</code> returns immediately. If <i>wait_for</i> is NULL, then <code>_lwp_wait()</code> waits for any undetached LWP in the current process. If <i>wait_for</i> is not NULL, then it must specify an undetached LWP in the current process. If <i>departed_lwp</i> is not NULL, then it points to location where the ID of the exited LWP is stored (see <code>_lwp_exit(2)</code>). When an LWP exits and there are one or more LWPs in this process waiting for this specific LWP to exit, then one of the waiting LWPs is unblocked and it returns from <code>_lwp_wait()</code> successfully. Any other LWPs waiting for this same LWP to exit are also unblocked, however, they return from <code>_lwp_wait()</code> with an error (ESRCH) indicating the waited for LWP no longer exists. If there are no LWPs in this process waiting for this specific LWP to exit but there are one or more LWPs waiting for any LWP to exit, then one of the waiting LWPs is unblocked and it returns from <code>_lwp_wait()</code> successfully. The ID of an LWP that has exited may be reused via <code>_lwp_create()</code> after the LWP has been successfully waited for.						
RETURN VALUES	Upon successful completion, 0 is returned. A non-zero value indicates an error.						
ERRORS	If any of the following conditions are detected, <code>_lwp_wait()</code> fails and returns the corresponding value: <table><tr><td>EINTR</td><td>The <code>_lwp_wait()</code> function was interrupted by a signal.</td></tr><tr><td>EDEADLK</td><td>All LWPs in this process would be blocked waiting for LWPs to terminate, or the calling LWP is attempting to wait for itself.</td></tr></table> If any of the following conditions occur, <code>_lwp_wait()</code> fails and returns the corresponding value: <table><tr><td>ESRCH</td><td>The <i>wait_for</i> argument cannot be found in the current process or it was detached.</td></tr></table>	EINTR	The <code>_lwp_wait()</code> function was interrupted by a signal.	EDEADLK	All LWPs in this process would be blocked waiting for LWPs to terminate, or the calling LWP is attempting to wait for itself.	ESRCH	The <i>wait_for</i> argument cannot be found in the current process or it was detached.
EINTR	The <code>_lwp_wait()</code> function was interrupted by a signal.						
EDEADLK	All LWPs in this process would be blocked waiting for LWPs to terminate, or the calling LWP is attempting to wait for itself.						
ESRCH	The <i>wait_for</i> argument cannot be found in the current process or it was detached.						
SEE ALSO	<code>_lwp_create(2)</code> , <code>_lwp_exit(2)</code>						

NAME	memcntl – memory management control																
SYNOPSIS	<pre>#include <sys/types.h> #include <sys/mman.h> int memcntl(caddr_t addr, size_t len, int cmd, caddr_t arg, int attr, int mask);</pre>																
DESCRIPTION	The <code>memcntl()</code> function allows the calling process to apply a variety of control operations over the address space identified by the mappings established for the address range <code>[addr, addr + len)</code>. The <code>addr</code> argument must be a multiple of the pagesize as returned by <code>sysconf(3C)</code>. The scope of the control operations can be further defined with additional selection criteria (in the form of attributes) according to the bit pattern contained in <code>attr</code>. The following attributes specify page mapping selection criteria: <table> <tr> <td>SHARED</td><td>Page is mapped shared.</td></tr> <tr> <td>PRIVATE</td><td>Page is mapped private.</td></tr> </table> The following attributes specify page protection selection criteria. The selection criteria are constructed by a bitwise OR operation on the attribute bits and must match exactly. <table> <tr> <td>PROT_READ</td><td>Page can be read.</td></tr> <tr> <td>PROT_WRITE</td><td>Page can be written.</td></tr> <tr> <td>PROT_EXEC</td><td>Page can be executed.</td></tr> </table> The following criteria may also be specified: <table> <tr> <td>PROC_TEXT</td><td>Process text.</td></tr> <tr> <td>PROC_DATA</td><td>Process data.</td></tr> </table> The <code>PROC_TEXT</code> attribute specifies all privately mapped segments with read and execute permission, and the <code>PROC_DATA</code> attribute specifies all privately mapped segments with write permission. Selection criteria can be used to describe various abstract memory objects within the address space on which to operate. If an operation shall not be constrained by the selection criteria, <code>attr</code> must have the value 0. The operation to be performed is identified by the argument <code>cmd</code>. The symbolic names for the operations are defined in <code><sys/mman.h></code> as follows: <table> <tr> <td>MC_LOCK</td><td>Lock in memory all pages in the range with attributes <code>attr</code>. A given page may be locked multiple times through different mappings; however, within a given mapping, page locks do not nest. Multiple lock operations on the same address in the same process will all be removed with a single unlock operation. A page locked in one</td></tr> </table>	SHARED	Page is mapped shared.	PRIVATE	Page is mapped private.	PROT_READ	Page can be read.	PROT_WRITE	Page can be written.	PROT_EXEC	Page can be executed.	PROC_TEXT	Process text.	PROC_DATA	Process data.	MC_LOCK	Lock in memory all pages in the range with attributes <code>attr</code> . A given page may be locked multiple times through different mappings; however, within a given mapping, page locks do not nest. Multiple lock operations on the same address in the same process will all be removed with a single unlock operation. A page locked in one
SHARED	Page is mapped shared.																
PRIVATE	Page is mapped private.																
PROT_READ	Page can be read.																
PROT_WRITE	Page can be written.																
PROT_EXEC	Page can be executed.																
PROC_TEXT	Process text.																
PROC_DATA	Process data.																
MC_LOCK	Lock in memory all pages in the range with attributes <code>attr</code> . A given page may be locked multiple times through different mappings; however, within a given mapping, page locks do not nest. Multiple lock operations on the same address in the same process will all be removed with a single unlock operation. A page locked in one																

memcntl(2)

	process and mapped in another (or visible through a different mapping in the locking process) is locked in memory as long as the locking process does neither an implicit nor explicit unlock operation. If a locked mapping is removed, or a page is deleted through file removal or truncation, an unlock operation is implicitly performed. If a writable MAP_PRIVATE page in the address range is changed, the lock will be transferred to the private page. At present <i>arg</i> is unused, but must be 0 to ensure compatibility with potential future enhancements.						
MC_LOCKAS	Lock in memory all pages mapped by the address space with attributes <i>attr</i>. At present <i>addr</i> and <i>len</i> are unused, but must be NULL and 0 respectively, to ensure compatibility with potential future enhancements. The <i>arg</i> argument is a bit pattern built from the flags: <table><tr><td>MCL_CURRENT</td><td>Lock current mappings</td></tr><tr><td>MCL_FUTURE</td><td>Lock future mappings</td></tr></table> The value of <i>arg</i> determines whether the pages to be locked are those currently mapped by the address space, those that will be mapped in the future, or both. If MCL_FUTURE is specified, then all mappings subsequently added to the address space will be locked, provided sufficient memory is available.	MCL_CURRENT	Lock current mappings	MCL_FUTURE	Lock future mappings		
MCL_CURRENT	Lock current mappings						
MCL_FUTURE	Lock future mappings						
MC_SYNC	Write to their backing storage locations all modified pages in the range with attributes <i>attr</i>. Optionally, invalidate cache copies. The backing storage for a modified MAP_SHARED mapping is the file the page is mapped to; the backing storage for a modified MAP_PRIVATE mapping is its swap area. The <i>arg</i> argument is a bit pattern built from the flags used to control the behavior of the operation: <table><tr><td>MS_ASYNC</td><td>perform asynchronous writes</td></tr><tr><td>MS_SYNC</td><td>perform synchronous writes</td></tr><tr><td>MS_INVALIDATE</td><td>invalidate mappings</td></tr></table> MS_ASYNC returns immediately once all write operations are scheduled; with MS_SYNC the function will not return until all write operations are completed. MS_INVALIDATE invalidates all cached copies of data in memory, so that further references to the pages will be obtained by the system from their backing storage locations. This operation should be used by applications that require a memory object to be in a known state.	MS_ASYNC	perform asynchronous writes	MS_SYNC	perform synchronous writes	MS_INVALIDATE	invalidate mappings
MS_ASYNC	perform asynchronous writes						
MS_SYNC	perform synchronous writes						
MS_INVALIDATE	invalidate mappings						

	MC_UNLOCK	Unlock all pages in the range with attributes <i>attr</i> . At present <i>arg</i> is unused, but must be 0 to ensure compatibility with potential future enhancements.				
	MC_UNLOCKAS	Remove address space memory locks, and locks on all pages in the address space with attributes <i>attr</i> . At present <i>addr</i> , <i>len</i> , and <i>arg</i> are unused, but must be NULL, 0 and 0 respectively, to ensure compatibility with potential future enhancements.				
	The <i>mask</i> argument must be 0; it is reserved for future use.					
	Locks established with the lock operations are not inherited by a child process after <code>fork(2)</code> . The <code>mlock(2)</code> function fails if it attempts to lock more memory than a system-specific limit.					
	Due to the potential impact on system resources, all operations, with the exception of MC_SYNC, are restricted to processes with super-user effective user ID.					
USAGE	The <code>mlock(2)</code> function subsumes the operations of <code>plock(3C)</code> and <code>mctl(3UCB)</code> .					
RETURN VALUES	Upon successful completion, <code>mlock(2)</code> returns 0; otherwise, it returns -1 and sets <code>errno</code> to indicate an error.					
ERRORS	The <code>mlock(2)</code> function will fail if:					
	EAGAIN	Some or all of the memory identified by the operation could not be locked when MC_LOCK or MC_LOCKAS was specified.				
	EBUSY	Some or all of the addresses in the range [<i>addr</i> , <i>addr</i> + <i>len</i>) are locked and MC_SYNC with the MS_INVALIDATE option was specified.				
	EINVAL	The <i>addr</i> argument specifies invalid selection criteria or is not a multiple of the page size as returned by <code>sysconf(3C)</code> ; the <i>addr</i> and/or <i>len</i> argument does not have the value 0 when MC_LOCKAS or MC_UNLOCKAS is specified; or the <i>arg</i> argument is not valid for the function specified.				
	ENOMEM	Some or all of the addresses in the range [<i>addr</i> , <i>addr</i> + <i>len</i>) are invalid for the address space of a process or specify one or more pages which are not mapped.				
	EPERM	The process's effective user ID is not super-user and MC_LOCK, MC_LOCKAS, MC_UNLOCK, or MC_UNLOCKAS was specified.				
ATTRIBUTES	See <code>attributes(5)</code> for descriptions of the following attributes:					
<table><tr><th>ATTRIBUTE TYPE</th><th>ATTRIBUTE VALUE</th></tr><tr><td>MT-Level</td><td>MT-Safe</td></tr></table>			ATTRIBUTE TYPE	ATTRIBUTE VALUE	MT-Level	MT-Safe
ATTRIBUTE TYPE	ATTRIBUTE VALUE					
MT-Level	MT-Safe					

memcntl(2)

SEE ALSO fork(2), mmap(2), mprotect(2), mctl(3UCB), mlock(3C), mlockall(3C), msync(3C),
plock(3C), sysconf(3C), attributes(5)

NAME	mincore – determine residency of memory pages						
SYNOPSIS	<pre>#include <sys/types.h> int mincore(caddr_t <i>addr</i>, size_t <i>len</i>, char *<i>vec</i>);</pre>						
DESCRIPTION	The <code>mincore()</code> function determines the residency of the memory pages in the address space covered by mappings in the range <code>[addr, addr + len]</code>. The status is returned as a character-per-page in the character array referenced by <i>vec</i> (which the system assumes to be large enough to encompass all the pages in the address range). The least significant bit of each character is set to 1 to indicate that the referenced page is in primary memory, and to 0 to indicate that it is not. The settings of other bits in each character are undefined and may contain other information in future implementations. Because the status of a page can change between the time <code>mincore()</code> checks and returns the information, returned information might be outdated. Only locked pages are guaranteed to remain in memory; see <code>mlock(3C)</code>.						
RETURN VALUES	Upon successful completion, <code>mincore()</code> returns 0. Otherwise, -1 is returned and <code>errno</code> is set to indicate the error.						
ERRORS	The <code>mincore()</code> function will fail if: <table> <tr> <td>EFAULT</td><td>The <i>vec</i> argument points to an illegal address.</td></tr> <tr> <td>EINVAL</td><td>The <i>addr</i> argument is not a multiple of the page size as returned by <code>sysconf(3C)</code>, or the <i>len</i> argument has a value less than or equal to 0.</td></tr> <tr> <td>ENOMEM</td><td>Addresses in the range <code>[addr, addr + len]</code> are invalid for the address space of a process or specify one or more pages which are not mapped.</td></tr> </table>	EFAULT	The <i>vec</i> argument points to an illegal address.	EINVAL	The <i>addr</i> argument is not a multiple of the page size as returned by <code>sysconf(3C)</code> , or the <i>len</i> argument has a value less than or equal to 0.	ENOMEM	Addresses in the range <code>[addr, addr + len]</code> are invalid for the address space of a process or specify one or more pages which are not mapped.
EFAULT	The <i>vec</i> argument points to an illegal address.						
EINVAL	The <i>addr</i> argument is not a multiple of the page size as returned by <code>sysconf(3C)</code> , or the <i>len</i> argument has a value less than or equal to 0.						
ENOMEM	Addresses in the range <code>[addr, addr + len]</code> are invalid for the address space of a process or specify one or more pages which are not mapped.						
SEE ALSO	<code>mmap(2)</code> , <code>mlock(3C)</code> , <code>sysconf(3C)</code>						

mkdir(2)

NAME	mkdir – make a directory												
SYNOPSIS	<pre>#include <sys/types.h> #include <sys/stat.h> int mkdir(const char *path, mode_t mode);</pre>												
DESCRIPTION	The <code>mkdir()</code> function creates a new directory named by the path name pointed to by <i>path</i>. The mode of the new directory is initialized from <i>mode</i> (see <code>chmod(2)</code> for values of mode). The protection part of the <i>mode</i> argument is modified by the process's file creation mask (see <code>umask(2)</code>). The directory's owner ID is set to the process's effective user ID. The directory's group ID is set to the process's effective group ID, or if the <code>S_ISGID</code> bit is set in the parent directory, then the group ID of the directory is inherited from the parent. The <code>S_ISGID</code> bit of the new directory is inherited from the parent directory. If <i>path</i> is a symbolic link, it is not followed. The newly created directory is empty with the exception of entries for itself (<code>.</code>) and its parent directory (<code>..</code>). Upon successful completion, <code>mkdir()</code> marks for update the <code>st_atime</code>, <code>st_ctime</code> and <code>st_mtime</code> fields of the directory. Also, the <code>st_ctime</code> and <code>st_mtime</code> fields of the directory that contains the new entry are marked for update.												
RETURN VALUES	Upon successful completion, 0 is returned. Otherwise, -1 is returned, no directory is created, and <code>errno</code> is set to indicate the error.												
ERRORS	The <code>mkdir()</code> function will fail if: <table> <tr> <td>EACCES</td><td>Either a component of the path prefix denies search permission or write permission is denied on the parent directory of the directory to be created.</td></tr> <tr> <td>EDQUOT</td><td>The directory where the new file entry is being placed cannot be extended because the user's quota of disk blocks on that file system has been exhausted; the new directory cannot be created because the user's quota of disk blocks on that file system has been exhausted; or the user's quota of inodes on the file system where the file is being created has been exhausted.</td></tr> <tr> <td>EEXIST</td><td>The named file already exists.</td></tr> <tr> <td>EFAULT</td><td>The <i>path</i> argument points to an illegal address.</td></tr> <tr> <td>EIO</td><td>An I/O error has occurred while accessing the file system.</td></tr> <tr> <td>ELOOP</td><td>Too many symbolic links were encountered in translating <i>path</i>.</td></tr> </table>	EACCES	Either a component of the path prefix denies search permission or write permission is denied on the parent directory of the directory to be created.	EDQUOT	The directory where the new file entry is being placed cannot be extended because the user's quota of disk blocks on that file system has been exhausted; the new directory cannot be created because the user's quota of disk blocks on that file system has been exhausted; or the user's quota of inodes on the file system where the file is being created has been exhausted.	EEXIST	The named file already exists.	EFAULT	The <i>path</i> argument points to an illegal address.	EIO	An I/O error has occurred while accessing the file system.	ELOOP	Too many symbolic links were encountered in translating <i>path</i> .
EACCES	Either a component of the path prefix denies search permission or write permission is denied on the parent directory of the directory to be created.												
EDQUOT	The directory where the new file entry is being placed cannot be extended because the user's quota of disk blocks on that file system has been exhausted; the new directory cannot be created because the user's quota of disk blocks on that file system has been exhausted; or the user's quota of inodes on the file system where the file is being created has been exhausted.												
EEXIST	The named file already exists.												
EFAULT	The <i>path</i> argument points to an illegal address.												
EIO	An I/O error has occurred while accessing the file system.												
ELOOP	Too many symbolic links were encountered in translating <i>path</i> .												

EMLINK	The maximum number of links to the parent directory would be exceeded.
ENAMETOOLONG	The length of the <i>path</i> argument exceeds <code>PATH_MAX</code> , or the length of a <i>path</i> component exceeds <code>NAME_MAX</code> while <code>_POSIX_NO_TRUNC</code> is in effect.
ENOENT	A component of the path prefix does not exist or is a null pathname.
ENOLINK	The <i>path</i> argument points to a remote machine and the link to that machine is no longer active.
ENOSPC	No free space is available on the device containing the directory.
ENOTDIR	A component of the path prefix is not a directory.
EROFS	The path prefix resides on a read-only file system.

ATTRIBUTES See `attributes(5)` for descriptions of the following attributes:

ATTRIBUTE TYPE	ATTRIBUTE VALUE
MT-Level	Async-Signal-Safe

SEE ALSO `chmod(2)`, `mknod(2)`, `umask(2)`, `stat(3HEAD)`, `attributes(5)`

mknod(2)

NAME	mknod – make a directory, or a special or ordinary file																																																									
SYNOPSIS	#include <sys/stat.h> int mknod (const char *path, mode_t mode, dev_t dev);																																																									
DESCRIPTION	The mknod() function creates a new file named by the path name pointed to by path. The file type and permissions of the new file are initialized from mode. The file type is specified in mode by the S_IFMT bits, which must be set to one of the following values: <table><tr><td>S_IFIFO</td><td>fifo special</td></tr><tr><td>S_IFCHR</td><td>character special</td></tr><tr><td>S_IFDIR</td><td>directory</td></tr><tr><td>S_IFBLK</td><td>block special</td></tr><tr><td>S_IFREG</td><td>ordinary file</td></tr></table> The file access permissions are specified in mode by the 0007777 bits, and may be constructed by a bitwise OR operation of the following values: <table><tr><td>S_ISUID</td><td>04000</td><td>Set user ID on execution.</td></tr><tr><td>S_ISGID</td><td>020#0</td><td>Set group ID on execution if # is 7, 5, 3, or 1. Enable mandatory file/record locking if # is 6, 4, 2, or 0</td></tr><tr><td>S_ISVTX</td><td>01000</td><td>On directories, restricted deletion flag; on regular files on a UFS file system, do not cache flag.</td></tr><tr><td>S_IRWXU</td><td>00700</td><td>Read, write, execute by owner.</td></tr><tr><td>S_IRUSR</td><td>00400</td><td>Read by owner.</td></tr><tr><td>S_IWUSR</td><td>00200</td><td>Write by owner.</td></tr><tr><td>S_IXUSR</td><td>00100</td><td>Execute (search if a directory) by owner.</td></tr><tr><td>S_IRWXG</td><td>00070</td><td>Read, write, execute by group.</td></tr><tr><td>S_IRGRP</td><td>00040</td><td>Read by group.</td></tr><tr><td>S_IWGRP</td><td>00020</td><td>Write by group.</td></tr><tr><td>S_IXGRP</td><td>00010</td><td>Execute by group.</td></tr><tr><td>S_IRWXO</td><td>00007</td><td>Read, write, execute (search) by others.</td></tr><tr><td>S_IROTH</td><td>00004</td><td>Read by others.</td></tr><tr><td>S_IWOTH</td><td>00002</td><td>Write by others</td></tr><tr><td>S_IXOTH</td><td>00001</td><td>Execute by others.</td></tr></table>			S_IFIFO	fifo special	S_IFCHR	character special	S_IFDIR	directory	S_IFBLK	block special	S_IFREG	ordinary file	S_ISUID	04000	Set user ID on execution.	S_ISGID	020#0	Set group ID on execution if # is 7, 5, 3, or 1. Enable mandatory file/record locking if # is 6, 4, 2, or 0	S_ISVTX	01000	On directories, restricted deletion flag; on regular files on a UFS file system, do not cache flag.	S_IRWXU	00700	Read, write, execute by owner.	S_IRUSR	00400	Read by owner.	S_IWUSR	00200	Write by owner.	S_IXUSR	00100	Execute (search if a directory) by owner.	S_IRWXG	00070	Read, write, execute by group.	S_IRGRP	00040	Read by group.	S_IWGRP	00020	Write by group.	S_IXGRP	00010	Execute by group.	S_IRWXO	00007	Read, write, execute (search) by others.	S_IROTH	00004	Read by others.	S_IWOTH	00002	Write by others	S_IXOTH	00001	Execute by others.
S_IFIFO	fifo special																																																									
S_IFCHR	character special																																																									
S_IFDIR	directory																																																									
S_IFBLK	block special																																																									
S_IFREG	ordinary file																																																									
S_ISUID	04000	Set user ID on execution.																																																								
S_ISGID	020#0	Set group ID on execution if # is 7, 5, 3, or 1. Enable mandatory file/record locking if # is 6, 4, 2, or 0																																																								
S_ISVTX	01000	On directories, restricted deletion flag; on regular files on a UFS file system, do not cache flag.																																																								
S_IRWXU	00700	Read, write, execute by owner.																																																								
S_IRUSR	00400	Read by owner.																																																								
S_IWUSR	00200	Write by owner.																																																								
S_IXUSR	00100	Execute (search if a directory) by owner.																																																								
S_IRWXG	00070	Read, write, execute by group.																																																								
S_IRGRP	00040	Read by group.																																																								
S_IWGRP	00020	Write by group.																																																								
S_IXGRP	00010	Execute by group.																																																								
S_IRWXO	00007	Read, write, execute (search) by others.																																																								
S_IROTH	00004	Read by others.																																																								
S_IWOTH	00002	Write by others																																																								
S_IXOTH	00001	Execute by others.																																																								

mknod(2)

The owner ID of the file is set to the effective user ID of the process. The group ID of the file is set to the effective group ID of the process. However, if the `S_ISGID` bit is set in the parent directory, then the group ID of the file is inherited from the parent. If the group ID of the new file does not match the effective group ID or one of the supplementary group IDs, the `S_ISGID` bit is cleared.

The access permission bits of *mode* are modified by the process's file mode creation mask: all bits set in the process's file mode creation mask are cleared (see `umask(2)`). If *mode* indicates a block or character special file, *dev* is a configuration-dependent specification of a character or block I/O device. If *mode* does not indicate a block special or character special device, *dev* is ignored. See `makedev(3C)`.

If *path* is a symbolic link, it is not followed.

RETURN VALUES Upon successful completion, `mknod()` returns 0. Otherwise, it returns -1, the new file is not created, and `errno` is set to indicate the error.

ERRORS The `mknod()` function will fail if:

EACCES	A component of the path prefix denies search permission, or write permission is denied on the parent directory.
EDQUOT	The directory where the new file entry is being placed cannot be extended because the user's quota of disk blocks on that file system has been exhausted, or the user's quota of inodes on the file system where the file is being created has been exhausted.
EEXIST	The named file exists.
EFAULT	The <i>path</i> argument points to an illegal address.
EINTR	A signal was caught during the execution of the <code>mknod()</code> function.
EINVAL	An invalid argument exists.
EIO	An I/O error occurred while accessing the file system.
ELOOP	Too many symbolic links were encountered in translating <i>path</i> .
ENAMETOOLONG	The length of the <i>path</i> argument exceeds <code>PATH_MAX</code> , or the length of a <i>path</i> component exceeds <code>NAME_MAX</code> while <code>_POSIX_NO_TRUNC</code> is in effect.
ENOENT	A component of the path prefix specified by <i>path</i> does not name an existing directory or <i>path</i> is an empty string.
ENOLINK	The <i>path</i> argument points to a remote machine and the link to that machine is no longer active.

mknod(2)

	ENOSPC	The directory that would contain the new file cannot be extended or the file system is out of file allocation resources.
	ENOTDIR	A component of the path prefix is not a directory.
	EPERM	The effective user of the calling process is not super-user.
	EROFS	The directory in which the file is to be created is located on a read-only file system.
	The <code>mknod()</code> function may fail if:	
	ENAMETOOLONG	Pathname resolution of a symbolic link produced an intermediate result whose length exceeds <code>PATH_MAX</code> .
USAGE	Normally, applications should use the <code>mkdir(2)</code> routine to make a directory, since the function <code>mknod()</code> may not establish directory entries for the directory itself (<code>.</code>) and the parent directory (<code>..</code>), and appropriate permissions are not required. Similarly, <code>mkfifo(3C)</code> should be used in place of <code>mknod()</code> in order to create FIFOs. The <code>mknod()</code> function may be invoked only by a privileged user for file types other than FIFO special.	
SEE ALSO	<code>chmod(2)</code> , <code>creat(2)</code> , <code>exec(2)</code> , <code>mkdir(2)</code> , <code>open(2)</code> , <code>stat(2)</code> , <code>umask(2)</code> , <code>makedev(3C)</code> , <code>mkfifo(3C)</code> , <code>stat(3HEAD)</code>	

NAME	mmap – map pages of memory								
SYNOPSIS	<pre>#include <sys/mman.h> void *mmap(void *addr, size_t len, int prot, int flags, int fildes, off_t off) ;</pre>								
DESCRIPTION	The <code>mmap()</code> function establishes a mapping between a process's address space and a file or shared memory object. The format of the call is as follows: <pre>pa = mmap(addr, len, prot, flags, fildes, off) ;</pre> The <code>mmap()</code> function establishes a mapping between the address space of the process at an address <i>pa</i> for <i>len</i> bytes to the memory object represented by the file descriptor <i>fildes</i> at offset <i>off</i> for <i>len</i> bytes. The value of <i>pa</i> is a function of the <i>addr</i> argument and values of <i>flags</i>, further described below. A successful <code>mmap()</code> call returns <i>pa</i> as its result. The address range starting at <i>pa</i> and continuing for <i>len</i> bytes will be legitimate for the possible (not necessarily current) address space of the process. The range of bytes starting at <i>off</i> and continuing for <i>len</i> bytes will be legitimate for the possible (not necessarily current) offsets in the file or shared memory object represented by <i>fildes</i>. The <code>mmap()</code> function allows [<i>pa</i>, <i>pa</i> + <i>len</i>) to extend beyond the end of the object both at the time of the <code>mmap()</code> and while the mapping persists, such as when the file is created prior to the <code>mmap()</code> call and has no contents, or when the file is truncated. Any reference to addresses beyond the end of the object, however, will result in the delivery of a SIGBUS or SIGSEGV signal. The <code>mmap()</code> function cannot be used to implicitly extend the length of files. The mapping established by <code>mmap()</code> replaces any previous mappings for those whole pages containing any part of the address space of the process starting at <i>pa</i> and continuing for <i>len</i> bytes. If the size of the mapped file changes after the call to <code>mmap()</code> as a result of some other operation on the mapped file, the effect of references to portions of the mapped region that correspond to added or removed portions of the file is unspecified. The <code>mmap()</code> function is supported for regular files and shared memory objects. Support for any other type of file is unspecified. The <i>prot</i> argument determines whether read, write, execute, or some combination of accesses are permitted to the data being mapped. The <i>prot</i> argument should be either <code>PROT_NONE</code> or the bitwise inclusive OR of one or more of the other flags in the following table, defined in the header <code><sys/mman.h></code>. <table> <tr> <td><code>PROT_READ</code></td><td>Data can be read.</td></tr> <tr> <td><code>PROT_WRITE</code></td><td>Data can be written.</td></tr> <tr> <td><code>PROT_EXEC</code></td><td>Data can be executed.</td></tr> <tr> <td><code>PROT_NONE</code></td><td>Data cannot be accessed.</td></tr> </table>	<code>PROT_READ</code>	Data can be read.	<code>PROT_WRITE</code>	Data can be written.	<code>PROT_EXEC</code>	Data can be executed.	<code>PROT_NONE</code>	Data cannot be accessed.
<code>PROT_READ</code>	Data can be read.								
<code>PROT_WRITE</code>	Data can be written.								
<code>PROT_EXEC</code>	Data can be executed.								
<code>PROT_NONE</code>	Data cannot be accessed.								

mmap(2)

If an implementation of `mmap()` for a specific platform cannot support the combination of access types specified by *prot*, the call to `mmap()` fails. An implementation may permit accesses other than those specified by *prot*; however, the implementation will not permit a write to succeed where `PROT_WRITE` has not been set or permit any access where `PROT_NONE` alone has been set. Each platform-specific implementation of `mmap()` supports the following values of *prot*: `PROT_NONE`, `PROT_READ`, `PROT_WRITE`, and the inclusive OR of `PROT_READ` and `PROT_WRITE`. On some platforms, the `PROT_WRITE` protection option is implemented as `PROT_READ | PROT_WRITE` and `PROT_EXEC` as `PROT_READ | PROT_EXEC`. The file descriptor *fd* is opened with read permission, regardless of the protection options specified. If `PROT_WRITE` is specified, the application must have opened the file descriptor *fd* with write permission unless `MAP_PRIVATE` is specified in the *flags* argument as described below.

The *flags* argument provides other information about the handling of the mapped data. The value of *flags* is the bitwise inclusive OR of these options, defined in `<sys/mman.h>`:

<code>MAP_SHARED</code>	Changes are shared.
<code>MAP_PRIVATE</code>	Changes are private.
<code>MAP_FIXED</code>	Interpret <i>addr</i> exactly.
<code>MAP_NORESERVE</code>	Do not reserve swap space.
<code>MAP_ANON</code>	Map anonymous memory.

The `MAP_SHARED` and `MAP_PRIVATE` options describe the disposition of write references to the underlying object. If `MAP_SHARED` is specified, write references will change the memory object. If `MAP_PRIVATE` is specified, the initial write reference will create a private copy of the memory object page and redirect the mapping to the copy. The private copy is not created until the first write; until then, other users who have the object mapped `MAP_SHARED` can change the object. Either `MAP_SHARED` or `MAP_PRIVATE` must be specified, but not both. The mapping type is retained across `fork(2)`.

When `MAP_FIXED` is set in the *flags* argument, the system is informed that the value of *pa* must be *addr*, exactly. If `MAP_FIXED` is set, `mmap()` may return `(void *)-1` and set `errno` to `EINVAL`. If a `MAP_FIXED` request is successful, the mapping established by `mmap()` replaces any previous mappings for the process's pages in the range $[pa, pa + len)$. The use of `MAP_FIXED` is discouraged, since it may prevent a system from making the most effective use of its resources.

When `MAP_FIXED` is set and the requested address is the same as previous mapping, the previous address is unmapped and the new mapping is created on top of the old one.

When `MAP_FIXED` is not set, the system uses *addr* to arrive at *pa*. The *pa* so chosen will be an area of the address space that the system deems suitable for a mapping of *len* bytes to the file. The `mmap()` function interprets an *addr* value of 0 as granting the system complete freedom in selecting *pa*, subject to constraints described below. A non-zero value of *addr* is taken to be a suggestion of a process address near which the mapping should be placed. When the system selects a value for *pa*, it will never place a mapping at address 0, nor will it replace any extant mapping, nor map into areas considered part of the potential data or stack “segments”.

The `MAP_NORESERVE` option specifies that no swap space be reserved for a mapping. Without this flag, the creation of a writable `MAP_PRIVATE` mapping reserves swap space equal to the size of the mapping; when the mapping is written into, the reserved space is employed to hold private copies of the data. A write into a `MAP_NORESERVE` mapping produces results which depend on the current availability of swap space in the system. If space is available, the write succeeds and a private copy of the written page is created; if space is not available, the write fails and a `SIGBUS` or `SIGSEGV` signal is delivered to the writing process. `MAP_NORESERVE` mappings are inherited across `fork()`; at the time of the `fork()`, swap space is reserved in the child for all private pages that currently exist in the parent; thereafter the child’s mapping behaves as described above.

When `MAP_ANON` is set in *flags*, and *fd* is set to -1, `mmap()` provides a direct path to return anonymous pages to the caller. This operation is equivalent to passing `mmap()` an open file descriptor on `/dev/zero` with `MAP_ANON` elided from the *flags* argument.

The *off* argument is constrained to be aligned and sized according to the value returned by `sysconf(3C)` when passed `_SC_PAGESIZE` or `_SC_PAGE_SIZE`. When `MAP_FIXED` is specified, the *addr* argument must also meet these constraints. The system performs mapping operations over whole pages. Thus, while the *len* argument need not meet a size or alignment constraint, the system will include, in any mapping operation, any partial page specified by the range [*pa*, *pa* + *len*).

The system will always zero-fill any partial page at the end of an object. Further, the system will never write out any modified portions of the last page of an object which are beyond its end. References to whole pages following the end of an object will result in the delivery of a `SIGBUS` or `SIGSEGV` signal. `SIGBUS` signals may also be delivered on various file system conditions, including quota exceeded errors.

The `mmap()` function adds an extra reference to the file associated with the file descriptor *fd* which is not removed by a subsequent `close(2)` on that file descriptor. This reference is removed when there are no more mappings to the file by a call to the `munmap(2)` function.

The `st_atime` field of the mapped file may be marked for update at any time between the `mmap()` call and the corresponding `munmap(2)` call. The initial read or write reference to a mapped region will cause the file’s `st_atime` field to be marked for update if it has not already been marked for update.

mmap(2)

	The <code>st_ctime</code> and <code>st_mtime</code> fields of a file that is mapped with <code>MAP_SHARED</code> and <code>PROT_WRITE</code>, will be marked for update at some point in the interval between a write reference to the mapped region and the next call to <code>msync(3C)</code> with <code>MS_ASYNC</code> or <code>MS_SYNC</code> for that portion of the file by any process. If there is no such call, these fields may be marked for update at any time after a write reference if the underlying file is modified as a result. If the process calls <code>mlockall(3C)</code> with the <code>MCL_FUTURE</code> flag, the pages mapped by all future calls to <code>mmap()</code> will be locked in memory. In this case, if not enough memory could be locked, <code>mmap()</code> fails and sets <code>errno</code> to <code>EAGAIN</code>.																				
RETURN VALUES	Upon successful completion, the <code>mmap()</code> function returns the address at which the mapping was placed (<i>pa</i>); otherwise, it returns a value of <code>MAP_FAILED</code> and sets <code>errno</code> to indicate the error. The symbol <code>MAP_FAILED</code> is defined in the header <code><sys/mman.h></code>. No successful return from <code>mmap()</code> will return the value <code>MAP_FAILED</code>. If <code>mmap()</code> fails for reasons other than <code>EBADF</code>, <code>EINVAL</code> or <code>ENOTSUP</code>, some of the mappings in the address range starting at <i>addr</i> and continuing for <i>len</i> bytes may have been unmapped.																				
ERRORS	The <code>mmap()</code> function will fail if: <table><tr><td>EACCES</td><td>The <i>fil-des</i> file descriptor is not open for read, regardless of the protection specified, or <i>fil-des</i> is not open for write and <code>PROT_WRITE</code> was specified for a <code>MAP_SHARED</code> type mapping.</td></tr><tr><td>EAGAIN</td><td>The mapping could not be locked in memory.</td></tr><tr><td></td><td>There was insufficient room to reserve swap space for the mapping.</td></tr><tr><td>EBADF</td><td>The <i>fil-des</i> file descriptor is not open (and <code>MAP_ANON</code> was not specified).</td></tr><tr><td>EINVAL</td><td>The arguments <i>addr</i> (if <code>MAP_FIXED</code> was specified) or <i>off</i> are not multiples of the page size as returned by <code>sysconf()</code>.</td></tr><tr><td></td><td>The field in <i>flags</i> is invalid (neither <code>MAP_PRIVATE</code> or <code>MAP_SHARED</code> is set).</td></tr><tr><td></td><td>The argument <i>len</i> has a value less than or equal to 0.</td></tr><tr><td></td><td><code>MAP_ANON</code> was specified, but the file descriptor was not <code>-1</code>.</td></tr><tr><td>EMFILE</td><td>The number of mapped regions would exceed an implementation-dependent limit (per process or per system).</td></tr><tr><td>ENODEV</td><td>The <i>fil-des</i> argument refers to an object for which <code>mmap()</code> is meaningless, such as a terminal.</td></tr></table>	EACCES	The <i>fil-des</i> file descriptor is not open for read, regardless of the protection specified, or <i>fil-des</i> is not open for write and <code>PROT_WRITE</code> was specified for a <code>MAP_SHARED</code> type mapping.	EAGAIN	The mapping could not be locked in memory.		There was insufficient room to reserve swap space for the mapping.	EBADF	The <i>fil-des</i> file descriptor is not open (and <code>MAP_ANON</code> was not specified).	EINVAL	The arguments <i>addr</i> (if <code>MAP_FIXED</code> was specified) or <i>off</i> are not multiples of the page size as returned by <code>sysconf()</code> .		The field in <i>flags</i> is invalid (neither <code>MAP_PRIVATE</code> or <code>MAP_SHARED</code> is set).		The argument <i>len</i> has a value less than or equal to 0.		<code>MAP_ANON</code> was specified, but the file descriptor was not <code>-1</code> .	EMFILE	The number of mapped regions would exceed an implementation-dependent limit (per process or per system).	ENODEV	The <i>fil-des</i> argument refers to an object for which <code>mmap()</code> is meaningless, such as a terminal.
EACCES	The <i>fil-des</i> file descriptor is not open for read, regardless of the protection specified, or <i>fil-des</i> is not open for write and <code>PROT_WRITE</code> was specified for a <code>MAP_SHARED</code> type mapping.																				
EAGAIN	The mapping could not be locked in memory.																				
	There was insufficient room to reserve swap space for the mapping.																				
EBADF	The <i>fil-des</i> file descriptor is not open (and <code>MAP_ANON</code> was not specified).																				
EINVAL	The arguments <i>addr</i> (if <code>MAP_FIXED</code> was specified) or <i>off</i> are not multiples of the page size as returned by <code>sysconf()</code> .																				
	The field in <i>flags</i> is invalid (neither <code>MAP_PRIVATE</code> or <code>MAP_SHARED</code> is set).																				
	The argument <i>len</i> has a value less than or equal to 0.																				
	<code>MAP_ANON</code> was specified, but the file descriptor was not <code>-1</code> .																				
EMFILE	The number of mapped regions would exceed an implementation-dependent limit (per process or per system).																				
ENODEV	The <i>fil-des</i> argument refers to an object for which <code>mmap()</code> is meaningless, such as a terminal.																				

	ENOMEM	The MAP_FIXED option was specified and the range <code>[addr, addr + len)</code> exceeds that allowed for the address space of a process. The MAP_FIXED option was <i>not</i> specified and there is insufficient room in the address space to effect the mapping. The mapping could not be locked in memory, if required by <code>mlockall(3C)</code>, because it would require more space than the system is able to supply. The composite size of <i>len</i> plus the lengths obtained from all previous calls to <code>mmap()</code> exceeds <code>RLIMIT_VMEM</code> (see <code>getrlimit(2)</code>).
	ENOTSUP	The system does not support the combination of accesses requested in the <i>prot</i> argument.
	ENXIO	Addresses in the range <code>[off, off + len)</code> are invalid for the object specified by <i>fildes</i> .
	EOVERFLOW	The MAP_FIXED option was specified in <i>flags</i> and the combination of <i>addr</i>, <i>len</i> and <i>off</i> is invalid for the object specified by <i>fildes</i>. The file is a regular file and the value of <i>off</i> plus <i>len</i> exceeds the offset maximum established in the open file description associated with <i>fildes</i>.
	The <code>mmap()</code> function may fail if:	
	EAGAIN	The file to be mapped is already locked using advisory or mandatory record locking. See <code>fcntl(2)</code> .
USAGE	Use of <code>mmap()</code> may reduce the amount of memory available to other memory allocation functions. Use of MAP_FIXED may result in unspecified behaviour in further use of <code>brk(2)</code>, <code>sbrk(2)</code>, <code>malloc(3C)</code>, and <code>shmat(2)</code>. The use of MAP_FIXED is discouraged, as it may prevent an implementation from making the most effective use of resources. The application must ensure correct synchronization when using <code>mmap()</code> in conjunction with any other file access method, such as <code>read(2)</code> and <code>write(2)</code>, standard input/output, and <code>shmat(2)</code>. The <code>mmap()</code> function has a transitional interface for 64-bit file offsets. See <code>lf64(5)</code>. The <code>mmap()</code> function allows access to resources using address space manipulations instead of the <code>read()/write()</code> interface. Once a file is mapped, all a process has to do to access it is use the data at the address to which the object was mapped. Consider the following pseudo-code:	

mmap(2)

```
fildes = open( . . . )
lseek(fildes, offset, whence)
read(fildes, buf, len)
/* use data in buf */
```

The following is a rewrite using `mmap()`:

```
fildes = open( . . . )
address = mmap((caddr_t) 0, len, (PROT_READ | PROT_WRITE),
              MAP_PRIVATE, fildes, offset)
/* use data at address */
```

SEE ALSO

`close(2)`, `exec(2)`, `fcntl(2)`, `fork(2)`, `getrlimit(2)`, `mprotect(2)`, `munmap(2)`,
`shmat(2)`, `lockf(3C)`, `mlockall(3C)`, `msync(3C)`, `plock(3C)`, `sysconf(3C)`, `lf64(5)`

NAME	mount – mount a file system				
SYNOPSIS	<pre>#include <sys/types.h> #include <sys/mount.h> #include <sys/mntent.h> int mount(const char *spec, const char *dir, int mflag, char *fstype, char *dataptr, int datalen, char *optptr, int optlen);</pre>				
DESCRIPTION	The <code>mount()</code> function requests that a removable file system contained on the block special file identified by <i>spec</i> be mounted on the directory identified by <i>dir</i>. The <i>spec</i> and <i>dir</i> arguments are pointers to path names. After a successful call to <code>mount()</code>, all references to the file <i>dir</i> refer to the root directory on the mounted file system. The mounted file system is inserted into the kernel list of all mounted file systems. This list can be examined through the mounted file system table (see <code>mnttab(4)</code>). The <i>fstype</i> argument is the file system type name. Standard file system names are defined with the prefix <code>MNTTYPE_</code> in <code><sys/mntent.h></code>. The <i>dataptr</i> argument is 0 if no file system-specific data is to be passed; otherwise it points to an area of size <i>datalen</i> that contains the file system-specific data for this mount and the <code>MS_DATA</code> flag should be set. If the <code>MS_OPTIONSTR</code> flag is set, then <i>optptr</i> points to a buffer containing the list of options to be used for this mount. The <i>optlen</i> argument specifies the length of the buffer. On completion of the <code>mount()</code> call, the options in effect for the mounted file system are returned in this buffer. If <code>MS_OPTIONSTR</code> is not specified, then the options for this mount will not appear in the mounted file systems table. The <i>mflag</i> argument is constructed by a bitwise-inclusive-OR of flags from the following list, defined in <code><sys/mount.h></code>. <table> <tr> <td><code>MS_DATA</code></td><td>If this flag is set, the <i>dataptr</i> and <i>datalen</i> arguments describe a block of file system-specific binary data at address <i>dataptr</i> of length <i>datalen</i>. This is interpreted by file system-specific code within the operating system and its format depends on the file system type. If a particular file system type does not require this data, <i>dataptr</i> and <i>datalen</i> should both be 0.</td></tr> <tr> <td><code>MS_OPTIONSTR</code></td><td>If this flag is set, the <i>optptr</i> and <i>optlen</i> arguments describe a character buffer at address <i>optptr</i> of size <i>optlen</i>. When calling <code>mount()</code>, the character buffer should contain a null-terminated string of options to be passed to the file system-specific code within the operating system. On a successful return, the file system-specific code will return the list of options recognized. Unrecognized options are ignored. The format of the string is a list of option names separated by commas. Options that have values (rather than binary options such as <code>suid</code> or <code>nosuid</code>), are separated by "=" such as <code>dev=2c4046c</code>. Standard option names are defined in <code><sys/mntent.h></code>. Only strings defined in the "C" locale are supported. The maximum length</td></tr> </table>	<code>MS_DATA</code>	If this flag is set, the <i>dataptr</i> and <i>datalen</i> arguments describe a block of file system-specific binary data at address <i>dataptr</i> of length <i>datalen</i> . This is interpreted by file system-specific code within the operating system and its format depends on the file system type. If a particular file system type does not require this data, <i>dataptr</i> and <i>datalen</i> should both be 0.	<code>MS_OPTIONSTR</code>	If this flag is set, the <i>optptr</i> and <i>optlen</i> arguments describe a character buffer at address <i>optptr</i> of size <i>optlen</i> . When calling <code>mount()</code> , the character buffer should contain a null-terminated string of options to be passed to the file system-specific code within the operating system. On a successful return, the file system-specific code will return the list of options recognized. Unrecognized options are ignored. The format of the string is a list of option names separated by commas. Options that have values (rather than binary options such as <code>suid</code> or <code>nosuid</code>), are separated by "=" such as <code>dev=2c4046c</code> . Standard option names are defined in <code><sys/mntent.h></code> . Only strings defined in the "C" locale are supported. The maximum length
<code>MS_DATA</code>	If this flag is set, the <i>dataptr</i> and <i>datalen</i> arguments describe a block of file system-specific binary data at address <i>dataptr</i> of length <i>datalen</i> . This is interpreted by file system-specific code within the operating system and its format depends on the file system type. If a particular file system type does not require this data, <i>dataptr</i> and <i>datalen</i> should both be 0.				
<code>MS_OPTIONSTR</code>	If this flag is set, the <i>optptr</i> and <i>optlen</i> arguments describe a character buffer at address <i>optptr</i> of size <i>optlen</i> . When calling <code>mount()</code> , the character buffer should contain a null-terminated string of options to be passed to the file system-specific code within the operating system. On a successful return, the file system-specific code will return the list of options recognized. Unrecognized options are ignored. The format of the string is a list of option names separated by commas. Options that have values (rather than binary options such as <code>suid</code> or <code>nosuid</code>), are separated by "=" such as <code>dev=2c4046c</code> . Standard option names are defined in <code><sys/mntent.h></code> . Only strings defined in the "C" locale are supported. The maximum length				

mount(2)

		option string that can be passed to or returned from a <code>mount ()</code> call is defined by the <code>MAX_MNTOPT_STR</code> constant. The buffer should be long enough to contain more options than were passed in, as the state of any default options that were not passed in the input option string may also be returned in the recognized options list that is returned.
	<code>MS_RDONLY</code>	The file system should be mounted for reading only. This flag should also be specified for file systems that are incapable of writing (for example, CDROM). Without this flag, writing is permitted according to individual file accessibility.
	<code>MS_NOSUID</code>	This option prevents programs that are marked set-user-ID or set-group-ID from executing (see <code>chmod(1)</code>). It also causes <code>open(2)</code> to return <code>ENXIO</code> when attempting to open block or character special files.
	<code>MS_REMOUNT</code>	Remounts a read-only file system as read-write.
	<code>MS_OVERLAY</code>	Allow the file system to be mounted over an existing file system mounted on <i>dir</i> , making the underlying file system inaccessible. If a mount is attempted on a pre-existing mount point without setting this flag, the mount will fail.
	<code>MS_GLOBAL</code>	Mount a file system globally if the system is configured and booted as part of a cluster (see <code>clinfo(1M)</code>).
RETURN VALUES	Upon successful completion, 0 is returned. Otherwise, -1 is returned and <code>errno</code> is set to indicate the error.	
ERRORS	The <code>mount ()</code> function will fail if:	
	<code>EBUSY</code>	The <i>dir</i> argument is currently mounted on, is someone's current working directory, or is otherwise busy; the device associated with <i>spec</i> is currently mounted; or there are no more mount table entries.
	<code>EFAULT</code>	The <i>spec</i> , <i>dir</i> , <i>fstype</i> , <i>dataptr</i> , or <i>optptr</i> argument points outside the allocated address space of the process.
	<code>EINVAL</code>	The super block has an invalid magic number or the <i>fstype</i> is invalid.
	<code>ELOOP</code>	Too many symbolic links were encountered in translating <i>spec</i> or <i>dir</i> .

	ENAMETOOLONG	The length of the <i>path</i> argument exceeds <code>PATH_MAX</code> , or the length of a <i>path</i> component exceeds <code>NAME_MAX</code> while <code>_POSIX_NO_TRUNC</code> is in effect.
	ENOENT	None of the named files exists or is a null pathname.
	ENOLINK	The <i>path</i> argument points to a remote machine and the link to that machine is no longer active.
	ENOSPC	The file system state in the super-block is not <code>FSOKAY</code> and <i>mflag</i> requests write permission.
	ENOTBLK	The <i>spec</i> argument is not a block special device.
	ENOTDIR	The <i>dir</i> argument is not a directory, or a component of a path prefix is not a directory.
	ENOTSUP	A global mount is attempted (the <code>MS_GLOBAL</code> flag is set in <i>mflag</i>) on a machine which is not booted as a cluster or a local mount is attempted and <i>dir</i> is within a globally mounted file system.
	ENXIO	The device associated with <i>spec</i> does not exist.
	EOVERFLOW	The length of the option string to be returned in the <i>dataptr</i> argument exceeds the size of the buffer specified by <i>datalen</i> .
	EPERM	The effective user ID is not super-user.
	EREMOTE	The <i>spec</i> argument is remote and cannot be mounted.
	EROFS	The <i>spec</i> argument is write protected and <i>mflag</i> requests write permission.
USAGE	The <code>mount ()</code> function can be invoked only by processes with super-user privileges.	
SEE ALSO	<code>mount(1M)</code> , <code>umount(2)</code> , <code>mnttab(4)</code>	
NOTES	<code>MS_OPTIONSTR</code>-type option strings should be used. Some flag bits set file system options that can also be passed in an option string. Options are first set from the option string with the last setting of an option in the string determining the value to be set by the option string. Any options controlled by flags are then applied, overriding any value set by the option string.	

mprotect(2)

NAME	mprotect – set protection of memory mapping								
SYNOPSIS	<pre>#include <sys/mman.h> int mprotect(void *<i>addr</i>, size_t <i>len</i>, int <i>prot</i>);</pre>								
DESCRIPTION	The <code>mprotect()</code> function changes the access protections on the mappings specified by the range <code>[addr, addr + len)</code>, rounding <code>len</code> up to the next multiple of the page size as returned by <code>sysconf(3C)</code>, to be that specified by <code>prot</code>. Legitimate values for <code>prot</code> are the same as those permitted for <code>mmap(2)</code> and are defined in <code><sys/mman.h></code> as: <pre>PROT_READ /* page can be read */ PROT_WRITE /* page can be written */ PROT_EXEC /* page can be executed */ PROT_NONE /* page can not be accessed */</pre> When <code>mprotect()</code> fails for reasons other than <code>EINVAL</code>, the protections on some of the pages in the range <code>[addr, addr + len)</code> may have been changed. If the error occurs on some page at <code>addr2</code>, then the protections of all whole pages in the range <code>[addr, addr2]</code> will have been modified.								
RETURN VALUES	Upon successful completion, <code>mprotect()</code> returns 0. Otherwise, it returns <code>-1</code> and sets <code>errno</code> to indicate the error.								
ERRORS	The <code>mprotect()</code> function will fail if: <table><tr><td>EACCES</td><td>The <code>prot</code> argument specifies a protection that violates the access permission the process has to the underlying memory object.</td></tr><tr><td>EINVAL</td><td>The <code>len</code> argument has a value equal to 0, or <code>addr</code> is not a multiple of the page size as returned by <code>sysconf(3C)</code>.</td></tr><tr><td>ENOMEM</td><td>Addresses in the range <code>[addr, addr + len)</code> are invalid for the address space of a process, or specify one or more pages which are not mapped.</td></tr></table> The <code>mprotect()</code> function may fail if: <table><tr><td>EAGAIN</td><td>The address range <code>[addr, addr + len)</code> includes one or more pages that have been locked in memory and that were mapped <code>MAP_PRIVATE</code>; <code>prot</code> includes <code>PROT_WRITE</code>; and the system has insufficient resources to reserve memory for the private pages that may be created. These private pages may be created by store operations in the now-writable address range.</td></tr></table>	EACCES	The <code>prot</code> argument specifies a protection that violates the access permission the process has to the underlying memory object.	EINVAL	The <code>len</code> argument has a value equal to 0, or <code>addr</code> is not a multiple of the page size as returned by <code>sysconf(3C)</code> .	ENOMEM	Addresses in the range <code>[addr, addr + len)</code> are invalid for the address space of a process, or specify one or more pages which are not mapped.	EAGAIN	The address range <code>[addr, addr + len)</code> includes one or more pages that have been locked in memory and that were mapped <code>MAP_PRIVATE</code> ; <code>prot</code> includes <code>PROT_WRITE</code> ; and the system has insufficient resources to reserve memory for the private pages that may be created. These private pages may be created by store operations in the now-writable address range.
EACCES	The <code>prot</code> argument specifies a protection that violates the access permission the process has to the underlying memory object.								
EINVAL	The <code>len</code> argument has a value equal to 0, or <code>addr</code> is not a multiple of the page size as returned by <code>sysconf(3C)</code> .								
ENOMEM	Addresses in the range <code>[addr, addr + len)</code> are invalid for the address space of a process, or specify one or more pages which are not mapped.								
EAGAIN	The address range <code>[addr, addr + len)</code> includes one or more pages that have been locked in memory and that were mapped <code>MAP_PRIVATE</code> ; <code>prot</code> includes <code>PROT_WRITE</code> ; and the system has insufficient resources to reserve memory for the private pages that may be created. These private pages may be created by store operations in the now-writable address range.								
SEE ALSO	<code>mmap(2)</code> , <code>plock(3C)</code> , <code>mlock(3C)</code> , <code>mlockall(3C)</code> , <code>sysconf(3C)</code>								

NAME	msgctl – message control operations
SYNOPSIS	<pre>#include <sys/msg.h> int msgctl(int <i>msqid</i>, int <i>cmd</i>, struct msqid_ds *<i>buf</i>);</pre>
DESCRIPTION	The <code>msgctl()</code> function provides a variety of message control operations as specified by <i>cmd</i>. The following <i>cmds</i> are available: IPC_STAT Place the current value of each member of the data structure associated with <i>msqid</i> into the structure pointed to by <i>buf</i>. The contents of this structure are defined in <code>intro(2)</code>. IPC_SET Set the value of the following members of the data structure associated with <i>msqid</i> to the corresponding value found in the structure pointed to by <i>buf</i>: <pre>msg_perm.uid msg_perm.gid msg_perm.mode /* access permission bits only */ msg_qbytes</pre> This <i>cmd</i> can only be executed by a process that has an effective user ID equal to either that of super-user, or to the value of <code>msg_perm.cuid</code> or <code>msg_perm.uid</code> in the data structure associated with <i>msqid</i>. Only super-user can raise the value of <code>msg_qbytes</code>. IPC_RMID Remove the message queue identifier specified by <i>msqid</i> from the system and destroy the message queue and data structure associated with it. This <i>cmd</i> can only be executed by a process that has an effective user ID equal to either that of super-user, or to the value of <code>msg_perm.cuid</code> or <code>msg_perm.uid</code> in the data structure associated with <i>msqid</i>. The <i>buf</i> argument is ignored.
RETURN VALUES	Upon successful completion, <code>msgctl()</code> returns 0. Otherwise, it returns -1 and sets <code>errno</code> to indicate the error.
ERRORS	The <code>msgctl()</code> function will fail if: EACCES The <i>cmd</i> argument is IPC_STAT and operation permission is denied to the calling process (see <code>intro(2)</code>). EFAULT The <i>buf</i> argument points to an illegal address. EINVAL The <i>msqid</i> argument is not a valid message queue identifier; or the <i>cmd</i> argument is not a valid command or is IPC_SET and <code>msg_perm.uid</code> or <code>msg_perm.gid</code> is not valid. EPERM The <i>cmd</i> argument is IPC_RMID or IPC_SET and the effective user ID of the calling process is not super-user and is not equal to the value of <code>msg_perm.cuid</code> or <code>msg_perm.uid</code> in the data structure associated with <i>msqid</i>.

msgctl(2)

EPERM	The <i>cmd</i> argument is <code>IPC_SET</code> , an attempt is being made to increase to the value of <code>msg_qbytes</code> , and the effective user ID of the calling process is not super-user.
E_OVERFLOW	The <i>cmd</i> argument is <code>IPC_STAT</code> and <i>uid</i> or <i>gid</i> is too large to be stored in the structure pointed to by <i>buf</i> .

SEE ALSO `intro(2)`, `msgget(2)`, `msgrcv(2)`, `msgsnd(2)`

NAME	msgget – get message queue								
SYNOPSIS	<pre>#include <sys/msg.h> int msgget(key_t <i>key</i>, int <i>msgflg</i>);</pre>								
DESCRIPTION	The <code>msgget()</code> argument returns the message queue identifier associated with <i>key</i>. A message queue identifier and associated message queue and data structure (see intro(3)) are created for <i>key</i> if one of the following are true: ■ <i>key</i> is <code>IPC_PRIVATE</code>. ■ <i>key</i> does not already have a message queue identifier associated with it, and <code>(msgflg&IPC_CREAT)</code> is true. On creation, the data structure associated with the new message queue identifier is initialized as follows: ■ <code>msg_perm.cuid</code>, <code>msg_perm.uid</code>, <code>msg_perm.cgid</code>, and <code>msg_perm.gid</code> are set to the effective user ID and effective group ID, respectively, of the calling process. ■ The low-order 9 bits of <code>msg_perm.mode</code> are set to the low-order 9 bits of <i>msgflg</i>. ■ <code>msg_qnum</code>, <code>msg_lspid</code>, <code>msg_lrpid</code>, <code>msg_stime</code>, and <code>msg_rtime</code> are set to 0. ■ <code>msg_ctime</code> is set to the current time. ■ <code>msg_qbytes</code> is set to the system limit.								
RETURN VALUES	Upon successful completion, a non-negative integer representing a message queue identifier is returned. Otherwise, <code>-1</code> is returned and <code>errno</code> is set to indicate the error.								
ERRORS	The <code>msgget()</code> function will fail if: <table> <tr> <td>EACCES</td><td>A message queue identifier exists for <i>key</i>, but operation permission (see intro(3)) as specified by the low-order 9 bits of <i>msgflg</i> would not be granted.</td></tr> <tr> <td>EEXIST</td><td>A message queue identifier exists for <i>key</i> but <code>(msgflg&IPC_CREAT)</code> and <code>(msgflg&IPC_EXCL)</code> are both true.</td></tr> <tr> <td>ENOENT</td><td>A message queue identifier does not exist for <i>key</i> and <code>(msgflg&IPC_CREAT)</code> is false.</td></tr> <tr> <td>ENOSPC</td><td>A message queue identifier is to be created but the system-imposed limit on the maximum number of allowed message queue identifiers system wide would be exceeded.</td></tr> </table>	EACCES	A message queue identifier exists for <i>key</i> , but operation permission (see intro(3)) as specified by the low-order 9 bits of <i>msgflg</i> would not be granted.	EEXIST	A message queue identifier exists for <i>key</i> but <code>(msgflg&IPC_CREAT)</code> and <code>(msgflg&IPC_EXCL)</code> are both true.	ENOENT	A message queue identifier does not exist for <i>key</i> and <code>(msgflg&IPC_CREAT)</code> is false.	ENOSPC	A message queue identifier is to be created but the system-imposed limit on the maximum number of allowed message queue identifiers system wide would be exceeded.
EACCES	A message queue identifier exists for <i>key</i> , but operation permission (see intro(3)) as specified by the low-order 9 bits of <i>msgflg</i> would not be granted.								
EEXIST	A message queue identifier exists for <i>key</i> but <code>(msgflg&IPC_CREAT)</code> and <code>(msgflg&IPC_EXCL)</code> are both true.								
ENOENT	A message queue identifier does not exist for <i>key</i> and <code>(msgflg&IPC_CREAT)</code> is false.								
ENOSPC	A message queue identifier is to be created but the system-imposed limit on the maximum number of allowed message queue identifiers system wide would be exceeded.								
SEE ALSO	intro(3) , msgctl(2) , msgrcv(2) , msgsnd(2) , ftok(3C)								

msgids(2)

NAME	msgids – discover all message queue identifiers
SYNOPSIS	<pre>#include <sys/msg.h> int msgids(int *buf, uint_t nids, uint_t *pnids);</pre>
DESCRIPTION	The <code>msgids()</code> function copies all active message queue identifiers from the system into the user-defined buffer specified by <i>buf</i>, provided that the number of such identifiers is not greater than the number of integers the buffer can contain, as specified by <i>nids</i>. If the size of the buffer is insufficient to contain all of the active message queue identifiers in the system, none are copied. Whether or not the size of the buffer is sufficient to contain all of them, the number of active message queue identifiers in the system is copied into the unsigned integer pointed to by <i>pnids</i>. If <i>nids</i> is 0 or less than the number of active message queue identifiers in the system, <i>buf</i> is ignored.
RETURN VALUES	Upon successful completion, <code>msgids()</code> returns 0. Otherwise, -1 is returned and <code>errno</code> is set to indicate the error.
ERRORS	The <code>msgids()</code> function will fail if: EFAULT The <i>buf</i> or <i>pnids</i> argument points to an illegal address.
USAGE	The <code>msgids()</code> function returns a snapshot of all the active message queue identifiers in the system. More may be added and some may be removed before they can be used by the caller.
EXAMPLES	EXAMPLE 1 <code>msgids()</code> example This is sample C code indicating how to use the <code>msgids()</code> function (see <code>msgsnap(2)</code>): <pre>void examine_queues() { int *ids = NULL; uint_t nids = 0; uint_t n; int i; for (;;) { if (msgids(ids, nids, &n) != 0) { perror("msgids"); exit(1); } if (n <= nids) /* we got them all */ break; /* we need a bigger buffer */ ids = realloc(ids, (nids = n) * sizeof (int)); } for (i = 0; i < n; i++)</pre>

EXAMPLE 1 msgids() example (Continued)

```
        process_msgid(ids[i]);  
    free(ids);  
}
```

ATTRIBUTES See attributes(5) for descriptions of the following attributes:

ATTRIBUTE TYPE	ATTRIBUTE VALUE
MT-Level	Async-Signal-Safe

SEE ALSO ipcrm(1), ipcs(1), intro(2), msgctl(2), msgget(2), msgsnap(2), msgrcv(2), msgsnd(2), attributes(5)

msgrcv(2)

NAME	msgrcv – message receive operation
SYNOPSIS	<pre>#include <sys/msg.h> ssize_t msgrcv(int <i>msqid</i>, void *<i>msgp</i>, size_t <i>msgsz</i>, long int <i>msgtyp</i>, int <i>msgflg</i>);</pre>
DESCRIPTION	The <code>msgrcv()</code> function reads a message from the queue associated with the message queue identifier specified by <i>msqid</i> and places it in the user-defined buffer pointed to by <i>msgp</i>. The <i>msgp</i> argument points to a user-defined buffer that must contain first a field of type <code>long int</code> that will specify the type of the message, and then a data portion that will hold the data bytes of the message. The structure below is an example of what this user-defined buffer might look like: <pre>struct mymsg { long int mtype; /* message type */ char mtext[1]; /* message text */ }</pre> The <code>mtype</code> member is the received message's type as specified by the sending process. The <code>mtext</code> member is the text of the message. The <i>msgsz</i> argument specifies the size in bytes of <code>mtext</code>. The received message is truncated to <i>msgsz</i> bytes if it is larger than <i>msgsz</i> and (<i>msgflg</i> & <code>MSG_NOERROR</code>) is non-zero. The truncated part of the message is lost and no indication of the truncation is given to the calling process. The <i>msgtyp</i> argument specifies the type of message requested as follows: ■ If <i>msgtyp</i> is 0, the first message on the queue is received. ■ If <i>msgtyp</i> is greater than 0, the first message of type <i>msgtyp</i> is received. ■ If <i>msgtyp</i> is less than 0, the first message of the lowest type that is less than or equal to the absolute value of <i>msgtyp</i> is received. The <i>msgflg</i> argument specifies which of the following actions is to be taken if a message of the desired type is not on the queue: ■ If (<i>msgflg</i> & <code>IPC_NOWAIT</code>) is non-zero, the calling process will return immediately with a return value of -1 and <code>errno</code> set to <code>ENOMSG</code>. ■ If (<i>msgflg</i> & <code>IPC_NOWAIT</code>) is 0, the calling process will suspend execution until one of the following occurs: ■ A message of the desired type is placed on the queue. ■ The message queue identifier <i>msqid</i> is removed from the system (see <code>msgctl(2)</code>); when this occurs, <code>errno</code> is set equal to <code>EIDRM</code> and -1 is returned. ■ The calling process receives a signal that is to be caught; in this case a message is not received and the calling process resumes execution in the manner prescribed in <code>sigaction(2)</code>.

	Upon successful completion, the following actions are taken with respect to the data structure associated with <i>msqid</i> (see <code>intro(2)</code>):														
	■ <code>msg_qnum</code> is decremented by 1. ■ <code>msg_lrpid</code> is set equal to the process ID of the calling process. ■ <code>msg_rtime</code> is set equal to the current time.														
RETURN VALUES	Upon successful completion, <code>msgrcv()</code> returns a value equal to the number of bytes actually placed into the buffer <i>mtext</i> . Otherwise, <code>-1</code> is returned, no message is received, and <code>errno</code> is set to indicate the error.														
ERRORS	The <code>msgrcv()</code> function will fail if: <table> <tr> <td><code>E2BIG</code></td><td>The value of <i>mtext</i> is greater than <i>msgsz</i> and (<i>msgflg</i>&<code>MSG_NOERROR</code>) is 0.</td></tr> <tr> <td><code>EACCES</code></td><td>Operation permission is denied to the calling process. See <code>intro(2)</code>.</td></tr> <tr> <td><code>EIDRM</code></td><td>The message queue identifier <i>msqid</i> is removed from the system.</td></tr> <tr> <td><code>EINTR</code></td><td>The <code>msgrcv()</code> function was interrupted by a signal.</td></tr> <tr> <td><code>EINVAL</code></td><td>The <i>msqid</i> argument is not a valid message queue identifier.</td></tr> <tr> <td><code>ENOMSG</code></td><td>The queue does not contain a message of the desired type and (<i>msgflg</i>&<code>IPC_NOWAIT</code>) is non-zero.</td></tr> </table> The <code>msgrcv()</code> function may fail if: <table> <tr> <td><code>EFAULT</code></td><td>The <i>msgp</i> argument points to an illegal address.</td></tr> </table>	<code>E2BIG</code>	The value of <i>mtext</i> is greater than <i>msgsz</i> and (<i>msgflg</i> & <code>MSG_NOERROR</code>) is 0.	<code>EACCES</code>	Operation permission is denied to the calling process. See <code>intro(2)</code> .	<code>EIDRM</code>	The message queue identifier <i>msqid</i> is removed from the system.	<code>EINTR</code>	The <code>msgrcv()</code> function was interrupted by a signal.	<code>EINVAL</code>	The <i>msqid</i> argument is not a valid message queue identifier.	<code>ENOMSG</code>	The queue does not contain a message of the desired type and (<i>msgflg</i> & <code>IPC_NOWAIT</code>) is non-zero.	<code>EFAULT</code>	The <i>msgp</i> argument points to an illegal address.
<code>E2BIG</code>	The value of <i>mtext</i> is greater than <i>msgsz</i> and (<i>msgflg</i> & <code>MSG_NOERROR</code>) is 0.														
<code>EACCES</code>	Operation permission is denied to the calling process. See <code>intro(2)</code> .														
<code>EIDRM</code>	The message queue identifier <i>msqid</i> is removed from the system.														
<code>EINTR</code>	The <code>msgrcv()</code> function was interrupted by a signal.														
<code>EINVAL</code>	The <i>msqid</i> argument is not a valid message queue identifier.														
<code>ENOMSG</code>	The queue does not contain a message of the desired type and (<i>msgflg</i> & <code>IPC_NOWAIT</code>) is non-zero.														
<code>EFAULT</code>	The <i>msgp</i> argument points to an illegal address.														
USAGE	The value passed as the <i>msgp</i> argument should be converted to type <code>void *</code> .														
SEE ALSO	<code>intro(2)</code> , <code>msgctl(2)</code> , <code>msgget(2)</code> , <code>msgsnd(2)</code> , <code>sigaction(2)</code>														

msgsnap(2)

NAME	msgsnap – message queue snapshot operation		
SYNOPSIS	<pre>#include <sys/msg.h> msgsnap(int <i>msqid</i>, void *<i>buf</i>, size_t <i>bufsz</i>, long <i>msgtyp</i>);</pre>		
DESCRIPTION	The <code>msgsnap()</code> function reads all of the messages of type <i>msgtyp</i> from the queue associated with the message queue identifier specified by <i>msqid</i> and places them in the user-defined buffer pointed to by <i>buf</i>. The <i>buf</i> argument points to a user-defined buffer that on return will contain first a buffer header structure: <pre>struct msgsnap_head { size_t msgsnap_size; /* bytes used/required in the buffer */ size_t msgsnap_nmsg; /* number of messages in the buffer */ };</pre> followed by <code>msgsnap_nmsg</code> messages, each of which starts with a message header: <pre>struct msgsnap_mhead { size_t msgsnap_mlen; /* number of bytes in the message */ long msgsnap_mtype; /* message type */ };</pre> and followed by <code>msgsnap_mlen</code> bytes containing the message contents. Each subsequent message header is located at the first byte following the previous message contents, rounded up to a <code>sizeof(size_t)</code> boundary. The <i>bufsz</i> argument specifies the size of <i>buf</i> in bytes. If <i>bufsz</i> is less than <code>sizeof(msgsnap_head)</code>, <code>msgsnap()</code> fails with <code>EINVAL</code>. If <i>bufsz</i> is insufficient to contain all of the requested messages, <code>msgsnap()</code> succeeds but returns with <code>msgsnap_nmsg</code> set to 0 and with <code>msgsnap_size</code> set to the required size of the buffer in bytes. The <i>msgtyp</i> argument specifies the types of messages requested as follows: ■ If <i>msgtyp</i> is 0, all of the messages on the queue are read. ■ If <i>msgtyp</i> is greater than 0, all messages of type <i>msgtyp</i> are read. ■ If <i>msgtyp</i> is less than 0, all messages with type less than or equal to the absolute value of <i>msgtyp</i> are read. The <code>msgsnap()</code> function is a non-destructive operation. Upon completion, no changes are made to the data structures associated with <i>msqid</i>.		
RETURN VALUES	Upon successful completion, <code>msgsnap()</code> returns 0. Otherwise, -1 is returned and <code>errno</code> is set to indicate the error.		
ERRORS	The <code>msgsnap()</code> function will fail if: <table><tr><td>EACCES</td><td>Operation permission is denied to the calling process. See <code>intro(2)</code>.</td></tr></table>	EACCES	Operation permission is denied to the calling process. See <code>intro(2)</code> .
EACCES	Operation permission is denied to the calling process. See <code>intro(2)</code> .		

EINVAL The *msqid* argument is not a valid message queue identifier or the value of *bufsz* is less than `sizeof(struct msgsnap_head)`.

EFAULT The *buf* argument points to an illegal address.

USAGE The `msgsnap()` function returns a snapshot of messages on a message queue at one point in time. The queue contents can change immediately following return from `msgsnap()`.

EXAMPLES **EXAMPLE 1** `msgsnap()` example

This is sample C code indicating how to use the `msgsnap` function (see `msgids(2)`).

```
void
process_msgid(int msqid)
{
    size_t bufsize;
    struct msgsnap_head *buf;
    struct msgsnap_mhead *mhead;
    int i;

    /* allocate a minimum-size buffer */
    buf = malloc(bufsize = sizeof(struct msgsnap_head));

    /* read all of the messages from the queue */
    for (;;) {
        if (msgsnap(msqid, buf, bufsize, 0) != 0) {
            perror("msgsnap");
            free(buf);
            return;
        }
        if (bufsize >= buf->msgsnap_size) /* we got them all */
            break;
        /* we need a bigger buffer */
        buf = realloc(buf, bufsize = buf->msgsnap_size);
    }

    /* process each message in the queue (there may be none) */
    mhead = (struct msgsnap_mhead *) (buf + 1); /* first message */
    for (i = 0; i < buf->msgsnap_nmsg; i++) {
        size_t mlen = mhead->msgsnap_mlen;

        /* process the message contents */
        process_message(mhead->msgsnap_mtype, (char *) (mhead+1), mlen);

        /* advance to the next message header */
        mhead = (struct msgsnap_mhead *)
            ((char *) mhead + sizeof(struct msgsnap_mhead) +
             (mlen + sizeof(size_t) - 1) & ~(sizeof(size_t) - 1));
    }

    free(buf);
}
```

msgsnap(2)

ATTRIBUTES See attributes(5) for descriptions of the following attributes:

ATTRIBUTE TYPE	ATTRIBUTE VALUE
MT-Level	Async-Signal-Safe

SEE ALSO ipcrm(1), ipcs(1), intro(2), msgctl(2), msgget(2), msgids(2), msgrcv(2), msgsnd(2), attributes(5)

NAME	msgsnd – message send operation
SYNOPSIS	<pre>#include <sys/msg.h> int msgsnd(int <i>msqid</i>, const void *<i>msgp</i>, size_t <i>msgsz</i>, int <i>msgflg</i>);</pre>
DESCRIPTION	The <code>msgsnd()</code> function is used to send a message to the queue associated with the message queue identifier specified by <i>msqid</i>. The <i>msgp</i> argument points to a user-defined buffer that must contain first a field of type <code>long int</code> that will specify the type of the message, and then a data portion that will hold the data bytes of the message. The structure below is an example of what this user-defined buffer might look like: <pre>struct mymsg { long mtype; /* message type */ char mtext[1]; /* message text */ }</pre> The <code>mtype</code> member is a non-zero positive type <code>long int</code> that can be used by the receiving process for message selection. The <code>mtext</code> member is any text of length <i>msgsz</i> bytes. The <i>msgsz</i> argument can range from 0 to a system-imposed maximum. The <i>msgflg</i> argument specifies the action to be taken if one or more of the following are true: ■ The number of bytes already on the queue is equal to <code>msg_qbytes</code>; see <code>intro(2)</code>. ■ The total number of messages on all queues system-wide is equal to the system-imposed limit. These actions are as follows: ■ If (<i>msgflg</i> & <code>IPC_NOWAIT</code>) is non-zero, the message will not be sent and the calling process will return immediately. ■ If (<i>msgflg</i> & <code>IPC_NOWAIT</code>) is 0, the calling process will suspend execution until one of the following occurs: ■ The condition responsible for the suspension no longer exists, in which case the message is sent. ■ The message queue identifier <i>msqid</i> is removed from the system (see <code>msgctl(2)</code>); when this occurs, <code>errno</code> is set equal to <code>EIDRM</code> and <code>-1</code> is returned. ■ The calling process receives a signal that is to be caught; in this case the message is not sent and the calling process resumes execution in the manner prescribed in <code>sigaction(2)</code>. Upon successful completion, the following actions are taken with respect to the data structure associated with <i>msqid</i> (see <code>intro(2)</code>): ■ <code>msg_qnum</code> is incremented by 1. ■ <code>msg_lspid</code> is set equal to the process ID of the calling process.

msgsnd(2)

	■ <code>msg_stime</code> is set equal to the current time.												
RETURN VALUES	Upon successful completion, 0 is returned. Otherwise, -1 is returned, no message is sent, and <code>errno</code> is set to indicate the error.												
ERRORS	The <code>msgsnd()</code> function will fail if: <table><tr><td>EACCES</td><td>Operation permission is denied to the calling process. See <code>intro(2)</code>.</td></tr><tr><td>EAGAIN</td><td>The message cannot be sent for one of the reasons cited above and (<code>msgflg</code>&<code>IPC_NOWAIT</code>) is non-zero.</td></tr><tr><td>EIDRM</td><td>The message queue identifier <i>msgid</i> is removed from the system.</td></tr><tr><td>EINTR</td><td>The <code>msgsnd()</code> function was interrupted by a signal.</td></tr><tr><td>EINVAL</td><td>The value of <i>msgid</i> is not a valid message queue identifier, or the value of <code>mtype</code> is less than 1; or the value of <i>msgsz</i> is less than 0 or greater than the system-imposed limit.</td></tr></table> The <code>msgsnd()</code> function may fail if: <table><tr><td>EFAULT</td><td>The <i>msgp</i> argument points to an illegal address.</td></tr></table>	EACCES	Operation permission is denied to the calling process. See <code>intro(2)</code> .	EAGAIN	The message cannot be sent for one of the reasons cited above and (<code>msgflg</code> & <code>IPC_NOWAIT</code>) is non-zero.	EIDRM	The message queue identifier <i>msgid</i> is removed from the system.	EINTR	The <code>msgsnd()</code> function was interrupted by a signal.	EINVAL	The value of <i>msgid</i> is not a valid message queue identifier, or the value of <code>mtype</code> is less than 1; or the value of <i>msgsz</i> is less than 0 or greater than the system-imposed limit.	EFAULT	The <i>msgp</i> argument points to an illegal address.
EACCES	Operation permission is denied to the calling process. See <code>intro(2)</code> .												
EAGAIN	The message cannot be sent for one of the reasons cited above and (<code>msgflg</code> & <code>IPC_NOWAIT</code>) is non-zero.												
EIDRM	The message queue identifier <i>msgid</i> is removed from the system.												
EINTR	The <code>msgsnd()</code> function was interrupted by a signal.												
EINVAL	The value of <i>msgid</i> is not a valid message queue identifier, or the value of <code>mtype</code> is less than 1; or the value of <i>msgsz</i> is less than 0 or greater than the system-imposed limit.												
EFAULT	The <i>msgp</i> argument points to an illegal address.												
USAGE	The value passed as the <i>msgp</i> argument should be converted to type <code>void *</code> .												
SEE ALSO	<code>intro(2)</code> , <code>msgctl(2)</code> , <code>msgget(2)</code> , <code>msgrcv(2)</code> , <code>sigaction(2)</code>												

NAME	munmap – unmap pages of memory		
SYNOPSIS	<pre>#include <sys/mman.h> int munmap(void *addr, size_t len);</pre>		
DESCRIPTION	The <code>munmap()</code> function removes the mappings for pages in the range <code>[addr, addr + len)</code>, rounding the <code>len</code> argument up to the next multiple of the page size as returned by <code>sysconf(3C)</code>. If <code>addr</code> is not the address of a mapping established by a prior call to <code>mmap(2)</code>, the behavior is undefined. After a successful call to <code>munmap()</code> and before any subsequent mapping of the unmapped pages, further references to these pages will result in the delivery of a <code>SIGBUS</code> or <code>SIGSEGV</code> signal to the process. The <code>mmap(2)</code> function often performs an implicit <code>munmap()</code>.		
RETURN VALUES	Upon successful completion, <code>munmap()</code> returns 0; otherwise, it returns -1 and sets <code>errno</code> to indicate an error.		
ERRORS	The <code>munmap()</code> function will fail if: <table> <tr> <td><code>EINVAL</code></td><td>The <code>addr</code> argument is not a multiple of the page size as returned by <code>sysconf(3C)</code>; addresses in the range <code>[addr, addr + len)</code> are outside the valid range for the address space of a process; or the <code>len</code> argument has a value less than or equal to 0.</td></tr> </table>	<code>EINVAL</code>	The <code>addr</code> argument is not a multiple of the page size as returned by <code>sysconf(3C)</code> ; addresses in the range <code>[addr, addr + len)</code> are outside the valid range for the address space of a process; or the <code>len</code> argument has a value less than or equal to 0.
<code>EINVAL</code>	The <code>addr</code> argument is not a multiple of the page size as returned by <code>sysconf(3C)</code> ; addresses in the range <code>[addr, addr + len)</code> are outside the valid range for the address space of a process; or the <code>len</code> argument has a value less than or equal to 0.		
SEE ALSO	<code>mmap(2)</code> , <code>sysconf(3C)</code>		

nice(2)

NAME	nice – change priority of a process				
SYNOPSIS	<pre>#include <unistd.h> int nice(int incr);</pre>				
DESCRIPTION	The <code>nice()</code> function allows a process to change its priority. The invoking process must be in a scheduling class that supports the <code>nice()</code>. The <code>nice()</code> function adds the value of <code>incr</code> to the nice value of the calling process. A process's nice value is a non-negative number for which a greater positive value results in lower CPU priority. A maximum nice value of $(2 * NZERO) - 1$ and a minimum nice value of 0 are imposed by the system. <code>NZERO</code> is defined in <code><limits.h></code> with a default value of 20. Requests for values above or below these limits result in the nice value being set to the corresponding limit. A nice value of 40 is treated as 39. Only a process with super-user privileges can lower the nice value.				
RETURN VALUES	Upon successful completion, <code>nice()</code> returns the new nice value minus <code>NZERO</code> . Otherwise, <code>-1</code> is returned, the process's nice value is not changed, and <code>errno</code> is set to indicate the error.				
ERRORS	The <code>nice()</code> function will fail if: <table><tr><td>EINVAL</td><td>The <code>nice()</code> function is called by a process in a scheduling class other than time-sharing.</td></tr><tr><td>EPERM</td><td>The <code>inc</code> argument is negative or greater than 40 and the effective user ID of the calling process is not super-user.</td></tr></table>	EINVAL	The <code>nice()</code> function is called by a process in a scheduling class other than time-sharing.	EPERM	The <code>inc</code> argument is negative or greater than 40 and the effective user ID of the calling process is not super-user.
EINVAL	The <code>nice()</code> function is called by a process in a scheduling class other than time-sharing.				
EPERM	The <code>inc</code> argument is negative or greater than 40 and the effective user ID of the calling process is not super-user.				
USAGE	The <code>prctl(2)</code> function is a more general interface to scheduler functions. Since <code>-1</code> is a permissible return value in a successful situation, an application wishing to check for error situations should set <code>errno</code> to 0, then call <code>nice()</code>, and if it returns <code>-1</code>, check to see if <code>errno</code> is non-zero.				
SEE ALSO	<code>nice(1)</code> , <code>exec(2)</code> , <code>prctl(2)</code>				

NAME	ntp_adjtime – adjust local clock parameters						
SYNOPSIS	<pre>#include <sys/timex.h> int ntp_adjtime(struct timex *tptr);</pre>						
DESCRIPTION	The <code>ntp_adjtime()</code> function adjusts the parameters used to discipline the local clock, according to the values in the struct <code>timex</code> pointed to by <i>tptr</i>. Before returning, it fills in the structure with the most recent values kept in the kernel. The adjustment is effected in part by speeding up or slowing down the clock, as necessary, and in part by phase-locking onto a once-per second pulse (PPS) provided by a driver, if available. <pre>struct timex { uint32_t modes; /* clock mode bits (w) */ int32_t offset; /* time offset (us) (rw) */ int32_t freq; /* frequency offset (scaled ppm) (rw) */ int32_t maxerror; /* maximum error (us) (rw) */ int32_t esterror; /* estimated error (us) (rw) */ int32_t status; /* clock status bits (rw) */ int32_t constant; /* pll time constant (rw) */ int32_t precision; /* clock precision (us) (r) */ int32_t tolerance; /* clock frequency tolerance (scaled ppm) (r) */ int32_t ppsfreq; /* pps frequency (scaled ppm) (r) */ int32_t jitter; /* pps jitter (us) (r) */ int32_t shift; /* interval duration (s) (shift) (r) */ int32_t stabil; /* pps stability (scaled ppm) (r) */ int32_t jitcnt; /* jitter limit exceeded (r) */ int32_t calcnt; /* calibration intervals (r) */ int32_t errcnt; /* calibration errors (r) */ int32_t stbcnt; /* stability limit exceeded (r) */ };</pre>						
RETURN VALUES	Upon successful completion, <code>ntp_adjtime()</code> returns the current clock state (see <code><sys/timex.h></code>). Otherwise, it returns <code>-1</code> and sets <code>errno</code> to indicate the error.						
ERRORS	The <code>ntp_adjtime()</code> function will fail if: <table> <tr> <td>EFAULT</td><td>The <i>tptr</i> argument is an invalid pointer.</td></tr> <tr> <td>EINVAL</td><td>The constant member of the structure pointed to by <i>tptr</i> is less than 0 or greater than 30.</td></tr> <tr> <td>EPERM</td><td>The user is not super-user.</td></tr> </table>	EFAULT	The <i>tptr</i> argument is an invalid pointer.	EINVAL	The constant member of the structure pointed to by <i>tptr</i> is less than 0 or greater than 30.	EPERM	The user is not super-user.
EFAULT	The <i>tptr</i> argument is an invalid pointer.						
EINVAL	The constant member of the structure pointed to by <i>tptr</i> is less than 0 or greater than 30.						
EPERM	The user is not super-user.						
SEE ALSO	<code>xntpd(1M)</code> , <code>ntp_gettime(2)</code>						

ntp_gettime(2)

NAME	ntp_gettime – get local clock values				
SYNOPSIS	<pre>#include <sys/timex.h> int ntp_gettime(struct ntptimeval *tpr) ;</pre>				
DESCRIPTION	The <code>ntp_gettime()</code> function reads the local clock value and dispersion, returning the information in <i>tpr</i>. The <code>ntptimeval</code> structure contains the following members: <pre>struct ntptimeval { struct timeval time; /* current time (ro) */ int32_t maxerror; /* maximum error (us) (ro) */ int32_t esterror; /* estimated error (us) (ro) */ };</pre>				
RETURN VALUES	Upon successful completion, <code>ntp_gettime()</code> returns the current clock state (see <code><sys/timex.h></code>). Otherwise, it returns <code>-1</code> and sets <code>errno</code> to indicate the error.				
ERRORS	The <code>ntp_gettime()</code> function will fail if: <table><tr><td>EFAULT</td><td>The <i>tpr</i> argument points to an invalid address.</td></tr></table> The <code>ntp_gettime()</code> function will fail for 32-bit interfaces if: <table><tr><td>E_OVERFLOW</td><td>The size of the <code>time.tv_sec</code> member of the <code>ntptimeval</code> structure pointed to by <i>tpr</i> is too small to contain the correct number of seconds.</td></tr></table>	EFAULT	The <i>tpr</i> argument points to an invalid address.	E_OVERFLOW	The size of the <code>time.tv_sec</code> member of the <code>ntptimeval</code> structure pointed to by <i>tpr</i> is too small to contain the correct number of seconds.
EFAULT	The <i>tpr</i> argument points to an invalid address.				
E_OVERFLOW	The size of the <code>time.tv_sec</code> member of the <code>ntptimeval</code> structure pointed to by <i>tpr</i> is too small to contain the correct number of seconds.				
SEE ALSO	<code>xntpd(1M)</code> , <code>ntp_adjtime(2)</code>				

NAME	open – open a file										
SYNOPSIS	<pre>#include <sys/types.h> #include <sys/stat.h> #include <fcntl.h> int open(const char *path, int oflag, /* mode_t mode */...);</pre>										
DESCRIPTION	The <code>open()</code> function establishes the connection between a file and a file descriptor. It creates an open file description that refers to a file and a file descriptor that refers to that open file description. The file descriptor is used by other I/O functions to refer to that file. The <i>path</i> argument points to a pathname naming the file. The <code>open()</code> function returns a file descriptor for the named file that is the lowest file descriptor not currently open for that process. The open file description is new, and therefore the file descriptor does not share it with any other process in the system. The <code>FD_CLOEXEC</code> file descriptor flag associated with the new file descriptor is cleared. The file offset used to mark the current position within the file is set to the beginning of the file. The file status flags and file access modes of the open file description are set according to the value of <i>oflag</i>. The <i>mode</i> argument is used only when <code>O_CREAT</code> is specified (see below.) Values for <i>oflag</i> are constructed by a bitwise-inclusive-OR of flags from the following list, defined in <code><fcntl.h></code>. Applications must specify exactly one of the first three values (file access modes) below in the value of <i>oflag</i>: <table> <tr> <td><code>O_RDONLY</code></td><td>Open for reading only.</td></tr> <tr> <td><code>O_WRONLY</code></td><td>Open for writing only.</td></tr> <tr> <td><code>O_RDWR</code></td><td>Open for reading and writing. The result is undefined if this flag is applied to a FIFO.</td></tr> </table> Any combination of the following may be used: <table> <tr> <td><code>O_APPEND</code></td><td>If set, the file offset is set to the end of the file prior to each write.</td></tr> <tr> <td><code>O_CREAT</code></td><td>Create the file if it does not exist. This flag requires that the <i>mode</i> argument be specified.</td></tr> </table> If the file exists, this flag has no effect except as noted under <code>O_EXCL</code> below. Otherwise, the file is created with the user ID of the file set to the effective user ID of the process. The group ID of the file is set to the effective group IDs of the process, or if the <code>S_ISGID</code> bit is set in the directory in which the file is being created, the file's group ID is set to the group ID of its parent directory. If the group ID of the new file does not match the effective group ID or one of the supplementary groups IDs, the <code>S_ISGID</code> bit is cleared. The access permission bits (see	<code>O_RDONLY</code>	Open for reading only.	<code>O_WRONLY</code>	Open for writing only.	<code>O_RDWR</code>	Open for reading and writing. The result is undefined if this flag is applied to a FIFO.	<code>O_APPEND</code>	If set, the file offset is set to the end of the file prior to each write.	<code>O_CREAT</code>	Create the file if it does not exist. This flag requires that the <i>mode</i> argument be specified.
<code>O_RDONLY</code>	Open for reading only.										
<code>O_WRONLY</code>	Open for writing only.										
<code>O_RDWR</code>	Open for reading and writing. The result is undefined if this flag is applied to a FIFO.										
<code>O_APPEND</code>	If set, the file offset is set to the end of the file prior to each write.										
<code>O_CREAT</code>	Create the file if it does not exist. This flag requires that the <i>mode</i> argument be specified.										

open(2)

	<sys/stat.h>) of the file mode are set to the value of <i>mode</i>, modified as follows (see <code>creat(2)</code>): a bitwise-AND is performed on the file-mode bits and the corresponding bits in the complement of the process's file mode creation mask. Thus, all bits set in the process's file mode creation mask (see <code>umask(2)</code>) are correspondingly cleared in the file's permission mask. The "save text image after execution bit" of the mode is cleared (see <code>chmod(2)</code>). <code>O_SYNC</code> Write I/O operations on the file descriptor complete as defined by synchronized I/O file integrity completion (see <code>fcntl(3HEAD)</code> definition of <code>O_SYNC</code>.) When bits other than the file permission bits are set, the effect is unspecified. The <i>mode</i> argument does not affect whether the file is open for reading, writing or for both.
<code>O_DSYNC</code>	Write I/O operations on the file descriptor complete as defined by synchronized I/O data integrity completion.
<code>O_EXCL</code>	If <code>O_CREAT</code> and <code>O_EXCL</code> are set, <code>open()</code> fails if the file exists. The check for the existence of the file and the creation of the file if it does not exist is atomic with respect to other processes executing <code>open()</code> naming the same filename in the same directory with <code>O_EXCL</code> and <code>O_CREAT</code> set. If <code>O_CREAT</code> is not set, the effect is undefined.
<code>O_LARGEFILE</code>	If set, the offset maximum in the open file description is the largest value that can be represented correctly in an object of type <code>off64_t</code> .
<code>O_NOCTTY</code>	If set and <i>path</i> identifies a terminal device, <code>open()</code> does not cause the terminal device to become the controlling terminal for the process.
<code>O_NONBLOCK</code> or <code>O_NDELAY</code>	These flags may affect subsequent reads and writes (see <code>read(2)</code> and <code>write(2)</code>). If both <code>O_NDELAY</code> and <code>O_NONBLOCK</code> are set, <code>O_NONBLOCK</code> takes precedence. When opening a FIFO with <code>O_RDONLY</code> or <code>O_WRONLY</code> set: If <code>O_NONBLOCK</code> or <code>O_NDELAY</code> is set: An <code>open()</code> for reading only returns without delay. An <code>open()</code> for writing only returns an error if no process currently has the file open for reading. If <code>O_NONBLOCK</code> and <code>O_NDELAY</code> are clear: An <code>open()</code> for reading only blocks until a process opens the file for writing. An <code>open()</code> for writing only blocks until a process opens the file for reading.

open(2)

After both ends of a FIFO have been opened, there is no guarantee that further calls to `open()` `O_RDONLY` (`O_WRONLY`) will synchronize with later calls to `open()` `O_WRONLY` (`O_RDONLY`) until both ends of the FIFO have been closed by all readers and writers. Any data written into a FIFO will be lost if both ends of the FIFO are closed before the data is read.

When opening a block special or character special file that supports non-blocking opens:

If `O_NONBLOCK` or `O_NDELAY` is set:

The `open()` function returns without blocking for the device to be ready or available. Subsequent behavior of the device is device-specific.

If `O_NONBLOCK` and `O_NDELAY` are clear:

The `open()` function blocks until the device is ready or available before returning.

Otherwise, the behavior of `O_NONBLOCK` and `O_NDELAY` is unspecified.

<code>O_RSYNC</code>	Read I/O operations on the file descriptor complete at the same level of integrity as specified by the <code>O_DSYNC</code> and <code>O_SYNC</code> flags. If both <code>O_DSYNC</code> and <code>O_RSYNC</code> are set in <i>oflag</i> , all I/O operations on the file descriptor complete as defined by synchronized I/O data integrity completion. If both <code>O_SYNC</code> and <code>O_RSYNC</code> are set in <i>oflag</i> , all I/O operations on the file descriptor complete as defined by synchronized I/O file integrity completion.
<code>O_SYNC</code>	If <code>O_SYNC</code> is set on a regular file, writes to that file cause the process to block until the data is delivered to the underlying hardware.
<code>O_TRUNC</code>	If the file exists and is a regular file, and the file is successfully opened <code>O_RDWR</code> or <code>O_WRONLY</code> , its length is truncated to 0 and the mode and owner are unchanged. It has no effect on FIFO special files or terminal device files. Its effect on other file types is implementation-dependent. The result of using <code>O_TRUNC</code> with <code>O_RDONLY</code> is undefined.

If `O_CREAT` is set and the file did not previously exist, upon successful completion, `open()` marks for update the `st_atime`, `st_ctime`, and `st_mtime` fields of the file and the `st_ctime` and `st_mtime` fields of the parent directory.

If `O_TRUNC` is set and the file did previously exist, upon successful completion, `open()` marks for update the `st_ctime` and `st_mtime` fields of the file.

open(2)

If *path* refers to a STREAMS file, *oflag* may be constructed from `O_NONBLOCK` or `O_NODELAY` OR-ed with either `O_RDONLY`, `O_WRONLY`, or `O_RDWR`. Other flag values are not applicable to STREAMS devices and have no effect on them. The values `O_NONBLOCK` and `O_NODELAY` affect the operation of STREAMS drivers and certain functions (see `read(2)`, `getmsg(2)`, `putmsg(2)`, and `write(2)`) applied to file descriptors associated with STREAMS files. For STREAMS drivers, the implementation of `O_NONBLOCK` and `O_NODELAY` is device-specific.

When `open()` is invoked to open a named stream, and the `connld` module (see `connld(7M)`) has been pushed on the pipe, `open()` blocks until the server process has issued an `I_RECVFD ioctl()` (see `streamio(7I)`) to receive the file descriptor.

If *path* names the master side of a pseudo-terminal device, then it is unspecified whether `open()` locks the slave side so that it cannot be opened. Portable applications must call `unlockpt(3C)` before opening the slave side.

If *path* is a symbolic link and `O_CREAT` and `O_EXCL` are set, the link is not followed.

Certain flag values can be set following `open()` as described in `fcntl(2)`.

The largest value that can be represented correctly in an object of type `off_t` is established as the offset maximum in the open file description.

RETURN VALUES

Upon successful completion, the `open()` function opens the file and return a non-negative integer representing the lowest numbered unused file descriptor. Otherwise, `-1` is returned, `errno` is set to indicate the error, and no files are created or modified.

ERRORS

The `open()` function will fail if:

EACCES	Search permission is denied on a component of the path prefix, or the file exists and the permissions specified by <i>oflag</i> are denied, or the file does not exist and write permission is denied for the parent directory of the file to be created, or <code>O_TRUNC</code> is specified and write permission is denied.
EDQUOT	The file does not exist, <code>O_CREAT</code> is specified, and either the directory where the new file entry is being placed cannot be extended because the user's quota of disk blocks on that file system has been exhausted, or the user's quota of inodes on the file system where the file is being created has been exhausted.
EEXIST	The <code>O_CREAT</code> and <code>O_EXCL</code> flags are set, and the named file exists.
EINTR	A signal was caught during <code>open()</code> .
EFAULT	The <i>path</i> argument points to an illegal address.
EIO	The <i>path</i> argument names a STREAMS file and a hangup or error occurred during the <code>open()</code> .

EISDIR	The named file is a directory and <i>oflag</i> includes <code>O_WRONLY</code> or <code>O_RDWR</code> .
ELOOP	Too many symbolic links were encountered in resolving <i>path</i> .
EMFILE	<code>OPEN_MAX</code> file descriptors are currently open in the calling process.
EMULTIHOP	Components of <i>path</i> require hopping to multiple remote machines and the file system does not allow it.
ENAMETOOLONG	The length of the <i>path</i> argument exceeds <code>PATH_MAX</code> or a pathname component is longer than <code>NAME_MAX</code> .
ENFILE	The maximum allowable number of files is currently open in the system.
ENOENT	The <code>O_CREAT</code> flag is not set and the named file does not exist; or the <code>O_CREAT</code> flag is set and either the path prefix does not exist or the <i>path</i> argument points to an empty string.
ENOLINK	The <i>path</i> argument points to a remote machine, and the link to that machine is no longer active.
ENOSR	The <i>path</i> argument names a STREAMS-based file and the system is unable to allocate a STREAM.
ENOSPC	The directory or file system that would contain the new file cannot be expanded, the file does not exist, and <code>O_CREAT</code> is specified.
ENOTDIR	A component of the path prefix is not a directory.
ENXIO	The <code>O_NONBLOCK</code> flag is set, the named file is a FIFO, the <code>O_WRONLY</code> flag is set, and no process has the file open for reading; or the named file is a character special or block special file and the device associated with this special file does not exist.
EOPNOTSUPP	An attempt was made to open a path that corresponds to a <code>AF_UNIX</code> socket.
EOVERFLOW	The named file is a regular file and either <code>O_LARGEFILE</code> is not set and the size of the file cannot be represented correctly in an object of type <code>off_t</code> or <code>O_LARGEFILE</code> is set and the size of the file cannot be represented correctly in an object of type <code>off64_t</code> .
EROFS	The named file resides on a read-only file system and either <code>O_WRONLY</code> , <code>O_RDWR</code> , <code>O_CREAT</code> (if file does not exist), or <code>O_TRUNC</code> is set in the <i>oflag</i> argument.
The <code>open()</code> function may fail if:	
EAGAIN	The <i>path</i> argument names the slave side of a pseudo-terminal device that is locked.
EINVAL	The value of the <i>oflag</i> argument is not valid.

open(2)

ENAMETOOLONG	Pathname resolution of a symbolic link produced an intermediate result whose length exceeds <code>PATH_MAX</code> .
ENOMEM	The <i>path</i> argument names a STREAMS file and the system is unable to allocate resources.
ETXTBSY	The file is a pure procedure (shared text) file that is being executed and <i>oflag</i> is <code>O_WRONLY</code> or <code>O_RDWR</code> .

USAGE The `open()` function has a transitional interface for 64-bit file offsets. See `lf64(5)`. Note that using `open64()` is equivalent to using `open()` with `O_LARGEFILE` set in *oflag*.

ATTRIBUTES See `attributes(5)` for descriptions of the following attributes:

ATTRIBUTE TYPE	ATTRIBUTE VALUE
MT-Level	Async-Signal-Safe

SEE ALSO `intro(3)`, `chmod(2)`, `close(2)`, `creat(2)`, `dup(2)`, `exec(2)`, `fcntl(2)`, `getmsg(2)`, `getrlimit(2)`, `lseek(2)`, `putmsg(2)`, `read(2)`, `stat(2)`, `umask(2)`, `write(2)`, `unlockpt(3C)`, `attributes(5)`, `fcntl(3HEAD)`, `lf64(5)`, `stat(3HEAD)`, `connlld(7M)`, `streamio(7I)`

NOTES Hierarchical Storage Management (HSM) file systems can sometimes cause long delays when opening a file, since HSM files must be recalled from secondary storage.

NAME	pause – suspend process until signal				
SYNOPSIS	<pre>#include <unistd.h> int pause(void);</pre>				
DESCRIPTION	The <code>pause()</code> function suspends the calling process until it receives a signal. The signal must be one that is not currently set to be ignored by the calling process. If the signal causes termination of the calling process, <code>pause()</code> does not return. If the signal is caught by the calling process and control is returned from the signal-catching function (see <code>signal(3C)</code>), the calling process resumes execution from the point of suspension.				
RETURN VALUES	Since <code>pause()</code> suspends thread execution indefinitely unless interrupted by a signal, there is no successful completion return value. If interrupted, it returns <code>-1</code> and sets <code>errno</code> to indicate the error.				
ERRORS	The <code>pause()</code> function will fail if: <table> <tr> <td><code>EINTR</code></td><td>A signal is caught by the calling process and control is returned from the signal-catching function.</td></tr> </table>	<code>EINTR</code>	A signal is caught by the calling process and control is returned from the signal-catching function.		
<code>EINTR</code>	A signal is caught by the calling process and control is returned from the signal-catching function.				
ATTRIBUTES	See <code>attributes(5)</code> for descriptions of the following attributes: <table> <tr> <th>ATTRIBUTE TYPE</th><th>ATTRIBUTE VALUE</th></tr> <tr> <td>MT-Level</td><td>Async-Signal-Safe</td></tr> </table>	ATTRIBUTE TYPE	ATTRIBUTE VALUE	MT-Level	Async-Signal-Safe
ATTRIBUTE TYPE	ATTRIBUTE VALUE				
MT-Level	Async-Signal-Safe				
SEE ALSO	<code>alarm(2)</code> , <code>kill(2)</code> , <code>wait(2)</code> , <code>signal(3C)</code> , <code>attributes(5)</code>				

pcsample(2)

NAME	pcsample – program execution time profile						
SYNOPSIS	<pre>#include <pcsample.h> long pcsample(uintptr_t samples[], long nsamples);</pre>						
DESCRIPTION	The <code>pcsample()</code> function provides CPU-use statistics by profiling the amount of CPU time expended by a program. For profiling dynamically-linked programs and 64-bit programs, it is superior to the <code>profil(2)</code> function, which assumes that the entire program is contained in a small, contiguous segment of the address space, divides this segment into “bins”, and on each clock tick increments the counter in the bin where the program is currently executing. With shared libraries creating discontinuous program segments spread throughout the address space, and with 64-bit address spaces so large that the size of “bins” would be measured in megabytes, the <code>profil()</code> function is of limited value. The <code>pcsample()</code> function is passed an array <i>samples</i> containing <i>nsamples</i> pointer-sized elements. During program execution, the kernel samples the program counter of the process, storing unadulterated values in the array on each clock tick. The kernel stops writing to the array when it is full, which occurs after <i>nsamples</i> / HZ seconds of process virtual time. The HZ value is obtained by invoking the call <code>sysconf(_SC_CLK_TCK)</code>. See <code>sysconf(3C)</code>. The sampling can be stopped by a subsequent call to <code>pcsample()</code> with the <i>nsamples</i> argument set to 0. Like <code>profil()</code>, sampling continues across a call to <code>fork(2)</code>, but is disabled by a call to one of the <code>exec</code> family of functions (see <code>exec(2)</code>). It is also disabled if an update of the <i>samples[]</i> array causes a memory fault.						
RETURN VALUES	The <code>pcsample()</code> function always returns 0 the first time it is called. On subsequent calls, it returns the number of samples that were stored during the previous invocation. If <i>nsamples</i> is invalid, it returns -1 and sets <code>errno</code> to indicate the error.						
ERRORS	The <code>pcsample()</code> function will fail if: EINVAL The value of <i>nsamples</i> is not valid.						
ATTRIBUTES	See <code>attributes(5)</code> for descriptions of the following attributes:						
	<table><tr><th>ATTRIBUTE TYPE</th><th>ATTRIBUTE VALUE</th></tr><tr><td>MT-Level</td><td>Async-Signal-Safe</td></tr><tr><td>Interface Stability</td><td>Stable</td></tr></table>	ATTRIBUTE TYPE	ATTRIBUTE VALUE	MT-Level	Async-Signal-Safe	Interface Stability	Stable
ATTRIBUTE TYPE	ATTRIBUTE VALUE						
MT-Level	Async-Signal-Safe						
Interface Stability	Stable						
SEE ALSO	<code>exec(2)</code> , <code>fork(2)</code> , <code>profil(2)</code> , <code>sysconf(3C)</code> , <code>attributes(5)</code>						

NAME	pipe – create an interprocess channel				
SYNOPSIS	<pre>#include <unistd.h> int pipe(int <i>fildes</i>[2]);</pre>				
DESCRIPTION	The <code>pipe()</code> function creates an I/O mechanism called a pipe and returns two file descriptors, <i>fildes</i>[0] and <i>fildes</i>[1]. The files associated with <i>fildes</i>[0] and <i>fildes</i>[1] are streams and are both opened for reading and writing. The <code>O_NDELAY</code> and <code>O_NONBLOCK</code> flags are cleared. A read from <i>fildes</i>[0] accesses the data written to <i>fildes</i>[1] on a first-in-first-out (FIFO) basis and a read from <i>fildes</i>[1] accesses the data written to <i>fildes</i>[0] also on a FIFO basis. The <code>FD_CLOEXEC</code> flag will be clear on both file descriptors. Upon successful completion <code>pipe()</code> marks for update the <code>st_atime</code>, <code>st_ctime</code>, and <code>st_mtime</code> fields of the pipe.				
RETURN VALUES	Upon successful completion, 0 is returned. Otherwise, -1 is returned and <code>errno</code> is set to indicate the error.				
ERRORS	The <code>pipe()</code> function will fail if: <table> <tr> <td><code>EMFILE</code></td><td>There are <code>OPEN_MAX-1</code> or more file descriptors currently open for this process.</td></tr> <tr> <td><code>ENFILE</code></td><td>A file table entry could not be allocated.</td></tr> </table>	<code>EMFILE</code>	There are <code>OPEN_MAX-1</code> or more file descriptors currently open for this process.	<code>ENFILE</code>	A file table entry could not be allocated.
<code>EMFILE</code>	There are <code>OPEN_MAX-1</code> or more file descriptors currently open for this process.				
<code>ENFILE</code>	A file table entry could not be allocated.				
ATTRIBUTES	See <code>attributes(5)</code> for descriptions of the following attributes: <table> <tr> <th>ATTRIBUTE TYPE</th><th>ATTRIBUTE VALUE</th></tr> <tr> <td>MT-Level</td><td>Async-Signal-Safe</td></tr> </table>	ATTRIBUTE TYPE	ATTRIBUTE VALUE	MT-Level	Async-Signal-Safe
ATTRIBUTE TYPE	ATTRIBUTE VALUE				
MT-Level	Async-Signal-Safe				
SEE ALSO	<code>sh(1)</code> , <code>fcntl(2)</code> , <code>fstat(2)</code> , <code>getmsg(2)</code> , <code>poll(2)</code> , <code>putmsg(2)</code> , <code>read(2)</code> , <code>write(2)</code> , <code>attributes(5)</code> , <code>streamio(7I)</code>				
NOTES	Since a pipe is bi-directional, there are two separate flows of data. Therefore, the size (<code>st_size</code>) returned by a call to <code>fstat(2)</code> with argument <i>fildes</i>[0] or <i>fildes</i>[1] is the number of bytes available for reading from <i>fildes</i>[0] or <i>fildes</i>[1] respectively. Previously, the size (<code>st_size</code>) returned by a call to <code>fstat()</code> with argument <i>fildes</i>[1] (the write-end) was the number of bytes available for reading from <i>fildes</i>[0] (the read-end).				

poll(2)

NAME	poll – input/output multiplexing																
SYNOPSIS	<pre>#include <poll.h> int poll(struct pollfd <i>fds</i>[], nfds_t <i>nfds</i>, int <i>timeout</i>);</pre>																
DESCRIPTION	The <code>poll()</code> function provides applications with a mechanism for multiplexing input/output over a set of file descriptors. For each member of the array pointed to by <i>fds</i>, <code>poll()</code> examines the given file descriptor for the event(s) specified in <i>events</i>. The number of <code>pollfd</code> structures in the <i>fds</i> array is specified by <i>nfds</i>. The <code>poll()</code> function identifies those file descriptors on which an application can read or write data, or on which certain events have occurred. The <i>fds</i> argument specifies the file descriptors to be examined and the events of interest for each file descriptor. It is a pointer to an array with one member for each open file descriptor of interest. The array's members are <code>pollfd</code> structures, which contain the following members: <pre>int fd; /* file descriptor */ short events; /* requested events */ short revents; /* returned events */</pre> The <i>fd</i> member specifies an open file descriptor and the <i>events</i> and <i>revents</i> members are bitmasks constructed by a logical OR operation of any combination of the following event flags: <table> <tr> <td>POLLIN</td><td>Data other than high priority data may be read without blocking. For STREAMS, this flag is set in <i>revents</i> even if the message is of zero length.</td></tr> <tr> <td>POLLRDNORM</td><td>Normal data (priority band equals 0) may be read without blocking. For STREAMS, this flag is set in <i>revents</i> even if the message is of zero length.</td></tr> <tr> <td>POLLRDBAND</td><td>Data from a non-zero priority band may be read without blocking. For STREAMS, this flag is set in <i>revents</i> even if the message is of zero length.</td></tr> <tr> <td>POLLPRI</td><td>High priority data may be received without blocking. For STREAMS, this flag is set in <i>revents</i> even if the message is of zero length.</td></tr> <tr> <td>POLLOUT</td><td>Normal data (priority band equals 0) may be written without blocking.</td></tr> <tr> <td>POLLWRNORM</td><td>The same as POLLOUT.</td></tr> <tr> <td>POLLWRBAND</td><td>Priority data (priority band > 0) may be written. This event only examines bands that have been written to at least once.</td></tr> <tr> <td>POLLERR</td><td>An error has occurred on the device or stream. This flag is only valid in the <i>revents</i> bitmask; it is not used in the <i>events</i> member.</td></tr> </table>	POLLIN	Data other than high priority data may be read without blocking. For STREAMS, this flag is set in <i>revents</i> even if the message is of zero length.	POLLRDNORM	Normal data (priority band equals 0) may be read without blocking. For STREAMS, this flag is set in <i>revents</i> even if the message is of zero length.	POLLRDBAND	Data from a non-zero priority band may be read without blocking. For STREAMS, this flag is set in <i>revents</i> even if the message is of zero length.	POLLPRI	High priority data may be received without blocking. For STREAMS, this flag is set in <i>revents</i> even if the message is of zero length.	POLLOUT	Normal data (priority band equals 0) may be written without blocking.	POLLWRNORM	The same as POLLOUT.	POLLWRBAND	Priority data (priority band > 0) may be written. This event only examines bands that have been written to at least once.	POLLERR	An error has occurred on the device or stream. This flag is only valid in the <i>revents</i> bitmask; it is not used in the <i>events</i> member.
POLLIN	Data other than high priority data may be read without blocking. For STREAMS, this flag is set in <i>revents</i> even if the message is of zero length.																
POLLRDNORM	Normal data (priority band equals 0) may be read without blocking. For STREAMS, this flag is set in <i>revents</i> even if the message is of zero length.																
POLLRDBAND	Data from a non-zero priority band may be read without blocking. For STREAMS, this flag is set in <i>revents</i> even if the message is of zero length.																
POLLPRI	High priority data may be received without blocking. For STREAMS, this flag is set in <i>revents</i> even if the message is of zero length.																
POLLOUT	Normal data (priority band equals 0) may be written without blocking.																
POLLWRNORM	The same as POLLOUT.																
POLLWRBAND	Priority data (priority band > 0) may be written. This event only examines bands that have been written to at least once.																
POLLERR	An error has occurred on the device or stream. This flag is only valid in the <i>revents</i> bitmask; it is not used in the <i>events</i> member.																

POLLHUP	A hangup has occurred on the stream. This event and POLLOUT are mutually exclusive; a stream can never be writable if a hangup has occurred. However, this event and POLLIN, POLLRDNORM, POLLRDBAND, or POLLPRI are not mutually exclusive. This flag is only valid in the <code>revents</code> bitmask; it is not used in the <code>events</code> member.
POLLNVAL	The specified <code>fd</code> value does not belong to an open file. This flag is only valid in the <code>revents</code> member; it is not used in the <code>events</code> member.

If the value `fd` is less than zero, `events` is ignored and `revents` is set to 0 in that entry on return from `poll()`.

The results of the `poll()` query are stored in the `revents` member in the `pollfd` structure. Bits are set in the `revents` bitmask to indicate which of the requested events are true. If none are true, none of the specified bits are set in `revents` when the `poll()` call returns. The event flags `POLLHUP`, `POLLERR`, and `POLLNVAL` are always set in `revents` if the conditions they indicate are true; this occurs even though these flags were not present in `events`.

If none of the defined events have occurred on any selected file descriptor, `poll()` waits at least *timeout* milliseconds for an event to occur on any of the selected file descriptors. On a computer where millisecond timing accuracy is not available, *timeout* is rounded up to the nearest legal value available on that system. If the value *timeout* is 0, `poll()` returns immediately. If the value of *timeout* is `INFTIM` (or `-1`), `poll()` blocks until a requested event occurs or until the call is interrupted. The `poll()` function is not affected by the `O_NDELAY` and `O_NONBLOCK` flags.

The `poll()` function supports regular files, terminal and pseudo-terminal devices, STREAMS-based files, FIFOs and pipes. The behavior of `poll()` on elements of *fds* that refer to other types of file is unspecified.

The `poll()` function supports sockets.

A file descriptor for a socket that is listening for connections will indicate that it is ready for reading, once connections are available. A file descriptor for a socket that is connecting asynchronously will indicate that it is ready for writing, once a connection has been established.

Regular files always poll `TRUE` for reading and writing.

RETURN VALUES

Upon successful completion, a non-negative value is returned. A positive value indicates the total number of file descriptors that has been selected (that is, file descriptors for which the `revents` member is non-zero). A value of 0 indicates that the call timed out and no file descriptors have been selected. Upon failure, `-1` is returned and `errno` is set to indicate the error.

ERRORS

The `poll()` function will fail if:

poll(2)

	EAGAIN	Allocation of internal data structures failed, but the request may be attempted again.
	EFAULT	Some argument points to an illegal address.
	EINTR	A signal was caught during the <code>poll()</code> function.
	EINVAL	The argument <i>nfds</i> is greater than <code>{OPEN_MAX}</code> , or one of the <code>fd</code> members refers to a STREAM or multiplexer that is linked (directly or indirectly) downstream from a multiplexer.
SEE ALSO	intro(3), getmsg(2), getrlimit(2), putmsg(2), read(2), write(2), select(3C), chpoll(9E) <i>STREAMS Programming Guide</i>	
NOTES	Non-STREAMS drivers use <code>chpoll(9E)</code> to implement <code>poll()</code> on these devices.	

NAME	p_online – return or change processor operational status				
SYNOPSIS	<pre>#include <sys/types.h> #include <sys/processor.h> int p_online(processorid_t processorid, int flag);</pre>				
DESCRIPTION	The <code>p_online()</code> function changes or returns the operational status of processors. The state of the processor specified by the <i>processorid</i> argument is changed to the state represented by the <i>flag</i> argument. Legal values for <i>flag</i> are <code>P_STATUS</code>, <code>P_ONLINE</code>, <code>P_OFFLINE</code>, and <code>P_NOINTR</code>. When <i>flag</i> is <code>P_STATUS</code>, no processor status change occurs, but the current processor status is returned. The <code>P_ONLINE</code>, <code>P_OFFLINE</code>, and <code>P_NOINTR</code> values for <i>flag</i> refer to valid processor states. A processor in the <code>P_ONLINE</code> state is allowed to process LWPs (lightweight processes) and perform system activities. The processor is also interruptible by I/O devices attached to the system. A processor in the <code>P_OFFLINE</code> state is not allowed to process LWPs. The processor is as inactive as possible. If the hardware supports such a feature, the processor is not interruptible by attached I/O devices. A processor in the <code>P_NOINTR</code> state is allowed to process LWPs, but it is not interruptible by attached I/O devices. Typically, interrupts, when they occur are routed to other processors in the system. Not all systems support putting a processor into the <code>P_NOINTR</code> state. It is not permitted to put all the processors of a system into the <code>P_NOINTR</code> state. At least one processor must always be available to service system clock interrupts. Processor numbers are integers, greater than or equal to 0, and are defined by the hardware platform. Processor numbers are not necessarily contiguous, but “not too sparse.” Processor numbers should always be printed in decimal. The number of processors present can be determined by calling <code>sysconf(_SC_NPROCESSORS_CONF)</code>. The list of valid processor numbers can be determined by calling <code>p_online()</code> with <i>processorid</i> values starting at 0 until all processors have been found. The <code>EINVAL</code> error is returned for invalid processor numbers. See EXAMPLES below.				
RETURN VALUES	On successful completion, the value returned is the previous state of the processor, <code>P_ONLINE</code> , <code>P_OFFLINE</code> , <code>P_NOINTR</code> , or <code>P_POWEROFF</code> . Otherwise, <code>-1</code> is returned and <code>errno</code> is set to indicate the error.				
ERRORS	The <code>p_online()</code> function will fail if: <table> <tr> <td><code>EPERM</code></td><td>The effective user of the calling process is not super-user.</td></tr> <tr> <td><code>EINVAL</code></td><td>A non-existent processor ID was specified or <i>flag</i> was invalid.</td></tr> </table>	<code>EPERM</code>	The effective user of the calling process is not super-user.	<code>EINVAL</code>	A non-existent processor ID was specified or <i>flag</i> was invalid.
<code>EPERM</code>	The effective user of the calling process is not super-user.				
<code>EINVAL</code>	A non-existent processor ID was specified or <i>flag</i> was invalid.				

p_online(2)

EBUSY	The <i>flag</i> was P_OFFLINE and the specified processor is the only on-line processor, there are currently LWPs bound to the processor, or the processor performs some essential function that cannot be performed by another processor.
EBUSY	The <i>flag</i> was P_NOINTR and the specified processor is the only interruptible processor in the system, or it handles interrupts that cannot be handled by another processor.
EBUSY	The specified processor is powered off and cannot be powered on because some platform- specific resource is not available.
ENOTSUP	The specified processor is powered off, and the platform does not support power on of individual processors.

EXAMPLES

EXAMPLE 1 List the legal processor numbers.

The following code sample will list the legal processor numbers:

```
#include <sys/unistd.h>
#include <sys/processor.h>
#include <sys/types.h>
#include <stdio.h>
#include <errno.h>

int
main()
{
    processorid_t i;
    int    status;
    int    n = sysconf(_SC_NPROCESSORS_ONLN);
    for (i = 0; n > 0; i++) {
        status = p_online(i, P_STATUS);
        if (status == -1 && errno == EINVAL)
            continue;
        printf("processor %d present\n", i);
        n--;
    }
    return (0);
}
```

ATTRIBUTES

See attributes(5) for descriptions of the following attributes:

ATTRIBUTE TYPE	ATTRIBUTE VALUE
MT-Level	MT-Safe

SEE ALSO

psradm(1M), psrinfo(1M), processor_bind(2), processor_info(2), pset_create(2), sysconf(3C), attributes(5)

NAME	prioset – process scheduler control										
SYNOPSIS	<pre>#include <sys/types.h> #include <sys/prioset.h> #include <sys/rtprioset.h> #include <sys/tprioset.h> long prioset(idtype_t idtype, id_t id, int cmd, /* arg */ ...);</pre>										
DESCRIPTION	The prioset() function provides for control over the scheduling of an active light weight process (LWP). LTWs fall into distinct classes with a separate scheduling policy applied to each class. The two classes currently supported are the realtime class and the time-sharing class. The characteristics of these classes are described under the corresponding headings below. The class attribute of an LWP is inherited across the fork(2) and _lwp_create(2) functions and the exec family of functions (see exec(2)). The prioset() function can be used to dynamically change the class and other scheduling parameters associated with a running LWP or set of LTWs given the appropriate permissions as explained below. In the default configuration, a runnable realtime LWP runs before any other LWP. Therefore, inappropriate use of realtime LWP can have a dramatic negative impact on system performance. The prioset() function provides an interface for specifying a process, set of processes or an LWP to which the function is to apply. The priosetset(2) function provides the same functions as prioset(), but allows a more general interface for specifying the set of LTWs to which the function is to apply. For prioset(), the idtype and id arguments are used together to specify the set of LTWs. The interpretation of id depends on the value of idtype. The possible values for idtype and corresponding interpretations of id are as follows: <table> <tr> <td>P_LWPID</td><td>The id argument is an LWP ID. The prioset function applies to the LWP with the specified ID within the calling process.</td></tr> <tr> <td>P_PID</td><td>The id argument is a process ID specifying a single process. The prioset() function applies to all LTWs currently associated with the specified process.</td></tr> <tr> <td>P_PPID</td><td>The id argument is a parent process ID. The prioset() function applies to all LTWs currently associated with processes with the specified parent process ID.</td></tr> <tr> <td>P_PGID</td><td>The id argument is a process group ID. The prioset() function applies to all LTWs currently associated with processes in the specified process group.</td></tr> <tr> <td>P_SID</td><td>The id argument is a session ID. The prioset() function applies to all LTWs currently associated with processes in the specified session.</td></tr> </table>	P_LWPID	The id argument is an LWP ID. The prioset function applies to the LWP with the specified ID within the calling process.	P_PID	The id argument is a process ID specifying a single process. The prioset() function applies to all LTWs currently associated with the specified process.	P_PPID	The id argument is a parent process ID. The prioset() function applies to all LTWs currently associated with processes with the specified parent process ID.	P_PGID	The id argument is a process group ID. The prioset() function applies to all LTWs currently associated with processes in the specified process group.	P_SID	The id argument is a session ID. The prioset() function applies to all LTWs currently associated with processes in the specified session.
P_LWPID	The id argument is an LWP ID. The prioset function applies to the LWP with the specified ID within the calling process.										
P_PID	The id argument is a process ID specifying a single process. The prioset() function applies to all LTWs currently associated with the specified process.										
P_PPID	The id argument is a parent process ID. The prioset() function applies to all LTWs currently associated with processes with the specified parent process ID.										
P_PGID	The id argument is a process group ID. The prioset() function applies to all LTWs currently associated with processes in the specified process group.										
P_SID	The id argument is a session ID. The prioset() function applies to all LTWs currently associated with processes in the specified session.										

priocntl(2)

P_TASKID	The <i>id</i> argument is a task ID. The <code>priocntl()</code> function applies to all LWPs currently associated with processes in the specified task.
P_CID	The <i>id</i> argument is a class ID (returned by the <code>priocntl()</code> <code>PC_GETCID</code> command as explained below). The <code>priocntl()</code> function applies to all LWPs in the specified class.
P_UID	The <i>id</i> argument is a user ID. The <code>priocntl()</code> function applies to all LWPs with this effective user ID.
P_GID	The <i>id</i> argument is a group ID. The <code>priocntl()</code> function applies to all LWPs with this effective group ID.
P_PROJID	The <i>id</i> argument is a project ID. The <code>priocntl()</code> function applies to all LWPs with this project ID.
P_ALL	The <code>priocntl()</code> function applies to all existing LWPs. The value of <i>id</i> is ignored. The permission restrictions described below still apply.

An *id* value of `P_MYID` can be used in conjunction with the *idtype* value to specify the calling LWP's LWP ID, parent process ID, process group ID, session ID, task ID, class ID, user ID, group ID, or project ID.

In order to change the scheduling parameters of an LWP (using the `PC_SETPARMS` command as explained below) the real or effective user ID of the LWP calling `priocntl()` must match the real or effective user ID of the receiving LWP or the effective user ID of the calling LWP must be super-user. These are the minimum permission requirements enforced for all classes. An individual class may impose additional permissions requirements when setting LWPs to that class and/or when setting class-specific scheduling parameters.

A special `sys` scheduling class exists for the purpose of scheduling the execution of certain special system processes (such as the swapper process). It is not possible to change the class of any LWP to `sys`. In addition, any processes in the `sys` class that are included in a specified set of processes are disregarded by `priocntl()`. For example, an *idtype* of `P_UID` and an *id* value of 0 would specify all processes with a user ID of 0 except processes in the `sys` class and (if changing the parameters using `PC_SETPARMS`) the `init(1M)` process.

The *init* process is a special case. In order for a `priocntl()` call to change the class or other scheduling parameters of the *init* process (process ID 1), it must be the only process specified by *idtype* and *id*. The *init* process may be assigned to any class configured on the system, but the time-sharing class is almost always the appropriate choice. (Other choices may be highly undesirable; see the *System Administration Guide, Volume 1* for more information.)

The data type and value of *arg* are specific to the type of command specified by *cmd*.

A structure with the following members is used by the `PC_GETCID` and `PC_GETCLINFO` commands.


```
id_t   pc_cid;                /* Class id */
char   pc_clname[PC_CLNMSZ];  /* Class name */
int     pc_clinfo[PC_CLINFOSZ]; /* Class information */
```

The `pc_cid` member is a class ID returned by the `prctl()` `PC_GETCID` command. The `pc_clname` member is a buffer of size `PC_CLNMSZ` (defined in `<sys/prctl.h>`) used to hold the class name (RT for realtime or TS for time-sharing).

The `pc_clinfo` member is a buffer of size `PC_CLINFOSZ` (defined in `<sys/prctl.h>`) used to return data describing the attributes of a specific class. The format of this data is class-specific and is described under the appropriate heading (REALTIME CLASS or TIME-SHARING CLASS) below.

A structure with the following elements is used by the `PC_SETPARMS` and `PC_GETPARMS` commands.

```
id_t   pc_cid;                /* LWP class */
int     pc_clparms[PC_CLPARMSZ]; /* Class-specific params */
```

The `pc_cid` member is a class ID (returned by the `prctl()` `PC_GETCID` command). The special class ID `PC_CLNULL` can also be assigned to `pc_cid` when using the `PC_SETPARMS` command as explained below.

The `pc_clparms` buffer holds class-specific scheduling parameters. The format of this parameter data for a particular class is described under the appropriate heading below. `PC_CLPARMSZ` is the length of the `pc_clparms` buffer and is defined in `<sys/prctl.h>`.

COMMANDS

Available `prctl()` commands are:

PC_GETCID

Get class ID and class attributes for a specific class given class name. The *idtype* and *id* arguments are ignored. If *arg* is non-null, it points to a structure of type `pcinfo_t`. The `pc_clname` buffer contains the name of the class whose attributes you are getting.

On success, the class ID is returned in `pc_cid`, the class attributes are returned in the `pc_clinfo` buffer, and the `prctl()` call returns the total number of classes configured in the system (including the `sys` class). If the class specified by `pc_clname` is invalid or is not currently configured the `prctl()` call returns `-1` with `errno` set to `EINVAL`. The format of the attribute data returned for a given class is defined in the `<sys/rtpriocntl.h>` or `<sys/tspriocntl.h>` header and described under the appropriate heading below.

If *arg* is a null pointer, no attribute data is returned but the `prctl()` call still returns the number of configured classes.

priocntl(2)

PC_GETCLINFO

Get class name and class attributes for a specific class given class ID. The *idtype* and *id* arguments are ignored. If *arg* is non-null, it points to a structure of type `pcinfo_t`. The `pc_cid` member is the class ID of the class whose attributes you are getting.

On success, the class name is returned in the `pc_clname` buffer, the class attributes are returned in the `pc_clinfo` buffer, and the `priocntl()` call returns the total number of classes configured in the system (including the `sys` class). The format of the attribute data returned for a given class is defined in the `<sys/rtpriocntl.h>` or `<sys/tspriocntl.h>` header file and described under the appropriate heading below.

If *arg* is a null pointer, no attribute data is returned but the `priocntl()` call still returns the number of configured classes.

PC_SETPARMS

Set the class and class-specific scheduling parameters of the specified LWP(s) associated with the specified process(es). When this command is used with the *idtype* of `P_LWPID`, it will set the class and class-specific scheduling parameters of the LWP. The *arg* argument points to a structure of type `pcparms_t`. The `pc_cid` member specifies the class you are setting and the `pc_clparms` buffer contains the class-specific parameters you are setting. The format of the class-specific parameter data is defined in the `<sys/rtpriocntl.h>` or `<sys/tspriocntl.h>` header and described under the appropriate class heading below.

When setting parameters for a set of LWPs, `priocntl()` acts on the LWPs in the set in an implementation-specific order. If `priocntl()` encounters an error for one or more of the target processes, it may or may not continue through the set of LWPs, depending on the nature of the error. If the error is related to permissions (`EPERM`), `priocntl()` continues through the LWP set, resetting the parameters for all target LWPs for which the calling LWP has appropriate permissions. The `priocntl()` function then returns `-1` with `errno` set to `EPERM` to indicate that the operation failed for one or more of the target LWPs. If `priocntl()` encounters an error other than permissions, it does not continue through the set of target LWPs but returns the error immediately.

PC_GETPARMS

Get the class and/or class-specific scheduling parameters of an LWP. The *arg* member points to a structure of type `pcparms_t`.

If `pc_cid` specifies a configured class and a single LWP belonging to that class is specified by the *idtype* and *id* values or the `procset` structure, then the scheduling parameters of that LWP are returned in the `pc_clparms` buffer. If the LWP specified does not exist or does not belong to the specified class, the `priocntl()` call returns `-1` with `errno` set to `ESRCH`.

If `pc_cid` specifies a configured class and a set of LWPs is specified, the scheduling parameters of one of the specified LWP belonging to the specified class are returned in the `pc_clparms` buffer and the `priocntl()` call returns the process ID of the

REALTIME CLASS

selected LWP. The criteria for selecting an LWP to return in this case is class dependent. If none of the specified LWPs exist or none of them belong to the specified class the `priocntl()` call returns `-1` with `errno` set to `ESRCH`.

If `pc_cid` is `PC_CLNULL` and a single LWP is specified the class of the specified LWP is returned in `pc_cid` and its scheduling parameters are returned in the `pc_clparms` buffer.

PC_ADMIN

This command provides functionality needed for the implementation of the `dispadm(1M)` utility. It is not intended for general use by other applications.

The realtime class provides a fixed priority preemptive scheduling policy for those LWPs requiring fast and deterministic response and absolute user/application control of scheduling priorities. If the realtime class is configured in the system it should have exclusive control of the highest range of scheduling priorities on the system. This ensures that a runnable realtime LWP is given CPU service before any LWP belonging to any other class.

The realtime class has a range of realtime priority (`rt_pri`) values that may be assigned to an LWP within the class. Real-time priorities range from 0 to x , where the value of x is configurable and can be determined for a specific installation by using the `priocntl() PC_GETCID` or `PC_GETCLINFO` command.

The realtime scheduling policy is a fixed priority policy. The scheduling priority of a realtime LWP is never changed except as the result of an explicit request by the user/application to change the `rt_pri` value of the LWP.

For an LWP in the realtime class, the `rt_pri` value is, for all practical purposes, equivalent to the scheduling priority of the LWP. The `rt_pri` value completely determines the scheduling priority of a realtime LWP relative to other LWPs within its class. Numerically higher `rt_pri` values represent higher priorities. Since the realtime class controls the highest range of scheduling priorities in the system it is guaranteed that the runnable realtime LWP with the highest `rt_pri` value is always selected to run before any other LWPs in the system.

In addition to providing control over priority, `priocntl()` provides for control over the length of the time quantum allotted to the LWP in the realtime class. The time quantum value specifies the maximum amount of time an LWP may run assuming that it does not complete or enter a resource or event wait state (*sleep*). Note that if another LWP becomes runnable at a higher priority, the currently running LWP may be preempted before receiving its full time quantum.

The system's process scheduler keeps the runnable realtime LWPs on a set of scheduling queues. There is a separate queue for each configured realtime priority and all realtime LWPs with a given `rt_pri` value are kept together on the appropriate queue. The LWPs on a given queue are ordered in FIFO order (that is, the LWP at the front of the queue has been waiting longest for service and receives the CPU first). Real-time LWPs that wake up after sleeping, LWPs which change to the realtime class from some other class, LWPs which have used their full time quantum, and runnable

priocntl(2)

LWPs whose priority is reset by `priocntl()` are all placed at the back of the appropriate queue for their priority. An LWP that is preempted by a higher priority LWP remains at the front of the queue (with whatever time is remaining in its time quantum) and runs before any other LWP at this priority. Following a `fork(2)` or `_lwp_create(2)` function call by a realtime LWP, the parent LWP continues to run while the child LWP (which inherits its parent's `rt_pri` value) is placed at the back of the queue.

A structure with the following members (defined in `<sys/rtpriocntl.h>`) defines the format used for the attribute data for the realtime class.

```
short    rt_maxpri;        /* Maximum realtime priority */
```

The `priocntl()` `PC_GETCID` and `PC_GETCLINFO` commands return realtime class attributes in the `pc_clinfo` buffer in this format.

The `rt_maxpri` member specifies the configured maximum `rt_pri` value for the realtime class (if `rt_maxpri` is *x*, the valid realtime priorities range from 0 to *x*).

A structure with the following members (defined in `<sys/rtpriocntl.h>`) defines the format used to specify the realtime class-specific scheduling parameters of an LWP.

```
short    rt_pri;           /* Real-Time priority */
uint_t   rt_tqsecs;        /* Seconds in time quantum */
int       rt_tqnsecs;       /* Additional nanoseconds in quantum */
```

When using the `priocntl()` `PC_SETPARMS` or `PC_GETPARMS` commands, if `pc_cid` specifies the realtime class, the data in the `pc_clparms` buffer is in this format.

The above commands can be used to set the realtime priority to the specified value or get the current `rt_pri` value. Setting the `rt_pri` value of an LWP that is currently running or runnable (not sleeping) causes the LWP to be placed at the back of the scheduling queue for the specified priority. The LWP is placed at the back of the appropriate queue regardless of whether the priority being set is different from the previous `rt_pri` value of the LWP. Note that a running LWP can voluntarily release the CPU and go to the back of the scheduling queue at the same priority by resetting its `rt_pri` value to its current realtime priority value. In order to change the time quantum of an LWP without setting the priority or affecting the LWP's position on the queue, the `rt_pri` member should be set to the special value `RT_NOCHANGE` (defined in `<sys/rtpriocntl.h>`). Specifying `RT_NOCHANGE` when changing the class of an LWP to realtime from some other class results in the realtime priority being set to 0.

For the `priocntl()` `PC_GETPARMS` command, if `pc_cid` specifies the realtime class and more than one realtime LWP is specified, the scheduling parameters of the realtime LWP with the highest `rt_pri` value among the specified LWPs are returned and the LWP ID of this LWP is returned by the `priocntl()` call. If there is more than one LWP sharing the highest priority, the one returned is implementation-dependent.

The `rt_tqsecs` and `rt_tqnsecs` members are used for getting or setting the time quantum associated with an LWP or group of LWPs. `rt_tqsecs` is the number of seconds in the time quantum and `rt_tqnsecs` is the number of additional

priocntl(2)

nanoseconds in the quantum. For example setting `rt_tqsecs` to 2 and `rt_tqnsecs` to 500,000,000 (decimal) would result in a time quantum of two and one-half seconds. Specifying a value of 1,000,000,000 or greater in the `rt_tqnsecs` member results in an error return with `errno` set to `EINVAL`. Although the resolution of the `tq_nsecs` member is very fine, the specified time quantum length is rounded up by the system to the next integral multiple of the system clock's resolution. The maximum time quantum that can be specified is implementation-specific and equal to `LONG_MAX` ticks (defined in `<limits.h>`). Requesting a quantum greater than this maximum results in an error return with `errno` set to `ERANGE` (although infinite quanta may be requested using a special value as explained below). Requesting a time quantum of 0 (setting both `rt_tqsecs` and `rt_tqnsecs` to 0) results in an error return with `errno` set to `EINVAL`.

The `rt_tqnsecs` member can also be set to one of the following special values (defined in `<sys/rtpriocntl.h>`), in which case the value of `rt_tqsecs` is ignored:

<code>RT_TQINF</code>	Set an infinite time quantum.
<code>RT_TQDEF</code>	Set the time quantum to the default for this priority (see <code>rt_dptbl(4)</code>).
<code>RT_NOCHANGE</code>	Do not set the time quantum. This value is useful when you wish to change the realtime priority of an LWP without affecting the time quantum. Specifying this value when changing the class of an LWP to realtime from some other class is equivalent to specifying <code>RT_TQDEF</code> .

In order to change the class of an LWP to realtime (from any other class) the LWP invoking `priocntl()` must have super-user privileges. In order to change the priority or time quantum setting of a realtime LWP, the LWP invoking `priocntl()` must have super-user privileges or must itself be a realtime LWP whose real or effective user ID matches the real or effective user ID of the target LWP.

The realtime priority and time quantum are inherited across `fork(2)` and the `exec` family of functions (see `exec(2)`).

TIME-SHARING CLASS

The time-sharing scheduling policy provides for a fair and effective allocation of the CPU resource among LWPs with varying CPU consumption characteristics. The objectives of the time-sharing policy are to provide good response time to interactive LWPs and good throughput to CPU-bound jobs while providing a degree of user/application control over scheduling.

The time-sharing class has a range of time-sharing user priority (see `ts_upri` below) values that may be assigned to LWPs within the class. A `ts_upri` value of 0 is defined as the default base priority for the time-sharing class. User priorities range from $-x$ to $+x$ where the value of x is configurable and can be determined for a specific installation by using the `priocntl()` `PC_GETCID` or `PC_GETCLINFO` command.

priocntl(2)

The purpose of the user priority is to provide some degree of user/application control over the scheduling of LWPs in the time-sharing class. Raising or lowering the `ts_upri` value of an LWP in the time-sharing class raises or lowers the scheduling priority of the LWP. It is not guaranteed, however, that an LWP with a higher `ts_upri` value will run before one with a lower `ts_upri` value. This is because the `ts_upri` value is just one factor used to determine the scheduling priority of a time-sharing LWP. The system may dynamically adjust the internal scheduling priority of a time-sharing LWP based on other factors such as recent CPU usage.

In addition to the system-wide limits on user priority (returned by the `PC_GETCID` and `PC_GETCLINFO` commands) there is a per LWP user priority limit (see `ts_uprim` below), which specifies the maximum `ts_upri` value that may be set for a given LWP; by default, `ts_uprim` is 0.

A structure with the following members (defined in `<sys/tspriocntl.h>`) defines the format used for the attribute data for the time-sharing class.

```
short    ts_maxupri;    /* Limits of user priority range */
```

The `priocntl()` `PC_GETCID` and `PC_GETCLINFO` commands return time-sharing class attributes in the `pc_clinfo` buffer in this format.

`ts_maxupri` specifies the configured maximum user priority value for the time-sharing class. If `ts_maxupri` is x , the valid range for both user priorities and user priority limits is from $-x$ to $+x$.

A structure with the following members (defined in `<sys/tspriocntl.h>`) defines the format used to specify the time-sharing class-specific scheduling parameters of an LWP.

```
short    ts_uprim;      /* Time-Sharing user priority limit */
short    ts_upri;       /* Time-Sharing user priority */
```

When using the `priocntl()` `PC_SETPARMS` or `PC_GETPARMS` commands, if `pc_cid` specifies the time-sharing class, the data in the `pc_clparms` buffer is in this format.

For the `priocntl()` `PC_GETPARMS` command, if `pc_cid` specifies the time-sharing class and more than one time-sharing LWP is specified, the scheduling parameters of the time-sharing LWP with the highest `ts_upri` value among the specified LWPs is returned and the LWP ID of this LWP is returned by the `priocntl()` call. If there is more than one LWP sharing the highest user priority, the one returned is implementation-dependent.

Any time-sharing LWP may lower its own `ts_uprim` (or that of another LWP with the same user ID). Only a time-sharing LWP with super-user privileges may raise a `ts_uprim`. When changing the class of an LWP to time-sharing from some other class, super-user privileges are required in order to set the initial `ts_uprim` to a value greater than 0. Attempts by a non-super-user LWP to raise a `ts_uprim` or set an initial `ts_uprim` greater than 0 fail with a return value of `-1` and `errno` set to `EPERM`.

	priocntl(2)														
	Any time-sharing LWP may set its own <code>ts_upri</code> (or that of another LWP with the same user ID) to any value less than or equal to the LWP's <code>ts_uprilm</code>. Attempts to set the <code>ts_upri</code> above the <code>ts_uprilm</code> (and/or set the <code>ts_uprilm</code> below the <code>ts_upri</code>) result in the <code>ts_upri</code> being set equal to the <code>ts_uprilm</code>. Either of the <code>ts_uprilm</code> or <code>ts_upri</code> members may be set to the special value <code>TS_NOCHANGE</code> (defined in <code><sys/tspriocntl.h></code>) in order to set one of the values without affecting the other. Specifying <code>TS_NOCHANGE</code> for the <code>ts_upri</code> when the <code>ts_uprilm</code> is being set to a value below the current <code>ts_upri</code> causes the <code>ts_upri</code> to be set equal to the <code>ts_uprilm</code> being set. Specifying <code>TS_NOCHANGE</code> for a parameter when changing the class of an LWP to time-sharing (from some other class) causes the parameter to be set to a default value. The default value for the <code>ts_uprilm</code> is 0 and the default for the <code>ts_upri</code> is to set it equal to the <code>ts_uprilm</code> which is being set. The time-sharing user priority and user priority limit are inherited across <code>fork()</code> and the <i>exec</i> family of functions.														
RETURN VALUES	Unless otherwise noted above, <code>priocntl()</code> returns a value of 0 on success. On failure, <code>priocntl()</code> returns -1 and sets <code>errno</code> to indicate the error.														
ERRORS	The <code>priocntl()</code> function fails if: <table> <tr> <td>EAGAIN</td><td>An attempt to change the class of an LWP failed because of insufficient resources other than memory (for example, class-specific kernel data structures).</td></tr> <tr> <td>EFAULT</td><td>One of the arguments points to an illegal address.</td></tr> <tr> <td>EINVAL</td><td>The argument <i>cmd</i> was invalid, an invalid or unconfigured class was specified, or one of the parameters specified was invalid.</td></tr> <tr> <td>ENOMEM</td><td>An attempt to change the class of an LWP failed because of insufficient memory.</td></tr> <tr> <td>EPERM</td><td>The effective user of the calling LWP is not super-user.</td></tr> <tr> <td>ERANGE</td><td>The requested time quantum is out of range.</td></tr> <tr> <td>ESRCH</td><td>None of the specified LWPs exist.</td></tr> </table>	EAGAIN	An attempt to change the class of an LWP failed because of insufficient resources other than memory (for example, class-specific kernel data structures).	EFAULT	One of the arguments points to an illegal address.	EINVAL	The argument <i>cmd</i> was invalid, an invalid or unconfigured class was specified, or one of the parameters specified was invalid.	ENOMEM	An attempt to change the class of an LWP failed because of insufficient memory.	EPERM	The effective user of the calling LWP is not super-user.	ERANGE	The requested time quantum is out of range.	ESRCH	None of the specified LWPs exist.
EAGAIN	An attempt to change the class of an LWP failed because of insufficient resources other than memory (for example, class-specific kernel data structures).														
EFAULT	One of the arguments points to an illegal address.														
EINVAL	The argument <i>cmd</i> was invalid, an invalid or unconfigured class was specified, or one of the parameters specified was invalid.														
ENOMEM	An attempt to change the class of an LWP failed because of insufficient memory.														
EPERM	The effective user of the calling LWP is not super-user.														
ERANGE	The requested time quantum is out of range.														
ESRCH	None of the specified LWPs exist.														
SEE ALSO	<code>priocntl(1)</code>, <code>dispadmin(1M)</code>, <code>init(1M)</code>, <code>_lwp_create(2)</code>, <code>exec(2)</code>, <code>fork(2)</code>, <code>nice(2)</code>, <code>priocntlset(2)</code>, <code>rt_dptbl(4)</code> <i>System Administration Guide, Volume 1</i> <i>System Interface Guide</i>														

priocntlset(2)

NAME	priocntlset – generalized process scheduler control								
SYNOPSIS	<pre>#include <sys/types.h> #include <sys/procset.h> #include <sys/priocntl.h> #include <sys/rtpriocntl.h> #include <sys/tspriocntl.h> long priocntlset(procset_t *psp, int cmd, /* arg */ ...);</pre>								
DESCRIPTION	The <code>priocntlset()</code> function changes the scheduling properties of running processes. <code>priocntlset()</code> has the same functions as the <code>priocntl()</code> function, but a more general way of specifying the set of processes whose scheduling properties are to be changed. <i>cmd</i> specifies the function to be performed. <i>arg</i> is a pointer to a structure whose type depends on <i>cmd</i>. See <code>priocntl(2)</code> for the valid values of <i>cmd</i> and the corresponding <i>arg</i> structures. <i>psp</i> is a pointer to a <code>procset</code> structure, which <code>priocntlset()</code> uses to specify the set of processes whose scheduling properties are to be changed. The <code>procset</code> structure contains the following members: <pre>idop_t p_op; /* operator connecting left/right sets */ idtype_t p_lidtype; /* left set ID type */ id_t p_lid; /* left set ID */ idtype_t p_ridtype; /* right set ID type */ id_t p_rid; /* right set ID */</pre> The <code>p_lidtype</code> and <code>p_lid</code> members specify the ID type and ID of one (“left”) set of processes; the <code>p_ridtype</code> and <code>p_rid</code> members specify the ID type and ID of a second (“right”) set of processes. ID types and IDs are specified just as for the <code>priocntl()</code> function. The <code>p_op</code> member specifies the operation to be performed on the two sets of processes to get the set of processes the function is to apply to. The valid values for <code>p_op</code> and the processes they specify are: <table><tr><td>POP_DIFF</td><td>Set difference: processes in left set and not in right set.</td></tr><tr><td>POP_AND</td><td>Set intersection: processes in both left and right sets.</td></tr><tr><td>POP_OR</td><td>Set union: processes in either left or right sets or both.</td></tr><tr><td>POP_XOR</td><td>Set exclusive-or: processes in left or right set but not in both.</td></tr></table> The following macro, which is defined in <code><procset.h></code>, offers a convenient way to initialize a <code>procset</code> structure: <pre>#define setprocset(psp, op, ltype, lid, rtype, rid) \ (psp)⇒p_op = (op), \ (psp)⇒p_lidtype = (ltype), \ (psp)⇒p_lid = (lid), \ (psp)⇒p_ridtype = (rtype), \ (psp)⇒p_rid = (rid),</pre>	POP_DIFF	Set difference: processes in left set and not in right set.	POP_AND	Set intersection: processes in both left and right sets.	POP_OR	Set union: processes in either left or right sets or both.	POP_XOR	Set exclusive-or: processes in left or right set but not in both.
POP_DIFF	Set difference: processes in left set and not in right set.								
POP_AND	Set intersection: processes in both left and right sets.								
POP_OR	Set union: processes in either left or right sets or both.								
POP_XOR	Set exclusive-or: processes in left or right set but not in both.								

RETURN VALUES Unless otherwise noted above, `priocntlset()` returns 0 on success. Otherwise, it returns -1 and sets `errno` to indicate the error.

ERRORS The `priocntlset()` function will fail if:

- | | |
|---------------------|--|
| <code>EAGAIN</code> | An attempt to change the class of a process failed because of insufficient resources other than memory (for example, class-specific kernel data structures). |
| <code>EFAULT</code> | One of the arguments points to an illegal address. |
| <code>EINVAL</code> | The argument <i>cmd</i> was invalid, an invalid or unconfigured class was specified, or one of the parameters specified was invalid. |
| <code>ENOMEM</code> | An attempt to change the class of a process failed because of insufficient memory. |
| <code>EPERM</code> | The effective user of the calling process is not super-user. |
| <code>ERANGE</code> | The requested time quantum is out of range. |
| <code>ESRCH</code> | None of the specified processes exist. |

SEE ALSO `priocntl(1)`, `priocntl(2)`

processor_bind(2)

NAME	processor_bind – bind LWP's to a processor								
SYNOPSIS	<pre>#include <sys/types.h> #include <sys/processor.h> #include <sys/procset.h> int processor_bind(idtype_t <i>idtype</i>, id_t <i>id</i>, processorid_t <i>processorid</i>, processorid_t *<i>obind</i>);</pre>								
DESCRIPTION	The <code>processor_bind()</code> function binds the LWP (lightweight process) or set of LWP's specified by <i>idtype</i> and <i>id</i> to the processor specified by <i>processorid</i>. If <i>obind</i> is not NULL, this function also sets the <code>processorid_t</code> variable pointed to by <i>obind</i> to the previous binding of one of the specified LWP's, or to <code>PBIND_NONE</code> if the selected LWP was not bound. If <i>idtype</i> is <code>P_PID</code>, the binding affects all LWP's of the process with process ID (PID) <i>id</i>. If <i>idtype</i> is <code>P_LWPID</code>, the binding affects the LWP of the current process with LWP ID <i>id</i>. If <i>idtype</i> is <code>P_TASKID</code>, the binding affects all LWP's of all processes with task ID <i>id</i>. If <i>id</i> is <code>P_MYID</code>, the specified LWP, process, or task is the current one. If <i>processorid</i> is <code>PBIND_NONE</code>, the processor bindings of the specified LWP's are cleared. If <i>processorid</i> is <code>PBIND_QUERY</code>, the processor bindings are not changed. The effective user of the calling process must be superuser, or its real or effective user ID must match the real or effective user ID of the LWP's being bound. If the calling process does not have permission to change all of the specified LWP's, the bindings of the LWP's for which it does have permission will be changed even though an error is returned.								
RETURN VALUES	Upon successful completion, 0 is returned. Otherwise, -1 is returned and <code>errno</code> is set to indicate the error.								
ERRORS	The <code>processor_bind()</code> function will fail if: <table><tr><td>EFAULT</td><td>The location pointed to by <i>obind</i> was not NULL and not writable by the user.</td></tr><tr><td>EINVAL</td><td>The specified processor is not on-line, or the <i>idtype</i> argument was not <code>P_PID</code>, <code>P_LWPID</code>, or <code>P_TASKID</code>.</td></tr><tr><td>EPERM</td><td>The effective user of the calling process is not superuser, and its real or effective user ID does not match the real or effective user ID of one of the LWP's being bound.</td></tr><tr><td>ESRCH</td><td>No processes, LWP's, or tasks were found to match the criteria specified by <i>idtype</i> and <i>id</i>.</td></tr></table>	EFAULT	The location pointed to by <i>obind</i> was not NULL and not writable by the user.	EINVAL	The specified processor is not on-line, or the <i>idtype</i> argument was not <code>P_PID</code> , <code>P_LWPID</code> , or <code>P_TASKID</code> .	EPERM	The effective user of the calling process is not superuser, and its real or effective user ID does not match the real or effective user ID of one of the LWP's being bound.	ESRCH	No processes, LWP's, or tasks were found to match the criteria specified by <i>idtype</i> and <i>id</i> .
EFAULT	The location pointed to by <i>obind</i> was not NULL and not writable by the user.								
EINVAL	The specified processor is not on-line, or the <i>idtype</i> argument was not <code>P_PID</code> , <code>P_LWPID</code> , or <code>P_TASKID</code> .								
EPERM	The effective user of the calling process is not superuser, and its real or effective user ID does not match the real or effective user ID of one of the LWP's being bound.								
ESRCH	No processes, LWP's, or tasks were found to match the criteria specified by <i>idtype</i> and <i>id</i> .								

processor_bind(2)

SEE ALSO psradm(1M), psrinfo(1M), p_online(2), pset_bind(2), sysconf(3C)

processor_info(2)

NAME	processor_info – determine type and status of a processor				
SYNOPSIS	<pre>#include <sys/types.h> #include <sys/processor.h> int processor_info(processorid_t <i>processorid</i>, processor_info_t *<i>infop</i>) ;</pre>				
DESCRIPTION	The <code>processor_info()</code> function returns the status of the processor specified by <i>processorid</i> in the <code>processor_info_t</code> structure pointed to by <i>infop</i>. The structure <code>processor_info_t</code> contains the following members: <pre>int pi_state; char pi_processor_type[PI_TYPELEN]; char pi_fputypes[PI_FPUTYPE]; int pi_clock;</pre> The <code>pi_state</code> member is the current state of the processor, either <code>P_ONLINE</code>, <code>P_OFFLINE</code>, or <code>P_POWEROFF</code>. The <code>pi_processor_type</code> member is a null-terminated ASCII string specifying the type of the processor. The <code>pi_fputypes</code> member is a null-terminated ASCII string containing the comma-separated types of floating-point units (FPUs) attached to the processor. This string will be empty if no FPU is attached. The <code>pi_clock</code> member is the processor clock frequency rounded to the nearest megahertz. It may be 0 if not known.				
RETURN VALUES	Upon successful completion, 0 is returned. Otherwise, -1 is returned and <code>errno</code> is set to indicate the error.				
ERRORS	The <code>processor_info()</code> function will fail if: <table><tr><td>EINVAL</td><td>An non-existent processor ID was specified.</td></tr><tr><td>EFAULT</td><td>The <code>processor_info_t</code> structure pointed to by <i>infop</i> was not writable by the user.</td></tr></table>	EINVAL	An non-existent processor ID was specified.	EFAULT	The <code>processor_info_t</code> structure pointed to by <i>infop</i> was not writable by the user.
EINVAL	An non-existent processor ID was specified.				
EFAULT	The <code>processor_info_t</code> structure pointed to by <i>infop</i> was not writable by the user.				
SEE ALSO	<code>psradm(1M)</code> , <code>psrinfo(1M)</code> , <code>p_online(2)</code> , <code>sysconf(3C)</code>				

NAME	profil – execution time profile
SYNOPSIS	<pre>#include <unistd.h> void profil(unsigned short *<i>buff</i>, unsigned int <i>bufsiz</i>, unsigned int <i>offset</i>, unsigned int <i>scale</i>);</pre>
DESCRIPTION	The <code>profil()</code> function provides CPU-use statistics by profiling the amount of CPU time expended by a program. The <code>profil()</code> function generates the statistics by creating an execution histogram for a current process. The histogram is defined for a specific region of program code to be profiled, and the identified region is logically broken up into a set of equal size subdivisions, each of which corresponds to a count in the histogram. With each clock tick, the current subdivision is identified and its corresponding histogram count is incremented. These counts establish a relative measure of how much time is being spent in each code subdivision. The resulting histogram counts for a profiled region can be used to identify those functions that consume a disproportionately high percentage of CPU time. The <i>buff</i> argument is a buffer of <i>bufsiz</i> bytes in which the histogram counts are stored in an array of unsigned short int. The <i>offset</i>, <i>scale</i>, and <i>bufsiz</i> arguments specify the region to be profiled. The <i>offset</i> argument is effectively the start address of the region to be profiled. The <i>scale</i> argument is a contraction factor that indicates how much smaller the histogram buffer is than the region to be profiled. More precisely, <i>scale</i> is interpreted as an unsigned 16-bit fixed-point fraction with the decimal point implied on the left. Its value is the reciprocal of the number of bytes in a subdivision, per byte of histogram buffer. Since there are two bytes per histogram counter, the effective ratio of subdivision bytes per counter is one half the scale. The values of <i>scale</i> are as follows: ■ the maximum value of <i>scale</i>, 0xffff (approximately 1), maps subdivisions 2 bytes long to each counter. ■ the minimum value of <i>scale</i> (for which profiling is performed), 0x0002 (1/32,768), maps subdivision 65,536 bytes long to each counter. ■ the default value of <i>scale</i> (currently used by <code>cc -qp</code>), 0x4000, maps subdivisions 8 bytes long to each counter. The values are used within the kernel as follows: when the process is interrupted for a clock tick, the value of <i>offset</i> is subtracted from the current value of the program counter (<i>pc</i>), and the remainder is multiplied by <i>scale</i> to derive a result. That result is used as an index into the histogram array to locate the cell to be incremented. Therefore, the cell count represents the number of times that the process was executing code in the subdivision associated with that cell when the process was interrupted.

profil(2)

The value of *scale* can be computed as $(RATIO * 0200000L)$, where *RATIO* is the desired ratio of *bufsiz* to profiled region size, and has a value between 0 and 1. Qualitatively speaking, the closer *RATIO* is to 1, the higher the resolution of the profile information.

The value of *bufsiz* can be computed as $(size_of_region_to_be_profiled * RATIO)$.

Profiling is turned off by giving a *scale* value of 0 or 1, and is rendered ineffective by giving a *bufsiz* value of 0. Profiling is turned off when one of the *exec* family of functions (see *exec(2)*) is executed, but remains on in both child and parent processes after a *fork(2)*. Profiling is turned off if a *buff* update would cause a memory fault.

USAGE The *pcsample(2)* function should be used when profiling dynamically-linked programs and 64-bit programs.

SEE ALSO *exec(2)*, *fork(2)*, *pcsample(2)*, *times(2)*, *monitor(3C)*, *prof(5)*

NOTES In Solaris releases prior to 2.6, calling *profil()* in a multithreaded program would impact only the calling LWP; the profile state was not inherited at LWP creation time. To profile a multithreaded program with a global profile buffer, each thread needed to issue a call to *profil()* at threads start-up time, and each thread had to be a bound thread. This was cumbersome and did not easily support dynamically turning profiling on and off. In Solaris 2.6, the *profil()* system call for multithreaded processes has global impact — that is, a call to *profil()* impacts all LWPs/threads in the process. This may cause applications that depend on the previous per-LWP semantic to break, but it is expected to improve multithreaded programs that wish to turn profiling on and off dynamically at runtime.

NAME	pset_bind – bind LWPs to a set of processors						
SYNOPSIS	<pre>#include <sys/pset.h> int pset_bind(psetid_t <i>pset</i>, idtype_t <i>idtype</i>, id_t <i>id</i>, psetid_t *<i>opset</i>) ;</pre>						
DESCRIPTION	The <code>pset_bind()</code> function binds the LWP or set of LWPs specified by <i>idtype</i> and <i>id</i> to the processor set specified by <i>pset</i>. If <i>obind</i> is not NULL, <code>pset_bind()</code> sets the <code>psetid_t</code> variable pointed to by <i>opset</i> to the previous processor set binding of one of the specified LWP, or to <code>PS_NONE</code> if the selected LWP was not bound. If <i>idtype</i> is <code>P_PID</code>, the binding affects all LWPs of the process with process ID (PID) <i>id</i>. If <i>idtype</i> is <code>P_LWPID</code>, the binding affects the LWP of the current process with LWP ID <i>id</i>. If <i>idtype</i> is <code>P_TASKID</code>, the binding affects all LWPs of all processes with task ID <i>id</i>. If <i>id</i> is <code>P_MYID</code>, the specified LWP or process is the current one. If <i>pset</i> is <code>PS_NONE</code>, the processor set bindings of the specified LWPs are cleared. If <i>pset</i> is <code>PS_QUERY</code>, the processor set bindings are not changed. The effective user of the calling process must be super-user, or its real or effective user ID must match the real or effective user ID of the LWPs being bound, or <i>pset</i> must be <code>PS_QUERY</code>. If the calling process does not have permission to change all of the specified LWPs, the bindings of the LWPs for which it does have permission will be changed even though an error is returned. If the processor set type of <i>pset</i> is <code>PS_PRIVATE</code> (see <code>pset_info(2)</code>), the effective user of the calling process must be super-user. LWPs that have been bound to a processor with <code>processor_bind(2)</code> may also be bound to a processor set if the processor is part of the processor set. If this occurs, the binding to the processor remains in effect. If the processor binding is later removed, the processor set binding becomes effective.						
RETURN VALUES	Upon successful completion, 0 is returned. Otherwise, -1 is returned and <code>errno</code> is set to indicate the error.						
ERRORS	The <code>pset_bind()</code> function will fail if: <table> <tr> <td>EBUSY</td><td>One of the LWPs is bound to a processor, and the specified processor set does not include that processor.</td></tr> <tr> <td>EFAULT</td><td>The location pointed to by <i>opset</i> was not NULL and not writable by the user.</td></tr> <tr> <td>EINVAL</td><td>An invalid processor set ID was specified; or <i>idtype</i> was not <code>P_PID</code>, <code>P_LWPID</code>, or <code>P_TASKID</code>.</td></tr> </table>	EBUSY	One of the LWPs is bound to a processor, and the specified processor set does not include that processor.	EFAULT	The location pointed to by <i>opset</i> was not NULL and not writable by the user.	EINVAL	An invalid processor set ID was specified; or <i>idtype</i> was not <code>P_PID</code> , <code>P_LWPID</code> , or <code>P_TASKID</code> .
EBUSY	One of the LWPs is bound to a processor, and the specified processor set does not include that processor.						
EFAULT	The location pointed to by <i>opset</i> was not NULL and not writable by the user.						
EINVAL	An invalid processor set ID was specified; or <i>idtype</i> was not <code>P_PID</code> , <code>P_LWPID</code> , or <code>P_TASKID</code> .						

pset_bind(2)

EPERM	The effective user of the calling process is not super-user, and either the processor set type of <i>pset</i> is PS_USER, or the real or effective user ID of the calling process does not match the real or effective user ID of one of the LWPs being bound.
ESRCH	No processes, LWPs, or tasks were found to match the criteria specified by <i>idtype</i> and <i>id</i> .

SEE ALSO pbind(1M), psrset(1M), processor_bind(2), pset_create(2), pset_info(2)

NAME	pset_create, pset_destroy, pset_assign – manage sets of processors				
SYNOPSIS	<pre>#include <sys/pset.h> int pset_create(psetid_t *newpset); int pset_destroy(psetid_t pset); int pset_assign(psetid_t pset, processorid_t cpu, psetid_t *opset);</pre>				
DESCRIPTION	These functions control the creation and management of sets of processors. Processor sets allow a subset of the system's processors to be set aside for exclusive use by specified LWPs and processes. The binding of LWPs and processes to processor sets is controlled by pset_bind(2). The pset_create() function creates an empty processor set that contains no processors. On successful return, <i>newpset</i> will contain the ID of the new processor set. Only a limited number of processor sets may be active (created and not destroyed) at a given time. This limit will always be greater than the number of processors in the system. If pset_create() is called when the maximum number of processor sets is already active, the function will return -1 and errno will be set to ENOMEM. The pset_destroy() function destroys the processor set <i>pset</i>, releasing its constituent processors and processes. The pset_assign() function assigns the processor <i>cpu</i> to the processor set <i>pset</i>. A processor that has been assigned to a processor set will run only LWPs and processes that have been explicitly bound to that processor set, unless another LWP requires a resource that is only available on that processor. On successful return, if <i>opset</i> is non-null, <i>opset</i> will contain the processor set ID of the former processor set of the processor. If <i>pset</i> is PS_NONE, pset_assign() releases processor <i>cpu</i> from its current processor set. If <i>pset</i> is PS_QUERY, pset_assign() makes no change to processor sets, but returns the current processor set ID of processor <i>cpu</i> in <i>opset</i>. These functions are restricted to super-user use, except for pset_assign() when <i>pset</i> is PS_QUERY.				
RETURN VALUES	Upon successful completion, these functions return 0. Otherwise, -1 is returned and errno is set to indicate the error.				
ERRORS	These functions will fail if: <table> <tr> <td>EBUSY</td><td>The processor could not be moved to the specified processor set.</td></tr> <tr> <td>EFAULT</td><td>The location pointed to by <i>newpset</i> was not writable by the user, or the location pointed to by <i>opset</i> was not NULL and not writable by the user.</td></tr> </table>	EBUSY	The processor could not be moved to the specified processor set.	EFAULT	The location pointed to by <i>newpset</i> was not writable by the user, or the location pointed to by <i>opset</i> was not NULL and not writable by the user.
EBUSY	The processor could not be moved to the specified processor set.				
EFAULT	The location pointed to by <i>newpset</i> was not writable by the user, or the location pointed to by <i>opset</i> was not NULL and not writable by the user.				

pset_create(2)

EINVAL	The specified processor does not exist, the specified processor is not on-line, or an invalid processor set was specified.
ENOMEM	There was insufficient space for pset_create to create a new processor set.
EPERM	The effective user of the calling process is not super-user.

SEE ALSO psradm(1M), psrinfo(1M), psrset(1M), p_online(2), processor_bind(2), pset_bind(2), pset_info(2)

NOTES Processors belonging to different processor sets of type PS_SYSTEM (see pset_info(2)) cannot be assigned to the same processor set of type PS_PRIVATE. If this is attempted, pset_assign() will fail and set errno to EINVAL.

Processors with LWPs bound to them using processor_bind(2) cannot be assigned to a new processor set. If this is attempted, pset_assign() will fail and set errno to EBUSY.

NAME	pset_info – get information about a processor set				
SYNOPSIS	<pre>#include <sys/pset.h> int pset_info(psetid_t <i>pset</i>, int *<i>type</i>, uint_t *<i>numcpus</i>, processorid_t *<i>cpulist</i>);</pre>				
DESCRIPTION	The <code>pset_info()</code> function returns information on the processor set <i>pset</i>. If <i>type</i> is non-null, then on successful completion the type of the processor set will be stored in the location pointed to by <i>type</i>. Processor set types can have the following values: <table> <tr> <td>PS_SYSTEM</td><td>The processor set was created by the system. Processor sets of this type cannot be modified or removed by the user, but LWPs and processes can be bound to them using <code>pset_bind(2)</code>.</td></tr> <tr> <td>PS_PRIVATE</td><td>The processor set was created by <code>pset_create(2)</code> and can be modified by <code>pset_assign(2)</code> and removed by <code>pset_destroy(2)</code>. LWPs and processes can also be bound to this processor set using <code>pset_bind()</code>.</td></tr> </table> If <i>numcpus</i> is non-null, then on successful completion the number of processors in the processor set will be stored in the location pointed to by <i>numcpus</i>. If <i>numcpus</i> and <i>cpulist</i> are both non-null, then <i>cpulist</i> points to a buffer where a list of processors assigned to the processor set is to be stored, and <i>numcpus</i> points to the maximum number of processor IDs the buffer can hold. On successful completion, the list of processors up to the maximum buffer size is stored in the buffer pointed to by <i>cpulist</i>.	PS_SYSTEM	The processor set was created by the system. Processor sets of this type cannot be modified or removed by the user, but LWPs and processes can be bound to them using <code>pset_bind(2)</code> .	PS_PRIVATE	The processor set was created by <code>pset_create(2)</code> and can be modified by <code>pset_assign(2)</code> and removed by <code>pset_destroy(2)</code> . LWPs and processes can also be bound to this processor set using <code>pset_bind()</code> .
PS_SYSTEM	The processor set was created by the system. Processor sets of this type cannot be modified or removed by the user, but LWPs and processes can be bound to them using <code>pset_bind(2)</code> .				
PS_PRIVATE	The processor set was created by <code>pset_create(2)</code> and can be modified by <code>pset_assign(2)</code> and removed by <code>pset_destroy(2)</code> . LWPs and processes can also be bound to this processor set using <code>pset_bind()</code> .				
RETURN VALUES	Upon successful completion, 0 is returned. Otherwise, -1 is returned and <code>errno</code> is set to indicate the error.				
ERRORS	The <code>pset_info()</code> function will fail if: <table> <tr> <td>EFAULT</td><td>The location pointed to by <i>type</i>, <i>numcpus</i>, or <i>cpulist</i> was not null and not writable by the user.</td></tr> <tr> <td>EINVAL</td><td>An invalid processor set ID was specified.</td></tr> </table>	EFAULT	The location pointed to by <i>type</i> , <i>numcpus</i> , or <i>cpulist</i> was not null and not writable by the user.	EINVAL	An invalid processor set ID was specified.
EFAULT	The location pointed to by <i>type</i> , <i>numcpus</i> , or <i>cpulist</i> was not null and not writable by the user.				
EINVAL	An invalid processor set ID was specified.				
SEE ALSO	<code>psrinfo(1M)</code> , <code>psrset(1M)</code> , <code>processor_info(2)</code> , <code>pset_assign(2)</code> , <code>pset_bind(2)</code> , <code>pset_create(2)</code> , <code>pset_destroy(2)</code>				

ptrace(2)

NAME	ptrace – allows a parent process to control the execution of a child process
SYNOPSIS	<pre>#include <unistd.h> #include <sys/types.h> int ptrace(int <i>request</i>, pid_t <i>pid</i>, int <i>addr</i>, int <i>data</i>);</pre>
DESCRIPTION	The <code>ptrace()</code> function allows a parent process to control the execution of a child process. Its primary use is for the implementation of breakpoint debugging. The child process behaves normally until it encounters a signal (see <code>signal(3HEAD)</code>), at which time it enters a stopped state and its parent is notified via the <code>wait(2)</code> function. When the child is in the stopped state, its parent can examine and modify its “core image” using <code>ptrace()</code>. Also, the parent can cause the child either to terminate or continue, with the possibility of ignoring the signal that caused it to stop. The <i>request</i> argument determines the action to be taken by <code>ptrace()</code> and is one of the following: 0 This request must be issued by the child process if it is to be traced by its parent. It turns on the child’s trace flag that stipulates that the child should be left in a stopped state on receipt of a signal rather than the state specified by <i>func</i> (see <code>signal(3C)</code>). The <i>pid</i>, <i>addr</i>, and <i>data</i> arguments are ignored, and a return value is not defined for this request. Peculiar results ensue if the parent does not expect to trace the child. The remainder of the requests can only be used by the parent process. For each, <i>pid</i> is the process ID of the child. The child must be in a stopped state before these requests are made. 1, 2 With these requests, the word at location <i>addr</i> in the address space of the child is returned to the parent process. If instruction and data space are separated, request 1 returns a word from instruction space, and request 2 returns a word from data space. If instruction and data space are not separated, either request 1 or request 2 may be used with equal results. The <i>data</i> argument is ignored. These two requests fail if <i>addr</i> is not the start address of a word, in which case <code>-1</code> is returned to the parent process and the parent’s <code>errno</code> is set to <code>EIO</code>. 3 With this request, the word at location <i>addr</i> in the child’s user area in the system’s address space (see <code><sys/user.h></code>) is returned to the parent process. The <i>data</i> argument is ignored. This request fails if <i>addr</i> is not the start address of a word or is outside the user area, in which case <code>-1</code> is returned to the parent process and the parent’s <code>errno</code> is set to <code>EIO</code>. 4, 5 With these requests, the value given by the <i>data</i> argument is written into the address space of the child at location <i>addr</i>. If instruction and data space are separated, request 4 writes a word into instruction space, and request 5 writes a word into data space. If instruction and data space are not separated, either request 4 or request 5 may be used with equal results. On success, the value written into the address space of the child is returned to

the parent. These two requests fail if *addr* is not the start address of a word. On failure `-1` is returned to the parent process and the parent's `errno` is set to `EIO`.

- | | |
|---|---|
| 6 | With this request, a few entries in the child's user area can be written. <i>data</i> gives the value that is to be written and <i>addr</i> is the location of the entry. The few entries that can be written are the general registers and the condition codes of the Processor Status Word. |
| 7 | This request causes the child to resume execution. If the <i>data</i> argument is 0, all pending signals including the one that caused the child to stop are canceled before it resumes execution. If the <i>data</i> argument is a valid signal number, the child resumes execution as if it had incurred that signal, and any other pending signals are canceled. The <i>addr</i> argument must be equal to 1 for this request. On success, the value of <i>data</i> is returned to the parent. This request fails if <i>data</i> is not 0 or a valid signal number, in which case <code>-1</code> is returned to the parent process and the parent's <code>errno</code> is set to <code>EIO</code> . |
| 8 | This request causes the child to terminate with the same consequences as <code>exit(2)</code> . |
| 9 | This request sets the trace bit in the Processor Status Word of the child and then executes the same steps as listed above for request 7. The trace bit causes an interrupt on completion of one machine instruction. This effectively allows single stepping of the child. |

To forestall possible fraud, `ptrace()` inhibits the set-user-ID facility on subsequent calls to one of the `exec` family of functions (see `exec(2)`). If a traced process calls one of the `exec` functions, it stops before executing the first instruction of the new image showing signal `SIGTRAP`.

ERRORS The `ptrace()` function will fail if:

<code>EIO</code>	The <i>request</i> argument is an illegal number.
<code>EPERM</code>	The effective user of the calling process is not super-user.
<code>ESRCH</code>	The <i>pid</i> argument identifies a child that does not exist or has not executed a <code>ptrace()</code> call with request 0.

SEE ALSO `exec(2)`, `exit(2)`, `wait(2)`, `signal(3C)`, `signal(3HEAD)`

putmsg(2)

NAME	putmsg, putpmsg – send a message on a stream
SYNOPSIS	<pre>#include <stropts.h> int putmsg(int fildes, const struct strbuf *ctlptr, const struct strbuf *dataptr, int flags); int putpmsg(int fildes, const struct strbuf *ctlptr, const struct strbuf *dataptr, int band, int flags);</pre>
DESCRIPTION	The <code>putmsg()</code> function creates a message from user-specified buffer(s) and sends the message to a STREAMS file. The message may contain either a data part, a control part, or both. The data and control parts to be sent are distinguished by placement in separate buffers, as described below. The semantics of each part is defined by the STREAMS module that receives the message. The <code>putpmsg()</code> function does the same thing as <code>putmsg()</code>, but provides the user the ability to send messages in different priority bands. Except where noted, all information pertaining to <code>putmsg()</code> also pertains to <code>putpmsg()</code>. The <i>fildes</i> argument specifies a file descriptor referencing an open stream. The <i>ctlptr</i> and <i>dataptr</i> arguments each point to a <code>strbuf</code> structure, which contains the following members: <pre>int maxlen; /* not used here */ int len; /* length of data */ void *buf; /* ptr to buffer */</pre> The <i>ctlptr</i> argument points to the structure describing the control part, if any, to be included in the message. The <i>buf</i> member in the <code>strbuf</code> structure points to the buffer where the control information resides, and the <i>len</i> member indicates the number of bytes to be sent. The <i>maxlen</i> member is not used in <code>putmsg()</code> (see <code>getmsg(2)</code>). In a similar manner, <i>dataptr</i> specifies the data, if any, to be included in the message. The <i>flags</i> argument indicates what type of message should be sent and is described later. To send the data part of a message, <i>dataptr</i> must not be <code>NULL</code>, and the <i>len</i> member of <i>dataptr</i> must have a value of 0 or greater. To send the control part of a message, the corresponding values must be set for <i>ctlptr</i>. No data (control) part is sent if either <i>dataptr</i> (<i>ctlptr</i>) is <code>NULL</code> or the <i>len</i> member of <i>dataptr</i> (<i>ctlptr</i>) is negative. For <code>putmsg()</code>, if a control part is specified, and <i>flags</i> is set to <code>RS_HIPRI</code>, a high priority message is sent. If no control part is specified, and <i>flags</i> is set to <code>RS_HIPRI</code>, <code>putmsg()</code> fails and sets <code>errno</code> to <code>EINVAL</code>. If <i>flags</i> is set to 0, a normal (non-priority) message is sent. If no control part and no data part are specified, and <i>flags</i> is set to 0, no message is sent, and 0 is returned. The stream head guarantees that the control part of a message generated by <code>putmsg()</code> is at least 64 bytes in length.

For `putpmsg()`, the flags are different. The *flags* argument is a bitmask with the following mutually-exclusive flags defined: `MSG_HIPRI` and `MSG_BAND`. If *flags* is set to 0, `putpmsg()` fails and sets `errno` to `EINVAL`. If a control part is specified and *flags* is set to `MSG_HIPRI` and *band* is set to 0, a high-priority message is sent. If *flags* is set to `MSG_HIPRI` and either no control part is specified or *band* is set to a non-zero value, `putpmsg()` fails and sets `errno` to `EINVAL`. If *flags* is set to `MSG_BAND`, then a message is sent in the priority band specified by *band*. If a control part and data part are not specified and *flags* is set to `MSG_BAND`, no message is sent and 0 is returned.

Normally, `putmsg()` will block if the stream write queue is full due to internal flow control conditions. For high-priority messages, `putmsg()` does not block on this condition. For other messages, `putmsg()` does not block when the write queue is full and `O_NDELAY` or `O_NONBLOCK` is set. Instead, it fails and sets `errno` to `EAGAIN`.

The `putmsg()` or `putpmsg()` function also blocks, unless prevented by lack of internal resources, waiting for the availability of message blocks in the stream, regardless of priority or whether `O_NDELAY` or `O_NONBLOCK` has been specified. No partial message is sent.

RETURN VALUES Upon successful completion, 0 is returned. Otherwise, -1 is returned and `errno` is set to indicate the error.

ERRORS The `putmsg()` and `putpmsg()` functions will fail if:

<code>EAGAIN</code>	A non-priority message was specified, the <code>O_NDELAY</code> or <code>O_NONBLOCK</code> flag is set and the stream write queue is full due to internal flow control conditions.
<code>EBADF</code>	The <i>fildev</i> argument is not a valid file descriptor open for writing.
<code>EFAULT</code>	The <i>ctlptr</i> or <i>dataptr</i> argument points to an illegal address.
<code>EINTR</code>	A signal was caught during the execution of the <code>putmsg()</code> function.
<code>EINVAL</code>	An undefined value was specified in <i>flags</i> ; <i>flags</i> is set to <code>MSG_HIPRI</code> and no control part was supplied; or the stream referenced by <i>fildev</i> is linked below a multiplexor.
<code>ENOSR</code>	Buffers could not be allocated for the message that was to be created due to insufficient STREAMS memory resources.
<code>ENOSTR</code>	The <i>fildev</i> argument is not associated with a STREAM.
<code>ENXIO</code>	A hangup condition was generated downstream for the specified stream, or the other end of the pipe is closed.
<code>EPIPE</code> or <code>EIO</code>	The <i>fildev</i> argument refers to a STREAMS-based pipe and the other end of the pipe is closed. A <code>SIGPIPE</code> signal is generated for the calling process. This error condition occurs only with SUS-compliant applications. See <code>standards(5)</code> .

putmsg(2)

ERANGE

The size of the data part of the message does not fall within the range specified by the maximum and minimum packet sizes of the topmost stream module. This value is also returned if the control part of the message is larger than the maximum configured size of the control part of a message, or if the data part of a message is larger than the maximum configured size of the data part of a message.

In addition, `putmsg()` and `putpmsg()` will fail if the STREAM head had processed an asynchronous error before the call. In this case, the value of `errno` does not reflect the result of `putmsg()` or `putpmsg()` but reflects the prior error.

The `putpmsg()` function will fail if:

EINVAL

The *flags* argument is set to `MSG_HIPRI` and *band* is non-zero.

SEE ALSO

`intro(2)`, `getmsg(2)`, `poll(2)`, `read(2)`, `write(2)`, `standards(5)`

STREAMS Programming Guide

NAME	read, readv, pread – read from file
SYNOPSIS	<pre>#include <unistd.h> ssize_t read(int fildes, void *buf, size_t nbyte); ssize_t pread(int fildes, void *buf, size_t nbyte, off_t offset); #include <sys/uio.h> ssize_t readv(int fildes, const struct iovec *iov, int iovcnt);</pre>
DESCRIPTION	The <code>read()</code> function attempts to read <i>nbyte</i> bytes from the file associated with the open file descriptor, <i>fildes</i>, into the buffer pointed to by <i>buf</i>. If <i>nbyte</i> is 0, <code>read()</code> will return 0 and have no other results. On files that support seeking (for example, a regular file), the <code>read()</code> starts at a position in the file given by the file offset associated with <i>fildes</i>. The file offset is incremented by the number of bytes actually read. Files that do not support seeking (for example, terminals) always read from the current position. The value of a file offset associated with such a file is undefined. If <i>fildes</i> refers to a socket, <code>read()</code> is equivalent to <code>recv(3SOCKET)</code> with no flags set. No data transfer will occur past the current end-of-file. If the starting position is at or after the end-of-file, 0 will be returned. If the file refers to a device special file, the result of subsequent <code>read()</code> requests is implementation-dependent. If the value of <i>nbyte</i> is greater than <code>SSIZE_MAX</code>, the result is implementation-dependent. When attempting to read from a regular file with mandatory file/record locking set (see <code>chmod(2)</code>), and there is a write lock owned by another process on the segment of the file to be read: ■ If <code>O_NDELAY</code> or <code>O_NONBLOCK</code> is set, <code>read()</code> returns -1 and sets <code>errno</code> to <code>EAGAIN</code>. ■ If <code>O_NDELAY</code> and <code>O_NONBLOCK</code> are clear, <code>read()</code> sleeps until the blocking record lock is removed. When attempting to read from an empty pipe (or FIFO): ■ If no process has the pipe open for writing, <code>read()</code> returns 0 to indicate end-of-file. ■ If some process has the pipe open for writing and <code>O_NDELAY</code> is set, <code>read()</code> returns 0. ■ If some process has the pipe open for writing and <code>O_NONBLOCK</code> is set, <code>read()</code> returns -1 and sets <code>errno</code> to <code>EAGAIN</code>.

read(2)

- If `O_NDELAY` and `O_NONBLOCK` are clear, `read()` blocks until data is written to the pipe or the pipe is closed by all processes that had opened the pipe for writing.

When attempting to read a file associated with a terminal that has no data currently available:

- If `O_NDELAY` is set, `read()` returns 0.
- If `O_NONBLOCK` is set, `read()` returns -1 and sets `errno` to `EAGAIN`.
- If `O_NDELAY` and `O_NONBLOCK` are clear, `read()` blocks until data become available.

When attempting to read a file associated with a socket or a stream that is not a pipe, a FIFO, or a terminal, and the file has no data currently available:

- If `O_NDELAY` or `O_NONBLOCK` is set, `read()` returns -1 and sets `errno` to `EAGAIN`.
- If `O_NDELAY` and `O_NONBLOCK` are clear, `read()` blocks until data becomes available.

The `read()` function reads data previously written to a file. If any portion of a regular file prior to the end-of-file has not been written, `read()` returns bytes with value 0. For example, `lseek(2)` allows the file offset to be set beyond the end of existing data in the file. If data is later written at this point, subsequent reads in the gap between the previous end of data and the newly written data will return bytes with value 0 until data is written into the gap.

For regular files, no data transfer will occur past the offset maximum established in the open file description associated with *filides*.

Upon successful completion, where *nbyte* is greater than 0, `read()` will mark for update the `st_atime` field of the file, and return the number of bytes read. This number will never be greater than *nbyte*. The value returned may be less than *nbyte* if the number of bytes left in the file is less than *nbyte*, if the `read()` request was interrupted by a signal, or if the file is a pipe or FIFO or special file and has fewer than *nbyte* bytes immediately available for reading. For example, a `read()` from a file associated with a terminal may return one typed line of data.

If a `read()` is interrupted by a signal before it reads any data, it will return -1 with `errno` set to `EINTR`.

If a `read()` is interrupted by a signal after it has successfully read some data, it will return the number of bytes read.

A `read()` from a STREAMS file can read data in three different modes: byte-stream mode, message-nondiscard mode, and message-discard mode. The default is byte-stream mode. This can be changed using the `I_SRDOPT` `ioctl(2)` request, and can be tested with the `I_GRDOPT` `ioctl()`. In byte-stream mode, `read()` retrieves data from the STREAM until as many bytes as were requested are transferred, or until there is no more data to be retrieved. Byte-stream mode ignores message boundaries.

read(2)

In STREAMS message-nondiscard mode, `read()` retrieves data until as many bytes as were requested are transferred, or until a message boundary is reached. If `read()` does not retrieve all the data in a message, the remaining data is left on the STREAM, and can be retrieved by the next `read()` call. Message-discard mode also retrieves data until as many bytes as were requested are transferred, or a message boundary is reached. However, unread data remaining in a message after the `read()` returns is discarded, and is not available for a subsequent `read()`, `readv()` or `getmsg(2)` call.

How `read()` handles zero-byte STREAMS messages is determined by the current read mode setting. In byte-stream mode, `read()` accepts data until it has read *nbyte* bytes, or until there is no more data to read, or until a zero-byte message block is encountered. The `read()` function then returns the number of bytes read, and places the zero-byte message back on the STREAM to be retrieved by the next `read()`, `readv()` or `getmsg(2)`. In message-nondiscard mode or message-discard mode, a zero-byte message returns 0 and the message is removed from the STREAM. When a zero-byte message is read as the first message on a STREAM, the message is removed from the STREAM and 0 is returned, regardless of the read mode.

A `read()` from a STREAMS file returns the data in the message at the front of the STREAM head read queue, regardless of the priority band of the message.

By default, STREAMS are in control-normal mode, in which a `read()` from a STREAMS file can only process messages that contain a data part but do not contain a control part. The `read()` fails if a message containing a control part is encountered at the STREAM head. This default action can be changed by placing the STREAM in either control-data mode or control-discard mode with the `I_SRDOPT` `ioctl()` command. In control-data mode, `read()` converts any control part to data and passes it to the application before passing any data part originally present in the same message. In control-discard mode, `read()` discards message control parts but returns to the process any data part in the message.

In addition, `read()` and `readv()` will fail if the STREAM head had processed an asynchronous error before the call. In this case, the value of `errno` does not reflect the result of `read()` or `readv()` but reflects the prior error. If a hangup occurs on the STREAM being read, `read()` continues to operate normally until the STREAM head read queue is empty. Thereafter, it returns 0.

readv()

The `readv()` function is equivalent to `read()`, but places the input data into the *iovcnt* buffers specified by the members of the *iov* array: *iov0*, *iov1*, ..., *iov[iovcnt-1]*. The *iovcnt* argument is valid if greater than 0 and less than or equal to `IOV_MAX`.

The `iovec` structure contains the following members:

```
caddr_t    iov_base;  
int         iov_len;
```

Each `iovec` entry specifies the base address and length of an area in memory where data should be placed. The `readv()` function always fills an area completely before proceeding to the next.

read(2)

	Upon successful completion, <code>readv()</code> marks for update the <code>st_atime</code> field of the file.																						
<code>pread()</code>	The <code>pread()</code> function performs the same action as <code>read()</code> , except that it reads from a given position in the file without changing the file pointer. The first three arguments to <code>pread()</code> are the same as <code>read()</code> with the addition of a fourth argument <i>offset</i> for the desired position inside the file. <code>pread()</code> will read up to the maximum offset value that can be represented in an <code>off_t</code> for regular files. An attempt to perform a <code>pread()</code> on a file that is incapable of seeking results in an error.																						
RETURN VALUES	Upon successful completion, <code>read()</code> and <code>readv()</code> return a non-negative integer indicating the number of bytes actually read. Otherwise, the functions return <code>-1</code> and set <code>errno</code> to indicate the error.																						
ERRORS	The <code>read()</code>, <code>readv()</code>, and <code>pread()</code> functions will fail if: <table><tr><td><code>EAGAIN</code></td><td>Mandatory file/record locking was set, <code>O_NDELAY</code> or <code>O_NONBLOCK</code> was set, and there was a blocking record lock; total amount of system memory available when reading using raw I/O is temporarily insufficient; no data is waiting to be read on a file associated with a tty device and <code>O_NONBLOCK</code> was set; or no message is waiting to be read on a stream and <code>O_NDELAY</code> or <code>O_NONBLOCK</code> was set.</td></tr><tr><td><code>EBADF</code></td><td>The <i>fildevs</i> argument is not a valid file descriptor open for reading.</td></tr><tr><td><code>EBADMSG</code></td><td>Message waiting to be read on a stream is not a data message.</td></tr><tr><td><code>EDEADLK</code></td><td>The read was going to go to sleep and cause a deadlock to occur.</td></tr><tr><td><code>EFAULT</code></td><td>The <i>buf</i> argument points to an illegal address.</td></tr><tr><td><code>EINTR</code></td><td>A signal was caught during the read operation and no data was transferred.</td></tr><tr><td><code>EINVAL</code></td><td>An attempt was made to read from a stream linked to a multiplexor.</td></tr><tr><td><code>EIO</code></td><td>A physical I/O error has occurred, or the process is in a background process group and is attempting to read from its controlling terminal, and either the process is ignoring or blocking the <code>SIGTTIN</code> signal or the process group of the process is orphaned.</td></tr><tr><td><code>EISDIR</code></td><td>The <i>fildevs</i> argument refers to a directory on a file system type that does not support read operations on directories.</td></tr><tr><td><code>ENOLCK</code></td><td>The system record lock table was full, so the <code>read()</code> or <code>readv()</code> could not go to sleep until the blocking record lock was removed.</td></tr><tr><td><code>ENOLINK</code></td><td>The <i>fildevs</i> argument is on a remote machine and the link to that machine is no longer active.</td></tr></table>	<code>EAGAIN</code>	Mandatory file/record locking was set, <code>O_NDELAY</code> or <code>O_NONBLOCK</code> was set, and there was a blocking record lock; total amount of system memory available when reading using raw I/O is temporarily insufficient; no data is waiting to be read on a file associated with a tty device and <code>O_NONBLOCK</code> was set; or no message is waiting to be read on a stream and <code>O_NDELAY</code> or <code>O_NONBLOCK</code> was set.	<code>EBADF</code>	The <i>fildevs</i> argument is not a valid file descriptor open for reading.	<code>EBADMSG</code>	Message waiting to be read on a stream is not a data message.	<code>EDEADLK</code>	The read was going to go to sleep and cause a deadlock to occur.	<code>EFAULT</code>	The <i>buf</i> argument points to an illegal address.	<code>EINTR</code>	A signal was caught during the read operation and no data was transferred.	<code>EINVAL</code>	An attempt was made to read from a stream linked to a multiplexor.	<code>EIO</code>	A physical I/O error has occurred, or the process is in a background process group and is attempting to read from its controlling terminal, and either the process is ignoring or blocking the <code>SIGTTIN</code> signal or the process group of the process is orphaned.	<code>EISDIR</code>	The <i>fildevs</i> argument refers to a directory on a file system type that does not support read operations on directories.	<code>ENOLCK</code>	The system record lock table was full, so the <code>read()</code> or <code>readv()</code> could not go to sleep until the blocking record lock was removed.	<code>ENOLINK</code>	The <i>fildevs</i> argument is on a remote machine and the link to that machine is no longer active.
<code>EAGAIN</code>	Mandatory file/record locking was set, <code>O_NDELAY</code> or <code>O_NONBLOCK</code> was set, and there was a blocking record lock; total amount of system memory available when reading using raw I/O is temporarily insufficient; no data is waiting to be read on a file associated with a tty device and <code>O_NONBLOCK</code> was set; or no message is waiting to be read on a stream and <code>O_NDELAY</code> or <code>O_NONBLOCK</code> was set.																						
<code>EBADF</code>	The <i>fildevs</i> argument is not a valid file descriptor open for reading.																						
<code>EBADMSG</code>	Message waiting to be read on a stream is not a data message.																						
<code>EDEADLK</code>	The read was going to go to sleep and cause a deadlock to occur.																						
<code>EFAULT</code>	The <i>buf</i> argument points to an illegal address.																						
<code>EINTR</code>	A signal was caught during the read operation and no data was transferred.																						
<code>EINVAL</code>	An attempt was made to read from a stream linked to a multiplexor.																						
<code>EIO</code>	A physical I/O error has occurred, or the process is in a background process group and is attempting to read from its controlling terminal, and either the process is ignoring or blocking the <code>SIGTTIN</code> signal or the process group of the process is orphaned.																						
<code>EISDIR</code>	The <i>fildevs</i> argument refers to a directory on a file system type that does not support read operations on directories.																						
<code>ENOLCK</code>	The system record lock table was full, so the <code>read()</code> or <code>readv()</code> could not go to sleep until the blocking record lock was removed.																						
<code>ENOLINK</code>	The <i>fildevs</i> argument is on a remote machine and the link to that machine is no longer active.																						

ENXIO The device associated with *fildev* is a block special or character special file and the value of the file pointer is out of range.

The `read()` and `readv()` functions will fail if:

EOVERFLOW The file is a regular file, *nbyte* is greater than 0, the starting position is before the end-of-file, and the starting position is greater than or equal to the offset maximum established in the open file description associated with *fildev*.

The `readv()` function may fail if:

EFAULT The *iov* argument points outside the allocated address space.

EINVAL The *iovcnt* argument was less than or equal to 0, or greater than or equal to `{IOV_MAX}`. (See `intro(3)` for a definition of `{IOV_MAX}`).

EINVAL The sum of the *iov_len* values in the *iov* array overflowed an int.

The `pread()` function will fail and the file pointer remain unchanged if:

ESPIPE The *fildev* argument is associated with a pipe or FIFO.

USAGE The `pread()` function has a transitional interface for 64-bit file offsets. See `lf64(5)`.

ATTRIBUTES See `attributes(5)` for descriptions of the following attributes:

ATTRIBUTE TYPE	ATTRIBUTE VALUE
MT-Level	<code>read()</code> is Async-Signal-Safe

SEE ALSO `intro(3)`, `chmod(2)`, `creat(2)`, `dup(2)`, `fcntl(2)`, `getmsg(2)`, `ioctl(2)`, `lseek(2)`, `open(2)`, `pipe(2)`, `recv(3SOCKET)`, `attributes(5)`, `lf64(5)`, `streamio(7I)`, `termio(7I)`

readlink(2)

NAME	readlink – read the contents of a symbolic link																						
SYNOPSIS	<pre>#include <unistd.h> int readlink(const char *path, char *buf, size_t bufsiz);</pre>																						
DESCRIPTION	The <code>readlink()</code> function places the contents of the symbolic link referred to by <i>path</i> in the buffer <i>buf</i> which has size <i>bufsiz</i> . If the number of bytes in the symbolic link is less than <i>bufsiz</i> , the contents of the remainder of <i>buf</i> are unspecified.																						
RETURN VALUES	Upon successful completion, <code>readlink()</code> returns the count of bytes placed in the buffer. Otherwise, it returns <code>-1</code> , leaves the buffer unchanged, and sets <code>errno</code> to indicate the error.																						
ERRORS	The <code>readlink()</code> function will fail if: <table> <tr> <td>EACCES</td><td>Search permission is denied for a component of the path prefix of <i>path</i>.</td></tr> <tr> <td>EFAULT</td><td><i>path</i> or <i>buf</i> points to an illegal address.</td></tr> <tr> <td>EINVAL</td><td>The <i>path</i> argument names a file that is not a symbolic link.</td></tr> <tr> <td>EIO</td><td>An I/O error occurred while reading from the file system.</td></tr> <tr> <td>ENOENT</td><td>A component of <i>path</i> does not name an existing file or <i>path</i> is an empty string.</td></tr> <tr> <td>ELOOP</td><td>Too many symbolic links were encountered in resolving <i>path</i>.</td></tr> <tr> <td>ENAMETOOLONG</td><td>The length of <i>path</i> exceeds <code>PATH_MAX</code>, or a pathname component is longer than <code>NAME_MAX</code> while <code>_POSIX_NO_TRUNC</code> is in effect.</td></tr> <tr> <td>ENOTDIR</td><td>A component of the path prefix is not a directory.</td></tr> <tr> <td>ENOSYS</td><td>The file system does not support symbolic links.</td></tr> </table> The <code>readlink()</code> function may fail if: <table> <tr> <td>EACCES</td><td>Read permission is denied for the directory.</td></tr> <tr> <td>ENAMETOOLONG</td><td>Path name resolution of a symbolic link produced an intermediate result whose length exceeds <code>PATH_MAX</code>.</td></tr> </table>	EACCES	Search permission is denied for a component of the path prefix of <i>path</i> .	EFAULT	<i>path</i> or <i>buf</i> points to an illegal address.	EINVAL	The <i>path</i> argument names a file that is not a symbolic link.	EIO	An I/O error occurred while reading from the file system.	ENOENT	A component of <i>path</i> does not name an existing file or <i>path</i> is an empty string.	ELOOP	Too many symbolic links were encountered in resolving <i>path</i> .	ENAMETOOLONG	The length of <i>path</i> exceeds <code>PATH_MAX</code> , or a pathname component is longer than <code>NAME_MAX</code> while <code>_POSIX_NO_TRUNC</code> is in effect.	ENOTDIR	A component of the path prefix is not a directory.	ENOSYS	The file system does not support symbolic links.	EACCES	Read permission is denied for the directory.	ENAMETOOLONG	Path name resolution of a symbolic link produced an intermediate result whose length exceeds <code>PATH_MAX</code> .
EACCES	Search permission is denied for a component of the path prefix of <i>path</i> .																						
EFAULT	<i>path</i> or <i>buf</i> points to an illegal address.																						
EINVAL	The <i>path</i> argument names a file that is not a symbolic link.																						
EIO	An I/O error occurred while reading from the file system.																						
ENOENT	A component of <i>path</i> does not name an existing file or <i>path</i> is an empty string.																						
ELOOP	Too many symbolic links were encountered in resolving <i>path</i> .																						
ENAMETOOLONG	The length of <i>path</i> exceeds <code>PATH_MAX</code> , or a pathname component is longer than <code>NAME_MAX</code> while <code>_POSIX_NO_TRUNC</code> is in effect.																						
ENOTDIR	A component of the path prefix is not a directory.																						
ENOSYS	The file system does not support symbolic links.																						
EACCES	Read permission is denied for the directory.																						
ENAMETOOLONG	Path name resolution of a symbolic link produced an intermediate result whose length exceeds <code>PATH_MAX</code> .																						
USAGE	Portable applications should not assume that the returned contents of the symbolic link are null-terminated.																						
SEE ALSO	<code>stat(2)</code> , <code>symlink(2)</code>																						

NAME	rename – change the name of a file
SYNOPSIS	<pre>#include <stdio.h> int rename(const char *old, const char *new);</pre>
DESCRIPTION	The <code>rename()</code> function changes the name of a file. The <i>old</i> argument points to the pathname of the file to be renamed. The <i>new</i> argument points to the new pathname of the file. If <i>old</i> and <i>new</i> both refer to the same existing file, the <code>rename()</code> function returns successfully and performs no other action. If <i>old</i> points to the pathname of a file that is not a directory, <i>new</i> must not point to the pathname of a directory. If the link named by <i>new</i> exists, it will be removed and <i>old</i> will be renamed to <i>new</i>. In this case, a link named <i>new</i> must remain visible to other processes throughout the renaming operation and will refer to either the file referred to by <i>new</i> or the file referred to as <i>old</i> before the operation began. If <i>old</i> points to the pathname of a directory, <i>new</i> must not point to the pathname of a file that is not a directory. If the directory named by <i>new</i> exists, it will be removed and <i>old</i> will be renamed to <i>new</i>. In this case, a link named <i>new</i> will exist throughout the renaming operation and will refer to either the file referred to by <i>new</i> or the file referred to as <i>old</i> before the operation began. Thus, if <i>new</i> names an existing directory, it must be an empty directory. The <i>new</i> pathname must not contain a path prefix that names <i>old</i>. Write access permission is required for both the directory containing <i>old</i> and the directory containing <i>new</i>. If <i>old</i> points to the pathname of a directory, write access permission is required for the directory named by <i>old</i>, and, if it exists, the directory named by <i>new</i>. If the directory containing <i>old</i> has the sticky bit set, at least one of the following conditions listed below must be true: ■ the user must own <i>old</i> ■ the user must own the directory containing <i>old</i> ■ <i>old</i> must be writable by the user ■ the user must be a privileged user If <i>new</i> exists, and the directory containing <i>new</i> is writable and has the sticky bit set, at least one of the following conditions must be true: ■ the user must own <i>new</i> ■ the user must own the directory containing <i>new</i> ■ <i>new</i> must be writable by the user ■ the user must be a privileged user

rename(2)

	If the link named by <i>new</i> exists, the file's link count becomes zero when it is removed, and no process has the file open, then the space occupied by the file will be freed and the file will no longer be accessible. If one or more processes have the file open when the last link is removed, the link will be removed before <code>rename()</code> returns, but the removal of the file contents will be postponed until all references to the file have been closed. Upon successful completion, the <code>rename()</code> function will mark for update the <code>st_ctime</code> and <code>st_mtime</code> fields of the parent directory of each file.																							
RETURN VALUES	Upon successful completion, 0 is returned. Otherwise, <code>-1</code> is returned and <code>errno</code> is set to indicate an error.																							
ERRORS	The <code>rename()</code> function will fail if: <table><tr><td>EACCES</td><td>A component of either path prefix denies search permission; one of the directories containing <i>old</i> and <i>new</i> denies write permissions; or write permission is denied by a directory pointed to by <i>old</i> or <i>new</i>.</td></tr><tr><td>EBUSY</td><td>The <i>new</i> argument is a directory and the mount point for a mounted file system.</td></tr><tr><td>EDQUOT</td><td>The directory where the new name entry is being placed cannot be extended because the user's quota of disk blocks on that file system has been exhausted.</td></tr><tr><td>EEXIST</td><td>The link named by <i>new</i> is a directory containing entries other than <code>'.'</code> (the directory itself) and <code>'..'</code> (the parent directory).</td></tr><tr><td>EINVAL</td><td>The <i>new</i> argument directory pathname contains a path prefix that names the <i>old</i> directory.</td></tr><tr><td>EISDIR</td><td>The <i>new</i> argument points to a directory but <i>old</i> points to a file that is not a directory.</td></tr><tr><td>ELOOP</td><td>Too many symbolic links were encountered in translating the pathname.</td></tr><tr><td>ENAMETOOLONG</td><td>The length of <i>old</i> or <i>new</i> exceeds <code>PATH_MAX</code>, or a pathname component is longer than <code>NAME_MAX</code> while <code>_POSIX_NO_TRUNC</code> is in effect.</td></tr><tr><td>EMLINK</td><td>The file named by <i>old</i> is a directory, and the link count of the parent directory of <i>new</i> would exceed <code>LINK_MAX</code>.</td></tr><tr><td>ENOENT</td><td>The link named by <i>old</i> does not exist, or either <i>old</i> or <i>new</i> points to an empty string.</td></tr><tr><td>ENOSPC</td><td>The directory that would contain <i>new</i> cannot be extended.</td></tr></table>		EACCES	A component of either path prefix denies search permission; one of the directories containing <i>old</i> and <i>new</i> denies write permissions; or write permission is denied by a directory pointed to by <i>old</i> or <i>new</i> .	EBUSY	The <i>new</i> argument is a directory and the mount point for a mounted file system.	EDQUOT	The directory where the new name entry is being placed cannot be extended because the user's quota of disk blocks on that file system has been exhausted.	EEXIST	The link named by <i>new</i> is a directory containing entries other than <code>'.'</code> (the directory itself) and <code>'..'</code> (the parent directory).	EINVAL	The <i>new</i> argument directory pathname contains a path prefix that names the <i>old</i> directory.	EISDIR	The <i>new</i> argument points to a directory but <i>old</i> points to a file that is not a directory.	ELOOP	Too many symbolic links were encountered in translating the pathname.	ENAMETOOLONG	The length of <i>old</i> or <i>new</i> exceeds <code>PATH_MAX</code> , or a pathname component is longer than <code>NAME_MAX</code> while <code>_POSIX_NO_TRUNC</code> is in effect.	EMLINK	The file named by <i>old</i> is a directory, and the link count of the parent directory of <i>new</i> would exceed <code>LINK_MAX</code> .	ENOENT	The link named by <i>old</i> does not exist, or either <i>old</i> or <i>new</i> points to an empty string.	ENOSPC	The directory that would contain <i>new</i> cannot be extended.
EACCES	A component of either path prefix denies search permission; one of the directories containing <i>old</i> and <i>new</i> denies write permissions; or write permission is denied by a directory pointed to by <i>old</i> or <i>new</i> .																							
EBUSY	The <i>new</i> argument is a directory and the mount point for a mounted file system.																							
EDQUOT	The directory where the new name entry is being placed cannot be extended because the user's quota of disk blocks on that file system has been exhausted.																							
EEXIST	The link named by <i>new</i> is a directory containing entries other than <code>'.'</code> (the directory itself) and <code>'..'</code> (the parent directory).																							
EINVAL	The <i>new</i> argument directory pathname contains a path prefix that names the <i>old</i> directory.																							
EISDIR	The <i>new</i> argument points to a directory but <i>old</i> points to a file that is not a directory.																							
ELOOP	Too many symbolic links were encountered in translating the pathname.																							
ENAMETOOLONG	The length of <i>old</i> or <i>new</i> exceeds <code>PATH_MAX</code> , or a pathname component is longer than <code>NAME_MAX</code> while <code>_POSIX_NO_TRUNC</code> is in effect.																							
EMLINK	The file named by <i>old</i> is a directory, and the link count of the parent directory of <i>new</i> would exceed <code>LINK_MAX</code> .																							
ENOENT	The link named by <i>old</i> does not exist, or either <i>old</i> or <i>new</i> points to an empty string.																							
ENOSPC	The directory that would contain <i>new</i> cannot be extended.																							

rename(2)

ENOTDIR	A component of either path prefix is not a directory, or <i>old</i> names a directory and <i>new</i> names a nondirectory file.
EROFS	The requested operation requires writing in a directory on a read-only file system.
EXDEV	The links named by <i>old</i> and <i>new</i> are on different file systems.
EIO	An I/O error occurred while making or updating a directory entry.

ATTRIBUTES See `attributes(5)` for descriptions of the following attributes:

ATTRIBUTE TYPE	ATTRIBUTE VALUE
MT-Level	Async-Signal-Safe

SEE ALSO `chmod(2)`, `link(2)`, `unlink(2)`, `attributes(5)`

NOTES The system can deadlock if there is a loop in the file system graph. Such a loop can occur if there is an entry in directory *a*, *a/name1*, that is a hard link to directory *b*, and an entry in directory *b*, *b/name2*, that is a hard link to directory *a*. When such a loop exists and two separate processes attempt to rename *a/name1* to *b/name2* and *b/name2* to *a/name1*, the system may deadlock attempting to lock both directories for modification. Use symbolic links instead of hard links for directories.

resolvepath(2)

NAME	resolvepath – resolve all symbolic links of a path name														
SYNOPSIS	<pre>#include <unistd.h> int resolvepath(const char *path, char *buf, size_t bufsiz);</pre>														
DESCRIPTION	The <code>resolvepath()</code> function fully resolves all symbolic links in the path name <i>path</i> into a resulting path name free of symbolic links and places the resulting path name in the buffer <i>buf</i> which has size <i>bufsiz</i> . The resulting path name names the same file or directory as the original path name. All "." components are eliminated and every non-leading "." component is eliminated together with its preceding directory component. If leading "." components reach to the root directory, they are replaced by "/". If the number of bytes in the resulting path name is less than <i>bufsiz</i> , the contents of the remainder of <i>buf</i> are unspecified.														
RETURN VALUES	Upon successful completion, <code>resolvepath()</code> returns the count of bytes placed in the buffer. Otherwise, it returns -1, leaves the buffer unchanged, and sets <code>errno</code> to indicate the error.														
ERRORS	The <code>resolvepath()</code> function will fail if: <table><tr><td>EACCES</td><td>Search permission is denied for a component of the path prefix of <i>path</i> or for a path prefix component resulting from the resolution of a symbolic link.</td></tr><tr><td>EFAULT</td><td>The <i>path</i> or <i>buf</i> argument points to an illegal address.</td></tr><tr><td>EIO</td><td>An I/O error occurred while reading from the file system.</td></tr><tr><td>ENOENT</td><td>The <i>path</i> argument is an empty string or a component of <i>path</i> or a path name component produced by resolving a symbolic link does not name an existing file.</td></tr><tr><td>ELOOP</td><td>Too many symbolic links were encountered in resolving <i>path</i>.</td></tr><tr><td>ENAMETOOLONG</td><td>The length of <i>path</i> exceeds <code>PATH_MAX</code>, or a path name component is longer than <code>NAME_MAX</code>. Path name resolution of a symbolic link produced an intermediate result whose length exceeds <code>PATH_MAX</code> or a component whose length exceeds <code>NAME_MAX</code>.</td></tr><tr><td>ENOTDIR</td><td>A component of the path prefix of <i>path</i> or of a path prefix component resulting from the resolution of a symbolic link is not a directory.</td></tr></table>	EACCES	Search permission is denied for a component of the path prefix of <i>path</i> or for a path prefix component resulting from the resolution of a symbolic link.	EFAULT	The <i>path</i> or <i>buf</i> argument points to an illegal address.	EIO	An I/O error occurred while reading from the file system.	ENOENT	The <i>path</i> argument is an empty string or a component of <i>path</i> or a path name component produced by resolving a symbolic link does not name an existing file.	ELOOP	Too many symbolic links were encountered in resolving <i>path</i> .	ENAMETOOLONG	The length of <i>path</i> exceeds <code>PATH_MAX</code> , or a path name component is longer than <code>NAME_MAX</code> . Path name resolution of a symbolic link produced an intermediate result whose length exceeds <code>PATH_MAX</code> or a component whose length exceeds <code>NAME_MAX</code> .	ENOTDIR	A component of the path prefix of <i>path</i> or of a path prefix component resulting from the resolution of a symbolic link is not a directory.
EACCES	Search permission is denied for a component of the path prefix of <i>path</i> or for a path prefix component resulting from the resolution of a symbolic link.														
EFAULT	The <i>path</i> or <i>buf</i> argument points to an illegal address.														
EIO	An I/O error occurred while reading from the file system.														
ENOENT	The <i>path</i> argument is an empty string or a component of <i>path</i> or a path name component produced by resolving a symbolic link does not name an existing file.														
ELOOP	Too many symbolic links were encountered in resolving <i>path</i> .														
ENAMETOOLONG	The length of <i>path</i> exceeds <code>PATH_MAX</code> , or a path name component is longer than <code>NAME_MAX</code> . Path name resolution of a symbolic link produced an intermediate result whose length exceeds <code>PATH_MAX</code> or a component whose length exceeds <code>NAME_MAX</code> .														
ENOTDIR	A component of the path prefix of <i>path</i> or of a path prefix component resulting from the resolution of a symbolic link is not a directory.														
USAGE	No more than <code>PATH_MAX</code> bytes will be placed in the buffer. Applications should not assume that the returned contents of the buffer are null-terminated.														
SEE ALSO	<code>readlink(2)</code> , <code>realpath(3C)</code>														

NAME	rmdir – remove a directory																	
SYNOPSIS	#include <unistd.h> int rmdir (const char *path);																	
DESCRIPTION	The <code>rmdir()</code> function removes the directory named by the path name pointed to by <i>path</i>. The directory must not have any entries other than “.” and “..”. If the directory’s link count becomes zero and no process has the directory open, the space occupied by the directory is freed and the directory is no longer accessible. If one or more processes have the directory open when the last link is removed, the “.” and “..” entries, if present, are removed before <code>rmdir()</code> returns and no new entries may be created in the directory, but the directory is not removed until all references to the directory have been closed. Upon successful completion <code>rmdir()</code> marks for update the <code>st_ctime</code> and <code>st_mtime</code> fields of the parent directory.																	
RETURN VALUES	Upon successful completion, 0 is returned. Otherwise, -1 is returned, <code>errno</code> is set to indicate the error, and the named directory is not changed.																	
ERRORS	The <code>rmdir()</code> function will fail if: <table><tr><td>EACCES</td><td>Search permission is denied for a component of the path prefix; write permission is denied on the directory containing the directory to be removed; the parent directory has the <code>S_ISVTX</code> variable set and is not owned by the user; the directory is not owned by the user and is not writable by the user; or the user is not a super-user.</td></tr><tr><td>EBUSY</td><td>The directory to be removed is the mount point for a mounted file system.</td></tr><tr><td>EEXIST</td><td>The directory contains entries other than those for “.” and “..”.</td></tr><tr><td>EFAULT</td><td>The <i>path</i> argument points to an illegal address.</td></tr><tr><td>EINVAL</td><td>The directory to be removed is the current directory, or the final component of <i>path</i> is “..”.</td></tr><tr><td>EIO</td><td>An I/O error occurred while accessing the file system.</td></tr><tr><td>ELOOP</td><td>Too many symbolic links were encountered in translating <i>path</i>.</td></tr><tr><td>ENAMETOOLONG</td><td>The length of the <i>path</i> argument exceeds <code>PATH_MAX</code>, or the length of a <i>path</i> component exceeds <code>NAME_MAX</code> while <code>_POSIX_NO_TRUNC</code> is in effect.</td></tr></table>		EACCES	Search permission is denied for a component of the path prefix; write permission is denied on the directory containing the directory to be removed; the parent directory has the <code>S_ISVTX</code> variable set and is not owned by the user; the directory is not owned by the user and is not writable by the user; or the user is not a super-user.	EBUSY	The directory to be removed is the mount point for a mounted file system.	EEXIST	The directory contains entries other than those for “.” and “..”.	EFAULT	The <i>path</i> argument points to an illegal address.	EINVAL	The directory to be removed is the current directory, or the final component of <i>path</i> is “..”.	EIO	An I/O error occurred while accessing the file system.	ELOOP	Too many symbolic links were encountered in translating <i>path</i> .	ENAMETOOLONG	The length of the <i>path</i> argument exceeds <code>PATH_MAX</code> , or the length of a <i>path</i> component exceeds <code>NAME_MAX</code> while <code>_POSIX_NO_TRUNC</code> is in effect.
EACCES	Search permission is denied for a component of the path prefix; write permission is denied on the directory containing the directory to be removed; the parent directory has the <code>S_ISVTX</code> variable set and is not owned by the user; the directory is not owned by the user and is not writable by the user; or the user is not a super-user.																	
EBUSY	The directory to be removed is the mount point for a mounted file system.																	
EEXIST	The directory contains entries other than those for “.” and “..”.																	
EFAULT	The <i>path</i> argument points to an illegal address.																	
EINVAL	The directory to be removed is the current directory, or the final component of <i>path</i> is “..”.																	
EIO	An I/O error occurred while accessing the file system.																	
ELOOP	Too many symbolic links were encountered in translating <i>path</i> .																	
ENAMETOOLONG	The length of the <i>path</i> argument exceeds <code>PATH_MAX</code> , or the length of a <i>path</i> component exceeds <code>NAME_MAX</code> while <code>_POSIX_NO_TRUNC</code> is in effect.																	

rmkdir(2)

ENOENT	The named directory does not exist or is the null pathname.
ENOLINK	The <i>path</i> argument points to a remote machine, and the connection to that machine is no longer active.
ENOTDIR	A component of the path prefix is not a directory.
EROFS	The directory entry to be removed is part of a read-only file system.

ATTRIBUTES See `attributes(5)` for descriptions of the following attributes:

ATTRIBUTE TYPE	ATTRIBUTE VALUE
MT-Level	Async-Signal-Safe

SEE ALSO `mkdir(1)`, `rm(1)`, `mkdir(2)`, `attributes(5)`

NAME	semctl – semaphore control operations														
SYNOPSIS	<pre>#include <sys/types.h> #include <sys/ipc.h> #include <sys/sem.h> int semctl(int <i>semid</i>, int <i>semnum</i>, int <i>cmd</i>, ...);</pre>														
DESCRIPTION	The <code>semctl()</code> function provides a variety of semaphore control operations as specified by <i>cmd</i>. The fourth argument is optional, depending upon the operation requested. If required, it is of type union <code>semun</code>, which must be explicitly declared by the application program. <pre>union semun { int val; struct semid_ds *buf; ushort_t *array; } arg ;</pre> The permission required for a semaphore operation is given as <i>{token}</i>, where <i>token</i> is the type of permission needed. The types of permission are interpreted as follows: <pre>00400 READ by user 00200 ALTER by user 00040 READ by group 00020 ALTER by group 00004 READ by others 00002 ALTER by others</pre> See the Semaphore Operation Permissions subsection of the DEFINITIONS section of <code>intro(2)</code> for more information. The following semaphore operations as specified by <i>cmd</i> are executed with respect to the semaphore specified by <i>semid</i> and <i>semnum</i>. <table> <tr> <td>GETVAL</td><td>Return the value of <code>semval</code> (see <code>intro(2)</code>). {READ}</td></tr> <tr> <td>SETVAL</td><td>Set the value of <code>semval</code> to <i>arg.val</i>. {ALTER} When this command is successfully executed, the <code>semadj</code> value corresponding to the specified semaphore in all processes is cleared.</td></tr> <tr> <td>GETPID</td><td>Return the value of <code>(int) sempid</code>. {READ}</td></tr> <tr> <td>GETNCNT</td><td>Return the value of <code>semncnt</code>. {READ}</td></tr> <tr> <td>GETZCNT</td><td>Return the value of <code>semzcnt</code>. {READ}</td></tr> </table> The following operations return and set, respectively, every <code>semval</code> in the set of semaphores. <table> <tr> <td>GETALL</td><td>Place <code>semvals</code> into array pointed to by <i>arg.array</i>. {READ}</td></tr> <tr> <td>SETALL</td><td>Set <code>semvals</code> according to the array pointed to by <i>arg.array</i>. {ALTER}. When this <i>cmd</i> is successfully executed, the <code>semadj</code> values corresponding to each specified semaphore in all processes are cleared.</td></tr> </table>	GETVAL	Return the value of <code>semval</code> (see <code>intro(2)</code>). {READ}	SETVAL	Set the value of <code>semval</code> to <i>arg.val</i> . {ALTER} When this command is successfully executed, the <code>semadj</code> value corresponding to the specified semaphore in all processes is cleared.	GETPID	Return the value of <code>(int) sempid</code> . {READ}	GETNCNT	Return the value of <code>semncnt</code> . {READ}	GETZCNT	Return the value of <code>semzcnt</code> . {READ}	GETALL	Place <code>semvals</code> into array pointed to by <i>arg.array</i> . {READ}	SETALL	Set <code>semvals</code> according to the array pointed to by <i>arg.array</i> . {ALTER}. When this <i>cmd</i> is successfully executed, the <code>semadj</code> values corresponding to each specified semaphore in all processes are cleared.
GETVAL	Return the value of <code>semval</code> (see <code>intro(2)</code>). {READ}														
SETVAL	Set the value of <code>semval</code> to <i>arg.val</i> . {ALTER} When this command is successfully executed, the <code>semadj</code> value corresponding to the specified semaphore in all processes is cleared.														
GETPID	Return the value of <code>(int) sempid</code> . {READ}														
GETNCNT	Return the value of <code>semncnt</code> . {READ}														
GETZCNT	Return the value of <code>semzcnt</code> . {READ}														
GETALL	Place <code>semvals</code> into array pointed to by <i>arg.array</i> . {READ}														
SETALL	Set <code>semvals</code> according to the array pointed to by <i>arg.array</i> . {ALTER}. When this <i>cmd</i> is successfully executed, the <code>semadj</code> values corresponding to each specified semaphore in all processes are cleared.														

semctl(2)

The following operations are also available.

IPC_STAT Place the current value of each member of the data structure associated with *semid* into the structure pointed to by *arg.buf*. The contents of this structure are defined in *intro(2)*. {READ}

IPC_SET Set the value of the following members of the data structure associated with *semid* to the corresponding value found in the structure pointed to by *arg.buf*:

```
sem_perm.uid
sem_perm.gid
sem_perm.mode    /* access permission bits only */
```

This command can be executed only by a process that has an effective user ID equal to either that of super-user, or to the value of *sem_perm.cuid* or *sem_perm.uid* in the data structure associated with *semid*.

IPC_RMID Remove the semaphore identifier specified by *semid* from the system and destroy the set of semaphores and data structure associated with it. This command can only be executed by a process that has an effective user ID equal to either that of super-user, or to the value of *sem_perm.cuid* or *sem_perm.uid* in the data structure associated with *semid*.

RETURN VALUES Upon successful completion, the value returned depends on *cmd* as follows:

GETVAL the value of *semval*

GETPID the value of (int) *sempid*

GETNCNT the value of *semncnt*

GETZCNT the value of *semzcnt*

All other successful completions return 0; otherwise, -1 is returned and *errno* is set to indicate the error.

ERRORS The *semctl()* function will fail if:

EACCES Operation permission is denied to the calling process (see *intro(2)*).

EINVAL The *semid* argument is not a valid semaphore identifier; the *semnum* argument is less than 0 or greater than *sem_nsems* - 1; or the *cmd* argument is not a valid command or is *IPC_SET* and *sem_perm.uid* or *sem_perm.gid* is not valid.

EPERM The *cmd* argument is equal to *IPC_RMID* or *IPC_SET* and the effective user of the calling process is not super-user, or *cmd* is equal to the value of *sem_perm.cuid* or *sem_perm.uid* in the data structure associated with *semid*.

EOVERFLOW	The <i>cmd</i> argument is <code>IPC_STAT</code> and <i>uid</i> or <i>gid</i> is too large to be stored in the structure pointed to by <i>arg.buf</i> .
ERANGE	The <i>cmd</i> argument is <code>SETVAL</code> or <code>SETALL</code> and the value to which <i>semval</i> is to be set is greater than the system imposed maximum.

SEE ALSO `ipcs(1)`, `intro(2)`, `semget(2)`, `semop(2)`

semget(2)

NAME	semget – get set of semaphores										
SYNOPSIS	<pre>#include <sys/types.h> #include <sys/ipc.h> #include <sys/sem.h> int semget(key_t key, int nsems, int semflg);</pre>										
DESCRIPTION	The <code>semget()</code> function returns the semaphore identifier associated with <i>key</i>. A semaphore identifier and associated data structure and set containing <i>nsems</i> semaphores (see <code>intro(3)</code>) are created for <i>key</i> if one of the following is true: ■ <i>key</i> is equal to <code>IPC_PRIVATE</code>. ■ <i>key</i> does not already have a semaphore identifier associated with it, and (<i>semflg</i> & <code>IPC_CREAT</code>) is true. On creation, the data structure associated with the new semaphore identifier is initialized as follows: ■ <code>sem_perm.cuid</code>, <code>sem_perm.uid</code>, <code>sem_perm.cgid</code>, and <code>sem_perm.gid</code> are set equal to the effective user ID and effective group ID, respectively, of the calling process. ■ The access permission bits of <code>sem_perm.mode</code> are set equal to the access permission bits of <i>semflg</i>. ■ <code>sem_nsems</code> is set equal to the value of <i>nsems</i>. ■ <code>sem_otime</code> is set equal to 0 and <code>sem_ctime</code> is set equal to the current time.										
RETURN VALUES	Upon successful completion, a non-negative integer representing a semaphore identifier is returned. Otherwise, -1 is returned and <code>errno</code> is set to indicate the error.										
ERRORS	The <code>semget()</code> function will fail if: <table><tr><td>EACCES</td><td>A semaphore identifier exists for <i>key</i>, but operation permission (see <code>intro(3)</code>) as specified by the low-order 9 bits of <i>semflg</i> would not be granted.</td></tr><tr><td>EEXIST</td><td>A semaphore identifier exists for <i>key</i> but both (<i>semflg</i> & <code>IPC_CREAT</code>) and (<i>semflg</i> & <code>IPC_EXCL</code>) are both true.</td></tr><tr><td>EINVAL</td><td>The <i>nsems</i> argument is either less than or equal to 0 or greater than the system-imposed limit; or a semaphore identifier exists for <i>key</i>, but the number of semaphores in the set associated with it is less than <i>nsems</i> and <i>nsems</i> is not equal to 0.</td></tr><tr><td>ENOENT</td><td>A semaphore identifier does not exist for <i>key</i> and (<i>semflg</i> & <code>IPC_CREAT</code>) is false.</td></tr><tr><td>ENOSPC</td><td>A semaphore identifier is to be created but the system-imposed limit on the maximum number of allowed semaphores or semaphore identifiers system-wide would be exceeded.</td></tr></table>	EACCES	A semaphore identifier exists for <i>key</i> , but operation permission (see <code>intro(3)</code>) as specified by the low-order 9 bits of <i>semflg</i> would not be granted.	EEXIST	A semaphore identifier exists for <i>key</i> but both (<i>semflg</i> & <code>IPC_CREAT</code>) and (<i>semflg</i> & <code>IPC_EXCL</code>) are both true.	EINVAL	The <i>nsems</i> argument is either less than or equal to 0 or greater than the system-imposed limit; or a semaphore identifier exists for <i>key</i> , but the number of semaphores in the set associated with it is less than <i>nsems</i> and <i>nsems</i> is not equal to 0.	ENOENT	A semaphore identifier does not exist for <i>key</i> and (<i>semflg</i> & <code>IPC_CREAT</code>) is false.	ENOSPC	A semaphore identifier is to be created but the system-imposed limit on the maximum number of allowed semaphores or semaphore identifiers system-wide would be exceeded.
EACCES	A semaphore identifier exists for <i>key</i> , but operation permission (see <code>intro(3)</code>) as specified by the low-order 9 bits of <i>semflg</i> would not be granted.										
EEXIST	A semaphore identifier exists for <i>key</i> but both (<i>semflg</i> & <code>IPC_CREAT</code>) and (<i>semflg</i> & <code>IPC_EXCL</code>) are both true.										
EINVAL	The <i>nsems</i> argument is either less than or equal to 0 or greater than the system-imposed limit; or a semaphore identifier exists for <i>key</i> , but the number of semaphores in the set associated with it is less than <i>nsems</i> and <i>nsems</i> is not equal to 0.										
ENOENT	A semaphore identifier does not exist for <i>key</i> and (<i>semflg</i> & <code>IPC_CREAT</code>) is false.										
ENOSPC	A semaphore identifier is to be created but the system-imposed limit on the maximum number of allowed semaphores or semaphore identifiers system-wide would be exceeded.										

semget(2)

SEE ALSO ipcrm(1), ipcs(1), intro(3), semctl(2), semop(2), ftok(3C)

semids(2)

NAME	semids – discover all semaphore identifiers
SYNOPSIS	<pre>#include <sys/sem.h> int semids(int *buf, uint_t nids, uint_t *pnids);</pre>
DESCRIPTION	The <code>semids()</code> function copies all active semaphore identifiers from the system into the user-defined buffer specified by <i>buf</i>, provided that the number of such identifiers is not greater than the number of integers the buffer can contain, as specified by <i>nids</i>. If the size of the buffer is insufficient to contain all of the active semaphore identifiers in the system, none are copied. Whether or not the size of the buffer is sufficient to contain all of them, the number of active semaphore identifiers in the system is copied into the unsigned integer pointed to by <i>pnids</i>. If <i>nids</i> is 0 or less than the number of active semaphore identifiers in the system, <i>buf</i> is ignored.
RETURN VALUES	Upon successful completion, <code>semids()</code> returns 0. Otherwise, -1 is returned and <code>errno</code> is set to indicate the error.
ERRORS	The <code>semids()</code> function will fail if: EFAULT The <i>buf</i> or <i>pnids</i> argument points to an illegal address.
USAGE	The <code>semids()</code> function returns a snapshot of all the active semaphore identifiers in the system. More may be added and some may be removed before they can be used by the caller.
EXAMPLES	EXAMPLE 1 <code>semids()</code> example This is sample C code indicating how to use the <code>semids()</code> function. <pre>void examine_semids() { int *ids = NULL; uint_t nids = 0; uint_t n; int i; for (;;) { if (semids(ids, nids, &n) != 0) { perror("semids"); exit(1); } if (n <= nids) /* we got them all */ break; /* we need a bigger buffer */ ids = realloc(ids, (nids = n) * sizeof (int)); } for (i = 0; i < n; i++)</pre>

EXAMPLE 1 semids() example *(Continued)*

```
        process_semids(ids[i]);  
    free(ids);  
}
```

ATTRIBUTES See attributes(5) for descriptions of the following attributes:

ATTRIBUTE TYPE	ATTRIBUTE VALUE
MT-Level	Async-Signal-Safe

SEE ALSO ipcrm(1), ipcs(1), intro(2), semctl(2), semget(2), semop(2), attributes(5)

semop(2)

NAME	semop, semtimedop – semaphore operations
SYNOPSIS	<pre>#include <sys/types.h> #include <sys/ipc.h> #include <sys/sem.h> int semop(int <i>semid</i>, struct sembuf <i>*sops</i>, size_t <i>nsops</i>); int semtimedop(int <i>semid</i>, struct sembuf <i>*sops</i>, size_t <i>nsops</i>, const struct timespec <i>*timeout</i>);</pre>
DESCRIPTION	The <code>semop()</code> function is used to perform atomically an array of semaphore operations on the set of semaphores associated with the semaphore identifier specified by <i>semid</i>. The <i>sops</i> argument is a pointer to the array of semaphore-operation structures. The <i>nsops</i> argument is the number of such structures in the array. Each <code>sembuf</code> structure contains the following members: <pre>short sem_num; /* semaphore number */ short sem_op; /* semaphore operation */ short sem_flg; /* operation flags */</pre> Each semaphore operation specified by <code>sem_op</code> is performed on the corresponding semaphore specified by <i>semid</i> and <code>sem_num</code>. The permission required for a semaphore operation is given as <i>{token}</i>, where <i>token</i> is the type of permission needed. The types of permission are interpreted as follows: <pre>00400 READ by user 00200 ALTER by user 00040 READ by group 00020 ALTER by group 00004 READ by others 00002 ALTER by others</pre> See the <i>Semaphore Operation Permissions</i> section of <code>intro(3)</code> for more information. The <code>sem_op</code> member specifies one of three semaphore operations: 1. The <code>sem_op</code> member is a negative integer; {ALTER} ■ If <code>semval</code> (see <code>intro(3)</code>) is greater than or equal to the absolute value of <code>sem_op</code>, the absolute value of <code>sem_op</code> is subtracted from <code>semval</code>. Also, if (<code>sem_flg&SEM_UNDO</code>) is true, the absolute value of <code>sem_op</code> is added to the calling process's <code>semadj</code> value (see <code>exit(2)</code>) for the specified semaphore. ■ If <code>semval</code> is less than the absolute value of <code>sem_op</code> and (<code>sem_flg&IPC_NOWAIT</code>) is true, <code>semop()</code> returns immediately. ■ If <code>semval</code> is less than the absolute value of <code>sem_op</code> and (<code>sem_flg&IPC_NOWAIT</code>) is false, <code>semop()</code> increments the <code>semncnt</code> associated with the specified semaphore and suspends execution of the calling process until one of the following conditions occur:

semop(2)

- The value of `semval` becomes greater than or equal to the absolute value of `sem_op`. When this occurs, the value of `semncnt` associated with the specified semaphore is decremented, the absolute value of `sem_op` is subtracted from `semval` and, if (`sem_flg&SEM_UNDO`) is true, the absolute value of `sem_op` is added to the calling process's `semadj` value for the specified semaphore.
 - The `semid` for which the calling process is awaiting action is removed from the system (see `semctl(2)`). When this occurs, `errno` is set to `EIDRM` and `-1` is returned.
 - The calling process receives a signal that is to be caught. When this occurs, the value of `semncnt` associated with the specified semaphore is decremented, and the calling process resumes execution in the manner prescribed in `signal(3C)`.
2. The `sem_op` member is a positive integer; {ALTER}
- The value of `sem_op` is added to `semval` and, if (`sem_flg&SEM_UNDO`) is true, the value of `sem_op` is subtracted from the calling process's `semadj` value for the specified semaphore.
3. The `sem_op` member is 0; {READ}
- If `semval` is 0, `semop()` returns immediately.
 - If `semval` is not equal to 0 and (`sem_flg&IPC_NOWAIT`) is true, `semop()` returns immediately.
 - If `semval` is not equal to 0 and (`sem_flg&IPC_NOWAIT`) is false, `semop()` increments the `semzcnt` associated with the specified semaphore and suspends execution of the calling process until one of the following occurs:
 - The value of `semval` becomes 0, at which time the value of `semzcnt` associated with the specified semaphore is decremented.
 - The `semid` for which the calling process is awaiting action is removed from the system. When this occurs, `errno` is set to `EIDRM` and `-1` is returned.
 - The calling process receives a signal that is to be caught. When this occurs, the value of `semzcnt` associated with the specified semaphore is decremented, and the calling process resumes execution in the manner prescribed in `signal(3C)`.

Upon successful completion, the value of `sempid` for each semaphore specified in the array pointed to by `sops` is set to the process ID of the calling process.

The `semtimedop()` function behaves as `semop()` except when it must suspend execution of the calling process to complete its operation. If `semtimedop()` must suspend the calling process after the time interval specified in `timeout` expires, or if the timeout expires while the process is suspended, `semtimedop()` returns with an error. If the `timespec` structure pointed to by `timeout` is zero-valued and `semtimedop()` needs to suspend the calling process to complete the requested operation(s), it returns immediately with an error. If `timeout` is the `NULL` pointer, the behavior of `semtimedop()` is identical to that of `semop()`.

semop(2)

RETURN VALUES	Upon successful completion, 0 is returned. Otherwise, -1 is returned and <code>errno</code> is set to indicate the error.																										
ERRORS	The <code>semop()</code> and <code>semtimedop()</code> functions will fail if: <table><tr><td>E2BIG</td><td>The <i>nsops</i> argument is greater than the system-imposed maximum.</td></tr><tr><td>EACCES</td><td>Operation permission is denied to the calling process (see <code>intro(3)</code>).</td></tr><tr><td>EAGAIN</td><td>The operation would result in suspension of the calling process but (<i>sem_flg</i> & <code>IPC_NOWAIT</code>) is true.</td></tr><tr><td>EFAULT</td><td>The <i>sops</i> argument points to an illegal address.</td></tr><tr><td>EFBIG</td><td>The value of <i>sem_num</i> is less than 0 or greater than or equal to the number of semaphores in the set associated with <i>semid</i>.</td></tr><tr><td>EIDRM</td><td>A <i>semid</i> was removed from the system.</td></tr><tr><td>EINTR</td><td>A signal was received.</td></tr><tr><td>EINVAL</td><td>The <i>semid</i> argument is not a valid semaphore identifier, or the number of individual semaphores for which the calling process requests a <code>SEM_UNDO</code> would exceed the limit.</td></tr><tr><td>ENOSPC</td><td>The limit on the number of individual processes requesting an <code>SEM_UNDO</code> would be exceeded.</td></tr><tr><td>ERANGE</td><td>An operation would cause a <i>semval</i> or a <i>semadj</i> value to overflow the system-imposed limit.</td></tr></table> The <code>semtimedop()</code> function will fail if: <table><tr><td>EAGAIN</td><td>The timeout expired before the requested operation could be completed.</td></tr></table> The <code>semtimedop()</code> function will fail if one of the following is detected: <table><tr><td>EFAULT</td><td>The <i>timeout</i> argument points to an illegal address.</td></tr><tr><td>EINVAL</td><td>The <i>timeout</i> argument specified a <i>tv_sec</i> or <i>tv_nsec</i> value less than 0, or a <i>tv_nsec</i> value greater than or equal to 1000 million.</td></tr></table>	E2BIG	The <i>nsops</i> argument is greater than the system-imposed maximum.	EACCES	Operation permission is denied to the calling process (see <code>intro(3)</code>).	EAGAIN	The operation would result in suspension of the calling process but (<i>sem_flg</i> & <code>IPC_NOWAIT</code>) is true.	EFAULT	The <i>sops</i> argument points to an illegal address.	EFBIG	The value of <i>sem_num</i> is less than 0 or greater than or equal to the number of semaphores in the set associated with <i>semid</i> .	EIDRM	A <i>semid</i> was removed from the system.	EINTR	A signal was received.	EINVAL	The <i>semid</i> argument is not a valid semaphore identifier, or the number of individual semaphores for which the calling process requests a <code>SEM_UNDO</code> would exceed the limit.	ENOSPC	The limit on the number of individual processes requesting an <code>SEM_UNDO</code> would be exceeded.	ERANGE	An operation would cause a <i>semval</i> or a <i>semadj</i> value to overflow the system-imposed limit.	EAGAIN	The timeout expired before the requested operation could be completed.	EFAULT	The <i>timeout</i> argument points to an illegal address.	EINVAL	The <i>timeout</i> argument specified a <i>tv_sec</i> or <i>tv_nsec</i> value less than 0, or a <i>tv_nsec</i> value greater than or equal to 1000 million.
E2BIG	The <i>nsops</i> argument is greater than the system-imposed maximum.																										
EACCES	Operation permission is denied to the calling process (see <code>intro(3)</code>).																										
EAGAIN	The operation would result in suspension of the calling process but (<i>sem_flg</i> & <code>IPC_NOWAIT</code>) is true.																										
EFAULT	The <i>sops</i> argument points to an illegal address.																										
EFBIG	The value of <i>sem_num</i> is less than 0 or greater than or equal to the number of semaphores in the set associated with <i>semid</i> .																										
EIDRM	A <i>semid</i> was removed from the system.																										
EINTR	A signal was received.																										
EINVAL	The <i>semid</i> argument is not a valid semaphore identifier, or the number of individual semaphores for which the calling process requests a <code>SEM_UNDO</code> would exceed the limit.																										
ENOSPC	The limit on the number of individual processes requesting an <code>SEM_UNDO</code> would be exceeded.																										
ERANGE	An operation would cause a <i>semval</i> or a <i>semadj</i> value to overflow the system-imposed limit.																										
EAGAIN	The timeout expired before the requested operation could be completed.																										
EFAULT	The <i>timeout</i> argument points to an illegal address.																										
EINVAL	The <i>timeout</i> argument specified a <i>tv_sec</i> or <i>tv_nsec</i> value less than 0, or a <i>tv_nsec</i> value greater than or equal to 1000 million.																										
SEE ALSO	<code>ipcs(1)</code> , <code>intro(3)</code> , <code>exec(2)</code> , <code>exit(2)</code> , <code>fork(2)</code> , <code>semctl(2)</code> , <code>semget(2)</code>																										

NAME	setpgid – set process group ID												
SYNOPSIS	<pre>#include <sys/types.h> #include <unistd.h> int setpgid(pid_t <i>pid</i>, pid_t <i>pgid</i>);</pre>												
DESCRIPTION	The <code>setpgid()</code> function sets the process group ID of the process with ID <i>pid</i> to <i>pgid</i>. If <i>pgid</i> is equal to <i>pid</i>, the process becomes a process group leader. See <code>intro(2)</code> for more information on session leaders and process group leaders. If <i>pgid</i> is not equal to <i>pid</i>, the process becomes a member of an existing process group. If <i>pid</i> is equal to 0, the process ID of the calling process is used. If <i>pgid</i> is equal to 0, the process specified by <i>pid</i> becomes a process group leader.												
RETURN VALUES	Upon successful completion, 0 is returned. Otherwise, -1 is returned and <code>errno</code> is set to indicate the error.												
ERRORS	The <code>setpgid()</code> function will fail if: <table> <tr> <td>EACCES</td><td>The <i>pid</i> argument matches the process ID of a child process of the calling process and the child process has successfully executed one of the <i>exec</i> family of functions (see <code>exec(2)</code>).</td></tr> <tr> <td>EINVAL</td><td>The <i>pgid</i> argument is less than (pid_t) 0 or greater than or equal to <code>PID_MAX</code>, or the calling process has a controlling terminal that does not support job control.</td></tr> <tr> <td>EPERM</td><td>The process indicated by the <i>pid</i> argument is a session leader.</td></tr> <tr> <td>EPERM</td><td>The <i>pid</i> argument matches the process ID of a child process of the calling process and the child process is not in the same session as the calling process.</td></tr> <tr> <td>EPERM</td><td>The <i>pgid</i> argument does not match the process ID of the process indicated by the <i>pid</i> argument, and there is no process with a process group ID that matches <i>pgid</i> in the same session as the calling process.</td></tr> <tr> <td>ESRCH</td><td>The <i>pid</i> argument does not match the process ID of the calling process or of a child process of the calling process.</td></tr> </table>	EACCES	The <i>pid</i> argument matches the process ID of a child process of the calling process and the child process has successfully executed one of the <i>exec</i> family of functions (see <code>exec(2)</code>).	EINVAL	The <i>pgid</i> argument is less than (pid_t) 0 or greater than or equal to <code>PID_MAX</code> , or the calling process has a controlling terminal that does not support job control.	EPERM	The process indicated by the <i>pid</i> argument is a session leader.	EPERM	The <i>pid</i> argument matches the process ID of a child process of the calling process and the child process is not in the same session as the calling process.	EPERM	The <i>pgid</i> argument does not match the process ID of the process indicated by the <i>pid</i> argument, and there is no process with a process group ID that matches <i>pgid</i> in the same session as the calling process.	ESRCH	The <i>pid</i> argument does not match the process ID of the calling process or of a child process of the calling process.
EACCES	The <i>pid</i> argument matches the process ID of a child process of the calling process and the child process has successfully executed one of the <i>exec</i> family of functions (see <code>exec(2)</code>).												
EINVAL	The <i>pgid</i> argument is less than (pid_t) 0 or greater than or equal to <code>PID_MAX</code> , or the calling process has a controlling terminal that does not support job control.												
EPERM	The process indicated by the <i>pid</i> argument is a session leader.												
EPERM	The <i>pid</i> argument matches the process ID of a child process of the calling process and the child process is not in the same session as the calling process.												
EPERM	The <i>pgid</i> argument does not match the process ID of the process indicated by the <i>pid</i> argument, and there is no process with a process group ID that matches <i>pgid</i> in the same session as the calling process.												
ESRCH	The <i>pid</i> argument does not match the process ID of the calling process or of a child process of the calling process.												
ATTRIBUTES	See <code>attributes(5)</code> for descriptions of the following attributes: <table> <tr> <th>ATTRIBUTE TYPE</th><th>ATTRIBUTE VALUE</th></tr> <tr> <td>MT-Level</td><td>Async-Signal-Safe</td></tr> </table>	ATTRIBUTE TYPE	ATTRIBUTE VALUE	MT-Level	Async-Signal-Safe								
ATTRIBUTE TYPE	ATTRIBUTE VALUE												
MT-Level	Async-Signal-Safe												
SEE ALSO	<code>intro(2)</code> , <code>exec(2)</code> , <code>exit(2)</code> , <code>fork(2)</code> , <code>getpid(2)</code> , <code>getsid(2)</code> , <code>attributes(5)</code>												

setpgrp(2)

NAME	setpgrp – set process group ID
SYNOPSIS	<pre>#include <sys/types.h> #include <unistd.h> pid_t setpgrp(void);</pre>
DESCRIPTION	If the calling process is not already a session leader, the <code>setpgrp()</code> function makes it one by setting its process group ID and session ID to the value of its process ID, and releases its controlling terminal. See <code>intro(3)</code> for more information on process group IDs and session leaders.
RETURN VALUES	The <code>setpgrp()</code> function returns the value of the new process group ID.
SEE ALSO	<code>intro(3)</code> , <code>exec(2)</code> , <code>fork(2)</code> , <code>getpid(2)</code> , <code>getsid(2)</code> , <code>kill(2)</code> , <code>signal(3C)</code>
NOTES	The <code>setpgrp()</code> function will be phased out in favor of the <code>setsid(2)</code> function.

NAME	setregid – set real and effective group IDs				
SYNOPSIS	<pre>#include <unistd.h> int setregid(gid_t rgid, gid_t egid);</pre>				
DESCRIPTION	The <code>setregid()</code> function is used to set the real and effective group IDs of the calling process. If <i>rgid</i> is <code>-1</code>, the real group ID is not changed; if <i>egid</i> is <code>-1</code>, the effective group ID is not changed. The real and effective group IDs may be set to different values in the same call. If the effective user ID of the calling process is super-user, the real group ID and the effective group ID can be set to any legal value. If the effective user ID of the calling process is not super-user, either the real group ID can be set to the saved set-group-ID from <code>execve(2)</code>, or the effective group ID can either be set to the saved set-group-ID or the real group ID. In either case, if the real group ID is being changed (that is, if <i>rgid</i> is not <code>-1</code>), or the effective group ID is being changed to a value not equal to the real group ID, the saved set-group-ID is set equal to the new effective group ID.				
RETURN VALUES	Upon successful completion, 0 is returned. Otherwise, <code>-1</code> is returned, <code>errno</code> is set to indicate the error, and neither of the group IDs will be changed.				
ERRORS	The <code>setregid()</code> function will fail if: <table> <tr> <td><code>EINVAL</code></td><td>The value of <i>rgid</i> or <i>egid</i> is less than 0 or greater than <code>UID_MAX</code> (defined in <code><limits.h></code>).</td></tr> <tr> <td><code>EPERM</code></td><td>The calling process's effective UID is not the super-user and a change other than changing the real group ID to the saved set-group-ID or changing the effective group ID to the real group ID or the saved group ID, was specified.</td></tr> </table>	<code>EINVAL</code>	The value of <i>rgid</i> or <i>egid</i> is less than 0 or greater than <code>UID_MAX</code> (defined in <code><limits.h></code>).	<code>EPERM</code>	The calling process's effective UID is not the super-user and a change other than changing the real group ID to the saved set-group-ID or changing the effective group ID to the real group ID or the saved group ID, was specified.
<code>EINVAL</code>	The value of <i>rgid</i> or <i>egid</i> is less than 0 or greater than <code>UID_MAX</code> (defined in <code><limits.h></code>).				
<code>EPERM</code>	The calling process's effective UID is not the super-user and a change other than changing the real group ID to the saved set-group-ID or changing the effective group ID to the real group ID or the saved group ID, was specified.				
USAGE	If a set-group-ID process sets its effective group ID to its real group ID, it can still set its effective group ID back to the saved set-group-ID.				
SEE ALSO	<code>execve(2)</code> , <code>getgid(2)</code> , <code>setreuid(2)</code> , <code>setuid(2)</code>				

setreuid(2)

NAME	setreuid – set real and effective user IDs				
SYNOPSIS	<pre>#include <unistd.h> int setreuid(uid_t <i>ruid</i>, uid_t <i>euid</i>);</pre>				
DESCRIPTION	The <code>setreuid()</code> function is used to set the real and effective user IDs of the calling process. If <i>ruid</i> is <code>-1</code>, the real user ID is not changed; if <i>euid</i> is <code>-1</code>, the effective user ID is not changed. The real and effective user IDs may be set to different values in the same call. If the effective user ID of the calling process is super-user, the real user ID and the effective user ID can be set to any legal value. If the effective user ID of the calling process is not super-user, either the real user ID can be set to the effective user ID, or the effective user ID can either be set to the saved set-user ID from <code>execve()</code> (see <code>exec(2)</code>) or the real user ID. In either case, if the real user ID is being changed (that is, if <i>ruid</i> is not <code>-1</code>), or the effective user ID is being changed to a value not equal to the real user ID, the saved set-user ID is set equal to the new effective user ID.				
RETURN VALUES	Upon successful completion, <code>0</code> is returned. Otherwise, <code>-1</code> is returned, <code>errno</code> is set to indicate the error, and neither of the user IDs will be changed.				
ERRORS	The <code>setreuid()</code> function will fail if: <table><tr><td>EINVAL</td><td>The value of <i>ruid</i> or <i>euid</i> is less than <code>0</code> or greater than <code>UID_MAX</code> (defined in <code><limits.h></code>).</td></tr><tr><td>EPERM</td><td>The calling process's effective user ID is not the super-user and a change other than changing the real user ID to the effective user ID, or changing the effective user ID to the real user ID or the saved set-user ID, was specified.</td></tr></table>	EINVAL	The value of <i>ruid</i> or <i>euid</i> is less than <code>0</code> or greater than <code>UID_MAX</code> (defined in <code><limits.h></code>).	EPERM	The calling process's effective user ID is not the super-user and a change other than changing the real user ID to the effective user ID, or changing the effective user ID to the real user ID or the saved set-user ID, was specified.
EINVAL	The value of <i>ruid</i> or <i>euid</i> is less than <code>0</code> or greater than <code>UID_MAX</code> (defined in <code><limits.h></code>).				
EPERM	The calling process's effective user ID is not the super-user and a change other than changing the real user ID to the effective user ID, or changing the effective user ID to the real user ID or the saved set-user ID, was specified.				
USAGE	If a set-user-ID process sets its effective user ID to its real user ID, it can still set its effective user ID back to the saved set-user ID.				
SEE ALSO	<code>exec(2)</code> , <code>getuid(2)</code> , <code>setregid(2)</code> , <code>setuid(2)</code>				

NAME	setsid – create session and set process group ID				
SYNOPSIS	<pre>#include <sys/types.h> #include <unistd.h> pid_t setsid(void);</pre>				
DESCRIPTION	The setsid() function creates a new session, if the calling process is not a process group leader. Upon return the calling process will be the session leader of this new session, will be the process group leader of a new process group, and will have no controlling terminal. The process group ID of the calling process will be set equal to the process ID of the calling process. The calling process will be the only process in the new process group and the only process in the new session.				
RETURN VALUES	Upon successful completion, setsid() returns the value of the process group ID of the calling process. Otherwise it returns (pid_t)-1 and sets errno to indicate the error.				
ERRORS	The setsid() function will fail if: EPERM The calling process is already a process group leader, or the process group ID of a process other than the calling process matches the process ID of the calling process.				
ATTRIBUTES	See attributes(5) for descriptions of the following attributes: <table border="1"> <thead> <tr> <th>ATTRIBUTE TYPE</th><th>ATTRIBUTE VALUE</th></tr> </thead> <tbody> <tr> <td>MT-Level</td><td>Async-Signal-Safe</td></tr> </tbody> </table>	ATTRIBUTE TYPE	ATTRIBUTE VALUE	MT-Level	Async-Signal-Safe
ATTRIBUTE TYPE	ATTRIBUTE VALUE				
MT-Level	Async-Signal-Safe				
SEE ALSO	getsid(2), setpgid(2), setpgrp(2), attributes(5)				
WARNINGS	A call to setsid() by a process that is a process group leader will fail. A process can become a process group leader by being the last member of a pipeline started by a job control shell. Thus, a process that expects to be part of a pipeline, and that calls setsid(), should always first fork; the parent should exit and the child should call setsid(). This will ensure that the calling process will work reliably when started by both job control shells and non-job control shells.				

settaskid(2)

NAME	settaskid, gettaskid, getprojid – set or get task or project IDs						
SYNOPSIS	<pre>#include <sys/types.h> #include <sys/task.h> #include <unistd.h> taskid_t settaskid(projid_t <i>project</i>, int <i>flags</i>) ; taskid_t gettaskid(void) ; projid_t getprojid(void) ;</pre>						
DESCRIPTION	The <code>settaskid()</code> function makes a request of the system to assign a new task ID to the calling process, changing the associated project ID to that specified. The calling process must have superuser privileges to perform this operation. The <i>flags</i> argument should be either <code>TASK_NORMAL</code> for a regular task, or <code>TASK_FINAL</code>, which disallows subsequent <code>settaskid()</code> calls by the created task. The <code>gettaskid()</code> function returns the task ID of the calling process. The <code>getprojid()</code> function returns the project ID of the calling process.						
RETURN VALUES	Upon successful completion, these functions return the appropriate task or project ID. Otherwise, <code>-1</code> is returned and <code>errno</code> is set to indicate the error.						
ERRORS	The <code>settaskid()</code> function will fail if: <table><tr><td>EACCES</td><td>The invoking task was created with the <code>TASK_FINAL</code> flag.</td></tr><tr><td>EPERM</td><td>The effective user of the calling process is not superuser.</td></tr><tr><td>EINVAL</td><td>The given project ID is not within the valid project ID range.</td></tr></table>	EACCES	The invoking task was created with the <code>TASK_FINAL</code> flag.	EPERM	The effective user of the calling process is not superuser.	EINVAL	The given project ID is not within the valid project ID range.
EACCES	The invoking task was created with the <code>TASK_FINAL</code> flag.						
EPERM	The effective user of the calling process is not superuser.						
EINVAL	The given project ID is not within the valid project ID range.						
ATTRIBUTES	See <code>attributes(5)</code> for descriptions of the following attributes: <table><thead><tr><th>ATTRIBUTE TYPE</th><th>ATTRIBUTE VALUE</th></tr></thead><tbody><tr><td>MT-Level</td><td>Async-Signal-Safe</td></tr></tbody></table>	ATTRIBUTE TYPE	ATTRIBUTE VALUE	MT-Level	Async-Signal-Safe		
ATTRIBUTE TYPE	ATTRIBUTE VALUE						
MT-Level	Async-Signal-Safe						
SEE ALSO	<code>setsid(2)</code> , <code>project(4)</code> , <code>attributes(5)</code>						

NAME	setuid, setegid, seteuid, setgid – set user and group IDs
SYNOPSIS	<pre>#include <sys/types.h> #include <unistd.h> int setuid(uid_t uid) ; int setegid(gid_t egid) ; int seteuid(uid_t euid) ; int setgid(gid_t gid) ;</pre>
DESCRIPTION	The <code>setuid()</code> function sets the real user ID, effective user ID, and saved user ID of the calling process. The <code>setgid()</code> function sets the real group ID, effective group ID, and saved group ID of the calling process. The <code>setegid()</code> and <code>seteuid()</code> functions set the effective group and user IDs respectively for the calling process. See <code>intro(2)</code> for more information on real, effective, and saved user and group IDs. At login time, the real user ID, effective user ID, and saved user ID of the login process are set to the login ID of the user responsible for the creation of the process. The same is true for the real, effective, and saved group IDs; they are set to the group ID of the user responsible for the creation of the process. When a process calls one of the <code>exec</code> family of functions (see <code>exec(2)</code>) to execute a file (program), the user and/or group identifiers associated with the process can change. If the file executed is a set-user-ID file, the effective and saved user IDs of the process are set to the owner of the file executed. If the file executed is a set-group-ID file, the effective and saved group IDs of the process are set to the group of the file executed. If the file executed is not a set-user-ID or set-group-ID file, the effective user ID, saved user ID, effective group ID, and saved group ID are not changed. If the effective user ID of the process calling <code>setuid()</code> is the super-user, the real, effective, and saved user IDs are set to the <i>uid</i> argument. If the effective user ID of the calling process is not the super-user, but <i>uid</i> is either the real user ID or the saved user ID of the calling process, the effective user ID is set to <i>uid</i>. If the effective user ID of the process calling <code>setgid()</code> is the super-user, the real, effective, and saved group IDs are set to the <i>gid</i> argument. If the effective user ID of the calling process is not the superuser, but <i>gid</i> is either the real group ID or the saved group ID of the calling process, the effective group ID is set to <i>gid</i>.
RETURN VALUES	Upon successful completion, 0 is returned. Otherwise, -1 is returned and <code>errno</code> is set to indicate the error.
ERRORS	The <code>setuid()</code> and <code>setgid()</code> functions will fail if: <code>EINVAL</code> The value of <i>uid</i> or <i>gid</i> is out of range.

setuid(2)

EPERM

For `setuid()` and `seteuid()` the effective user of the calling process is not super-user, and the *uid* argument does not match either the real or saved user IDs. For `setgid()` and `setegid()` the effective user of the calling process is not the superuser, and the *gid* argument does not match either the real or saved group IDs.

ATTRIBUTES

See `attributes(5)` for descriptions of the following attributes:

ATTRIBUTE TYPE	ATTRIBUTE VALUE
MT-Level	<code>setuid()</code> and <code>setgid()</code> and Async-Signal-Safe

SEE ALSO

`intro(2)`, `exec(2)`, `getgroups(2)`, `getuid(2)`, `stat(3HEAD)` `attributes(5)`

NAME	shmctl – shared memory control operations										
SYNOPSIS	<pre>#include <sys/types.h> #include <sys/ipc.h> #include <sys/shm.h> int shmctl(int shmid, int cmd, struct shmid_ds *buf);</pre>										
DESCRIPTION	The <code>shmctl()</code> function provides a variety of shared memory control operations as specified by <i>cmd</i>. The permission required for a shared memory control operation is given as <i>{token}</i>, where <i>token</i> is the type of permission needed. The types of permission are interpreted as follows: <pre>00400 READ by user 00200 WRITE by user 00040 READ by group 00020 WRITE by group 00004 READ by others 00002 WRITE by others</pre> See the <i>Shared Memory Operation Permissions</i> section of <code>intro(2)</code> for more information. The following operations require the specified tokens: <table> <tr> <td>IPC_STAT</td><td>Place the current value of each member of the data structure associated with <i>shmid</i> into the structure pointed to by <i>buf</i>. The contents of this structure are defined in <code>intro(2)</code>. {READ}</td></tr> <tr> <td>IPC_SET</td><td>Set the value of the following members of the data structure associated with <i>shmid</i> to the corresponding value found in the structure pointed to by <i>buf</i>: <pre>shm_perm.uid shm_perm.gid shm_perm.mode /* access permission bits only */</pre> This command can be executed only by a process that has an effective user ID equal to that of super-user, or to the value of <code>shm_perm.cuid</code> or <code>shm_perm.uid</code> in the data structure associated with <i>shmid</i>.</td></tr> <tr> <td>IPC_RMID</td><td>Remove the shared memory identifier specified by <i>shmid</i> from the system and destroy the shared memory segment and data structure associated with it. This command can be executed only by a process that has an effective user ID equal to that of super-user, or to the value of <code>shm_perm.cuid</code> or <code>shm_perm.uid</code> in the data structure associated with <i>shmid</i>.</td></tr> <tr> <td>SHM_LOCK</td><td>Lock the shared memory segment specified by <i>shmid</i> in memory. This command can be executed only by a process that has an effective user ID equal to super-user.</td></tr> <tr> <td>SHM_UNLOCK</td><td>Unlock the shared memory segment specified by <i>shmid</i>. This command can be executed only by a process that has an effective user ID equal to super-user.</td></tr> </table>	IPC_STAT	Place the current value of each member of the data structure associated with <i>shmid</i> into the structure pointed to by <i>buf</i> . The contents of this structure are defined in <code>intro(2)</code> . {READ}	IPC_SET	Set the value of the following members of the data structure associated with <i>shmid</i> to the corresponding value found in the structure pointed to by <i>buf</i> : <pre>shm_perm.uid shm_perm.gid shm_perm.mode /* access permission bits only */</pre> This command can be executed only by a process that has an effective user ID equal to that of super-user, or to the value of <code>shm_perm.cuid</code> or <code>shm_perm.uid</code> in the data structure associated with <i>shmid</i> .	IPC_RMID	Remove the shared memory identifier specified by <i>shmid</i> from the system and destroy the shared memory segment and data structure associated with it. This command can be executed only by a process that has an effective user ID equal to that of super-user, or to the value of <code>shm_perm.cuid</code> or <code>shm_perm.uid</code> in the data structure associated with <i>shmid</i> .	SHM_LOCK	Lock the shared memory segment specified by <i>shmid</i> in memory. This command can be executed only by a process that has an effective user ID equal to super-user.	SHM_UNLOCK	Unlock the shared memory segment specified by <i>shmid</i> . This command can be executed only by a process that has an effective user ID equal to super-user.
IPC_STAT	Place the current value of each member of the data structure associated with <i>shmid</i> into the structure pointed to by <i>buf</i> . The contents of this structure are defined in <code>intro(2)</code> . {READ}										
IPC_SET	Set the value of the following members of the data structure associated with <i>shmid</i> to the corresponding value found in the structure pointed to by <i>buf</i> : <pre>shm_perm.uid shm_perm.gid shm_perm.mode /* access permission bits only */</pre> This command can be executed only by a process that has an effective user ID equal to that of super-user, or to the value of <code>shm_perm.cuid</code> or <code>shm_perm.uid</code> in the data structure associated with <i>shmid</i> .										
IPC_RMID	Remove the shared memory identifier specified by <i>shmid</i> from the system and destroy the shared memory segment and data structure associated with it. This command can be executed only by a process that has an effective user ID equal to that of super-user, or to the value of <code>shm_perm.cuid</code> or <code>shm_perm.uid</code> in the data structure associated with <i>shmid</i> .										
SHM_LOCK	Lock the shared memory segment specified by <i>shmid</i> in memory. This command can be executed only by a process that has an effective user ID equal to super-user.										
SHM_UNLOCK	Unlock the shared memory segment specified by <i>shmid</i> . This command can be executed only by a process that has an effective user ID equal to super-user.										

shmctl(2)

	Shared memory segments must be explicitly removed after the last reference to them has been removed.														
RETURN VALUES	Upon successful completion, 0 is returned. Otherwise, -1 is returned and <code>errno</code> is set to indicate the error.														
ERRORS	The <code>shmctl()</code> function will fail if: <table><tr><td>EACCES</td><td>The <i>cmd</i> argument is equal to <code>IPC_STAT</code> and {<code>READ</code>} operation permission is denied to the calling process.</td></tr><tr><td>EFAULT</td><td>The <i>buf</i> argument points to an illegal address.</td></tr><tr><td>EINVAL</td><td>The <i>shmid</i> argument is not a valid shared memory identifier; or the <i>cmd</i> argument is not a valid command or is <code>IPC_SET</code> and <code>shm_perm.uid</code> or <code>shm_perm.gid</code> is not valid.</td></tr><tr><td>ENOMEM</td><td>The <i>cmd</i> argument is equal to <code>SHM_LOCK</code> and there is not enough memory.</td></tr><tr><td>EOVERFLOW</td><td>The <i>cmd</i> argument is <code>IPC_STAT</code> and <i>uid</i> or <i>gid</i> is too large to be stored in the structure pointed to by <i>buf</i>.</td></tr><tr><td>EPERM</td><td>The <i>cmd</i> argument is equal to <code>IPC_RMID</code> or <code>IPC_SET</code> and the effective user ID of the calling process is not super-user and it is not equal to the value of <code>shm_perm.cuid</code> or <code>shm_perm.uid</code> in the data structure associated with <i>shmid</i>.</td></tr><tr><td>EPERM</td><td>The <i>cmd</i> argument is equal to <code>SHM_LOCK</code> or <code>SHM_UNLOCK</code> and the effective user ID of the calling process is not equal to that of super-user.</td></tr></table>	EACCES	The <i>cmd</i> argument is equal to <code>IPC_STAT</code> and { <code>READ</code> } operation permission is denied to the calling process.	EFAULT	The <i>buf</i> argument points to an illegal address.	EINVAL	The <i>shmid</i> argument is not a valid shared memory identifier; or the <i>cmd</i> argument is not a valid command or is <code>IPC_SET</code> and <code>shm_perm.uid</code> or <code>shm_perm.gid</code> is not valid.	ENOMEM	The <i>cmd</i> argument is equal to <code>SHM_LOCK</code> and there is not enough memory.	EOVERFLOW	The <i>cmd</i> argument is <code>IPC_STAT</code> and <i>uid</i> or <i>gid</i> is too large to be stored in the structure pointed to by <i>buf</i> .	EPERM	The <i>cmd</i> argument is equal to <code>IPC_RMID</code> or <code>IPC_SET</code> and the effective user ID of the calling process is not super-user and it is not equal to the value of <code>shm_perm.cuid</code> or <code>shm_perm.uid</code> in the data structure associated with <i>shmid</i> .	EPERM	The <i>cmd</i> argument is equal to <code>SHM_LOCK</code> or <code>SHM_UNLOCK</code> and the effective user ID of the calling process is not equal to that of super-user.
EACCES	The <i>cmd</i> argument is equal to <code>IPC_STAT</code> and { <code>READ</code> } operation permission is denied to the calling process.														
EFAULT	The <i>buf</i> argument points to an illegal address.														
EINVAL	The <i>shmid</i> argument is not a valid shared memory identifier; or the <i>cmd</i> argument is not a valid command or is <code>IPC_SET</code> and <code>shm_perm.uid</code> or <code>shm_perm.gid</code> is not valid.														
ENOMEM	The <i>cmd</i> argument is equal to <code>SHM_LOCK</code> and there is not enough memory.														
EOVERFLOW	The <i>cmd</i> argument is <code>IPC_STAT</code> and <i>uid</i> or <i>gid</i> is too large to be stored in the structure pointed to by <i>buf</i> .														
EPERM	The <i>cmd</i> argument is equal to <code>IPC_RMID</code> or <code>IPC_SET</code> and the effective user ID of the calling process is not super-user and it is not equal to the value of <code>shm_perm.cuid</code> or <code>shm_perm.uid</code> in the data structure associated with <i>shmid</i> .														
EPERM	The <i>cmd</i> argument is equal to <code>SHM_LOCK</code> or <code>SHM_UNLOCK</code> and the effective user ID of the calling process is not equal to that of super-user.														
SEE ALSO	<code>ipcs(1)</code> , <code>intro(2)</code> , <code>shmget(2)</code> , <code>shmop(2)</code>														

NAME	shmget – get shared memory segment identifier								
SYNOPSIS	<pre>#include <sys/types.h> #include <sys/ipc.h> #include <sys/shm.h> int shmget(key_t <i>key</i>, size_t <i>size</i>, int <i>shmflg</i>);</pre>								
DESCRIPTION	The <code>shmget()</code> function returns the shared memory identifier associated with <i>key</i>. A shared memory identifier and associated data structure and shared memory segment of at least <i>size</i> bytes (see <code>intro(3)</code>) are created for <i>key</i> if one of the following are true: ■ The <i>key</i> argument is equal to <code>IPC_PRIVATE</code>. ■ The <i>key</i> argument does not already have a shared memory identifier associated with it, and $(shmflg \& IPC_CREAT)$ is true. Upon creation, the data structure associated with the new shared memory identifier is initialized as follows: ■ The values of <code>shm_perm.cuid</code>, <code>shm_perm.uid</code>, <code>shm_perm.cgid</code>, and <code>shm_perm.gid</code> are set equal to the effective user ID and effective group ID, respectively, of the calling process. ■ The access permission bits of <code>shm_perm.mode</code> are set equal to the access permission bits of <i>shmflg</i>. <code>shm_segsz</code> is set equal to the value of <i>size</i>. ■ The values of <code>shm_lpid</code>, <code>shm_nattch</code>, <code>shm_atime</code>, and <code>shm_dtime</code> are set equal to 0. ■ The <code>shm_ctime</code> is set equal to the current time. Shared memory segments must be explicitly removed after the last reference to them has been removed.								
RETURN VALUES	Upon successful completion, a non-negative integer representing a shared memory identifier is returned. Otherwise, <code>-1</code> is returned and <code>errno</code> is set to indicate the error.								
ERRORS	The <code>shmget()</code> function will fail if: <table> <tr> <td>EACCES</td><td>A shared memory identifier exists for <i>key</i> but operation permission (see <code>intro(3)</code>) as specified by the low-order 9 bits of <i>shmflg</i> would not be granted.</td></tr> <tr> <td>EEXIST</td><td>A shared memory identifier exists for <i>key</i> but both $(shmflg \& IPC_CREATE)$ and $(shmflg \& IPC_EXCL)$ are true.</td></tr> <tr> <td>EINVAL</td><td>The <i>size</i> argument is less than the system-imposed minimum or greater than the system-imposed maximum.</td></tr> <tr> <td>EINVAL</td><td>A shared memory identifier exists for <i>key</i> but the size of the segment associated with it is less than <i>size</i> and <i>size</i> is not equal to 0.</td></tr> </table>	EACCES	A shared memory identifier exists for <i>key</i> but operation permission (see <code>intro(3)</code>) as specified by the low-order 9 bits of <i>shmflg</i> would not be granted.	EEXIST	A shared memory identifier exists for <i>key</i> but both $(shmflg \& IPC_CREATE)$ and $(shmflg \& IPC_EXCL)$ are true.	EINVAL	The <i>size</i> argument is less than the system-imposed minimum or greater than the system-imposed maximum.	EINVAL	A shared memory identifier exists for <i>key</i> but the size of the segment associated with it is less than <i>size</i> and <i>size</i> is not equal to 0.
EACCES	A shared memory identifier exists for <i>key</i> but operation permission (see <code>intro(3)</code>) as specified by the low-order 9 bits of <i>shmflg</i> would not be granted.								
EEXIST	A shared memory identifier exists for <i>key</i> but both $(shmflg \& IPC_CREATE)$ and $(shmflg \& IPC_EXCL)$ are true.								
EINVAL	The <i>size</i> argument is less than the system-imposed minimum or greater than the system-imposed maximum.								
EINVAL	A shared memory identifier exists for <i>key</i> but the size of the segment associated with it is less than <i>size</i> and <i>size</i> is not equal to 0.								

shmget(2)

ENOENT	A shared memory identifier does not exist for <i>key</i> and (<i>shmflg</i> &IPC_CREATE) is false.
ENOMEM	A shared memory identifier and associated shared memory segment are to be created but the amount of available memory is not sufficient to fill the request.
ENOSPC	A shared memory identifier is to be created but the system-imposed limit on the maximum number of allowed shared memory identifiers system-wide would be exceeded.

SEE ALSO intro(3), shmctl(2), shmop(2), ftok(3C)

NAME	shmids – discover all shared memory identifiers
SYNOPSIS	<pre>#include <sys/shm.h> int shmids(int *buf, uint_t nids, uint_t *pnids);</pre>
DESCRIPTION	The <code>shmids()</code> function copies all active shared memory identifiers from the system into the user-defined buffer specified by <i>buf</i>, provided that the number of such identifiers is not greater than the number of integers the buffer can contain, as specified by <i>nids</i>. If the size of the buffer is insufficient to contain all of the active shared memory identifiers in the system, none are copied. Whether or not the size of the buffer is sufficient to contain all of them, the number of active shared memory identifiers in the system is copied into the unsigned integer pointed to by <i>pnids</i>. If <i>nids</i> is 0 or less than the number of active shared memory identifiers in the system, <i>buf</i> is ignored.
RETURN VALUES	Upon successful completion, <code>shmids()</code> returns 0. Otherwise, -1 is returned and <code>errno</code> is set to indicate the error.
ERRORS	The <code>shmids()</code> function will fail if: EFAULT The <i>buf</i> or <i>pnids</i> argument points to an illegal address.
USAGE	The <code>shmids()</code> function returns a snapshot of all the active shared memory identifiers in the system. More may be added and some may be removed before they can be used by the caller.
EXAMPLES	EXAMPLE 1 <code>shmids()</code> example This is sample C code indicating how to use the <code>shmids()</code> function. <pre>void examine_shmids() { int *ids = NULL; uint_t nids = 0; uint_t n; int i; for (;;) { if (shmids(ids, nids, &n) != 0) { perror("shmids"); exit(1); } if (n <= nids) /* we got them all */ break; /* we need a bigger buffer */ ids = realloc(ids, (nids = n) * sizeof (int)); } for (i = 0; i < n; i++)</pre>

shmids(2)

EXAMPLE 1 shmids() example *(Continued)*

```
        process_shmid(ids[i]);  
    free(ids);  
}
```

ATTRIBUTES See attributes(5) for descriptions of the following attributes:

ATTRIBUTE TYPE	ATTRIBUTE VALUE
MT-Level	Async-Signal-Safe

SEE ALSO ipcrm(1), ipcs(1), intro(2), shmctl(2), shmget(2), shmop(2), attributes(5)

NAME	shmop, shmat, shmdt – shared memory operations
SYNOPSIS	<pre>#include <sys/types.h> #include <sys/shm.h> void *shmat(int shmid, const void *shmaddr, int shmflg); int shmdt(char *shmaddr);</pre>
Default	int shmdt(const void *shmaddr);
Standard-conforming	
DESCRIPTION	The <code>shmat()</code> function attaches the shared memory segment associated with the shared memory identifier specified by <code>shmid</code> to the data segment of the calling process. The permission required for a shared memory control operation is given as <code>{token}</code>, where <code>token</code> is the type of permission needed. The types of permission are interpreted as follows: <pre>00400 READ by user 00200 WRITE by user 00040 READ by group 00020 WRITE by group 00004 READ by others 00002 WRITE by others</pre> See the <i>Shared Memory Operation Permissions</i> section of <code>intro(2)</code> for more information. When <code>(shmflg & SHM_SHARE_MMU)</code> is true, virtual memory resources in addition to shared memory itself are shared among processes that use the same shared memory. When <code>(shmflg & SHM_PAGEABLE)</code> is true, virtual memory resources are shared and the dynamic shared memory (DISM) framework is created. The dynamic shared memory can be resized dynamically within the specified size in <code>shmget(2)</code>. The DISM shared memory is pageable unless it is locked. The shared memory segment is attached to the data segment of the calling process at the address specified based on one of the following criteria: ■ If <code>shmaddr</code> is equal to <code>(void *) 0</code>, the segment is attached to the first available address as selected by the system. ■ If <code>shmaddr</code> is equal to <code>(void *) 0</code> and <code>(shmflg & SHM_SHARE_MMU)</code> or <code>(shmflg & SHM_PAGEABLE)</code> is true, then the segment is attached to the first available suitably aligned address. When <code>(shmflg & SHM_SHARE_MMU)</code> or <code>(shmflg & SHM_PAGEABLE)</code> is set, however, the permission given by <code>shmget()</code> determines whether the segment is attached for reading or reading and writing. ■ If <code>shmaddr</code> is not equal to <code>(void *) 0</code> and <code>(shmflg & SHM_RND)</code> is true, the segment is attached to the address given by <code>(shmaddr - (shmaddr modulus SHMLBA))</code>. ■ If <code>shmaddr</code> is not equal to <code>(void *) 0</code> and <code>(shmflg & SHM_RND)</code> is false, the segment is attached to the address given by <code>shmaddr</code>. ■ The segment is attached for reading if <code>(shmflg & SHM_RDONLY)</code> is true {READ}, otherwise it is attached for reading and writing {READ/WRITE}.

shmop(2)

	The <code>shmdt()</code> function detaches from the calling process's data segment the shared memory segment located at the address specified by <i>shmaddr</i>. If the application is standard-conforming (see <code>standards(5)</code>), the <i>shmaddr</i> argument is of type <code>const void *</code>. Otherwise it is of type <code>char *</code>. Shared memory segments must be explicitly removed after the last reference to them has been removed.																		
RETURN VALUES	Upon successful completion, <code>shmat()</code> returns the data segment start address of the attached shared memory segment; <code>shmdt()</code> returns 0. Otherwise, -1 is returned, the shared memory segment is not attached, and <code>errno</code> is set to indicate the error.																		
ERRORS	The <code>shmat()</code> function will fail if: <table><tr><td>EACCES</td><td>Operation permission is denied to the calling process (see <code>intro(2)</code>).</td></tr><tr><td>EINVAL</td><td>The <i>shmid</i> argument is not a valid shared memory identifier.</td></tr><tr><td>EINVAL</td><td>The <i>shmaddr</i> argument is not equal to 0, and the value of (<i>shmaddr</i> - (<i>shmaddr</i> modulus SHMLBA)) is an illegal address.</td></tr><tr><td>EINVAL</td><td>The <i>shmaddr</i> argument is not equal to 0, is an illegal address, and (<i>shmflg</i>&SHM_RND) is false.</td></tr><tr><td>EINVAL</td><td>The <i>shmaddr</i> argument is not equal to 0, is not properly aligned, and (<i>shmflg</i>&SHM_SHARE_MMU) is true.</td></tr><tr><td>EINVAL</td><td>SHM_SHARE_MMU is not supported in certain architectures.</td></tr><tr><td>EMFILE</td><td>The number of shared memory segments attached to the calling process would exceed the system-imposed limit.</td></tr><tr><td>ENOMEM</td><td>The available data space is not large enough to accommodate the shared memory segment.</td></tr></table> The <code>shmdt()</code> function will fail if: <table><tr><td>EINVAL</td><td>The <i>shmaddr</i> argument is not the data segment start address of a shared memory segment.</td></tr></table>	EACCES	Operation permission is denied to the calling process (see <code>intro(2)</code>).	EINVAL	The <i>shmid</i> argument is not a valid shared memory identifier.	EINVAL	The <i>shmaddr</i> argument is not equal to 0, and the value of (<i>shmaddr</i> - (<i>shmaddr</i> modulus SHMLBA)) is an illegal address.	EINVAL	The <i>shmaddr</i> argument is not equal to 0, is an illegal address, and (<i>shmflg</i> &SHM_RND) is false.	EINVAL	The <i>shmaddr</i> argument is not equal to 0, is not properly aligned, and (<i>shmflg</i> &SHM_SHARE_MMU) is true.	EINVAL	SHM_SHARE_MMU is not supported in certain architectures.	EMFILE	The number of shared memory segments attached to the calling process would exceed the system-imposed limit.	ENOMEM	The available data space is not large enough to accommodate the shared memory segment.	EINVAL	The <i>shmaddr</i> argument is not the data segment start address of a shared memory segment.
EACCES	Operation permission is denied to the calling process (see <code>intro(2)</code>).																		
EINVAL	The <i>shmid</i> argument is not a valid shared memory identifier.																		
EINVAL	The <i>shmaddr</i> argument is not equal to 0, and the value of (<i>shmaddr</i> - (<i>shmaddr</i> modulus SHMLBA)) is an illegal address.																		
EINVAL	The <i>shmaddr</i> argument is not equal to 0, is an illegal address, and (<i>shmflg</i> &SHM_RND) is false.																		
EINVAL	The <i>shmaddr</i> argument is not equal to 0, is not properly aligned, and (<i>shmflg</i> &SHM_SHARE_MMU) is true.																		
EINVAL	SHM_SHARE_MMU is not supported in certain architectures.																		
EMFILE	The number of shared memory segments attached to the calling process would exceed the system-imposed limit.																		
ENOMEM	The available data space is not large enough to accommodate the shared memory segment.																		
EINVAL	The <i>shmaddr</i> argument is not the data segment start address of a shared memory segment.																		
SEE ALSO	<code>intro(2)</code>, <code>exec(2)</code>, <code>exit(2)</code>, <code>fork(2)</code>, <code>shmctl(2)</code>, <code>shmget(2)</code>, <code>standards(5)</code>																		

NAME	sigaction – detailed signal management				
SYNOPSIS	<pre>#include <signal.h> int sigaction(int <i>sig</i>, const struct sigaction *<i>act</i>, struct sigaction *<i>oact</i>);</pre>				
DESCRIPTION	The <code>sigaction()</code> function allows the calling process to examine or specify the action to be taken on delivery of a specific signal. (See <code>signal(3HEAD)</code> for an explanation of general signal concepts.) The <i>sig</i> argument specifies the signal and can be assigned any of the signals specified in <code>signal(3HEAD)</code> except <code>SIGKILL</code> and <code>SIGSTOP</code>. In a multithreaded process, <i>sig</i> cannot be <code>SIGWAITING</code>, <code>SIGCANCEL</code>, or <code>SIGLWP</code>. If the argument <i>act</i> is not <code>NULL</code>, it points to a structure specifying the new action to be taken when delivering <i>sig</i>. If the argument <i>oact</i> is not <code>NULL</code>, it points to a structure where the action previously associated with <i>sig</i> is to be stored on return from <code>sigaction()</code>. The <code>sigaction</code> structure includes the following members: <pre>void (*sa_handler)(); void (*sa_sigaction)(int, siginfo_t *, void *); sigset_t sa_mask; int sa_flags;</pre> The <code>sa_handler</code> member identifies the action to be associated with the specified signal, if the <code>SA_SIGINFO</code> flag (see below) is cleared in the <code>sa_flags</code> field of the <code>sigaction</code> structure. It may take any of the values specified in <code>signal(3HEAD)</code> or that of a user specified signal handler. If the <code>SA_SIGINFO</code> flag is set in the <code>sa_flags</code> field, the <code>sa_sigaction</code> field specifies a signal-catching function. The <code>sa_mask</code> member specifies a set of signals to be blocked while the signal handler is active. On entry to the signal handler, that set of signals is added to the set of signals already being blocked when the signal is delivered. In addition, the signal that caused the handler to be executed will also be blocked, unless the <code>SA_NODEFER</code> flag has been specified. <code>SIGSTOP</code> and <code>SIGKILL</code> cannot be blocked (the system silently enforces this restriction). The <code>sa_flags</code> member specifies a set of flags used to modify the delivery of the signal. It is formed by a logical OR of any of the following values: <table> <tr> <td><code>SA_ONSTACK</code></td><td>If set and the signal is caught, and if the LWP that is chosen to process a delivered signal has an alternate signal stack declared with <code>sigaltstack(2)</code>, then it will process the signal on that stack. Otherwise, the signal is delivered on the LWP main stack. Unbound threads (see <code>thr_create(3THR)</code>) may not have alternate signal stacks.</td></tr> <tr> <td><code>SA_RESETHAND</code></td><td>If set and the signal is caught, the disposition of the signal is reset to <code>SIG_DFL</code> and the signal will not be blocked on entry to the</td></tr> </table>	<code>SA_ONSTACK</code>	If set and the signal is caught, and if the LWP that is chosen to process a delivered signal has an alternate signal stack declared with <code>sigaltstack(2)</code> , then it will process the signal on that stack. Otherwise, the signal is delivered on the LWP main stack. Unbound threads (see <code>thr_create(3THR)</code>) may not have alternate signal stacks.	<code>SA_RESETHAND</code>	If set and the signal is caught, the disposition of the signal is reset to <code>SIG_DFL</code> and the signal will not be blocked on entry to the
<code>SA_ONSTACK</code>	If set and the signal is caught, and if the LWP that is chosen to process a delivered signal has an alternate signal stack declared with <code>sigaltstack(2)</code> , then it will process the signal on that stack. Otherwise, the signal is delivered on the LWP main stack. Unbound threads (see <code>thr_create(3THR)</code>) may not have alternate signal stacks.				
<code>SA_RESETHAND</code>	If set and the signal is caught, the disposition of the signal is reset to <code>SIG_DFL</code> and the signal will not be blocked on entry to the				

sigaction(2)

	signal handler (SIGILL, SIGTRAP, and SIGPWR cannot be automatically reset when delivered; the system silently enforces this restriction).		
SA_NODEFER	If set and the signal is caught, the signal will not be automatically blocked by the kernel while it is being caught.		
SA_RESTART	If set and the signal is caught, functions that are interrupted by the execution of this signal's handler are transparently restarted by the system, namely <code>fcntl(2)</code> , <code>ioctl(2)</code> , <code>wait(2)</code> , <code>waitid(2)</code> , and the following functions on slow devices like terminals: <code>getmsg()</code> and <code>getpmsg()</code> (see <code>getmsg(2)</code>); <code>putmsg()</code> and <code>putpmsg()</code> (see <code>putmsg(2)</code>); <code>pread()</code> , <code>read()</code> , and <code>readv()</code> (see <code>read(2)</code>); <code>pwrite()</code> , <code>write()</code> , and <code>writv()</code> (see <code>write(2)</code>); <code>recv()</code> , <code>recvfrom()</code> , and <code>recvmsg()</code> (see <code>recv(3SOCKET)</code>); and <code>send()</code> , <code>sendto()</code> , and <code>sendmsg()</code> (see <code>send(3SOCKET)</code>). Otherwise, the function returns an <code>EINTR</code> error.		
SA_SIGINFO	If cleared and the signal is caught, <i>sig</i> is passed as the only argument to the signal-catching function. If set and the signal is caught, two additional arguments are passed to the signal-catching function. If the second argument is not equal to <code>NULL</code> , it points to a <code>siginfo_t</code> structure containing the reason why the signal was generated (see <code>siginfo(3HEAD)</code>); the third argument points to a <code>ucontext_t</code> structure containing the receiving process's context when the signal was delivered (see <code>ucontext(3HEAD)</code>).		
SA_NOCLDWAIT	If set and <i>sig</i> equals <code>SIGCHLD</code> , the system will not create zombie processes when children of the calling process exit. If the calling process subsequently issues a <code>wait(2)</code> , it blocks until all of the calling process's child processes terminate, and then returns <code>-1</code> with <code>errno</code> set to <code>ECHILD</code> .		
SA_NOCLDSTOP	If set and <i>sig</i> equals <code>SIGCHLD</code> , <code>SIGCHLD</code> will not be sent to the calling process when its child processes stop or continue.		
SA_WAITSIG	If set and <i>sig</i> equals <code>SIGWAITING</code> , enables generation of <code>SIGWAITING</code> signals. Reserved for use by the threads library.		
RETURN VALUES	Upon successful completion, 0 is returned. Otherwise, <code>-1</code> is returned, <code>errno</code> is set to indicate the error, and no new signal handler is installed.		
ERRORS	The <code>sigaction()</code> function will fail if: <table> <tr> <td>EINVAL</td><td>The value of the <i>sig</i> argument is not a valid signal number or is equal to <code>SIGKILL</code> or <code>SIGSTOP</code>. In addition, if in a multithreaded process, it is equal to <code>SIGWAITING</code>, <code>SIGCANCEL</code>, or <code>SIGLWP</code>.</td></tr> </table>	EINVAL	The value of the <i>sig</i> argument is not a valid signal number or is equal to <code>SIGKILL</code> or <code>SIGSTOP</code> . In addition, if in a multithreaded process, it is equal to <code>SIGWAITING</code> , <code>SIGCANCEL</code> , or <code>SIGLWP</code> .
EINVAL	The value of the <i>sig</i> argument is not a valid signal number or is equal to <code>SIGKILL</code> or <code>SIGSTOP</code> . In addition, if in a multithreaded process, it is equal to <code>SIGWAITING</code> , <code>SIGCANCEL</code> , or <code>SIGLWP</code> .		

ATTRIBUTES See attributes(5) for descriptions of the following attributes:

ATTRIBUTE TYPE	ATTRIBUTE VALUE
MT-Level	Async-Signal-Safe

SEE ALSO kill(1), intro(3), exit(2), fcntl(2), getmsg(2), ioctl(2), kill(2), pause(2), putmsg(2), read(2), sigaltstack(2), sigprocmask(2), sigsend(2), sigsuspend(2), wait(2), waitid(2), write(2), recv(3SOCKET), recv(3SOCKET), signal(3C), sigsetops(3C), thr_create(3THR), attributes(5), siginfo(3HEAD), signal(3HEAD), ucontext(3HEAD)

NOTES The handler routine can be declared:

```
void handler (int sig, siginfo_t *sip, ucontext_t *uap) ;
```

The *sig* argument is the signal number. The *sip* argument is a pointer (to space on the stack) to a `siginfo_t` structure, which provides additional detail about the delivery of the signal. The *uap* argument is a pointer (again to space on the stack) to a `ucontext_t` structure (defined in `<sys/ucontext.h>`) which contains the context from before the signal. It is not recommended that *uap* be used by the handler to restore the context from before the signal delivery.

sigaltstack(2)

NAME	sigaltstack – set or get signal alternate stack context
SYNOPSIS	<pre>#include <signal.h> int sigaltstack(const stack_t *ss, stack_t *oss);</pre>
DESCRIPTION	The <code>sigaltstack()</code> function allows an LWP to define and examine the state of an alternate stack area on which signals are processed. If <code>ss</code> is non-zero, it specifies a pointer to and the size of a stack area on which to deliver signals, and informs the system whether the LWP is currently executing on that stack. When a signal's action indicates its handler should execute on the alternate signal stack (specified with a <code>sigaction(2)</code> call), the system checks whether the LWP chosen to execute the signal handler is currently executing on that stack. If the LWP is not currently executing on the signal stack, the system arranges a switch to the alternate signal stack for the duration of the signal handler's execution. The <code>stack_t</code> structure includes the following members: <pre>int *ss_sp long ss_size int ss_flags</pre> If <code>ss</code> is not NULL, it points to a structure specifying the alternate signal stack that will take effect upon successful return from <code>sigaltstack()</code>. The <code>ss_sp</code> and <code>ss_size</code> members specify the new base and size of the stack, which is automatically adjusted for direction of growth and alignment. The <code>ss_flags</code> member specifies the new stack state and may be set to the following: <code>SS_DISABLE</code> The stack is to be disabled and <code>ss_sp</code> and <code>ss_size</code> are ignored. If <code>SS_DISABLE</code> is not set, the stack will be enabled. If <code>oss</code> is not NULL, it points to a structure specifying the alternate signal stack that was in effect prior to the call to <code>sigaltstack()</code>. The <code>ss_sp</code> and <code>ss_size</code> members specify the base and size of that stack. The <code>ss_flags</code> member specifies the stack's state, and may contain the following values: <code>SS_ONSTACK</code> The LWP is currently executing on the alternate signal stack. Attempts to modify the alternate signal stack while the LWP is executing on it will fail. <code>SS_DISABLE</code> The alternate signal stack is currently disabled.
RETURN VALUES	Upon successful completion, 0 is return. Otherwise, -1 is returned and <code>errno</code> is set to indicate the error.
ERRORS	The <code>sigaltstack()</code> function will fail if: <code>EFAULT</code> The <code>ss</code> or <code>oss</code> argument points to an illegal address. <code>EINVAL</code> The <code>ss</code> argument is not a null pointer, and the <code>ss_flags</code> member pointed to by <code>ss</code> contains flags other than <code>SS_DISABLE</code>. <code>ENOMEM</code> The size of the alternate stack area is less than <code>MINSIGSTKSZ</code>.

EPERM An attempt was made to modify an active stack.

SEE ALSO getcontext(2), sigaction(2), ucontext(3HEAD)

NOTES The value SIGSTKSZ is defined to be the number of bytes that would be used to cover the usual case when allocating an alternate stack area. The value MINSIGSTKSZ is defined to be the minimum stack size for a signal handler. In computing an alternate stack size, a program should add that amount to its stack requirements to allow for the operating system overhead.

The sigaltstack() function creates an alternate signal stack for the calling LWP, with the implication that in a multithreaded application, only bound threads (or in POSIX terminology, only threads whose scheduling contention scope is PTHREAD_SCOPE_SYSTEM, see thr_create(3THR) and pthread_create(3THR)) should call sigaltstack(). The behavior of an application that calls sigaltstack() from an unbound thread (that is, a POSIX thread whose scheduling contention scope is PTHREAD_SCOPE_PROCESS) is undefined. In a future release, sigaltstack() may return a new error value if it is called from an unbound thread.

The following code fragment is typically used to allocate an alternate stack:

```
if ((sigstk.ss_sp = (char *)malloc(SIGSTKSZ)) == NULL)
    /* error return */;

sigstk.ss_size = SIGSTKSZ;
sigstk.ss_flags = 0;
if (sigaltstack(&sigstk, (stack_t *)0) < 0)
    perror("sigaltstack");
```

`_signotifywait(2)`

NAME	<code>_signotifywait</code> , <code>_lwp_sigredirect</code> – deliver process signals to specific LWPs
SYNOPSIS	<pre>#include <sys/lwp.h> int _signotifywait(void); int _lwp_sigredirect(lwpid_t target_lwp, int signo, int *queued);</pre>
DESCRIPTION	In a multithreaded process, signals that are generated for a process are delivered to one of the threads that does not have that signal masked. If all of the application threads are masking that signal, its delivery waits until one of them unmask it. The disposition of the each thread's signal mask is unknown to the kernel when it generates signals for the process. The <code>_signotifywait()</code> and <code>_lwp_sigredirect()</code> functions provide a mechanism to direct instances of signals generated for the process to application-specified LWPs. Each process has a set of signals pending for the process, and for each LWP there is a set of signals pending for that LWP. If no signals are pending, these sets are empty. There is also a process-wide signal set, termed the <i>notification</i> set, manipulated by these functions. A signal generated for the process where the signal number is not in the notification set is called an <i>unnotified</i> signal. In a multithreaded program there is an <code>aslwp</code>, a special LWP endowed with powers to handle signals that are generated for a process. The <code>_signotifywait()</code> function is used to await signals generated for the process, and should be called only from the <code>aslwp</code>. In general, these functions are not to be called from the application-level. If there is a pending unnotified signal when <code>_signotifywait()</code> is called, that signal is selected and the call returns immediately. If there is not a signal pending, the call suspends the calling LWP until the generation of an unnotified signal; that signal then is selected and the function returns. In both cases, the selected signal number is set in the notification set and returned as the value of <code>_signotifywait()</code>. The signal remains pending for the process, and any associated <code>siginfo(3HEAD)</code> information remains queued at the process. The <code>_lwp_sigredirect()</code> function requests that a signal pending for the process be delivered to the LWP specified by <code>target_lwp</code>. If <code>target_lwp</code> is 0, the signal is discarded. It is an error if <code>signo</code> is not currently in the notification set of the process. The signal specified by <code>signo</code> is removed from pending for the process and is made pending for the <code>target_lwp</code>. If there is an associated <code>siginfo</code> information structure queued at the process, that <code>siginfo</code> is queued to the <code>target_lwp</code>. Whenever a signal is cleared from the set of signals pending for the process, the corresponding signal is cleared from the notification set. After a successful call to <code>_lwp_sigredirect()</code>, the signal <code>signo</code> is cleared from the notification set and from the set of signals pending for the process. If another instance of <code>signo</code> is queued for the process, the signal number is again set in the process pending mask, and if another

`_signotifywait(2)`

LWP is blocked in a call to `_signotifywait()`, its wait for an unnotified signal will be satisfied. The effects described in this paragraph also apply when the signal *signo* is returned by a call to `sigtimedwait()` and *signo* was not pending for the calling LWP.

When *queued* is non-NULL, and there is another instance of the signal (*signo*), a non-zero value will be place in *queued* to indicate that more than one instance of the signal is pending on the process.

RETURN VALUES The `_signotifywait()` function returns the signal number of the pending but hitherto unnotified signal. The `_lwp_sigredirect()` function returns 0 when successful. A non-zero value indicates an error.

ERRORS No error conditions are specified for `_signotifywait()`.

If the following conditions occurs, `_lwp_sigredirect()` fails and return the corresponding value:

`EINVAL` The signal *signo* was not pending for the process, or *signo* was not in the notification set.

`ESRCH` The *target_lwp* cannot be found in the current process.

SEE ALSO `_lwp_create(2)`, `_lwp_kill(2)`, `sigtimedwait(3RT)`, `siginfo(3HEAD)`, `signal(3HEAD)`

NOTES This mechanism for delivering signals to multithreaded processes is subject to change in future versions of Solaris. Any process with explicit knowledge of this mechanism may not be compatible from release to release.

sigpending(2)

NAME	sigpending – examine signals that are blocked and pending				
SYNOPSIS	<pre>#include <signal.h> int sigpending(sigset_t *set);</pre>				
DESCRIPTION	The <code>sigpending()</code> function retrieves those signals that have been sent to the calling process but are being blocked from delivery by the calling process's signal mask. The signals are stored in the space pointed to by the <i>set</i> argument.				
RETURN VALUES	Upon successful completion, 0 is returned. Otherwise, -1 is returned and <code>errno</code> is set to indicate the error.				
ERRORS	The <code>sigpending()</code> function will fail if: EFAULT The <i>set</i> argument points to an illegal address.				
ATTRIBUTES	See <code>attributes(5)</code> for descriptions of the following attributes: <table><thead><tr><th>ATTRIBUTE TYPE</th><th>ATTRIBUTE VALUE</th></tr></thead><tbody><tr><td>MT-Level</td><td>Async-Signal-Safe</td></tr></tbody></table>	ATTRIBUTE TYPE	ATTRIBUTE VALUE	MT-Level	Async-Signal-Safe
ATTRIBUTE TYPE	ATTRIBUTE VALUE				
MT-Level	Async-Signal-Safe				
SEE ALSO	<code>sigaction(2)</code> , <code>sigprocmask(2)</code> , <code>sigsetops(3C)</code> , <code>attributes(5)</code>				

NAME	sigprocmask – change and/or examine caller’s signal mask				
SYNOPSIS	<pre>#include <signal.h> int sigprocmask(int <i>how</i>, const sigset_t *<i>set</i>, sigset_t *<i>oset</i>);</pre>				
DESCRIPTION	The <code>sigprocmask()</code> function is used to examine and/or change the caller’s signal mask. If the value is <code>SIG_BLOCK</code>, the set pointed to by the <i>set</i> argument is added to the current signal mask. If the value is <code>SIG_UNBLOCK</code>, the set pointed to by the <i>set</i> argument is removed from the current signal mask. If the value is <code>SIG_SETMASK</code>, the current signal mask is replaced by the set pointed to by the <i>set</i> argument. If the <i>oset</i> argument is not <code>NULL</code>, the previous mask is stored in the space pointed to by <i>oset</i>. If the value of the <i>set</i> argument is <code>NULL</code>, the value <i>how</i> is not significant and the caller’s signal mask is unchanged; thus, the call can be used to inquire about currently blocked signals. If there are any pending unblocked signals after the call to <code>sigprocmask()</code>, at least one of those signals will be delivered before the call to <code>sigprocmask()</code> returns. It is not possible to block those signals that cannot be ignored this restriction is silently imposed by the system. See <code>sigaction(2)</code>. If <code>sigprocmask()</code> fails, the caller’s signal mask is not changed.				
RETURN VALUES	Upon successful completion, 0 is returned. Otherwise, -1 is returned and <code>errno</code> is set to indicate the error.				
ERRORS	The <code>sigprocmask()</code> function will fail if: <table> <tr> <td><code>EFAULT</code></td><td>The <i>set</i> or <i>oset</i> argument points to an illegal address.</td></tr> <tr> <td><code>EINVAL</code></td><td>The value of the <i>how</i> argument is not equal to one of the defined values.</td></tr> </table>	<code>EFAULT</code>	The <i>set</i> or <i>oset</i> argument points to an illegal address.	<code>EINVAL</code>	The value of the <i>how</i> argument is not equal to one of the defined values.
<code>EFAULT</code>	The <i>set</i> or <i>oset</i> argument points to an illegal address.				
<code>EINVAL</code>	The value of the <i>how</i> argument is not equal to one of the defined values.				
ATTRIBUTES	See <code>attributes(5)</code> for descriptions of the following attributes: <table> <tr> <th>ATTRIBUTE TYPE</th><th>ATTRIBUTE VALUE</th></tr> <tr> <td>MT-Level</td><td>Async-Signal-Safe</td></tr> </table>	ATTRIBUTE TYPE	ATTRIBUTE VALUE	MT-Level	Async-Signal-Safe
ATTRIBUTE TYPE	ATTRIBUTE VALUE				
MT-Level	Async-Signal-Safe				
SEE ALSO	<code>sigaction(2)</code> , <code>signal(3C)</code> , <code>sigsetops(3C)</code> , <code>thr_sigsetmask(3THR)</code> , <code>attributes(5)</code> , <code>signal(3HEAD)</code>				
NOTES	In a multithreaded program, the call to <code>sigpromask()</code> impacts only the calling thread’s signal mask and is therefore identical to a call to <code>thr_sigsetmask(3THR)</code>. Signals that are generated synchronously should not be masked. If such a signal is blocked and delivered, the receiving process is killed.				

sigsend(2)

NAME	sigsend, sigsendset – send a signal to a process or a group of processes
SYNOPSIS	<pre>#include <signal.h> int sigsend(idtype_t idtype, id_t id, int sig); int sigsendset(procset_t *psp, int sig);</pre>
DESCRIPTION	The <code>sigsend()</code> function sends a signal to the process or group of processes specified by <i>id</i> and <i>idtype</i>. The signal to be sent is specified by <i>sig</i> and is either 0 or one of the values listed in <code>signal(3HEAD)</code>. If <i>sig</i> is 0 (the null signal), error checking is performed but no signal is actually sent. This value can be used to check the validity of <i>id</i> and <i>idtype</i>. The real or effective user ID of the sending process must match the real or saved user ID of the receiving process, unless the effective user ID of the sending process is super-user, or <i>sig</i> is SIGCONT and the sending process has the same session ID as the receiving process. If <i>idtype</i> is P_PID, <i>sig</i> is sent to the process with process ID <i>id</i>. If <i>idtype</i> is P_PGID, <i>sig</i> is sent to all processes with process group ID <i>id</i>. If <i>idtype</i> is P_SID, <i>sig</i> is sent to all processes with session ID <i>id</i>. If <i>idtype</i> is P_TASKID, <i>sig</i> is sent to all processes with task ID <i>id</i>. If <i>idtype</i> is P_UID, <i>sig</i> is sent to any process with effective user ID <i>id</i>. If <i>idtype</i> is P_GID, <i>sig</i> is sent to any process with effective group ID <i>id</i>. If <i>idtype</i> is P_PROJID, <i>sig</i> is sent to any process with project ID <i>id</i>. If <i>idtype</i> is P_CID, <i>sig</i> is sent to any process with scheduler class ID <i>id</i> (see <code>prctl(2)</code>). If <i>idtype</i> is P_ALL, <i>sig</i> is sent to all processes and <i>id</i> is ignored. If <i>id</i> is P_MYID, the value of <i>id</i> is taken from the calling process. The process with a process ID of 0 is always excluded. The process with a process ID of 1 is excluded unless <i>idtype</i> is equal to P_PID. The <code>sigsendset()</code> function provides an alternate interface for sending signals to sets of processes. This function sends signals to the set of processes specified by <i>psp</i>. <i>psp</i> is a pointer to a structure of type <code>procset_t</code>, defined in <code><sys/procset.h></code>, which includes the following members: <pre>idop_t p_op; idtype_t p_lidtype; id_t p_lid; idtype_t p_ridtype; id_t p_rid;</pre>

sigsend(2)

The `p_lidtype` and `p_lid` members specify the ID type and ID of one (“left”) set of processes; the `p_ridtype` and `p_rid` members specify the ID type and ID of a second (“right”) set of processes. ID types and IDs are specified just as for the *idtype* and *id* arguments to `sigsend()`. The `p_op` member specifies the operation to be performed on the two sets of processes to get the set of processes the function is to apply to. The valid values for `p_op` and the processes they specify are:

<code>POP_DIFF</code>	Set difference: processes in left set and not in right set.
<code>POP_AND</code>	Set intersection: processes in both left and right sets.
<code>POP_OR</code>	Set union: processes in either left or right set or both.
<code>POP_XOR</code>	Set exclusive-or: processes in left or right set but not in both.

RETURN VALUES Upon successful completion, 0 is return. Otherwise, -1 is returned and `errno` is set to indicate the error.

ERRORS The `sigsend()` and `sigsendset()` functions will fail if:

<code>EINVAL</code>	The <i>sig</i> argument is not a valid signal number, or the <i>idtype</i> argument is not a valid idtype field.
<code>EINVAL</code>	The <i>sig</i> argument is <code>SIGKILL</code> , <i>idtype</i> is <code>P_PID</code> and <i>id</i> is 1 (<code>proc1</code>).
<code>EPERM</code>	The effective user of the calling process is not superuser and its real or effective user ID does not match the real or effective user ID of the receiving process, and the calling process is not sending <code>SIGCONT</code> to a process that shares the same session.
<code>ESRCH</code>	No process can be found corresponding to that specified by <i>id</i> and <i>idtype</i> .

The `sigsendset()` function will fail if:

<code>EFAULT</code>	The <i>psp</i> argument points to an illegal address.
---------------------	---

SEE ALSO `kill(1)`, `getpid(2)`, `kill(2)`, `prctl(2)`, `signal(3C)`, `signal(3HEAD)`

sigsuspend(2)

NAME	sigsuspend – install a signal mask and suspend caller until signal				
SYNOPSIS	<pre>#include <signal.h> int sigsuspend(const sigset_t *set);</pre>				
DESCRIPTION	The <code>sigsuspend()</code> function replaces the caller's signal mask with the set of signals pointed to by the <code>set</code> argument and suspends the caller until delivery of a signal whose action is either to execute a signal catching function or to terminate the process. If the action is to terminate the process, <code>sigsuspend()</code> does not return. If the action is to execute a signal catching function, <code>sigsuspend()</code> returns after the signal catching function returns. On return, the signal mask is restored to the set that existed before the call to <code>sigsuspend()</code>. See NOTES for the precise semantics of signal mask restoration in a multithreaded process. It is not possible to block those signals that cannot be ignored (see <code>signal(3HEAD)</code>); this restriction is silently imposed by the system.				
RETURN VALUES	Since <code>sigsuspend()</code> suspends process execution indefinitely, there is no successful completion return value. On failure, it returns <code>-1</code> and sets <code>errno</code> to indicate the error.				
ERRORS	The <code>sigsuspend()</code> function will fail if: <table><tr><td>EFAULT</td><td>The <code>set</code> argument points to an illegal address.</td></tr><tr><td>EINTR</td><td>A signal was caught by the calling process and control was returned from the signal catching function.</td></tr></table>	EFAULT	The <code>set</code> argument points to an illegal address.	EINTR	A signal was caught by the calling process and control was returned from the signal catching function.
EFAULT	The <code>set</code> argument points to an illegal address.				
EINTR	A signal was caught by the calling process and control was returned from the signal catching function.				
ATTRIBUTES	See <code>attributes(5)</code> for descriptions of the following attributes: <table><thead><tr><th>ATTRIBUTE TYPE</th><th>ATTRIBUTE VALUE</th></tr></thead><tbody><tr><td>MT-Level</td><td>Async-Signal-Safe</td></tr></tbody></table>	ATTRIBUTE TYPE	ATTRIBUTE VALUE	MT-Level	Async-Signal-Safe
ATTRIBUTE TYPE	ATTRIBUTE VALUE				
MT-Level	Async-Signal-Safe				
SEE ALSO	<code>sigaction(2)</code> , <code>sigprocmask(2)</code> , <code>sigwait(2)</code> , <code>signal(3C)</code> , <code>sigsetops(3C)</code> , <code>attributes(5)</code> , <code>signal(3HEAD)</code>				
NOTES	In a multithreaded application, the <code>sigwait(2)</code>, function should be used instead of <code>sigsuspend()</code>. Should <code>sigsuspend()</code> be used, however, its semantics of signal mask restoration are slightly different from those for a single-threaded process on return from the signal catching function, the signal mask is restored to the set that existed before the call to <code>sigsuspend()</code>. This action raises the following implications: ■ If a thread specifies two signals in the mask to <code>sigsuspend()</code>, both signals could interrupt its call to <code>sigsuspend()</code> simultaneously. In the traditional program that does not use threads, a call to <code>sigsuspend()</code> with two signals in the mask always returns with only one signal delivered. The other signal remains pending if masked earlier, unlike the MT case.				

`sigsuspend(2)`

- While a thread is executing the signal handler that interrupted its call to `sigsuspend()`, its signal mask is the one passed to `sigsuspend()`. It does not get restored to the previous mask until it returns from all the signal handlers that interrupted `sigsuspend()`.

sigwait(2)

NAME	sigwait – wait until a signal is posted				
SYNOPSIS					
Default	<pre>#include <signal.h> int sigwait(sigset_t *set);</pre>				
POSIX	<pre>cc [flag ...] file ... -D_POSIX_PTHREAD_SEMANTICS [library...] #include <signal.h> int sigwait(const sigset_t *set, int *sig);</pre>				
DESCRIPTION	The <code>sigwait()</code> function selects a signal in <i>set</i> that is pending on the calling thread (see <code>thr_create(3THR)</code>) or LWP. If no signal in <i>set</i> is pending, then <code>sigwait()</code> blocks until a signal in <i>set</i> becomes pending. The selected signal is cleared from the set of signals pending on the calling thread or LWP and the number of the signal is returned, or in the POSIX version (see <code>standards(5)</code>) placed in <i>sig</i>. The selection of a signal in <i>set</i> is independent of the signal mask of the calling thread or LWP. This means a thread or LWP can synchronously wait for signals that are being blocked by the signal mask of the calling thread or LWP. To ensure that only the caller receives the signals defined in <i>set</i>, all threads should have signals in <i>set</i> masked including the calling thread. If <code>sigwait()</code> is called on an ignored signal, then the occurrence of the signal will be ignored, unless <code>sigaction()</code> changes the disposition. If more than one thread or LWP waits for the same signal, only one is unblocked when the signal arrives.				
RETURN VALUES	Upon successful completion, the default version of <code>sigwait()</code> returns a signal number; the POSIX version returns 0 and stores the received signal number at the location pointed to by <i>sig</i>. Otherwise, -1 is returned and <code>errno</code> is set to indicate an error.				
ERRORS	The <code>sigwait()</code> function will fail if: <table> <tr> <td>EINVAL</td><td>The <i>set</i> argument contains an unsupported signal number.</td></tr> <tr> <td>EFAULT</td><td>The <i>set</i> argument points to an invalid address.</td></tr> </table>	EINVAL	The <i>set</i> argument contains an unsupported signal number.	EFAULT	The <i>set</i> argument points to an invalid address.
EINVAL	The <i>set</i> argument contains an unsupported signal number.				
EFAULT	The <i>set</i> argument points to an invalid address.				
EXAMPLES	EXAMPLE 1 Creating a thread to handle receipt of a signal The following sample C code creates a thread to handle the receipt of a signal. More specifically, it catches the asynchronously generated signal, SIGINT. <pre> /***** * * compile with -D_POSIX_PTHREAD_SEMANTICS switch; * required by sigwait() * * sigint thread handles delivery of signal. uses sigwait() to wait * for SIGINT signal. * *****/ #include <pthread.h></pre>				

EXAMPLE 1 Creating a thread to handle receipt of a signal (Continued)

```

#include <stdlib.h>
#include <stdio.h>
#include <string.h>
#include <unistd.h>
#include <signal.h>
#include <synch.h>

static void    *threadTwo(void *);
static void    *threadThree(void *);
static void    *sigint(void *);

sigset_t      signalSet;

void *
main(void)
{
    pthread_t    t;
    pthread_t    t2;
    pthread_t    t3;

    thr_setconcurrency(3);
    sigfillset ( &signalSet );
    /*
     * Block signals in initial thread. New threads will
     * inherit this signal mask.
     */
    pthread_sigmask ( SIG_BLOCK, &signalSet, NULL );

    printf("Creating threads\n");

    /* POSIX thread create arguments:
     * thr_id, attr, strt_func, arg
     */
    pthread_create(&t, NULL, sigint, NULL);
    pthread_create(&t2, NULL, threadTwo, NULL);
    pthread_create(&t3, NULL, threadThree, NULL);

    printf("#####\n");
    printf("press CTRL-C to deliver SIGINT to sigint thread\n");
    printf("#####\n");

    thr_exit((void *)0);
}

static void *
threadTwo(void *arg)
{
    printf("hello world, from threadTwo [tid: %d]\n",
           pthread_self( ));
    printf("threadTwo [tid: %d] is now complete and exiting\n",
           pthread_self( ));
    thr_exit((void *)0);
}

static void *

```

sigwait(2)

EXAMPLE 1 Creating a thread to handle receipt of a signal (Continued)

```
threadThree(void *arg)
{
    printf("hello world, from threadThree [tid: %d]\n",
           pthread_self( ));
    printf("threadThree [tid: %d] is now complete and exiting\n",
           pthread_self( ));
    thr_exit((void *)0);
}

void *
sigint(void *arg)
{
    int    sig;
    int    err;

    printf("thread sigint [tid: %d] awaiting SIGINT\n",
           pthread_self( ));

    /* use POSIX sigwait( ) -- 2 args
     * signal set, signum
     */
    err = sigwait ( &signalSet, &sig );

    /* test for SIGINT; could catch other signals */
    if (err || sig != SIGINT)
        abort( );

    printf("\nSIGINT signal %d caught by sigint thread [tid: %d]\n",
           sig, pthread_self( ));
    thr_exit((void *)0);
}
```

SEE ALSO sigaction(2), sigpending(2), sigprocmask(2), sigsuspend(2), thr_create(3THR), thr_sigsetmask(3THR), signal(3HEAD), standards(5)

NOTES The sigwait() function cannot be used to wait for signals that cannot be caught (see sigaction(2)). This restriction is silently imposed by the system.

Solaris 2.4 and earlier releases provided a sigwait() facility as specified in POSIX.1c Draft 6. The final POSIX.1c standard changed the interface as described above. Support for the Draft 6 interface is provided for compatibility only and may not be supported in future releases. New applications and libraries should use the POSIX standard interface.

In Solaris 2.4 and earlier releases, the call to sigwait() from a multithreaded process overrode the signal's ignore disposition; even if a signal's disposition was SIG_IGN, a call to sigwait() resulted in catching the signal, if generated. This is unspecified behavior from the standpoint of the POSIX 1003.1c specification.

sigwait(2)

In Solaris 2.5, the behavior of `sigwait()` was corrected, so that it does not override the signal's ignore disposition. This change can cause applications that rely on the old behavior to break. Applications should employ `sigwait()` as follows: Install a dummy signal handler, thereby changing the disposition from `SIG_IGN` to having a handler. Then, any calls to `sigwait()` for this signal would catch it upon generation.

__sparc_utrap_install(2)

NAME	__sparc_utrap_install – install a SPARC V9 user trap handler
SYNOPSIS	<pre>#include <sys/utrap.h> int __sparc_utrap_install(utrap_entry_t type, utrap_handler_t new_precise, utrap_handler_t new_deferred, utrap_handler_t *old_precise, utrap_handler_t *old_deferred);</pre>
DESCRIPTION	The <code>__sparc_utrap_install()</code> function establishes <i>new_precise</i> and <i>new_deferred</i> user trap handlers as the new values for the specified <i>type</i> and returns the existing user trap handler values in <i>*old_precise</i> and <i>*old_deferred</i> in a single atomic operation. A new handler address of <code>NULL</code> means no user handler of that type will be installed. A new handler address of <code>UTH_NOCHANGE</code> means that the user handler for that type should not be changed. An old handler pointer of <code>NULL</code> means that the user is not interested in the old handler address. A <i>precise trap</i> is caused by a specific instruction and occurs before any program-visible state has been changed by this instruction. When a precise trap occurs, the program counter (PC) saved in the Trap Program Counter (TPC) register points to the instruction that induced the trap; all instructions prior to this trapping instruction have been executed. The next program counter (nPC) saved in the Trap Next Program Counter (TnPC) register points to the next instruction following the trapping instruction, which has not yet been executed. A <i>deferred trap</i> is also caused by a particular instruction, but unlike a precise trap, a deferred trap may occur after the program-visible state has been changed. See the <i>SPARC Architecture Manual, Version 9</i> for further information on precise and deferred traps. The list that follows contains hardware traps and their corresponding user trap types. User trap types marked with a plus-sign (+) are required and must be provided by all ABI-conforming implementations. The others may not be present on every implementation; an attempt to install a user trap handler for those conditions will return <code>EINVAL</code>. User trap types marked with an asterisk (*) are implemented as precise traps only.

Trap Name	User Trap Type (utrap_entry_t)
illegal_instruction	UT_ILLTRAP_INSTRUCTION +* or UT_ILLEGAL_INSTRUCTION
fp_disabled	UT_FP_DISABLED +*
fp_exception_ieee_754	UT_FP_EXCEPTION_IEEE_754 +
fp_exception_other	UT_FP_EXCEPTION_OTHER
tag_overflow	UT_TAG_OVERFLOW +*
division_by_zero	UT_DIVISION_BY_ZERO +
mem_address_not_aligned	UT_MEM_ADDRESS_NOT_ALIGNED +

Trap Name	User Trap Type (<code>utrap_entry_t</code>)
<code>privileged_action</code>	<code>UT_PRIVILEGED_ACTION +</code>
<code>privileged_opcode</code>	<code>UT_PRIVILEGED_OPCODE</code>
<code>async_data_error</code>	<code>UT_ASYNC_DATA_ERROR</code>
<code>trap_instruction</code>	<code>UT_TRAP_INSTRUCTION_16</code> through <code>UT_TRAP_INSTRUCTION_31 +*</code>
<code>instruction_access_exception</code> <code>instruction_access_MMU_miss</code> <code>instruction_access_error</code>	<code>UT_INSTRUCTION_EXCEPTION</code> or <code>UT_INSTRUCTION_PROTECTION</code> or <code>UT_INSTRUCTION_ERROR</code>
<code>data_access_exception</code> <code>data_access_MMU_miss</code> <code>data_access_error</code> <code>data_access_protection</code>	<code>UT_DATA_EXCEPTION</code> or <code>UT_DATA_PROTECTION</code> or <code>UT_DATA_ERROR</code>

The following explanations are provided for those user trap types that are not self-explanatory.

`UT_ILLTRAP_INSTRUCTION`

This trap is raised by user execution of the `ILLTRAP INSTRUCTION`. It is always precise.

`UT_ILLEGAL_INSTRUCTION`

This trap will be raised by the execution of otherwise undefined opcodes. It is implementation-dependent as to what opcodes raise this trap; the ABI only specifies the interface. The trap may be precise or deferred.

`UT_PRIVILEGED_OPCODE`

All opcodes declared to be privileged in SPARC V9 will raise this trap. It is implementation-dependent whether other opcodes will raise it as well; the ABI only specifies the interface.

`UT_DATA_EXCEPTION`, `UT_INSTRUCTION_EXCEPTION`

No valid user mapping can be made to this address, for a data or instruction access, respectively.

`UT_DATA_PROTECTION`, `UT_INSTRUCTION_PROTECTION`

A valid mapping exists, and user privilege to it exists, but the type of access (read, write, or execute) is denied, for a data or instruction access, respectively.

`UT_DATA_ERROR`, `UT_INSTRUCTION_ERROR`

A valid mapping exists, and both user privilege and the type of access are allowed, but an unrecoverable error occurred in attempting the access, for a data or instruction access, respectively. `%11` will contain either `BUS_ADDRERR` or `BUS_OBJERR`.

__sparc_utrap_install(2)

UT_FP_DISABLED

This trap is raised when an application issues a floating point instruction (including load or store) and the SPARC V9 Floating Point Registers State (FPRS) FEF bit is 0.

If a user handler is installed for this trap, it will be given control. Otherwise the system will set FEF to one and retry the instruction.

For all traps, the handler executes in a new register window, where the *in* registers are the *out* registers of the previous frame and have the value they contained at the time of the trap, similar to a normal subroutine call after the *save* instruction. The *global* registers (including the special registers %ccr, %asi, and %y) and the *floating-point* registers have their values from the time of the trap. The stack pointer register %sp plus the BIAS will point to a properly-aligned 128-byte register save area; if the handler needs scratch space, it should decrement the stack pointer to obtain it. If the handler needs access to the previous frame's *in* registers or *local* registers, it should execute a FLUSHW instruction, and then access them off of the frame pointer. If the handler calls an ABI-conforming function, it must set the %asi register to ASI_PRIMARY_NOFAULT before the call.

On entry to a precise user trap handler %l6 contains the %pc and %l7 contains the %npc at the time of the trap. To return from a handler and reexecute the trapped instruction, the handler would execute:

```
jmp1 %l6, %g0 ! Trapped PC supplied to user trap handler
return %l7    ! Trapped nPC supplied to user trap handler
```

To return from a handler and skip the trapped instruction, the handler would execute:

```
jmp1 %l7, %g0 ! Trapped nPC supplied to user trap handler
return %l7 + 4 ! Trapped nPC + 4
```

On entry to a deferred trap handler %o0 contains the address of the instruction that caused the trap and %o1 contains the actual instruction (right-justified, zero-extended), if the information is available. Otherwise %o0 contains the value -1 and %o1 is undefined. Additional information may be made available for certain cases of deferred traps, as indicated in the following table.

Instructions	Additional Information
LD-type (LDSTUB)	%o2 contains the effective address (<i>rs1</i> + <i>rs2</i> <i>simmm13</i>).
ST-type (CAS, SWAP)	%o2 contains the effective address (<i>rs1</i> + <i>rs2</i> <i>simmm13</i>).
Integer arithmetic	%o2 contains the <i>rs1</i> value. %o3 contains the <i>rs2</i> <i>simmm13</i> value. %o4 contains the contents of the %y register.
Floating-point arithmetic	%o2 contains the address of <i>rs1</i> value. %o3 contains the address of <i>rs2</i> value.
Control-transfer	%o2 contains the target address (<i>rs1</i> + <i>rs2</i> <i>simmm13</i>).

Asynchronous data errors	%o2 contains the address that caused the error. %o3 contains the effective ASI, if available, else -1.
--------------------------	--

To return from a deferred trap, the trap handler issues:

```
ta    68    !ST_RETURN_FROM_DEFERRED_TRAP
```

The following pseudo-code explains how the operating system dispatches traps:

```
if (precise_trap) {
    if (precise_handler) {
        invoke(precise_handler);
        /* not reached */
    } else {
        convert_to_signal(precise_trap);
    }
} else if (deferred_trap) {
    invoke(deferred_handler);
    /* not reached */
} else {
    convert_to_signal(deferred_trap);
}
if (signal)
    send(signal);
```

User trap handlers must preserve all registers except the *locals* (%l0-7) and the *outs* (%o0-7), that is, %i0-7, %g1-7, %d0-d62, %asi, %fsr, %fprs, %ccr, and %y, except to the extent that modifying the registers is part of the desired functionality of the handler. For example, the handler for UT_FP_DISABLED may load floating-point registers.

RETURN VALUES

Upon successful completion, 0 is returned. Otherwise, a non-zero value is returned and `errno` is set to indicate the error.

ERRORS

The `__sparc_utrap_install()` function will fail if:

EINVAL	The <i>type</i> argument is not a supported user trap type; the new user trap handler address is not word aligned; the old user trap handler address cannot be returned; or the user program is not a 64-bit executable.
--------	--

EXAMPLES

EXAMPLE 1 A sample program using the `__sparc_utrap_install()` function

The `__sparc_utrap_install()` function is normally used by user programs that wish to provide their own tailored exception handlers as a faster alternative to `signal(3C)`s, or to handle exceptions that are not directly supported by the `signal()` interface, such as `fp_disabled`.

```
extern void *fpdis_trap_handler();
utrap_handler_t new_precise = (utrap_handler_t)fpdis_trap_handler;
double d;
```

__sparc_utrap_install(2)

EXAMPLE 1 A sample program using the `__sparc_utrap_install()` function
(Continued)

```
int err;
err = __sparc_utrap_install(UT_FP_DISABLED, new_precise,
    UTH_NOCHANGE, NULL, NULL);
if (err == EINVAL) {
    /* unexpected error, do something */
    exit (1);
}
d = 1.0e-300;
ENTRY(fpdis_trap_handler)
wr      %g0, FPRS_FEF, %fprs
jmpl    %l6, %g0
return  %l7
SET_SIZE(fpdis_trap_handler)
```

This example turns on bit 2, FEF, in the Floating-Point Registers State (FPRS) Register, after a floating-point instruction causes an `fp_disabled` trap. (Note that this example simulates part of the default system behavior; programs do not need such a handler. The example is for illustrative purposes only.)

ATTRIBUTES See `attributes(5)` for descriptions of the following attributes:

ATTRIBUTE TYPE	ATTRIBUTE VALUE
MT-Level	MT-Safe

SEE ALSO `signal(3C)`, `attributes(5)`

SPARC Architecture Manual, Version 9

Manufacturer's processor chip user manuals

NOTES The Exceptions and Interrupt Descriptions section of the SPARC V9 manual documents which hardware traps are mandatory or optional, and whether they can be implemented as precise or deferred traps, or both. The manufacturer's processor chip user manuals describe the details of the traps supported for the specific processor implementation.

NAME	stat, lstat, fstat – get file status
SYNOPSIS	<pre>#include <sys/types.h> #include <sys/stat.h> int stat(const char *path, struct stat *buf); int lstat(const char *path, struct stat *buf); int fstat(int <i>fildev</i>, struct stat *buf);</pre>
DESCRIPTION	The <code>stat()</code> function obtains information about the file pointed to by <i>path</i>. Read, write, or execute permission of the named file is not required, but all directories listed in the path name leading to the file must be searchable. The <code>lstat()</code> function obtains file attributes similar to <code>stat()</code>, except when the named file is a symbolic link; in that case <code>lstat()</code> returns information about the link, while <code>stat()</code> returns information about the file the link references. The <code>fstat()</code> function obtains information about an open file known by the file descriptor <i>fildev</i>, obtained from a successful <code>open(2)</code>, <code>creat(2)</code>, <code>dup(2)</code>, <code>fcntl(2)</code>, or <code>pipe(2)</code> function. The <i>buf</i> argument is a pointer to a <code>stat</code> structure into which information is placed concerning the file. A <code>stat</code> structure includes the following members: <pre>mode_t st_mode; /* File mode (see mknod(2)) */ ino_t st_ino; /* Inode number */ dev_t st_dev; /* ID of device containing */ dev_t st_rdev; /* a directory entry for this file */ /* ID of device */ /* This entry is defined only for */ /* char special or block special files */ nlink_t st_nlink; /* Number of links */ uid_t st_uid; /* User ID of the file's owner */ gid_t st_gid; /* Group ID of the file's group */ off_t st_size; /* File size in bytes */ time_t st_atime; /* Time of last access */ time_t st_mtime; /* Time of last data modification */ time_t st_ctime; /* Time of last file status change */ /* Times measured in seconds since */ /* 00:00:00 UTC, Jan. 1, 1970 */ long st_blksize; /* Preferred I/O block size */ blkcnt_t st_blocks; /* Number of 512 byte blocks allocated*/</pre> Descriptions of structure members are as follows: st_mode The mode of the file as described in <code>mknod(2)</code>. In addition to the modes described in <code>mknod()</code>, the mode of a file may also be <code>S_IFLNK</code> if the file is a symbolic link. <code>S_IFLNK</code> may only be returned by <code>lstat()</code>. st_ino This field uniquely identifies the file in a given file system. The pair <code>st_ino</code> and <code>st_dev</code> uniquely identifies regular files.

stat(2)

	<code>st_dev</code>	This field uniquely identifies the file system that contains the file. Its value may be used as input to the <code>ustat()</code> function to determine more information about this file system. No other meaning is associated with this value.
	<code>st_rdev</code>	This field should be used only by administrative commands. It is valid only for block special or character special files and only has meaning on the system where the file was configured.
	<code>st_nlink</code>	This field should be used only by administrative commands.
	<code>st_uid</code>	The user ID of the file's owner.
	<code>st_gid</code>	The group ID of the file's group.
	<code>st_size</code>	For regular files, this is the address of the end of the file. For block special or character special, this is not defined. See also <code>pipe(2)</code> .
	<code>st_atime</code>	Time when file data was last accessed. Changed by the following functions: <code>creat()</code> , <code>mknod()</code> , <code>pipe()</code> , <code>utime(2)</code> , and <code>read(2)</code> .
	<code>st_mtime</code>	Time when data was last modified. Changed by the following functions: <code>creat()</code> , <code>mknod()</code> , <code>pipe()</code> , <code>utime()</code> , and <code>write(2)</code> .
	<code>st_ctime</code>	Time when file status was last changed. Changed by the following functions: <code>chmod()</code> , <code>chown()</code> , <code>creat()</code> , <code>link(2)</code> , <code>mknod()</code> , <code>pipe()</code> , <code>unlink(2)</code> , <code>utime()</code> , and <code>write()</code> .
	<code>st_blksize</code>	A hint as to the "best" unit size for I/O operations. This field is not defined for block special or character special files.
	<code>st_blocks</code>	The total number of physical blocks of size 512 bytes actually allocated on disk. This field is not defined for block special or character special files.
RETURN VALUES	Upon successful completion, 0 is returned. Otherwise, -1 is returned and <code>errno</code> is set to indicate the error.	
ERRORS	The <code>stat()</code> , <code>fstat()</code> , and <code>lstat()</code> functions will fail if:	
	<code>EOVERFLOW</code>	The file size in bytes or the number of blocks allocated to the file or the file serial number cannot be represented correctly in the structure pointed to by <i>buf</i> .
	The <code>stat()</code> and <code>lstat()</code> functions will fail if:	
	<code>EACCES</code>	Search permission is denied for a component of the path prefix.
	<code>EFAULT</code>	The <i>buf</i> or <i>path</i> argument points to an illegal address.
	<code>EINTR</code>	A signal was caught during the execution of the <code>stat()</code> or <code>lstat()</code> function.

ELOOP	Too many symbolic links were encountered in translating <i>path</i> .
ENAMETOOLONG	The length of the <i>path</i> argument exceeds <code>PATH_MAX</code> , or the length of a <i>path</i> component exceeds <code>NAME_MAX</code> while <code>_POSIX_NO_TRUNC</code> is in effect.
ENOENT	The named file does not exist or is the null pathname.
ENOLINK	The <i>path</i> argument points to a remote machine and the link to that machine is no longer active.
ENOTDIR	A component of the path prefix is not a directory.
EOVERFLOW	A component is too large to store in the structure pointed to by <i>buf</i> .

The `fstat()` function will fail if:

EBADF	The <i>filides</i> argument is not a valid open file descriptor.
EFAULT	The <i>buf</i> argument points to an illegal address.
EINTR	A signal was caught during the execution of the <code>fstat()</code> function.
ENOLINK	The <i>filides</i> argument points to a remote machine and the link to that machine is no longer active.
EOVERFLOW	A component is too large to store in the structure pointed to by <i>buf</i> .

USAGE The `stat()`, `fstat()`, and `lstat()` functions have transitional interfaces for 64-bit file offsets. See `lf64(5)`.

ATTRIBUTES See `attributes(5)` for descriptions of the following attributes:

ATTRIBUTE TYPE	ATTRIBUTE VALUE
MT-Level	<code>stat()</code> and <code>fstat()</code> are Async-Signal-Safe

SEE ALSO `chmod(2)`, `chown(2)`, `creat(2)`, `link(2)`, `mknod(2)`, `pipe(2)`, `read(2)`, `time(2)`, `unlink(2)`, `utime(2)`, `write(2)`, `fattach(3C)`, `attributes(5)`, `lf64(5)`, `stat(3HEAD)`

NOTES If you use `chmod(2)` to change the file group owner permissions on a file with ACL entries, both the file group owner permissions and the ACL mask are changed to the new permissions. Be aware that the new ACL mask permissions may change the effective permissions for additional users and groups who have ACL entries on the file.

statvfs(2)

NAME	statvfs, fstatvfs – get file system information
SYNOPSIS	<pre>#include <sys/types.h> #include <sys/statvfs.h> int statvfs(const char *path, struct statvfs *buf); int fstatvfs(int fildes, struct statvfs *buf);</pre>
DESCRIPTION	The <code>statvfs()</code> function returns a “generic superblock” describing a file system; it can be used to acquire information about mounted file systems. The <i>buf</i> argument is a pointer to a structure (described below) that is filled by the function. The <i>path</i> argument should name a file that resides on that file system. The file system type is known to the operating system. Read, write, or execute permission for the named file is not required, but all directories listed in the path name leading to the file must be searchable. The <code>statvfs</code> structure pointed to by <i>buf</i> includes the following members: <pre>u_long f_bsize; /* preferred file system block size */ u_long f_frsize; /* fundamental filesystem block (size if supported) */ fsblkcnt_t f_blocks; /* total # of blocks on file system in units of f_frsize */ fsblkcnt_t f_bfree; /* total # of free blocks */ fsblkcnt_t f_bavail; /* # of free blocks avail to non-super-user */ fsfilcnt_t f_files; /* total # of file nodes (inodes) */ fsfilcnt_t f_ffree; /* total # of free file nodes */ fsfilcnt_t f_favail; /* # of inodes avail to non-super-user */ u_long f_fsid; /* file system id (dev for now) */ char f_basetype[FSTYPSZ]; /* target fs type name, null-terminated */ u_long f_flag; /* bit mask of flags */ u_long f_namemax; /* maximum file name length */ char f_fstr[32]; /* file system specific string */ u_long f_filler[16]; /* reserved for future expansion */</pre> The <code>f_basetype</code> member contains a null-terminated FSType name of the mounted target. The following values can be returned in the <code>f_flag</code> field: <pre>ST_RDONLY 0x01 /* read-only file system */ ST_NOSUID 0x02 /* does not support setuid/setgid semantics */ ST_NOTRUNC 0x04 /* does not truncate file names longer than NAME_MAX */</pre> The <code>fstatvfs()</code> function is similar to <code>statvfs()</code>, except that the file named by <i>path</i> in <code>statvfs()</code> is instead identified by an open file descriptor <i>fildes</i> obtained from a successful <code>open(2)</code>, <code>creat(2)</code>, <code>dup(2)</code>, <code>fcntl(2)</code>, or <code>pipe(2)</code> function call.
RETURN VALUES	Upon successful completion, 0 is returned. Otherwise, -1 is returned and <code>errno</code> is set to indicate the error.

ERRORS	The <code>statvfs()</code> and <code>fstatvfs()</code> functions will fail if:	
	<code>EOVERFLOW</code>	One of the values to be returned cannot be represented correctly in the structure pointed to by <i>buf</i> .
	The <code>statvfs()</code> function will fail if:	
	<code>EACCES</code>	Search permission is denied on a component of the path prefix.
	<code>EFAULT</code>	The <i>path</i> or <i>buf</i> argument points to an illegal address.
	<code>EINTR</code>	A signal was caught during the execution of the <code>statvfs()</code> function.
	<code>EIO</code>	An I/O error occurred while reading the file system.
	<code>ELOOP</code>	Too many symbolic links were encountered in translating <i>path</i> .
	<code>ENAMETOOLONG</code>	The length of a <i>path</i> component exceeds <code>NAME_MAX</code> characters, or the length of <i>path</i> exceeds <code>PATH_MAX</code> characters.
	<code>ENOENT</code>	Either a component of the path prefix or the file referred to by <i>path</i> does not exist.
	<code>ENOLINK</code>	The <i>path</i> argument points to a remote machine and the link to that machine is no longer active.
	<code>ENOTDIR</code>	A component of the path prefix of <i>path</i> is not a directory.
	The <code>fstatvfs()</code> function will fail if:	
	<code>EBADF</code>	The <i>fd</i> argument is not an open file descriptor.
	<code>EFAULT</code>	The <i>buf</i> argument points to an illegal address.
	<code>EINTR</code>	A signal was caught during the execution of the <code>fstatvfs()</code> function.
	<code>EIO</code>	An I/O error occurred while reading the file system.
USAGE	The <code>statvfs()</code> and <code>fstatvfs()</code> functions have transitional interfaces for 64-bit file offsets. See <code>lf64(5)</code> .	
SEE ALSO	<code>chmod(2)</code> , <code>chown(2)</code> , <code>creat(2)</code> , <code>dup(2)</code> , <code>fcntl(2)</code> , <code>link(2)</code> , <code>mknod(2)</code> , <code>open(2)</code> , <code>pipe(2)</code> , <code>read(2)</code> , <code>time(2)</code> , <code>unlink(2)</code> , <code>utime(2)</code> , <code>write(2)</code> , <code>lf64(5)</code>	
BUGS	The values returned for <code>f_files</code> , <code>f_ffree</code> , and <code>f_favail</code> may not be valid for NFS mounted file systems.	

stime(2)

NAME	stime – set system time and date
SYNOPSIS	<pre>#include <unistd.h> int stime(const time_t *tp);</pre>
DESCRIPTION	The <code>stime()</code> function sets the system's idea of the time and date. The <i>tp</i> argument points to the value of time as measured in seconds from 00:00:00 UTC January 1, 1970.
RETURN VALUES	Upon successful completion, 0 is returned. Otherwise, -1 is returned and <code>errno</code> is set to indicate the error.
ERRORS	The <code>stime()</code> function will fail if: EINVAL The <i>tp</i> argument points to an invalid (negative) time value. EPERM The effective user of the calling process is not super-user.
SEE ALSO	<code>time(2)</code>

NAME	swapctl – manage swap space
SYNOPSIS	<pre>#include <sys/stat.h> #include <sys/swap.h> int swapctl(int cmd, void *arg);</pre>
DESCRIPTION	The <code>swapctl()</code> function adds, deletes, or returns information about swap resources. <i>cmd</i> specifies one of the following options contained in <code><sys/swap.h></code>: <pre>SC_ADD /* add a resource for swapping */ SC_LIST /* list the resources for swapping */ SC_REMOVE /* remove a resource for swapping */ SC_GETNSWP /* return number of swap resources */</pre> When <code>SC_ADD</code> or <code>SC_REMOVE</code> is specified, <i>arg</i> is a pointer to a <code>swapres</code> structure containing the following members: <pre>char *sr_name; /* pathname of resource */ off_t sr_start; /* offset to start of swap area */ off_t sr_length; /* length of swap area */</pre> The <code>sr_start</code> and <code>sr_length</code> members are specified in 512-byte blocks. A swap resource can only be removed by specifying the same values for the <code>sr_start</code> and <code>sr_length</code> members as were specified when it was added. Swap resources need not be removed in the order in which they were added. When <code>SC_LIST</code> is specified, <i>arg</i> is a pointer to a <code>swaptable</code> structure containing the following members: <pre>int swt_n; /* number of swapents following */ struct swapent swt_ent[]; /* array of swt_n swapents */</pre> A <code>swapent</code> structure contains the following members: <pre>char *ste_path; /* name of the swap file */ off_t ste_start; /* starting block for swapping */ off_t ste_length; /* length of swap area */ long ste_pages; /* number of pages for swapping */ long ste_free; /* number of ste_pages free */ long ste_flags; /* ST_INDEL bit set if swap file */ /* is now being deleted */</pre> The <code>SC_LIST</code> function causes <code>swapctl()</code> to return at most <code>swt_n</code> entries. The return value of <code>swapctl()</code> is the number actually returned. The <code>ST_INDEL</code> bit is turned on in <code>ste_flags</code> if the swap file is in the process of being deleted. When <code>SC_GETNSWP</code> is specified, <code>swapctl()</code> returns as its value the number of swap resources in use. <i>arg</i> is ignored for this operation. The <code>SC_ADD</code> and <code>SC_REMOVE</code> functions will fail if calling process does not have appropriate privileges.

swapctl(2)

RETURN VALUES	Upon successful completion, the function <code>swapctl()</code> returns a value of 0 for <code>SC_ADD</code> or <code>SC_REMOVE</code> , the number of <code>struct swapent</code> entries actually returned for <code>SC_LIST</code> , or the number of swap resources in use for <code>SC_GETNSWP</code> . Upon failure, the function <code>swapctl()</code> returns a value of -1 and sets <code>errno</code> to indicate an error.
ERRORS	Under the following conditions, the function <code>swapctl()</code> fails and sets <code>errno</code> to:
<code>EEXIST</code>	Part of the range specified by <code>sr_start</code> and <code>sr_length</code> is already being used for swapping on the specified resource (<code>SC_ADD</code>).
<code>EFAULT</code>	Either <code>arg</code> , <code>sr_name</code> , or <code>ste_path</code> points to an illegal address.
<code>EINVAL</code>	The specified function value is not valid, the path specified is not a swap resource (<code>SC_REMOVE</code>), part of the range specified by <code>sr_start</code> and <code>sr_length</code> lies outside the resource specified (<code>SC_ADD</code>), or the specified swap area is less than one page (<code>SC_ADD</code>).
<code>EISDIR</code>	The path specified for <code>SC_ADD</code> is a directory.
<code>ELOOP</code>	Too many symbolic links were encountered in translating the pathname provided to <code>SC_ADD</code> or <code>SC_REMOVE</code> .
<code>ENAMETOOLONG</code>	The length of a component of the path specified for <code>SC_ADD</code> or <code>SC_REMOVE</code> exceeds <code>NAME_MAX</code> characters or the length of the path exceeds <code>PATH_MAX</code> characters and <code>_POSIX_NO_TRUNC</code> is in effect.
<code>ENOENT</code>	The pathname specified for <code>SC_ADD</code> or <code>SC_REMOVE</code> does not exist.
<code>ENOMEM</code>	An insufficient number of <code>struct swapent</code> structures were provided to <code>SC_LIST</code> , or there were insufficient system storage resources available during an <code>SC_ADD</code> or <code>SC_REMOVE</code> , or the system would not have enough swap space after an <code>SC_REMOVE</code> .
<code>ENOSYS</code>	The pathname specified for <code>SC_ADD</code> or <code>SC_REMOVE</code> is not a file or block special device.
<code>ENOTDIR</code>	Pathname provided to <code>SC_ADD</code> or <code>SC_REMOVE</code> contained a component in the path prefix that was not a directory.
<code>EPERM</code>	The effective user of the calling process is not super-user.
<code>EROFS</code>	The pathname specified for <code>SC_ADD</code> is a read-only file system.

Additionally, the `swapctl()` function will fail for 32-bit interfaces if:

EOVERFLOW The amount of swap space configured on the machine is too large to be represented by a 32-bit quantity.

EXAMPLES

EXAMPLE 1 The usage of the `SC_GETNSWP` and `SC_LIST` commands.

The following example demonstrates the usage of the `SC_GETNSWP` and `SC_LIST` commands.

```
#include <sys/stat.h>
#include <sys/swap.h>
#include <stdio.h>

#define MAXSTRSIZE 80

main(argc, argv)
    int      argc;
    char     *argv[];
{
    swaptbl_t *s;
    int      i, n, num;
    char     *strtab;    /* string table for path names */

again:
    if ((num = swapctl(SC_GETNSWP, 0)) == -1) {
        perror("swapctl: GETNSWP");
        exit(1);
    }
    if (num == 0) {
        fprintf(stderr, "No Swap Devices Configured\n");
        exit(2);
    }
    /* allocate swaptable for num+1 entries */
    if ((s = (swaptbl_t *)
        malloc(num * sizeof(swaptent_t) +
              sizeof(struct swaptable))) ==
        (void *) 0) {
        fprintf(stderr, "Malloc Failed\n");
        exit(3);
    }
    /* allocate num+1 string holders */
    if ((strtab = (char *)
        malloc((num + 1) * MAXSTRSIZE)) == (void *) 0) {
        fprintf(stderr, "Malloc Failed\n");
        exit(3);
    }
    /* initialize string pointers */
    for (i = 0; i < (num + 1); i++) {
        s->swt_ent[i].ste_path = strtab + (i * MAXSTRSIZE);
    }

    s->swt_n = num + 1;
    if ((n = swapctl(SC_LIST, s)) < 0) {
        perror("swapctl");
        exit(1);
    }
}
```

swapctl(2)

EXAMPLE 1 The usage of the SC_GETNSWP and SC_LIST commands. *(Continued)*

```
    }
    if (n > num) {          /* more were added */
        free(s);
        free(strtab);
        goto again;
    }
    for (i = 0; i < n; i++)
        printf("%s %ld\n",
            s->swt_ent[i].ste_path, s->swt_ent[i].ste_pages);
}
```

NAME	symlink – make a symbolic link to a file	
SYNOPSIS	<pre>#include <unistd.h> int symlink(const char *name1, const char *name2);</pre>	
DESCRIPTION	The <code>symlink()</code> function creates a symbolic link <i>name2</i> to the file <i>name1</i>. Either name may be an arbitrary pathname, the files need not be on the same file system, and <i>name1</i> may be nonexistent. The file to which the symbolic link points is used when an <code>open(2)</code> operation is performed on the link. A <code>stat()</code> operation performed on a symbolic link returns the linked-to file, while an <code>lstat()</code> operation returns information about the link itself. See <code>stat(2)</code>. Unexpected results may occur when a symbolic link is made to a directory. To avoid confusion in applications, the <code>readlink(2)</code> call can be used to read the contents of a symbolic link.	
RETURN VALUES	Upon successful completion, 0 is returned. Otherwise, -1 is returned, <code>errno</code> is set to indicate the error, and the symbolic link is not made.	
ERRORS	EACCES	Search permission is denied for a component of the path prefix of <i>name2</i> .
	EDQUOT	The directory where the entry for the new symbolic link is being placed cannot be extended because the user's quota of disk blocks on that file system has been exhausted; the new symbolic link cannot be created because the user's quota of disk blocks on that file system has been exhausted; or the user's quota of inodes on the file system where the file is being created has been exhausted.
	EEXIST	The file referred to by <i>name2</i> already exists.
	EFAULT	The <i>name1</i> or <i>name2</i> argument points to an illegal address.
	EIO	An I/O error occurs while reading from or writing to the file system.
	ELOOP	Too many symbolic links are encountered in translating <i>name2</i> .
	ENAMETOOLONG	The length of the <i>name2</i> argument exceeds <code>PATH_MAX</code> , or the length of a <i>name2</i> component exceeds <code>NAME_MAX</code> while <code>_POSIX_NO_TRUNC</code> is in effect.
	ENOENT	A component of the path prefix of <i>name2</i> does not exist.
	ENOSPC	The directory in which the entry for the new symbolic link is being placed cannot be extended because no space is left on the file system containing the directory; the new symbolic link cannot be created because no

symlink(2)

space is left on the file system which will contain the link; or there are no free inodes on the file system on which the file is being created.

ENOSYS

The file system does not support symbolic links

ENOTDIR

A component of the path prefix of *name2* is not a directory.

EROFS

The file *name2* would reside on a read-only file system.

SEE ALSO

cp(1), link(2), open(2), readlink(2), stat(2), unlink(2)

NAME	sync – update super block
SYNOPSIS	<pre>#include <unistd.h> void sync(void);</pre>
DESCRIPTION	The <code>sync()</code> function writes all information in memory that should be on disk, including modified super blocks, modified inodes, and delayed block I/O. Unlike <code>fsync(3C)</code>, which completes the writing before it returns, <code>sync()</code> schedules but does not necessarily complete the writing before returning.
USAGE	The <code>sync()</code> function should be used by applications that examine a file system, such as <code>fsck(1M)</code> , and <code>df(1M)</code> , and is mandatory before rebooting.
SEE ALSO	<code>df(1M)</code> , <code>fsck(1M)</code> , <code>fsync(3C)</code>

sysfs(2)

NAME	sysfs – get file system type information						
SYNOPSIS	<pre>#include <sys/fstyp.h> #include <sys/fsid.h> int sysfs(int <i>opcode</i>, const char *<i>fsname</i>); int sysfs(int <i>opcode</i>, int <i>fs_index</i>, char *<i>buf</i>); int sysfs(int <i>opcode</i>);</pre>						
DESCRIPTION	The <code>sysfs()</code> function returns information about the file system types configured in the system. The number of arguments accepted by <code>sysfs()</code> depends on the <i>opcode</i> argument, which can take the following values: <table> <tr> <td>GETFSIND</td><td>Translate <i>fsname</i>, a null-terminated file-system type identifier, into a file-system type index.</td></tr> <tr> <td>GETFSTYP</td><td>Translate <i>fs_index</i>, a file-system type index, into a null-terminated file-system type identifier and write it into the buffer pointed to by <i>buf</i>, which must be at least of size FSTYPSZ as defined in <code><sys/fstyp.h></code>.</td></tr> <tr> <td>GETNFSSTYP</td><td>Return the total number of file system types configured in the system.</td></tr> </table>	GETFSIND	Translate <i>fsname</i> , a null-terminated file-system type identifier, into a file-system type index.	GETFSTYP	Translate <i>fs_index</i> , a file-system type index, into a null-terminated file-system type identifier and write it into the buffer pointed to by <i>buf</i> , which must be at least of size FSTYPSZ as defined in <code><sys/fstyp.h></code> .	GETNFSSTYP	Return the total number of file system types configured in the system.
GETFSIND	Translate <i>fsname</i> , a null-terminated file-system type identifier, into a file-system type index.						
GETFSTYP	Translate <i>fs_index</i> , a file-system type index, into a null-terminated file-system type identifier and write it into the buffer pointed to by <i>buf</i> , which must be at least of size FSTYPSZ as defined in <code><sys/fstyp.h></code> .						
GETNFSSTYP	Return the total number of file system types configured in the system.						
RETURN VALUES	Upon successful completion, the value returned depends upon the <i>opcode</i> argument as follows: <table> <tr> <td>GETFSIND</td><td>the file-system type index</td></tr> <tr> <td>GETFSTYP</td><td>0</td></tr> <tr> <td>GETNFSSTYP</td><td>the number of file system types configured</td></tr> </table> Otherwise, <code>-1</code> is returned and <code>errno</code> is set to indicate the error.	GETFSIND	the file-system type index	GETFSTYP	0	GETNFSSTYP	the number of file system types configured
GETFSIND	the file-system type index						
GETFSTYP	0						
GETNFSSTYP	the number of file system types configured						
ERRORS	The <code>sysfs()</code> function will fail if: <table> <tr> <td>EFAULT</td><td>The <i>buf</i> or <i>fsname</i> argument points to an illegal address.</td></tr> <tr> <td>EINVAL</td><td>The <i>fsname</i> argument points to an invalid file-system identifier; the <i>fs_index</i> argument is 0 or invalid; or the <i>opcode</i> argument is invalid.</td></tr> </table>	EFAULT	The <i>buf</i> or <i>fsname</i> argument points to an illegal address.	EINVAL	The <i>fsname</i> argument points to an invalid file-system identifier; the <i>fs_index</i> argument is 0 or invalid; or the <i>opcode</i> argument is invalid.		
EFAULT	The <i>buf</i> or <i>fsname</i> argument points to an illegal address.						
EINVAL	The <i>fsname</i> argument points to an invalid file-system identifier; the <i>fs_index</i> argument is 0 or invalid; or the <i>opcode</i> argument is invalid.						

NAME	sysinfo – get and set system information strings										
SYNOPSIS	<pre>#include <sys/systeminfo.h> long sysinfo(int <i>command</i>, char *<i>buf</i>, long <i>count</i>);</pre>										
DESCRIPTION	The <code>sysinfo()</code> function copies information relating to the operating system on which the process is executing into the buffer pointed to by <i>buf</i>. It can also set certain information where appropriate <i>commands</i> are available. The <i>count</i> parameter indicates the size of the buffer. The POSIX P1003.1 interface (see <code>standards(5)</code>) <code>sysconf(3C)</code> provides a similar class of configuration information, but returns an integer rather than a string. The values for <i>command</i> are as follows: <table> <tr> <td>SI_SYSNAME</td><td>Copy into the array pointed to by <i>buf</i> the string that would be returned by <code>uname(2)</code> in the <i>sysname</i> field. This is the name of the implementation of the operating system, for example, SunOS or UTS.</td></tr> <tr> <td>SI_HOSTNAME</td><td>Copy into the array pointed to by <i>buf</i> a string that names the present host machine. This is the string that would be returned by <code>uname(2)</code> in the <i>nodename</i> field. This hostname or nodename is often the name the machine is known by locally. The <i>hostname</i> is the name of this machine as a node in some network. Different networks may have different names for the node, but presenting the nodename to the appropriate network directory or name-to-address mapping service should produce a transport end point address. The name may not be fully qualified. Internet host names may be up to 256 bytes in length (plus the terminating null).</td></tr> <tr> <td>SI_SET_HOSTNAME</td><td>Copy the null-terminated contents of the array pointed to by <i>buf</i> into the string maintained by the kernel whose value will be returned by succeeding calls to <code>sysinfo()</code> with the command <code>SI_HOSTNAME</code>. This command requires that the effective-user-id be super-user.</td></tr> <tr> <td>SI_RELEASE</td><td>Copy into the array pointed to by <i>buf</i> the string that would be returned by <code>uname(2)</code> in the <i>release</i> field. Typical values might be 5.2 or 4.1.</td></tr> <tr> <td>SI_VERSION</td><td>Copy into the array pointed to by <i>buf</i> the string that would be returned by <code>uname(2)</code> in the <i>version</i> field. The syntax and semantics of this string are defined by the system provider.</td></tr> </table>	SI_SYSNAME	Copy into the array pointed to by <i>buf</i> the string that would be returned by <code>uname(2)</code> in the <i>sysname</i> field. This is the name of the implementation of the operating system, for example, SunOS or UTS.	SI_HOSTNAME	Copy into the array pointed to by <i>buf</i> a string that names the present host machine. This is the string that would be returned by <code>uname(2)</code> in the <i>nodename</i> field. This hostname or nodename is often the name the machine is known by locally. The <i>hostname</i> is the name of this machine as a node in some network. Different networks may have different names for the node, but presenting the nodename to the appropriate network directory or name-to-address mapping service should produce a transport end point address. The name may not be fully qualified. Internet host names may be up to 256 bytes in length (plus the terminating null).	SI_SET_HOSTNAME	Copy the null-terminated contents of the array pointed to by <i>buf</i> into the string maintained by the kernel whose value will be returned by succeeding calls to <code>sysinfo()</code> with the command <code>SI_HOSTNAME</code> . This command requires that the effective-user-id be super-user.	SI_RELEASE	Copy into the array pointed to by <i>buf</i> the string that would be returned by <code>uname(2)</code> in the <i>release</i> field. Typical values might be 5.2 or 4.1.	SI_VERSION	Copy into the array pointed to by <i>buf</i> the string that would be returned by <code>uname(2)</code> in the <i>version</i> field. The syntax and semantics of this string are defined by the system provider.
SI_SYSNAME	Copy into the array pointed to by <i>buf</i> the string that would be returned by <code>uname(2)</code> in the <i>sysname</i> field. This is the name of the implementation of the operating system, for example, SunOS or UTS.										
SI_HOSTNAME	Copy into the array pointed to by <i>buf</i> a string that names the present host machine. This is the string that would be returned by <code>uname(2)</code> in the <i>nodename</i> field. This hostname or nodename is often the name the machine is known by locally. The <i>hostname</i> is the name of this machine as a node in some network. Different networks may have different names for the node, but presenting the nodename to the appropriate network directory or name-to-address mapping service should produce a transport end point address. The name may not be fully qualified. Internet host names may be up to 256 bytes in length (plus the terminating null).										
SI_SET_HOSTNAME	Copy the null-terminated contents of the array pointed to by <i>buf</i> into the string maintained by the kernel whose value will be returned by succeeding calls to <code>sysinfo()</code> with the command <code>SI_HOSTNAME</code> . This command requires that the effective-user-id be super-user.										
SI_RELEASE	Copy into the array pointed to by <i>buf</i> the string that would be returned by <code>uname(2)</code> in the <i>release</i> field. Typical values might be 5.2 or 4.1.										
SI_VERSION	Copy into the array pointed to by <i>buf</i> the string that would be returned by <code>uname(2)</code> in the <i>version</i> field. The syntax and semantics of this string are defined by the system provider.										

sysinfo(2)

SI_MACHINE	Copy into the array pointed to by <i>buf</i> the string that would be returned by <code>uname(2)</code> in the <i>machine</i> field, for example, <i>sun4c</i> , <i>sun4d</i> , or <i>sun4m</i> .
SI_ARCHITECTURE	Copy into the array pointed to by <i>buf</i> a string describing the basic instruction set architecture of the current system, for example, <i>sparc</i> , <i>mc68030</i> , <i>m32100</i> , or <i>i386</i> . These names may not match predefined names in the C language compilation system.
SI_ISALIST	Copy into the array pointed to by <i>buf</i> the names of the variant instruction set architectures executable on the current system. The names are space-separated and are ordered in the sense of best performance. That is, earlier-named instruction sets may contain more instructions than later-named instruction sets; a program that is compiled for an earlier-named instruction set will most likely run faster on this machine than the same program compiled for a later-named instruction set. Programs compiled for an instruction set that does not appear in the list will most likely experience performance degradation or not run at all on this machine. The instruction set names known to the system are listed in <code>isalist(5)</code>; these names may or may not match predefined names or compiler options in the C language compilation system.
SI_PLATFORM	Copy into the array pointed to by <i>buf</i> a string describing the specific model of the hardware platform, for example, <i>SUNW</i> , <i>Sun_4_75</i> , <i>SUNW</i> , <i>SPARCsystem-600</i> , or <i>i86pc</i> .
SI_HW_PROVIDER	Copies the name of the hardware manufacturer into the array pointed to by <i>buf</i> .
SI_HW_SERIAL	Copy into the array pointed to by <i>buf</i> a string which is the ASCII representation of the hardware-specific serial number of the physical machine on which the function is executed. Note that this may be implemented in Read-Only Memory, using software constants set when building the operating system, or by other means, and may contain non-numeric characters. It is anticipated that manufacturers will not issue the same "serial number" to more than one physical machine. The pair of strings returned by <code>SI_HW_PROVIDER</code> and

	SI_HW_SERIAL is likely to be unique across all vendor's SVR4 implementations.						
SI_SRPC_DOMAIN	Copies the Secure Remote Procedure Call domain name into the array pointed to by <i>buf</i> .						
SI_SET_SRPC_DOMAIN	Set the string to be returned by <code>sysinfo()</code> with the SI_SRPC_DOMAIN command to the value contained in the array pointed to by <i>buf</i> . This command requires that the effective-user-id be super-user.						
SI_DHCP_CACHE	Copy into the array pointed to by <i>buf</i> an ASCII string consisting of the ASCII hexadecimal encoding of the name of the interface configured by <code>boot(1M)</code> followed by the DHCPACK reply from the server. This command is intended for use only by the <code>dhcpgent(1M)</code> DHCP client daemon for the purpose of adopting the DHCP maintenance of the interface configured by <code>boot</code> .						
RETURN VALUES	Upon successful completion, the value returned indicates the buffer size in bytes required to hold the complete value and the terminating null character. If this value is no greater than the value passed in <i>count</i>, the entire string was copied. If this value is greater than <i>count</i>, the string copied into <i>buf</i> has been truncated to <i>count</i> - 1 bytes plus a terminating null character. Otherwise, -1 is returned and <code>errno</code> is set to indicate the error.						
ERRORS	The <code>sysinfo()</code> function will fail if: <table> <tr> <td>EFAULT</td><td>The <i>buf</i> argument does not point to a valid address.</td></tr> <tr> <td>EINVAL</td><td>The data for a SET command exceeds the limits established by the implementation.</td></tr> <tr> <td>EPERM</td><td>The effective user of the calling process is not super-user.</td></tr> </table>	EFAULT	The <i>buf</i> argument does not point to a valid address.	EINVAL	The data for a SET command exceeds the limits established by the implementation.	EPERM	The effective user of the calling process is not super-user.
EFAULT	The <i>buf</i> argument does not point to a valid address.						
EINVAL	The data for a SET command exceeds the limits established by the implementation.						
EPERM	The effective user of the calling process is not super-user.						
USAGE	In many cases there is no corresponding programming interface to set these values; such strings are typically settable only by the system administrator modifying entries in the <code>/etc/system</code> directory or the code provided by the particular OEM reading a serial number or code out of read-only memory, or hard-coded in the version of the operating system. A good estimation for <i>count</i> is 257, which is likely to cover all strings returned by this interface in typical installations.						
SEE ALSO	<code>boot(1M)</code> , <code>dhcpgent(1M)</code> , <code>uname(2)</code> , <code>gethostid(3C)</code> , <code>gethostname(3C)</code> , <code>sysconf(3C)</code> , <code>isalist(5)</code> , <code>standards(5)</code>						

time(2)

NAME	time – get time				
SYNOPSIS	<pre>#include <sys/types.h> #include <time.h> time_t time(time_t *tloc);</pre>				
DESCRIPTION	The <code>time()</code> function returns the value of time in seconds since 00:00:00 UTC, January 1, 1970. If <i>tloc</i> is non-zero, the return value is also stored in the location to which <i>tloc</i> points. If <i>tloc</i> points to an illegal address, <code>time()</code> fails and its actions are undefined.				
RETURN VALUES	Upon successful completion, <code>time()</code> returns the value of time. Otherwise, <code>(time_t)-1</code> is returned and <code>errno</code> is set to indicate the error.				
ATTRIBUTES	See <code>attributes(5)</code> for descriptions of the following attributes:				
	<table><tr><th>ATTRIBUTE TYPE</th><th>ATTRIBUTE VALUE</th></tr><tr><td>MT-Level</td><td>Async-Signal-Safe</td></tr></table>	ATTRIBUTE TYPE	ATTRIBUTE VALUE	MT-Level	Async-Signal-Safe
ATTRIBUTE TYPE	ATTRIBUTE VALUE				
MT-Level	Async-Signal-Safe				
SEE ALSO	<code>stime(2)</code> , <code>ctime(3C)</code> , <code>attributes(5)</code>				

NAME	times – get process and child process times
SYNOPSIS	<pre>#include <sys/times.h> #include <limits.h> clock_t times(struct tms *buffer);</pre>
DESCRIPTION	The <code>times()</code> function fills the <code>tms</code> structure pointed to by <i>buffer</i> with time-accounting information. The <code>tms</code> structure, defined in <code><sys/times.h></code>, contains the following members: <pre>clock_t tms_utime; clock_t tms_stime; clock_t tms_cutime; clock_t tms_cstime;</pre> All times are reported in clock ticks. The specific value for a clock tick is defined by the variable <code>CLK_TCK</code>, found in the header <code><limits.h></code>. The times of a terminated child process are included in the <code>tms_cutime</code> and <code>tms_cstime</code> members of the parent when <code>wait(2)</code> or <code>waitpid(2)</code> returns the process ID of this terminated child. If a child process has not waited for its children, their times will not be included in its times. The <code>tms_utime</code> member is the CPU time used while executing instructions in the user space of the calling process. The <code>tms_stime</code> member is the CPU time used by the system on behalf of the calling process. The <code>tms_cutime</code> member is the sum of the <code>tms_utime</code> and the <code>tms_cutime</code> of the child processes. The <code>tms_cstime</code> member is the sum of the <code>tms_stime</code> and the <code>tms_cstime</code> of the child processes.
RETURN VALUES	Upon successful completion, <code>times()</code> returns the elapsed real time, in clock ticks, since an arbitrary point in the past (for example, system start-up time). This point does not change from one invocation of <code>times()</code> within the process to another. The return value may overflow the possible range of type <code>clock_t</code>. If <code>times()</code> fails, <code>(clock_t)-1</code> is returned and <code>errno</code> is set to indicate the error.
ERRORS	The <code>times()</code> function will fail if: EFAULT The <i>buffer</i> argument points to an illegal address.
ATTRIBUTES	See <code>attributes(5)</code> for descriptions of the following attributes:

ATTRIBUTE TYPE	ATTRIBUTE VALUE
MT-Level	Async-Signal-Safe

times(2)

SEE ALSO time(1), timex(1), exec(2), fork(2), time(2), wait(2), waitid(2), waitpid(2), attributes(5)

NAME	uadmin – administrative control																								
SYNOPSIS	<pre>#include <sys/uadmin.h> int uadmin(int <i>cmd</i>, int <i>fcn</i>, uintptr_t <i>mdep</i>);</pre>																								
DESCRIPTION	The <code>uadmin()</code> function provides control for basic administrative functions. This function is tightly coupled to the system administrative procedures and is not intended for general use. The argument <i>mdep</i> is provided for machine-dependent use and is not defined here. As specified by <i>cmd</i>, the following commands are available: <table> <tr> <td>A_SHUTDOWN</td><td>The system is shut down. All user processes are killed, the buffer cache is flushed, and the root file system is unmounted. The action to be taken after the system has been shut down is specified by <i>fcn</i>. The functions are generic; the hardware capabilities vary on specific machines.</td></tr> <tr> <td>AD_HALT</td><td>Halt the processor(s).</td></tr> <tr> <td>AD_POWEROFF</td><td>Halt the processor(s) and turn off the power.</td></tr> <tr> <td>AD_BOOT</td><td>Reboot the system, using the kernel file.</td></tr> <tr> <td>AD_IBOOT</td><td>Interactive reboot; user is prompted for bootable program name.</td></tr> <tr> <td>A_REBOOT</td><td>The system stops immediately without any further processing. The action to be taken next is specified by <i>fcn</i> as above.</td></tr> <tr> <td>A_DUMP</td><td>The system is forced to panic immediately without any further processing and a crash dump is written to the dump device (see <code>dumpadm(1M)</code>). The action to be taken next is specified by <i>fcn</i> as above.</td></tr> <tr> <td>A_REMOUNT</td><td>The root file system is mounted again after having been fixed. This should be used only during the startup process.</td></tr> <tr> <td>A_FREEZE</td><td>Suspend the whole system. The system state is preserved in the state file. The following three subcommands are available.</td></tr> <tr> <td>AD_COMPRESS</td><td>Save the system state to the state file with compression of data.</td></tr> <tr> <td>AD_CHECK</td><td>Check if your system supports suspend and resume. Without performing a system suspend/resume, this command checks if this feature is currently available on your system.</td></tr> <tr> <td>AD_FORCE</td><td>Force AD_COMPRESS even when threads of drivers are not suspendable.</td></tr> </table>	A_SHUTDOWN	The system is shut down. All user processes are killed, the buffer cache is flushed, and the root file system is unmounted. The action to be taken after the system has been shut down is specified by <i>fcn</i> . The functions are generic; the hardware capabilities vary on specific machines.	AD_HALT	Halt the processor(s).	AD_POWEROFF	Halt the processor(s) and turn off the power.	AD_BOOT	Reboot the system, using the kernel file.	AD_IBOOT	Interactive reboot; user is prompted for bootable program name.	A_REBOOT	The system stops immediately without any further processing. The action to be taken next is specified by <i>fcn</i> as above.	A_DUMP	The system is forced to panic immediately without any further processing and a crash dump is written to the dump device (see <code>dumpadm(1M)</code>). The action to be taken next is specified by <i>fcn</i> as above.	A_REMOUNT	The root file system is mounted again after having been fixed. This should be used only during the startup process.	A_FREEZE	Suspend the whole system. The system state is preserved in the state file. The following three subcommands are available.	AD_COMPRESS	Save the system state to the state file with compression of data.	AD_CHECK	Check if your system supports suspend and resume. Without performing a system suspend/resume, this command checks if this feature is currently available on your system.	AD_FORCE	Force AD_COMPRESS even when threads of drivers are not suspendable.
A_SHUTDOWN	The system is shut down. All user processes are killed, the buffer cache is flushed, and the root file system is unmounted. The action to be taken after the system has been shut down is specified by <i>fcn</i> . The functions are generic; the hardware capabilities vary on specific machines.																								
AD_HALT	Halt the processor(s).																								
AD_POWEROFF	Halt the processor(s) and turn off the power.																								
AD_BOOT	Reboot the system, using the kernel file.																								
AD_IBOOT	Interactive reboot; user is prompted for bootable program name.																								
A_REBOOT	The system stops immediately without any further processing. The action to be taken next is specified by <i>fcn</i> as above.																								
A_DUMP	The system is forced to panic immediately without any further processing and a crash dump is written to the dump device (see <code>dumpadm(1M)</code>). The action to be taken next is specified by <i>fcn</i> as above.																								
A_REMOUNT	The root file system is mounted again after having been fixed. This should be used only during the startup process.																								
A_FREEZE	Suspend the whole system. The system state is preserved in the state file. The following three subcommands are available.																								
AD_COMPRESS	Save the system state to the state file with compression of data.																								
AD_CHECK	Check if your system supports suspend and resume. Without performing a system suspend/resume, this command checks if this feature is currently available on your system.																								
AD_FORCE	Force AD_COMPRESS even when threads of drivers are not suspendable.																								
RETURN VALUES	Upon successful completion, the value returned depends on <i>cmd</i> as follows:																								

uadmin(2)

A_SHUTDOWN	Never returns.
A_REBOOT	Never returns.
A_FREEZE	0 upon resume.
A_REMOUNT	0.

Otherwise, -1 is returned and `errno` is set to indicate the error.

ERRORS

The `uadmin()` function will fail if:

EPERM	The effective user of the calling process is not super-user.
ENOMEM	Suspend/resume ran out of physical memory.
ENOSPC	Suspend/resume could not allocate enough space on the root file system to store system information.
ENOTSUP	Suspend/resume not supported on this platform.
ENXIO	Unable to successfully suspend system.
EBUSY	Suspend already in progress.

SEE ALSO

`dumpadm(1M)`, `kernel(1M)`, `uadmin(1M)`

NAME	ulimit – get and set process limits								
SYNOPSIS	<pre>#include <ulimit.h> long ulimit(int <i>cmd</i>, /* <i>newlimit</i> */...);</pre>								
DESCRIPTION	The <code>ulimit()</code> function provides for control over process limits. It is effective in limiting the growth of regular files. Pipes are limited to <code>PIPE_MAX</code> bytes. The <i>cmd</i> values, defined in <code><ulimit.h></code>, include: <table> <tr> <td><code>UL_GETFSIZE</code></td><td>Return the soft file size limit of the process. The limit is in units of 512-byte blocks and is inherited by child processes. Files of any size can be read. The return value is the integer part of the soft file size limit divided by 512. If the result cannot be represented as a long int, the result is unspecified.</td></tr> <tr> <td><code>UL_SETFSIZE</code></td><td>Set the hard and soft file size limits for output operations of the process to the value of the second argument, taken as a long int. Any process may decrease its own hard limit, but only a process with appropriate privileges may increase the limit. The new file size limit is returned. The hard and soft file size limits are set to the specified value multiplied by 512. If the result would overflow an <code>rlimit_t</code>, the actual value set is unspecified.</td></tr> <tr> <td><code>UL_GMEMLIM</code></td><td>Get the maximum possible break value (see <code>brk(2)</code>).</td></tr> <tr> <td><code>UL_GDESLIM</code></td><td>Get the current value of the maximum number of open files per process configured in the system.</td></tr> </table>	<code>UL_GETFSIZE</code>	Return the soft file size limit of the process. The limit is in units of 512-byte blocks and is inherited by child processes. Files of any size can be read. The return value is the integer part of the soft file size limit divided by 512. If the result cannot be represented as a long int, the result is unspecified.	<code>UL_SETFSIZE</code>	Set the hard and soft file size limits for output operations of the process to the value of the second argument, taken as a long int. Any process may decrease its own hard limit, but only a process with appropriate privileges may increase the limit. The new file size limit is returned. The hard and soft file size limits are set to the specified value multiplied by 512. If the result would overflow an <code>rlimit_t</code> , the actual value set is unspecified.	<code>UL_GMEMLIM</code>	Get the maximum possible break value (see <code>brk(2)</code>).	<code>UL_GDESLIM</code>	Get the current value of the maximum number of open files per process configured in the system.
<code>UL_GETFSIZE</code>	Return the soft file size limit of the process. The limit is in units of 512-byte blocks and is inherited by child processes. Files of any size can be read. The return value is the integer part of the soft file size limit divided by 512. If the result cannot be represented as a long int, the result is unspecified.								
<code>UL_SETFSIZE</code>	Set the hard and soft file size limits for output operations of the process to the value of the second argument, taken as a long int. Any process may decrease its own hard limit, but only a process with appropriate privileges may increase the limit. The new file size limit is returned. The hard and soft file size limits are set to the specified value multiplied by 512. If the result would overflow an <code>rlimit_t</code> , the actual value set is unspecified.								
<code>UL_GMEMLIM</code>	Get the maximum possible break value (see <code>brk(2)</code>).								
<code>UL_GDESLIM</code>	Get the current value of the maximum number of open files per process configured in the system.								
RETURN VALUES	Upon successful completion, <code>ulimit()</code> returns the value of the requested limit. Otherwise, <code>-1</code> is returned, the limit is not changed, and <code>errno</code> is set to indicate the error.								
ERRORS	The <code>ulimit()</code> function will fail if: <table> <tr> <td><code>EINVAL</code></td><td>The <i>cmd</i> argument is not valid.</td></tr> <tr> <td><code>EPERM</code></td><td>A process not having appropriate privileges attempts to increase its file size limit.</td></tr> </table>	<code>EINVAL</code>	The <i>cmd</i> argument is not valid.	<code>EPERM</code>	A process not having appropriate privileges attempts to increase its file size limit.				
<code>EINVAL</code>	The <i>cmd</i> argument is not valid.								
<code>EPERM</code>	A process not having appropriate privileges attempts to increase its file size limit.								
USAGE	Since all return values are permissible in a successful situation, an application wishing to check for error situations should set <code>errno</code> to 0, then call <code>ulimit()</code>, and if it returns <code>-1</code>, check if <code>errno</code> is non-zero. The <code>getrlimit()</code> and <code>setrlimit()</code> functions provide a more general interface for controlling process limits, and are preferred over <code>ulimit()</code>. See <code>getrlimit(2)</code>.								
SEE ALSO	<code>brk(2)</code> , <code>getrlimit(2)</code> , <code>write(2)</code>								

umask(2)

NAME	umask – set and get file creation mask				
SYNOPSIS	<pre>#include <sys/types.h> #include <sys/stat.h> mode_t umask(mode_t <i>cmask</i>) ;</pre>				
DESCRIPTION	The <code>umask()</code> function sets the process's file mode creation mask to <i>cmask</i> and returns the previous value of the mask. Only the access permission bits of <i>cmask</i> and the file mode creation mask are used. The mask is inherited by child processes. See <code>intro(2)</code> for more information on masks.				
RETURN VALUES	The previous value of the file mode creation mask is returned.				
ATTRIBUTES	See <code>attributes(5)</code> for descriptions of the following attributes: <table><tr><th>ATTRIBUTE TYPE</th><th>ATTRIBUTE VALUE</th></tr><tr><td>MT-Level</td><td>Async-Signal-Safe</td></tr></table>	ATTRIBUTE TYPE	ATTRIBUTE VALUE	MT-Level	Async-Signal-Safe
ATTRIBUTE TYPE	ATTRIBUTE VALUE				
MT-Level	Async-Signal-Safe				
SEE ALSO	<code>mkdir(1)</code> , <code>sh(1)</code> , <code>intro(2)</code> , <code>chmod(2)</code> , <code>creat(2)</code> , <code>mknod(2)</code> , <code>open(2)</code> , <code>stat(3HEAD)</code> , <code>attributes(5)</code>				

NAME	umount, umount2 – unmount a file system																		
SYNOPSIS	<pre>#include <sys/mount.h> int umount(const char *file); int umount2(const char *file, int mflag);</pre>																		
DESCRIPTION	The <code>umount()</code> function requests that a previously mounted file system contained on a block special device or directory be unmounted. The <i>file</i> argument is a pointer to the absolute pathname of the file system to be unmounted. After unmounting the file system, the directory upon which the file system was mounted reverts to its ordinary interpretation. The <code>umount2()</code> function is identical to <code>umount()</code>, with the additional capability of unmounting file systems even if there are open files active. The <i>mflag</i> argument must contain one of the following values: <table> <tr> <td>0</td><td>Perform a normal unmount that is equivalent to <code>umount()</code>. The <code>umount2()</code> function returns <code>EBUSY</code> if there are open files active within the file system to be unmounted.</td></tr> <tr> <td><code>MS_FORCE</code></td><td>Unmount the file system, even if there are open files active. A forced unmount may resort in loss of data, so it should be used only when a regular unmount is unsuccessful. The <code>umount2()</code> function returns <code>ENOTSUP</code> if the specified file systems does not support <code>MS_FORCE</code>. Currently only <code>nfs</code>- and <code>ufs</code>-type file systems support <code>MS_FORCE</code>.</td></tr> </table>	0	Perform a normal unmount that is equivalent to <code>umount()</code> . The <code>umount2()</code> function returns <code>EBUSY</code> if there are open files active within the file system to be unmounted.	<code>MS_FORCE</code>	Unmount the file system, even if there are open files active. A forced unmount may resort in loss of data, so it should be used only when a regular unmount is unsuccessful. The <code>umount2()</code> function returns <code>ENOTSUP</code> if the specified file systems does not support <code>MS_FORCE</code> . Currently only <code>nfs</code> - and <code>ufs</code> -type file systems support <code>MS_FORCE</code> .														
0	Perform a normal unmount that is equivalent to <code>umount()</code> . The <code>umount2()</code> function returns <code>EBUSY</code> if there are open files active within the file system to be unmounted.																		
<code>MS_FORCE</code>	Unmount the file system, even if there are open files active. A forced unmount may resort in loss of data, so it should be used only when a regular unmount is unsuccessful. The <code>umount2()</code> function returns <code>ENOTSUP</code> if the specified file systems does not support <code>MS_FORCE</code> . Currently only <code>nfs</code> - and <code>ufs</code> -type file systems support <code>MS_FORCE</code> .																		
RETURN VALUES	Upon successful completion, 0 is returned. Otherwise, -1 is returned and <code>errno</code> is set to indicate the error.																		
ERRORS	The <code>umount()</code> and <code>umount2()</code> functions will fail if: <table> <tr> <td><code>EBUSY</code></td><td>A file on <i>file</i> is busy.</td></tr> <tr> <td><code>EFAULT</code></td><td>The file pointed to by <i>file</i> points to an illegal address.</td></tr> <tr> <td><code>EINVAL</code></td><td>The file pointed to by <i>file</i> is not mounted.</td></tr> <tr> <td><code>ENOENT</code></td><td>The file pointed to by <i>file</i> does not exist.</td></tr> <tr> <td><code>ELOOP</code></td><td>Too many symbolic links were encountered in translating the path pointed to by <i>file</i>.</td></tr> <tr> <td><code>ENAMETOOLONG</code></td><td>The length of the <i>file</i> argument exceeds <code>PATH_MAX</code>, or the length of a <i>file</i> component exceeds <code>NAME_MAX</code> while <code>_POSIX_NO_TRUNC</code> is in effect.</td></tr> <tr> <td><code>ENOLINK</code></td><td>The file pointed to by <i>file</i> is on a remote machine and the link to that machine is no longer active.</td></tr> <tr> <td><code>ENOTBLK</code></td><td>The file pointed to by <i>file</i> is not a block special device.</td></tr> <tr> <td><code>EPERM</code></td><td>The process's effective user ID is not superuser.</td></tr> </table>	<code>EBUSY</code>	A file on <i>file</i> is busy.	<code>EFAULT</code>	The file pointed to by <i>file</i> points to an illegal address.	<code>EINVAL</code>	The file pointed to by <i>file</i> is not mounted.	<code>ENOENT</code>	The file pointed to by <i>file</i> does not exist.	<code>ELOOP</code>	Too many symbolic links were encountered in translating the path pointed to by <i>file</i> .	<code>ENAMETOOLONG</code>	The length of the <i>file</i> argument exceeds <code>PATH_MAX</code> , or the length of a <i>file</i> component exceeds <code>NAME_MAX</code> while <code>_POSIX_NO_TRUNC</code> is in effect.	<code>ENOLINK</code>	The file pointed to by <i>file</i> is on a remote machine and the link to that machine is no longer active.	<code>ENOTBLK</code>	The file pointed to by <i>file</i> is not a block special device.	<code>EPERM</code>	The process's effective user ID is not superuser.
<code>EBUSY</code>	A file on <i>file</i> is busy.																		
<code>EFAULT</code>	The file pointed to by <i>file</i> points to an illegal address.																		
<code>EINVAL</code>	The file pointed to by <i>file</i> is not mounted.																		
<code>ENOENT</code>	The file pointed to by <i>file</i> does not exist.																		
<code>ELOOP</code>	Too many symbolic links were encountered in translating the path pointed to by <i>file</i> .																		
<code>ENAMETOOLONG</code>	The length of the <i>file</i> argument exceeds <code>PATH_MAX</code> , or the length of a <i>file</i> component exceeds <code>NAME_MAX</code> while <code>_POSIX_NO_TRUNC</code> is in effect.																		
<code>ENOLINK</code>	The file pointed to by <i>file</i> is on a remote machine and the link to that machine is no longer active.																		
<code>ENOTBLK</code>	The file pointed to by <i>file</i> is not a block special device.																		
<code>EPERM</code>	The process's effective user ID is not superuser.																		

umount(2)

EREMOTE

The file pointed to by *file* is remote.

The umount2 () function will fail if:

ENOTSUP

The file pointed to by *file* does not support this operation.

USAGE

The umount () and umount2 () functions may be invoked only by the superuser. Because it provides greater functionality, the umount2 () function is preferred.

SEE ALSO

mount(2)

NAME	uname – get name of current operating system				
SYNOPSIS	<pre>#include <sys/utsname.h> int uname(struct utsname *name);</pre>				
DESCRIPTION	The <code>uname()</code> function stores information identifying the current operating system in the structure pointed to by <i>name</i>. The <code>uname()</code> function uses the <code>utsname</code> structure, defined in <code><sys/utsname.h></code>, whose members include: <pre>char sysname[SYS_NMLN]; char nodename[SYS_NMLN]; char release[SYS_NMLN]; char version[SYS_NMLN]; char machine[SYS_NMLN];</pre> The <code>uname()</code> function returns a null-terminated character string naming the current operating system in the character array <code>sysname</code>. Similarly, the <code>nodename</code> member contains the name by which the system is known on a communications network. The <code>release</code> and <code>version</code> members further identify the operating system. The <code>machine</code> member contains a standard name that identifies the hardware on which the operating system is running.				
RETURN VALUES	Upon successful completion, a non-negative value is returned. Otherwise, <code>-1</code> is returned and <code>errno</code> is set to indicate the error.				
ERRORS	The <code>uname()</code> function will fail if: EFAULT The <i>name</i> argument points to an illegal address.				
ATTRIBUTES	See <code>attributes(5)</code> for descriptions of the following attributes:				
<table border="1"> <thead> <tr> <th>ATTRIBUTE TYPE</th><th>ATTRIBUTE VALUE</th></tr> </thead> <tbody> <tr> <td>MT-Level</td><td>Async-Signal-Safe</td></tr> </tbody> </table>		ATTRIBUTE TYPE	ATTRIBUTE VALUE	MT-Level	Async-Signal-Safe
ATTRIBUTE TYPE	ATTRIBUTE VALUE				
MT-Level	Async-Signal-Safe				
SEE ALSO	<code>uname(1)</code> , <code>sysinfo(2)</code> , <code>sysconf(3C)</code> , <code>attributes(5)</code>				

unlink(2)

NAME	unlink – remove directory entry														
SYNOPSIS	<pre>#include <unistd.h> int unlink(const char *path);</pre>														
DESCRIPTION	The <code>unlink()</code> function removes a link to a file. If <i>path</i> names a symbolic link, <code>unlink()</code> removes the symbolic link named by <i>path</i> and does not affect any file or directory named by the contents of the symbolic link. Otherwise, <code>unlink()</code> removes the link named by the pathname pointed to by <i>path</i> and decrements the link count of the file referenced by the link. When the file's link count becomes 0 and no process has the file open, the space occupied by the file will be freed and the file will no longer be accessible. If one or more processes have the file open when the last link is removed, the link will be removed before <code>unlink()</code> returns, but the removal of the file contents will be postponed until all references to the file are closed. The <i>path</i> argument must not name a directory unless the process has appropriate privileges and the implementation supports using <code>unlink()</code> on directories. Upon successful completion, <code>unlink()</code> will mark for update the <code>st_ctime</code> and <code>st_mtime</code> fields of the parent directory. If the file's link count is not 0, the <code>st_ctime</code> field of the file will be marked for update.														
RETURN VALUES	Upon successful completion, 0 is returned. Otherwise, -1 is returned, <code>errno</code> is set to indicate the error, and the file is not unlinked.														
ERRORS	The <code>unlink()</code> function will fail if: <table> <tr> <td>EACCES</td><td>Search permission is denied for a component of the <i>path</i> prefix; write permission is denied on the directory containing the link to be removed; the parent directory has the sticky bit set and the file is not writable by the user; or the user does not own the parent directory and the user does not own the file.</td></tr> <tr> <td>EBUSY</td><td>The entry to be unlinked is the mount point for a mounted file system.</td></tr> <tr> <td>EFAULT</td><td>The <i>path</i> argument points to an illegal address.</td></tr> <tr> <td>EINTR</td><td>A signal was caught during the execution of the <code>unlink()</code> function.</td></tr> <tr> <td>ELOOP</td><td>Too many symbolic links were encountered in translating <i>path</i>.</td></tr> <tr> <td>ENAMETOOLONG</td><td>The length of the <i>path</i> argument exceeds <code>PATH_MAX</code>, or the length of a <i>path</i> component exceeds <code>NAME_MAX</code> while <code>_POSIX_NO_TRUNC</code> is in effect.</td></tr> <tr> <td>ENOENT</td><td>The named file does not exist or is a null pathname.</td></tr> </table>	EACCES	Search permission is denied for a component of the <i>path</i> prefix; write permission is denied on the directory containing the link to be removed; the parent directory has the sticky bit set and the file is not writable by the user; or the user does not own the parent directory and the user does not own the file.	EBUSY	The entry to be unlinked is the mount point for a mounted file system.	EFAULT	The <i>path</i> argument points to an illegal address.	EINTR	A signal was caught during the execution of the <code>unlink()</code> function.	ELOOP	Too many symbolic links were encountered in translating <i>path</i> .	ENAMETOOLONG	The length of the <i>path</i> argument exceeds <code>PATH_MAX</code> , or the length of a <i>path</i> component exceeds <code>NAME_MAX</code> while <code>_POSIX_NO_TRUNC</code> is in effect.	ENOENT	The named file does not exist or is a null pathname.
EACCES	Search permission is denied for a component of the <i>path</i> prefix; write permission is denied on the directory containing the link to be removed; the parent directory has the sticky bit set and the file is not writable by the user; or the user does not own the parent directory and the user does not own the file.														
EBUSY	The entry to be unlinked is the mount point for a mounted file system.														
EFAULT	The <i>path</i> argument points to an illegal address.														
EINTR	A signal was caught during the execution of the <code>unlink()</code> function.														
ELOOP	Too many symbolic links were encountered in translating <i>path</i> .														
ENAMETOOLONG	The length of the <i>path</i> argument exceeds <code>PATH_MAX</code> , or the length of a <i>path</i> component exceeds <code>NAME_MAX</code> while <code>_POSIX_NO_TRUNC</code> is in effect.														
ENOENT	The named file does not exist or is a null pathname.														

- ENOLINK The *path* argument points to a remote machine and the link to that machine is no longer active.
 - ENOTDIR A component of the *path* prefix is not a directory.
 - EPERM The named file is a directory and the effective user of the calling process is not super-user.
 - EROFS The directory entry to be unlinked is part of a read-only file system.
- The unlink() function may fail if:
- ENAMETOOLONG Pathname resolution of a symbolic link produced an intermediate result whose length exceeds PATH_MAX.
 - ETXTBSY The entry to be unlinked is the last directory entry to a pure procedure (shared text) file that is being executed.

USAGE Applications should use rmdir(2) to remove a directory.

ATTRIBUTES See attributes(5) for descriptions of the following attributes:

ATTRIBUTE TYPE	ATTRIBUTE VALUE
MT-Level	Async-Signal-Safe

SEE ALSO rm(1), close(2), link(2), open(2), rmdir(2), remove(3C), attributes(5)

ustat(2)

NAME	ustat – get file system statistics												
SYNOPSIS	<pre>#include <sys/types.h> #include <ustat.h> int ustat(dev_t <i>dev</i>, struct ustat *<i>buf</i>);</pre>												
DESCRIPTION	The <code>ustat()</code> function returns information about a mounted file system. The <i>dev</i> argument is a device number identifying a device containing a mounted file system (see <code>makedev(3C)</code>). The <i>buf</i> argument is a pointer to a <code>ustat</code> structure that includes the following members: <pre>daddr_t f_tfree; /* Total free blocks */ ino_t f_tinode; /* Number of free inodes */ char f_fname[6]; /* Filsys name */ char f_fpack[6]; /* Filsys pack name */</pre> The <code>f_fname</code> and <code>f_fpack</code> members may not contain significant information on all systems; in this case, these members will contain the null character as the first character.												
RETURN VALUES	Upon successful completion, 0 is returned. Otherwise, -1 is returned and <code>errno</code> is set to indicate the error.												
ERRORS	The <code>ustat()</code> function will fail if: <table> <tr> <td>ECOMM</td><td>The <i>dev</i> argument is on a remote machine and the link to that machine is no longer active.</td></tr> <tr> <td>EFAULT</td><td>The <i>buf</i> argument points to an illegal address.</td></tr> <tr> <td>EINTR</td><td>A signal was caught during the execution of the <code>ustat()</code> function.</td></tr> <tr> <td>EINVAL</td><td>The <i>dev</i> argument is not the device number of a device containing a mounted file system.</td></tr> <tr> <td>ENOLINK</td><td>The <i>dev</i> argument refers to a device on a remote machine and the link to that machine is no longer active.</td></tr> <tr> <td>EOVERFLOW</td><td>One of the values returned cannot be represented in the structure pointed to by <i>buf</i>.</td></tr> </table>	ECOMM	The <i>dev</i> argument is on a remote machine and the link to that machine is no longer active.	EFAULT	The <i>buf</i> argument points to an illegal address.	EINTR	A signal was caught during the execution of the <code>ustat()</code> function.	EINVAL	The <i>dev</i> argument is not the device number of a device containing a mounted file system.	ENOLINK	The <i>dev</i> argument refers to a device on a remote machine and the link to that machine is no longer active.	EOVERFLOW	One of the values returned cannot be represented in the structure pointed to by <i>buf</i> .
ECOMM	The <i>dev</i> argument is on a remote machine and the link to that machine is no longer active.												
EFAULT	The <i>buf</i> argument points to an illegal address.												
EINTR	A signal was caught during the execution of the <code>ustat()</code> function.												
EINVAL	The <i>dev</i> argument is not the device number of a device containing a mounted file system.												
ENOLINK	The <i>dev</i> argument refers to a device on a remote machine and the link to that machine is no longer active.												
EOVERFLOW	One of the values returned cannot be represented in the structure pointed to by <i>buf</i> .												
USAGE	The <code>statvfs(2)</code> function should be used in favor of <code>ustat()</code> .												
SEE ALSO	<code>stat(2)</code> , <code>statvfs(2)</code> , <code>makedev(3C)</code>												
BUGS	The NFS revision 2 protocol does not permit the number of free files to be provided to the client; therefore, when <code>ustat()</code> has completed on an NFS file system, <code>f_tinode</code> is always -1.												

NAME	utime – set file access and modification times														
SYNOPSIS	<pre>#include <sys/types.h> #include <utime.h> int utime(const char *path, const struct utimbuf *times);</pre>														
DESCRIPTION	The <code>utime()</code> function sets the access and modification times of the file pointed to by <i>path</i>, and causes the time of the last file status change (<code>st_ctime</code>) to be updated. If <i>times</i> is NULL, the access and modification times of the file are set to the current time. A process must be the owner of the file or have write permission to use <code>utime()</code> in this manner. If <i>times</i> is not NULL, <i>times</i> is interpreted as a pointer to a <code>utimbuf</code> structure (defined in <code><utime.h></code>) and the access and modification times are set to the values contained in the designated structure. Only the owner of the file or the super-user may use <code>utime()</code> in this manner. The <code>utimbuf</code> structure contains the following members: <pre>time_t actime; /* access time */ time_t modtime; /* modification time */</pre> The times contained in the members of the <code>utimbuf</code> structure are measured in seconds since 00:00:00 UTC, January 1, 1970.														
RETURN VALUES	Upon successful completion, 0 is returned. Otherwise, -1 is returned and <code>errno</code> is set to indicate the error.														
ERRORS	The <code>utime()</code> function will fail if: <table> <tr> <td>EACCES</td><td>Search permission is denied by a component of the <i>path</i> prefix; or the effective user ID of the process is not super-user and not the owner of the file, write permission is denied for the file, and <i>times</i> is NULL.</td></tr> <tr> <td>EFAULT</td><td>The <i>path</i> argument points to an illegal address.</td></tr> <tr> <td>EINTR</td><td>A signal was caught during the execution of the <code>utime()</code> function.</td></tr> <tr> <td>EIO</td><td>An I/O error occurred while reading from or writing to the file system.</td></tr> <tr> <td>ELOOP</td><td>Too many symbolic links were encountered in translating <i>path</i>.</td></tr> <tr> <td>ENAMETOOLONG</td><td>The length of the <i>path</i> argument exceeds <code>PATH_MAX</code>, or the length of a <i>path</i> component exceeds <code>NAME_MAX</code> while <code>_POSIX_NO_TRUNC</code> is in effect.</td></tr> <tr> <td>ENOENT</td><td>The named file does not exist or is a null pathname.</td></tr> </table>	EACCES	Search permission is denied by a component of the <i>path</i> prefix; or the effective user ID of the process is not super-user and not the owner of the file, write permission is denied for the file, and <i>times</i> is NULL.	EFAULT	The <i>path</i> argument points to an illegal address.	EINTR	A signal was caught during the execution of the <code>utime()</code> function.	EIO	An I/O error occurred while reading from or writing to the file system.	ELOOP	Too many symbolic links were encountered in translating <i>path</i> .	ENAMETOOLONG	The length of the <i>path</i> argument exceeds <code>PATH_MAX</code> , or the length of a <i>path</i> component exceeds <code>NAME_MAX</code> while <code>_POSIX_NO_TRUNC</code> is in effect.	ENOENT	The named file does not exist or is a null pathname.
EACCES	Search permission is denied by a component of the <i>path</i> prefix; or the effective user ID of the process is not super-user and not the owner of the file, write permission is denied for the file, and <i>times</i> is NULL.														
EFAULT	The <i>path</i> argument points to an illegal address.														
EINTR	A signal was caught during the execution of the <code>utime()</code> function.														
EIO	An I/O error occurred while reading from or writing to the file system.														
ELOOP	Too many symbolic links were encountered in translating <i>path</i> .														
ENAMETOOLONG	The length of the <i>path</i> argument exceeds <code>PATH_MAX</code> , or the length of a <i>path</i> component exceeds <code>NAME_MAX</code> while <code>_POSIX_NO_TRUNC</code> is in effect.														
ENOENT	The named file does not exist or is a null pathname.														

utime(2)

ENOLINK	The <i>path</i> argument points to a remote machine and the link to that machine is no longer active.
ENOTDIR	A component of the <i>path</i> prefix is not a directory.
EPERM	The effective user of the calling process is not super-user and not the owner of the file, and <i>times</i> is not NULL.
EROFS	The file system containing the file is mounted read-only.

ATTRIBUTES See `attributes(5)` for descriptions of the following attributes:

ATTRIBUTE TYPE	ATTRIBUTE VALUE
MT-Level	Async-Signal-Safe

SEE ALSO `stat(2)`, `attributes(5)`

NAME	utimes – set file access and modification times																
SYNOPSIS	<pre>#include <sys/time.h> int utimes(const char *path, const struct timeval times[2]);</pre>																
DESCRIPTION	The <code>utimes()</code> function sets the access and modification times of the file pointed to by the <i>path</i> argument to the value of the <i>times</i> argument. It allows time specifications accurate to the microsecond. The <i>times</i> argument is an array of <code>timeval</code> structures. The first array member represents the date and time of last access, and the second member represents the date and time of last modification. The times in the <code>timeval</code> structure are measured in seconds and microseconds since the Epoch, although rounding toward the nearest second may occur. If the <i>times</i> argument is a null pointer, the access and modification times of the file are set to the current time. The effective user ID of the process must be the same as the owner of the file, or must have write access to the file or super-user privileges to use this call in this manner. Upon completion, <code>utimes()</code> will mark the time of the last file status change, <code>st_ctime</code>, for update.																
RETURN VALUES	Upon successful completion, 0 is returned. Otherwise, -1 is returned, <code>errno</code> is set to indicate the error, and the file times will not be affected.																
ERRORS	The <code>utimes()</code> function will fail if: <table> <tr> <td>EACCES</td><td>Search permission is denied by a component of the path prefix; or the <i>times</i> argument is a null pointer and the effective user ID of the process does not match the owner of the file and write access is denied.</td></tr> <tr> <td>EFAULT</td><td>The <i>path</i> or <i>times</i> argument points to an illegal address.</td></tr> <tr> <td>EINTR</td><td>A signal was caught during the execution of the <code>utimes()</code> function.</td></tr> <tr> <td>EINVAL</td><td>The number of microseconds specified in one or both of the <code>timeval</code> structures pointed to by <i>times</i> was greater than or equal to 1,000,000 or less than 0.</td></tr> <tr> <td>EIO</td><td>An I/O error occurred while reading from or writing to the file system.</td></tr> <tr> <td>ELOOP</td><td>Too many symbolic links were encountered in resolving <i>path</i>.</td></tr> <tr> <td>ENAMETOOLONG</td><td>The length of the <i>path</i> argument exceeds <code>PATH_MAX</code> or a pathname component is longer than <code>NAME_MAX</code>.</td></tr> <tr> <td>ENOLINK</td><td>The <i>path</i> argument points to a remote machine and the link to that machine is no longer active.</td></tr> </table>	EACCES	Search permission is denied by a component of the path prefix; or the <i>times</i> argument is a null pointer and the effective user ID of the process does not match the owner of the file and write access is denied.	EFAULT	The <i>path</i> or <i>times</i> argument points to an illegal address.	EINTR	A signal was caught during the execution of the <code>utimes()</code> function.	EINVAL	The number of microseconds specified in one or both of the <code>timeval</code> structures pointed to by <i>times</i> was greater than or equal to 1,000,000 or less than 0.	EIO	An I/O error occurred while reading from or writing to the file system.	ELOOP	Too many symbolic links were encountered in resolving <i>path</i> .	ENAMETOOLONG	The length of the <i>path</i> argument exceeds <code>PATH_MAX</code> or a pathname component is longer than <code>NAME_MAX</code> .	ENOLINK	The <i>path</i> argument points to a remote machine and the link to that machine is no longer active.
EACCES	Search permission is denied by a component of the path prefix; or the <i>times</i> argument is a null pointer and the effective user ID of the process does not match the owner of the file and write access is denied.																
EFAULT	The <i>path</i> or <i>times</i> argument points to an illegal address.																
EINTR	A signal was caught during the execution of the <code>utimes()</code> function.																
EINVAL	The number of microseconds specified in one or both of the <code>timeval</code> structures pointed to by <i>times</i> was greater than or equal to 1,000,000 or less than 0.																
EIO	An I/O error occurred while reading from or writing to the file system.																
ELOOP	Too many symbolic links were encountered in resolving <i>path</i> .																
ENAMETOOLONG	The length of the <i>path</i> argument exceeds <code>PATH_MAX</code> or a pathname component is longer than <code>NAME_MAX</code> .																
ENOLINK	The <i>path</i> argument points to a remote machine and the link to that machine is no longer active.																

utimes(2)

ENOENT	A component of <i>path</i> does not name an existing file or <i>path</i> is an empty string.
ENOTDIR	A component of the path prefix is not a directory.
EPERM	The <i>times</i> argument is not a null pointer and the calling process's effective user ID has write access to the file but does not match the owner of the file and the calling process does not have the appropriate privileges.
EROFS	The file system containing the file is read-only.
The <code>utimes()</code> function may fail if:	
ENAMETOOLONG	Path name resolution of a symbolic link produced an intermediate result whose length exceeds <code>PATH_MAX</code> .

SEE ALSO

`stat(2)`

NAME	vfork – spawn new process in a virtual memory efficient way				
SYNOPSIS	<pre>#include <unistd.h> pid_t vfork(void);</pre>				
DESCRIPTION	The <code>vfork()</code> function creates new processes without fully copying the address space of the old process. This function is useful in instances where the purpose of a <code>fork(2)</code> operation would be to create a new system context for an <code>execve()</code> operation (see <code>exec(2)</code>). Unlike with the <code>fork()</code> function, the child process borrows the parent's memory and thread of control until a call to <code>execve()</code> or an exit (either abnormally or by a call to <code>_exit()</code> (see <code>exit(2)</code>). The parent process is suspended while the child is using its resources. In a multithreaded application, <code>vfork()</code> borrows only the thread of control that called <code>vfork()</code> in the parent; that is, the child contains only one thread. In that sense, <code>vfork()</code> behaves like <code>fork()</code>. The <code>vfork()</code> function can normally be used the same way as <code>fork()</code>. The procedure that called <code>vfork()</code>, however, should not return while running in the child's context, since the eventual return from <code>vfork()</code> would be to a stack frame that no longer exists. The <code>_exit()</code> function should be used in favor of <code>exit(3C)</code> if unable to perform an <code>execve()</code> operation, since <code>exit()</code> will flush and close standard I/O channels, and thereby corrupt the parent process's standard I/O data structures. The <code>_exit()</code> function should be used even with <code>fork()</code> to avoid flushing the buffered data twice.				
RETURN VALUES	Upon successful completion, <code>vfork()</code> returns 0 to the child process and returns the process ID of the child process to the parent process. Otherwise, -1 is returned to the parent process, no child process is created, and <code>errno</code> is set to indicate the error.				
ERRORS	The <code>vfork()</code> function will fail if: <table> <tr> <td>EAGAIN</td><td>The system-imposed limit on the total number of processes under execution (either system-quality or by a single user) would be exceeded. This limit is determined when the system is generated.</td></tr> <tr> <td>ENOMEM</td><td>There is insufficient swap space for the new process.</td></tr> </table>	EAGAIN	The system-imposed limit on the total number of processes under execution (either system-quality or by a single user) would be exceeded. This limit is determined when the system is generated.	ENOMEM	There is insufficient swap space for the new process.
EAGAIN	The system-imposed limit on the total number of processes under execution (either system-quality or by a single user) would be exceeded. This limit is determined when the system is generated.				
ENOMEM	There is insufficient swap space for the new process.				
SEE ALSO	<code>exec(2)</code> , <code>exit(2)</code> , <code>fork(2)</code> , <code>ioctl(2)</code> , <code>wait(2)</code> , <code>exit(3C)</code>				
NOTES	The use of <code>vfork()</code> for any purpose other than as a prelude to an immediate call to a function from the <code>exec</code> family or to <code>_exit()</code> is not advised. The <code>vfork()</code> function is unsafe in multithreaded applications. This function will be eliminated in a future release. The memory sharing semantics of <code>vfork()</code> can be obtained through other mechanisms.				

`vfork(2)`

To avoid a possible deadlock situation, processes that are children in the middle of a `vfork()` are never sent `SIGTTOU` or `SIGTTIN` signals; rather, output or `ioctl`s are allowed and input attempts result in an EOF indication.

On some systems, the implementation of `vfork()` causes the parent to inherit register values from the child. This can create problems for certain optimizing compilers if `<unistd.h>` is not included in the source calling `vfork()`.

NAME	vhangup – virtually “hangup” the current controlling terminal
SYNOPSIS	<pre>#include <unistd.h> void vhangup(void);</pre>
DESCRIPTION	The <code>vhangup()</code> function is used by the initialization process <code>init(1M)</code> (among others) to ensure that users are given “clean” terminals at login by revoking access of the previous users’ processes to the terminal. To effect this, <code>vhangup()</code> searches the system tables for references to the controlling terminal of the invoking process and revokes access permissions on each instance of the terminal that it finds. Further attempts to access the terminal by the affected processes will yield I/O errors (EBADF or EIO). A <code>SIGHUP</code> (hangup signal) is sent to the process group of the controlling terminal.
SEE ALSO	<code>init(1M)</code>
BUGS	Access to the controlling terminal using <code>/dev/tty</code> is still possible. This call should be replaced by an automatic mechanism that takes place on process exit.

wait(2)

NAME	wait – wait for child process to stop or terminate		
SYNOPSIS	<pre>#include <sys/types.h> #include <sys/wait.h> pid_t wait(int *stat_loc);</pre>		
DESCRIPTION	The <code>wait()</code> function will suspend execution of the calling thread until status information for one of its terminated child processes is available, or until delivery of a signal whose action is either to execute a signal-catching function or to terminate the process. If more than one thread is suspended in <code>wait()</code> or <code>waitpid(2)</code> awaiting termination of the same process, exactly one thread will return the process status at the time of the target process termination. If status information is available prior to the call to <code>wait()</code>, return will be immediate. If <code>wait()</code> returns because the status of a child process is available, it returns the process ID of the child process. If the calling process specified a non-zero value for <i>stat_loc</i>, the status of the child process is stored in the location pointed to by <i>stat_loc</i>. That status may be evaluated with the macros described on the <code>wstat(3XFN)</code> manual page. In the following, <i>status</i> is the object pointed to by <i>stat_loc</i>: ■ If the child process stopped, the high order 8 bits of <i>status</i> will contain the number of the signal that caused the process to stop and the low order 8 bits will be set equal to <code>WSTOPFLG</code>. ■ If the child process terminated due to an <code>_exit()</code> call, the low order 8 bits of <i>status</i> will be 0 and the high order 8 bits will contain the low order 8 bits of the argument that the child process passed to <code>_exit()</code>; see <code>exit(2)</code>. ■ If the child process terminated due to a signal, the high order 8 bits of <i>status</i> will be 0 and the low order 8 bits will contain the number of the signal that caused the termination. In addition, if <code>WCOREFLG</code> is set, a “core image” will have been produced; see <code>signal(3HEAD)</code> and <code>wstat(3XFN)</code>. If the calling process has <code>SA_NOCLDWAIT</code> set or has <code>SIGCHLD</code> set to <code>SIG_IGN</code>, and the process has no unwaited children that were transformed into zombie processes, it will block until all of its children terminate, and <code>wait()</code> will fail and set <code>errno</code> to <code>ECHILD</code>. If a parent process terminates without waiting for its child processes to terminate, the parent process ID of each child process is set to 1, with the initialization process inheriting the child processes; see <code>intro(3)</code>.		
RETURN VALUES	When <code>wait()</code> returns due to a terminated child process, the process ID of the child is returned to the calling process. Otherwise, <code>-1</code> is returned and <code>errno</code> is set to indicate the error.		
ERRORS	The <code>wait()</code> function will fail if: <table><tr><td><code>ECHILD</code></td><td>The calling process has no existing unwaited-for child processes.</td></tr></table>	<code>ECHILD</code>	The calling process has no existing unwaited-for child processes.
<code>ECHILD</code>	The calling process has no existing unwaited-for child processes.		

wait(2)

EINTR The function was interrupted by a signal.

USAGE Since wait () blocks on a stopped child, a calling process wishing to see the return results of such a call should use waitid(2) or waitpid(2) instead of wait () .

ATTRIBUTES See attributes(5) for descriptions of the following attributes:

ATTRIBUTE TYPE	ATTRIBUTE VALUE
MT-Level	Async-Signal-Safe

SEE ALSO intro(3), exec(2), exit(2), fork(2), pause(2), ptrace(2), waitid(2), waitpid(2), signal(3C), attributes(5), signal(3HEAD), wstat(3XFN)

waitid(2)

NAME	waitid – wait for child process to change state												
SYNOPSIS	<pre>#include <wait.h> int waitid(idtype_t <i>idtype</i>, id_t <i>id</i>, siginfo_t *<i>infop</i>, int <i>options</i>);</pre>												
DESCRIPTION	The <code>waitid()</code> function suspends the calling process until one of its child processes changes state. It records the current state of a child in the structure pointed to by <i>infop</i>. It returns immediately if a child process changed state prior to the call. The <i>idtype</i> and <i>id</i> arguments specify which children <code>waitid()</code> is to wait for, as follows: ■ If <i>idtype</i> is <code>P_PID</code>, <code>waitid()</code> waits for the child with a process ID equal to <code>(pid_t)id</code>. ■ If <i>idtype</i> is <code>P_PGID</code>, <code>waitid()</code> waits for any child with a process group ID equal to <code>(pid_t)id</code>. ■ If <i>idtype</i> is <code>P_ALL</code>, <code>waitid()</code> waits for any child and <i>id</i> is ignored. The <i>options</i> argument is used to specify which state changes <code>waitid()</code> is to wait for. It is formed by bitwise OR operation of any of the following flags: <table> <tr> <td>WCONTINUED</td><td>Return the status for any child that was stopped and has been continued.</td></tr> <tr> <td>WEXITED</td><td>Wait for process(es) to exit.</td></tr> <tr> <td>WNOHANG</td><td>Return immediately.</td></tr> <tr> <td>WNOWAIT</td><td>Keep the process in a waitable state.</td></tr> <tr> <td>WSTOPPED</td><td>Wait for and return the process status of any child that has stopped upon receipt of a signal.</td></tr> <tr> <td>WTRAPPED</td><td>Wait for traced process(es) to become trapped or reach a breakpoint (see <code>ptrace(2)</code>).</td></tr> </table> The <i>infop</i> argument must point to a <code>siginfo_t</code> structure, as defined in <code>siginfo(3HEAD)</code>. If <code>waitid()</code> returns because a child process was found that satisfies the conditions indicated by the arguments <i>idtype</i> and <i>options</i>, then the structure pointed to by <i>infop</i> will be filled by the system with the status of the process. The <code>si_signo</code> member will always be equal to <code>SIGCHLD</code>.	WCONTINUED	Return the status for any child that was stopped and has been continued.	WEXITED	Wait for process(es) to exit.	WNOHANG	Return immediately.	WNOWAIT	Keep the process in a waitable state.	WSTOPPED	Wait for and return the process status of any child that has stopped upon receipt of a signal.	WTRAPPED	Wait for traced process(es) to become trapped or reach a breakpoint (see <code>ptrace(2)</code>).
WCONTINUED	Return the status for any child that was stopped and has been continued.												
WEXITED	Wait for process(es) to exit.												
WNOHANG	Return immediately.												
WNOWAIT	Keep the process in a waitable state.												
WSTOPPED	Wait for and return the process status of any child that has stopped upon receipt of a signal.												
WTRAPPED	Wait for traced process(es) to become trapped or reach a breakpoint (see <code>ptrace(2)</code>).												
RETURN VALUES	If <code>waitid()</code> returns due to a change of state of one of its children and <code>WNOHANG</code> was not used, 0 is returned. Otherwise, -1 is returned and <code>errno</code> is set to indicate the error. If <code>WNOHANG</code> was used, 0 can be returned (indicating no error); however, no children may have changed state if <code>info->si_pid</code> is 0.												
ERRORS	The <code>waitid()</code> function will fail if: <table> <tr> <td>ECHILD</td><td>The set of processes specified by <i>idtype</i> and <i>id</i> does not contain any unwaited processes.</td></tr> </table>	ECHILD	The set of processes specified by <i>idtype</i> and <i>id</i> does not contain any unwaited processes.										
ECHILD	The set of processes specified by <i>idtype</i> and <i>id</i> does not contain any unwaited processes.												

EFAULT	The <i>info</i> argument points to an illegal address.
EINTR	The <code>waitid()</code> function was interrupted due to the receipt of a signal by the calling process.
EINVAL	An invalid value was specified for <i>options</i> , or <i>idtype</i> and <i>id</i> specify an invalid set of processes.

USAGE With *idtype* equal to `P_ALL` and *options* equal to `WEXITED | WTRAPPED`, `waitid()` is equivalent to `wait(2)`.

SEE ALSO `intro(3)`, `exec(2)`, `exit(2)`, `fork(2)`, `pause(2)`, `ptrace(2)`, `sigaction(2)`, `wait(2)`, `signal(3C)`, `siginfo(3HEAD)`

waitpid(2)

NAME	waitpid – wait for child process to change state						
SYNOPSIS	<pre>#include <sys/types.h> #include <sys/wait.h> pid_t waitpid(pid_t pid, int *stat_loc, int options);</pre>						
DESCRIPTION	The <code>waitpid()</code> function will suspend execution of the calling thread until status information for one of its terminated child processes is available, or until delivery of a signal whose action is either to execute a signal-catching function or to terminate the process. If more than one thread is suspended in <code>waitpid()</code> or <code>wait(2)</code> awaiting termination of the same process, exactly one thread will return the process status at the time of the target process termination. If status information is available prior to the call to <code>waitpid()</code>, return will be immediate. The <i>pid</i> argument specifies a set of child processes for which status is requested, as follows: ■ If <i>pid</i> is equal to <code>(pid_t)-1</code>, status is requested for any child process. If <i>pid</i> is greater than <code>(pid_t)0</code>, it specifies the process ID of the child process for which status is requested. ■ If <i>pid</i> is equal to <code>(pid_t)0</code> status is requested for any child process whose process group ID is equal to that of the calling process. ■ If <i>pid</i> is less than <code>(pid_t)-1</code>, status is requested for any child process whose process group ID is equal to the absolute value of <i>pid</i>. If the calling process has <code>SA_NOCLDWAIT</code> set or has <code>SIGCHLD</code> set to <code>SIG_IGN</code> and the process has no unwaited children that were transformed into zombie processes, it will block until all of its children terminate, and <code>waitpid()</code> will fail and set <code>errno</code> to <code>ECHILD</code>. If <code>waitpid()</code> returns because the status of a child process is available, then that status may be evaluated with the macros defined by <code>wstat(3XFN)</code>. If the calling process had specified a non-zero value of <i>stat_loc</i>, the status of the child process will be stored in the location pointed to by <i>stat_loc</i>. The <i>options</i> argument is constructed from the bitwise inclusive OR of zero or more of the following flags, defined in the header <code><sys/wait.h></code>: <table><tr><td><code>WCONTINUED</code></td><td>The status of any continued child process specified by <i>pid</i>, whose status has not been reported since it continued, is also reported to the calling process.</td></tr><tr><td><code>WNOHANG</code></td><td><code>waitpid()</code> will not suspend execution of the calling process if status is not immediately available for one of the child processes specified by <i>pid</i>.</td></tr><tr><td><code>WNOWAIT</code></td><td>Keep the process whose status is returned in <i>stat_loc</i> in a waitable state. The process may be waited for again with identical results.</td></tr></table>	<code>WCONTINUED</code>	The status of any continued child process specified by <i>pid</i> , whose status has not been reported since it continued, is also reported to the calling process.	<code>WNOHANG</code>	<code>waitpid()</code> will not suspend execution of the calling process if status is not immediately available for one of the child processes specified by <i>pid</i> .	<code>WNOWAIT</code>	Keep the process whose status is returned in <i>stat_loc</i> in a waitable state. The process may be waited for again with identical results.
<code>WCONTINUED</code>	The status of any continued child process specified by <i>pid</i> , whose status has not been reported since it continued, is also reported to the calling process.						
<code>WNOHANG</code>	<code>waitpid()</code> will not suspend execution of the calling process if status is not immediately available for one of the child processes specified by <i>pid</i> .						
<code>WNOWAIT</code>	Keep the process whose status is returned in <i>stat_loc</i> in a waitable state. The process may be waited for again with identical results.						

	WUNTRACED	The status of any child processes specified by <i>pid</i> that are stopped, and whose status has not yet been reported since they stopped, is also reported to the calling process.				
RETURN VALUES	If <code>waitpid()</code> returns because the status of a child process is available, it returns a value equal to the process ID of the child process for which status is reported. If <code>waitpid()</code> returns due to the delivery of a signal to the calling process, <code>-1</code> is returned and <code>errno</code> is set to <code>EINTR</code> . If <code>waitpid()</code> was invoked with <code>WNOHANG</code> set in <i>options</i> , it has at least one child process specified by <i>pid</i> for which status is not available, and status is not available for any process specified by <i>pid</i> , then <code>0</code> is returned. Otherwise, <code>-1</code> is returned and <code>errno</code> is set to indicate the error.					
ERRORS	The <code>waitpid()</code> function will fail if:					
	ECHILD	The process or process group specified by <i>pid</i> does not exist or is not a child of the calling process or can never be in the states specified by <i>options</i> .				
	EINTR	The <code>waitpid()</code> function was interrupted due to the receipt of a signal sent by the calling process.				
	EINVAL	An invalid value was specified for <i>options</i> .				
USAGE	With <i>options</i> equal to <code>0</code> and <i>pid</i> equal to <code>(pid_t)-1</code> , <code>waitpid()</code> is identical to <code>wait(2)</code> .					
ATTRIBUTES	See <code>attributes(5)</code> for descriptions of the following attributes:					
	<table><tr><th>ATTRIBUTE TYPE</th><th>ATTRIBUTE VALUE</th></tr><tr><td>MT-Level</td><td>Async-Signal-Safe</td></tr></table>		ATTRIBUTE TYPE	ATTRIBUTE VALUE	MT-Level	Async-Signal-Safe
ATTRIBUTE TYPE	ATTRIBUTE VALUE					
MT-Level	Async-Signal-Safe					
SEE ALSO	<code>intro(3)</code> , <code>exec(2)</code> , <code>exit(2)</code> , <code>fork(2)</code> , <code>pause(2)</code> , <code>ptrace(2)</code> , <code>sigaction(2)</code> , <code>wait(2)</code> , <code>signal(3C)</code> , <code>attributes(5)</code> , <code>siginfo(3HEAD)</code> , <code>wstat(3XFN)</code>					

write(2)

NAME	write, pwrite, writev – write on a file
SYNOPSIS	<pre>#include <unistd.h> ssize_t write(int fildes, const void *buf, size_t nbyte); ssize_t pwrite(int fildes, const void *buf, size_t nbyte, off_t offset); #include <sys/uio.h> ssize_t writev(int fildes, const struct iovec *iov, int iovcnt);</pre>
DESCRIPTION	The <code>write()</code> function attempts to write <i>nbyte</i> bytes from the buffer pointed to by <i>buf</i> to the file associated with the open file descriptor, <i>fildes</i>. If <i>nbyte</i> is 0, <code>write()</code> will return 0 and have no other results if the file is a regular file; otherwise, the results are unspecified. On a regular file or other file capable of seeking, the actual writing of data proceeds from the position in the file indicated by the file offset associated with <i>fildes</i>. Before successful return from <code>write()</code>, the file offset is incremented by the number of bytes actually written. On a regular file, if this incremented file offset is greater than the length of the file, the length of the file will be set to this file offset. If the <code>O_SYNC</code> flag of the file status flags is set and <i>fildes</i> refers to a regular file, a successful <code>write()</code> does not return until the data is delivered to the underlying hardware. If <i>fildes</i> refers to a socket, <code>write()</code> is equivalent to <code>send(3SOCKET)</code> with no flags set. On a file not capable of seeking, writing always takes place starting at the current position. The value of a file offset associated with such a device is undefined. If the <code>O_APPEND</code> flag of the file status flags is set, the file offset will be set to the end of the file prior to each write and no intervening file modification operation will occur between changing the file offset and the write operation. For regular files, no data transfer will occur past the offset maximum established in the open file description with <i>fildes</i>. A <code>write()</code> to a regular file is blocked if mandatory file/record locking is set (see <code>chmod(2)</code>), and there is a record lock owned by another process on the segment of the file to be written: ■ If <code>O_NDELAY</code> or <code>O_NONBLOCK</code> is set, <code>write()</code> returns -1 and sets <code>errno</code> to <code>EAGAIN</code>. ■ If <code>O_NDELAY</code> and <code>O_NONBLOCK</code> are clear, <code>write()</code> sleeps until all blocking locks are removed or the <code>write()</code> is terminated by a signal. If a <code>write()</code> requests that more bytes be written than there is room for—for example, if the write would exceed the process file size limit (see <code>getrlimit(2)</code> and <code>ulimit(2)</code>), the system file size limit, or the free space on the device—only as many

bytes as there is room for will be written. For example, suppose there is space for 20 bytes more in a file before reaching a limit. A `write()` of 512-bytes returns 20. The next `write()` of a non-zero number of bytes gives a failure return (except as noted for pipes and FIFO below).

If `write()` is interrupted by a signal before it writes any data, it will return `-1` with `errno` set to `EINTR`.

If `write()` is interrupted by a signal after it successfully writes some data, it will return the number of bytes written.

If the value of *nbyte* is greater than `SSIZE_MAX`, the result is implementation-dependent.

After a `write()` to a regular file has successfully returned:

- Any successful `read(2)` from each byte position in the file that was modified by that write will return the data specified by the `write()` for that position until such byte positions are again modified.
- Any subsequent successful `write()` to the same byte position in the file will overwrite that file data.

Write requests to a pipe or FIFO are handled the same as a regular file with the following exceptions:

- There is no file offset associated with a pipe, hence each write request appends to the end of the pipe.
- Write requests of `{PIPE_BUF}` bytes or less are guaranteed not to be interleaved with data from other processes doing writes on the same pipe. Writes of greater than `{PIPE_BUF}` bytes may have data interleaved, on arbitrary boundaries, with writes by other processes, whether or not the `O_NONBLOCK` or `O_NDELAY` flags are set.
- If `O_NONBLOCK` and `O_NDELAY` are clear, a write request may cause the process to block, but on normal completion it returns *nbyte*.
- If `O_NONBLOCK` and `O_NDELAY` are set, `write()` does not block the process. If a `write()` request for `PIPE_BUF` or fewer bytes succeeds completely `write()` returns *nbyte*. Otherwise, if `O_NONBLOCK` is set, it returns `-1` and sets `errno` to `EAGAIN` or if `O_NDELAY` is set, it returns 0. A `write()` request for greater than `{PIPE_BUF}` bytes transfers what it can and returns the number of bytes written or it transfers no data and, if `O_NONBLOCK` is set, returns `-1` with `errno` set to `EAGAIN` or if `O_NDELAY` is set, it returns 0. Finally, if a request is greater than `PIPE_BUF` bytes and all data previously written to the pipe has been read, `write()` transfers at least `PIPE_BUF` bytes.

When attempting to write to a file descriptor (other than a pipe, a FIFO, a socket, or a STREAM) that supports nonblocking writes and cannot accept the data immediately:

- If `O_NONBLOCK` and `O_NDELAY` are clear, `write()` blocks until the data can be accepted.

write(2)

- If `O_NONBLOCK` or `O_NDELAY` is set, `write()` does not block the process. If some data can be written without blocking the process, `write()` writes what it can and returns the number of bytes written. Otherwise, if `O_NONBLOCK` is set, it returns `-1` and sets `errno` to `EAGAIN` or if `O_NDELAY` is set, it returns `0`.

Upon successful completion, where *nbyte* is greater than 0, `write()` will mark for update the `st_ctime` and `st_mtime` fields of the file, and if the file is a regular file, the `S_ISUID` and `S_ISGID` bits of the file mode may be cleared.

For STREAMS files (see `intro(3)` and `streamio(7I)`), the operation of `write()` is determined by the values of the minimum and maximum *nbyte* range ("packet size") accepted by the STREAM. These values are contained in the topmost STREAM module, and can not be set or tested from user level. If *nbyte* falls within the packet size range, *nbyte* bytes are written. If *nbyte* does not fall within the range and the minimum packet size value is zero, `write()` breaks the buffer into maximum packet size segments prior to sending the data downstream (the last segment may be smaller than the maximum packet size). If *nbyte* does not fall within the range and the minimum value is non-zero, `write()` fails and sets `errno` to `ERANGE`. Writing a zero-length buffer (*nbyte* is zero) to a STREAMS device sends a zero length message with zero returned. However, writing a zero-length buffer to a pipe or FIFO sends no message and zero is returned. The user program may issue the `I_SWROPT` `ioctl(2)` to enable zero-length messages to be sent across the pipe or FIFO (see `streamio(7I)`).

When writing to a STREAM, data messages are created with a priority band of zero. When writing to a socket or to a STREAM that is not a pipe or a FIFO:

- If `O_NDELAY` and `O_NONBLOCK` are not set, and the STREAM cannot accept data (the STREAM write queue is full due to internal flow control conditions), `write()` blocks until data can be accepted.
- If `O_NDELAY` or `O_NONBLOCK` is set and the STREAM cannot accept data, `write()` returns `-1` and sets `errno` to `EAGAIN`.
- If `O_NDELAY` or `O_NONBLOCK` is set and part of the buffer has already been written when a condition occurs in which the STREAM cannot accept additional data, `write()` terminates and returns the number of bytes written.

The `write()` and `writew()` functions will fail if the STREAM head had processed an asynchronous error before the call. In this case, the value of `errno` does not reflect the result of `write()` or `writew()` but reflects the prior error.

pwrite() The `pwrite()` function performs the same action as `write()`, except that it writes into a given position without changing the file pointer. The first three arguments to `pwrite()` are the same as `write()` with the addition of a fourth argument *offset* for the desired position inside the file.

writew() The `writew()` function performs the same action as `write()`, but gathers the output data from the *iovcnt* buffers specified by the members of the *iov* array: *iov*[0], *iov*[1], ..., *iov*[*iovcnt* - 1]. The *iovcnt* buffer is valid if greater than 0 and less than or equal to `{IOV_MAX}`. See `intro(3)` for a definition of `{IOV_MAX}`.

The `iovec` structure contains the following members:

```
caddr_t  iov_base;
int      iov_len;
```

Each `iovec` entry specifies the base address and length of an area in memory from which data should be written. The `writew()` function always writes all data from an area before proceeding to the next.

If *fildev* refers to a regular file and all of the `iov_len` members in the array pointed to by *iov* are 0, `writew()` will return 0 and have no other effect. For other file types, the behavior is unspecified.

If the sum of the `iov_len` values is greater than `SSIZE_MAX`, the operation fails and no data is transferred.

RETURN VALUES

Upon successful completion, `write()` returns the number of bytes actually written to the file associated with *fildev*. This number is never greater than *nbyte*. Otherwise, `-1` is returned, the file-pointer remains unchanged, and `errno` is set to indicate the error.

Upon successful completion, `writew()` returns the number of bytes actually written. Otherwise, it returns `-1`, the file-pointer remains unchanged, and `errno` is set to indicate an error.

ERRORS

The `write()`, `pwrite()`, and `writew()` functions will fail if:

EAGAIN	Mandatory file/record locking is set, <code>O_NDELAY</code> or <code>O_NONBLOCK</code> is set, and there is a blocking record lock; an attempt is made to write to a <code>STREAM</code> that can not accept data with the <code>O_NDELAY</code> or <code>O_NONBLOCK</code> flag set; or a write to a pipe or FIFO of <code>PIPE_BUF</code> bytes or less is requested and less than <i>nbytes</i> of free space is available.
EBADF	The <i>fildev</i> argument is not a valid file descriptor open for writing.
EDEADLK	The write was going to go to sleep and cause a deadlock situation to occur.
EDQUOT	The user's quota of disk blocks on the file system containing the file has been exhausted.
EFAULT	The <i>buf</i> argument points to an illegal address.
EFBIG	An attempt is made to write a file that exceeds the process's file size limit or the maximum file size (see <code>getrlimit(2)</code> and <code>ulimit(2)</code>).
EFBIG	The file is a regular file, <i>nbyte</i> is greater than 0, and the starting position is greater than or equal to the offset maximum established in the file description associated with <i>fildev</i> .
EINTR	A signal was caught during the write operation and no data was transferred.

write(2)

EIO	The process is in the background and is attempting to write to its controlling terminal whose TOSTOP flag is set, or the process is neither ignoring nor blocking SIGTTOU signals and the process group of the process is orphaned.
ENOLCK	Enforced record locking was enabled and {LOCK_MAX} regions are already locked in the system, or the system record lock table was full and the write could not go to sleep until the blocking record lock was removed.
ENOLINK	The <i>fildev</i> argument is on a remote machine and the link to that machine is no longer active.
ENOSPC	During a write to an ordinary file, there is no free space left on the device.
ENOSR	An attempt is made to write to a STREAMS with insufficient STREAMS memory resources available in the system.
ENXIO	A hangup occurred on the STREAM being written to.
EPIPE	An attempt is made to write to a pipe or a FIFO that is not open for reading by any process, or that has only one end open (or to a file descriptor created by <code>socket(3SOCKET)</code> , using type <code>SOCK_STREAM</code> that is no longer connected to a peer endpoint). A <code>SIGPIPE</code> signal will also be sent to the process. The process dies unless special provisions were taken to catch or ignore the signal.
ERANGE	The transfer request size was outside the range supported by the STREAMS file associated with <i>fildev</i> .

The `pwrite()` function fails and the file pointer remains unchanged if:

ESPIPE	The <i>fildev</i> argument is associated with a pipe or FIFO.
--------	---

The `writew()` function will fail if:

EINVAL	The sum of the <code>iov_len</code> values in the <i>iov</i> array would overflow an <code>ssize_t</code> .
--------	---

The `write()` and `writew()` functions may fail if:

EINVAL	The STREAM or multiplexer referenced by <i>fildev</i> is linked (directly or indirectly) downstream from a multiplexer.
ENXIO	A request was made of a non-existent device, or the request was outside the capabilities of the device.
ENXIO	A hangup occurred on the STREAM being written to.

A write to a STREAMS file may fail if an error message has been received at the STREAM head. In this case, `errno` is set to the value included in the error message.

The `writew()` function may fail if:

EINVAL The *iovcnt* argument was less than or equal to 0 or greater than {IOV_MAX}; one of the *iov_len* values in the *iov* array was negative; or the sum of the *iov_len* values in the *iov* array overflowed an int.

USAGE The `pwrite()` function has a transitional interface for 64-bit file offsets. See `lf64(5)`.

ATTRIBUTES See `attributes(5)` for descriptions of the following attributes:

ATTRIBUTE TYPE	ATTRIBUTE VALUE
MT-Level	<code>write()</code> is Async-Signal-Safe

SEE ALSO `Intro(3)`, `chmod(2)`, `creat(2)`, `dup(2)`, `fcntl(2)`, `getrlimit(2)`, `ioctl(2)`, `lseek(2)`, `open(2)`, `pipe(2)`, `ulimit(2)`, `send(3SOCKET)`, `socket(3SOCKET)`, `attributes(5)`, `lf64(5)`, `streamio(7I)`

yield(2)

NAME	yield – yield execution to another lightweight process
SYNOPSIS	<pre>#include <unistd.h> void yield(void);</pre>
DESCRIPTION	The <code>yield()</code> function causes the current lightweight process to yield its execution in favor of another lightweight process with the same or greater priority.
SEE ALSO	<code>thr_yield(3THR)</code>

Index

Numbers and Symbols

install a SPARC V9 user trap handler, 272

A

access — determine accessibility of a file, 40
access permission mode of file, change —
 chmod, 59
accounting, enable or disable process
 accounting — acct, 42
acct — enable or disable process accounting, 42
acl — get or set a file's Access Control List
 (ACL), 43
adjtime — correct the time to allow
 synchronization of the system clock, 45
adjust local clock parameters —
 ntp_adjtime, 175
alarm — set a process alarm clock, 46
audit — write an audit record, 48
auditon — manipulate auditing, 49
auditsvc() function, 53

B

bind LWPs to a processor —
 processor_bind, 202
bind LWPs to a set of processors —
 pset_bind, 207
brk — change the amount of space allocated for
 the calling process's data segment, 55

C

chdir — change working directory, 57
child processes
 allows a parent process to control the
 execution of a child process — ptrace, 212
 get time — times, 295
 wait for child process to change state —
 waitid, 316
 wait for child process to change state —
 waitpid, 318
 wait for child process to stop or terminate —
 wait, 314
chmod — change access permission mode of
 file, 59
chown — change owner and group of a file, 62
chroot — change root directory, 65
clock, get local clock values — ntp_gettime, 176
continue or suspend LWP execution
 — _lwp_continue, 139
 — _lwp_suspend, 139
CPU-use, process execution time profile —
 profil, 205
creat — create a new file or rewrite an existing
 one, 69
create a new process — fork, 89
 fork1, 89
create a new light-weight process —
 _lwp_create, 129
create session and set process group ID —
 setsid, 243

D

deliver process signals to specific LWPs
— `_lwp_sigredirect`, 260
— `_signotifywait`, 260
determine accessibility of a file — `access`, 40
devices, I/O control functions — `ioctl`, 116
directories
 change working directory — `chdir`, 57
 create a new one — `mknod`, 148
 get configurable pathname variables —
 `pathconf`, 92
 make a new one — `mkdir`, 146
 read directory entries and put in a file system
 independent format — `getdents`, 101
 remove — `rmdir`, 227
discover all message queue identifiers —
 `msgids`, 164
discover all shared memory identifiers —
 `shmids`, 251
`dup` — duplicate an open file descriptor, 71

E

effective group ID, set — `setregid()`, 241
effective user ID, set — `setreuid()`, 242
`exec` — execute a file, 72
`execl` — execute a file, 72
`execle` — execute a file, 72
`execvp` — execute a file, 72
`execv` — execute a file, 72
`execve` — execute a file, 72
`execvp` — execute a file, 72
`exit` — terminate process, 78
`_exit` — terminate process, 78

F

`facl` — get or set a file's Access Control List
 (ACL), 43
`fchdir` — change working directory, 57
`fchmod` — change access permission mode of
 file, 59
`fchown` — change owner and group of a file, 62
`fcntl` — file control, 81
file control — `fcntl`, 81

file descriptor, duplicate an open one —
 `dup`, 71
file pointer, read/write
 move — `lseek`, 122, 123
file status, get — `stat`, `lstat`, `fstat`, 277
file system
 get information — `statvfs`, `fstatvfs`, 280
 get statistics — `ustat`, 306
 make a symbolic link to a file —
 `symlink`, 287
 remove link — `unlink`, 304
 returns information about the file system
 types configured in the system —
 `sysfs`, 290
 update super block — `sync`, 289
files
 change access permission mode of file —
 `chmod`, 59
 change owner and group of a file —
 `chown`, 62
 change the name of a file — `rename`, 223
 create a new file or rewrite an existing one —
 `creat`, 69
 execute — `exec`, 72
 get configurable pathname variables —
 `pathconf`, 92
 link to a file — `link`, 120
 move read/write file pointer — `lseek`, 122,
 123
 set file access and modification times —
 `utime`, 307
`fork` — create a new process, 89
`fork`, spawn new process in a virtual memory
 efficient way — `vfork`, 311
`fork1` — create a new process, 89
`fpathconf` — get configurable pathname
 variables, 92
`fstat` — get status on open file known by file
 descriptor, 277
`fstatvfs` — get file system information, 280

G

get and set process audit information —
 `getaudit`, 97
get and set process audit information —
 `getaudit_addr`, 97

- get and set process audit information —
 setaudit, 97
- get and set process audit information —
 setaudit_addr, 97
- get, put, or write extended accounting data —
 getacct, 95
- get, put, or write extended accounting data —
 putacct, 95
- get, put, or write extended accounting data —
 wraacct, 95
- get and set process limits — ulimit, 299
- get information about a processor set —
 pset_info, 211
- get LWP identifier — _lwp_self, 136
- get or change processor operational status —
 p_online, 189
- get or set a file's Access Control List (ACL)
 — acl, 43
 — facl, 43
- get process group ID of session leader —
 getsid, 114
- getacct — get, put, or write extended
 accounting data, 95
- getaudit — get and set process audit
 information, 97
- getaudit_addr — get and set process audit
 information, 97
- getaudit — get user audit identity, 99
- getdents — read directory entries and put in a
 file system independent format, 101
- getegid — get effective group ID, 115
- geteuid — get effective user ID, 115
- getgid — get real group ID, 115
- getgroups — get supplementary group access
 list IDs, 102
- getitimer — get value of interval timer, 103
- getmsg — get next message off a stream, 107
- getpgid — get process group IDs, 110
- getpgrp — get process group IDs, 110
- getpid — get process IDs, 110
- getpmsg — get next message off a stream, 107
- getppid — get parent process IDs, 110
- getprojid — set or get task or project IDs, 244
- getrlimit — control maximum system resource
 consumption, 111
- getsid — get process group ID of session
 leader, 114
- gettaskid — set or get task or project IDs, 244

- getuid — get real user ID, 115
- group ID, set real and effective —
 setregid(), 241
- group IDs
 - get — getgid, getegid, 115
 - set — setgid, 245
 - supplementary group access list IDs —
 getgroups, setgroups, 102

H

- halt system, — uadmin, 297
- hangup signal, the current controlling terminal
 — vhangup, 313

I

- I/O
 - audit — audit, 48
 - multiplexing — poll, 186
- initialize an LWP context —
 _lwp_makecontext, 134
- interprocess communication, — pipe, 185
- interval timer, get or set value of interval timer
 — getitimer, setitimer, 103
- ioctl — control device, 116

K

- kill — send a signal to a process or a group of
 processes, 118

L

- lchown — change owner and group of a file, 62
- link — link to a file, 120
- link, remove — unlink, 304
- link, symbolic, make one to a file —
 symlink, 287
- lseek — move extended read/write file
 pointer, 122
- lseek — move read/write file pointer, 123
- lstat — get status on symbolic link file, 277
- LWP, scheduler control — priocntl, 191

- `_lwp_cond_broadcast` — signal a condition variable, 125
- `_lwp_cond_reltimedwait` — wait on a condition variable, 126
- `_lwp_cond_signal` — signal a condition variable, 125
- `_lwp_cond_timedwait` — wait on a condition variable, 126
- `_lwp_cond_wait` — wait on a condition variable, 126
- `_lwp_continue` — continue or suspend LWP execution, 139
- `_lwp_create` — create a new light-weight process, 129
- `_lwp_exit` — terminate the calling LWP, 131
- `_lwp_getprivate` — set/get LWP specific storage, 138
- `_lwp_info` — return the time-accounting information of a single LWP, 132
- `_lwp_kill` — send a signal to a LWP, 133
- `_lwp_makecontext` — initialize an LWP context, 134
- `_lwp_mutex_lock` — mutual exclusion, 135
- `_lwp_mutex_trylock` — mutual exclusion, 135
- `_lwp_mutex_unlock` — mutual exclusion, 135
- `_lwp_self` — get LWP identifier, 136
- `_lwp_sema_init` — semaphore operations, 137
- `_lwp_sema_post` — semaphore operations, 137
- `_lwp_sema_trywait` — semaphore operations, 137
- `_lwp_sema_wait` — semaphore operations, 137
- `_lwp_setprivate` — set/get LWP specific storage, 138
- `_lwp_sigredirect` — deliver process signals to specific LWPs, 260
- `_lwp_suspend` — continue or suspend LWP execution, 139
- `_lwp_wait` — wait for a LWP to terminate, 140

M

- make a directory, or a special or ordinary file — `mknod`, 148
- manage sets of processors
 - `pset_assign`, 209
 - `pset_create`, 209
 - `pset_destroy`, 209

- manipulate auditing — `auditon`, 49
- masks, set and get file creation mask — `umask`, 300
- `memcntl` — memory management control, 141
- memory, management control — `memcntl`, 141
- memory, shared
 - control operations — `shmctl`, 247
 - get segment identifier — `shmget`, 249
 - operations — `shmop`, 253
- memory management, change the amount of space allocated for the calling process's data segment — `brk`, `sbrk`, 55
- memory mapping, set protection — `mprotect`, 160
- memory pages
 - determine residency — `mincore`, 145
 - map — `mmap`, 151
 - unmap — `munmap`, 173
- message control operations, — `msgctl`, 161
- message queue, get — `msgget`, 163
- message queue snapshot operation — `msgsnap`, 168
- message receive operation — `msgrcv`, 166
- message send operation — `msgsnd`, 171
- messages, send a message on a stream — `putmsg`, 214
- `mincore` — determine residency of memory pages, 145
- `mkdir` — make a directory, 146
- `mknod` — make a directory, or a special or ordinary file, 148
- `mmap` — map pages of memory, 151
- `mount` — mount a file system, 157
- mount a file system — `mount`, 157
- `mprotect` — set protection of memory mapping, 160
- `msgctl` — message control operations, 161
- `msgget` — get message queue, 163
- `msgids` — discover all message queue identifiers, 164
- `msgrcv` — message receive operation, 166
- `msgsnap` — message queue snapshot operation, 168
- `msgsnd` — message send operation, 171
- `munmap` — unmap pages of memory, 173
- mutual exclusion
 - `_lwp_mutex_lock`, 135
 - `_lwp_mutex_trylock`, 135

mutual exclusion (Continued)
— `_lwp_mutex_unlock`, 135

N

`nice` — change priority of a time-sharing process, 174
`ntp_adjtime` — adjust local clock parameters, 175
`ntp_gettime` — get local clock values, 176

O

`open` — open a file, 177
`open` a file — `open`, 177
operating system, get name of current one — `uname`, 303
owner of file, change — `chown`, 62

P

`p_online` — get or change processor operational status, 189
`pathconf` — get configurable pathname variables, 92
`pathname`, get configurable variables — `pathconf`, 92
`pause` — suspend process until signal, 183
`pipe` — create an interprocess channel, 185
`poll` — input/output multiplexing, 186
`pread` — read from file, 217
`prctl` — process scheduler control, 191
`prctlset` — generalized process scheduler control, 200
process, time-sharing, change priority — `nice`, 174
process accounting, enable or disable — `acct`, 42
process alarm clock, set — `alarm`, 46
process group ID
 set — `setpgid`, 239
 set — `setpgrp`, 240
process scheduler
 control — `prctl`, 191
 generalized control — `prctlset`, 200

process statistics, process execution time profile
— `profil`, 205

processes

allows a parent process to control the
 execution of a child process — `ptrace`, 212
change priority of a time-sharing process —
 `nice`, 174
create a new one — `fork`, 89
create an interprocess channel — `pipe`, 185
execute a file — `exec`, 72
execution time profile — `profil`, 205
generalized scheduler control —
 `prctlset`, 200
get identification — `getpid`, `getpgrp`, `getppid`,
 `getpgid`, 110
get next message off a stream — `getmsg`, 107
get or set value of interval timer — `getitimer`,
 `setitimer`, 103
get real user, effective user, real group, and
 effective group IDs — `getuid`, `geteuid`,
 `getgid`, `getegid`, 115
get time — `times`, 295
read from file — `read`, 217
read directory entries and put in a file system
 independent format — `getdents`, 101
send a signal to a process or a group of
 processes — `kill`, 118
set a process alarm clock — `alarm`, 46
set and get file creation mask — `umask`, 300
set process group ID — `setpgid`, 239
set process group ID — `setpgrp`, 240
spawn new process in a virtual memory
 efficient way — `vfork`, 311
supplementary group access list IDs —
 `getgroups`, `setgroups`, 102
suspend process until signal — `pause`, 183
the current controlling terminal —
 `vhangup`, 313
wait for child process to change state —
 `waitid`, 316
wait for child process to change state —
 `waitpid`, 318
wait for child process to stop or terminate —
 `wait`, 314
processes and protection
 — `setregid()`, 241
 — `setreuid()`, 242

processor_bind — bind LWPs to a processor, 202
 processor_info — determine type and status of a processor, 204
 profil — process execution time profile, 205
 profiling utilities, execution time profile — profil, 205
 pset_assign — manage sets of processors, 209
 pset_bind — bind LWPs to a set of processors, 207
 pset_create — manage sets of processors, 209
 pset_destroy — manage sets of processors, 209
 pset_info — get information about a processor set, 211
 ptrace — allows a parent process to control the execution of a child process, 212
 putacct — get, put, or write extended accounting data, 95
 putmsg — send a message on a stream, 214
 putpmsg — send a message on a stream, 214
 pwrite — write on a file, 320

R

read from file — read, 217
 pread, 217
 readv, 217
 read the contents of a symbolic link — readlink, 222
 read/write file pointer
 move — lseek, 122, 123
 readlink — read the contents of a symbolic link, 222
 read — read from file, 217
 real group ID, set — setregid(), 241
 real user ID, set — setreuid(), 242
 reboot system, — uadmin, 297
 remount root file system, — uadmin, 297
 rename — change the name of a file, 223
 resolve all symbolic links of a path name — resolvepath, 226
 resolvepath — resolve all symbolic links of a path name, 226
 rmdir — remove a directory, 227
 root directory, change — chroot, 65

S

sbrk — change the amount of space allocated for the calling process's data segment, 55
 semaphore operations — semop, 236
 semaphore operations — semtimedop, 236
 semaphore operations
 — _lwp_sema_init, 137
 — _lwp_sema_post, 137
 — _lwp_sema_trywait, 137
 — _lwp_sema_wait, 137
 semaphores
 control operations — semctl, 229
 get a set — semget, 232
 semctl — semaphore control operations, 229
 semget — get set of semaphores, 232
 semop — semaphore operations, 236
 semtimedop — semaphore operations, 236
 send a signal to a LWP — _lwp_kill, 133
 set or get task or project IDs — getprojid, 244
 set or get task or project IDs — gettaskid, 244
 set or get task or project IDs — settaskid, 244
 set file access and modification times — utimes, 309
 set/get LWP specific storage
 — _lwp_getprivate, 138
 — _lwp_setprivate, 138
 setaudit — get and set process audit information, 97
 setaudit_addr — get and set process audit information, 97
 setaudit — set user audit identity, 99
 setegid — set effective group ID, 245
 seteuid — set effective user ID, 245
 setgid — set group ID, 245
 setgroups — set supplementary group access list IDs, 102
 setitimer — set value of interval timer, 103
 setpgid — set process group ID, 239
 setpgrp — set process group ID, 240
 setregid() — set real and effective group ID, 241
 setreuid() — set real and effective user IDs, 242
 setrlimit — control maximum system resource consumption, 111
 setsid — create session and set process group ID, 243
 settaskid — set or get task or project IDs, 244
 setuid — set user ID, 245

- shared memory
 - control operations — shmctl, 247
 - get segment identifier — shmget, 249
 - operations — shmop, 253
- shmctl — shared memory control
 - operations, 247
- shmget — get shared memory segment identifier, 249
- shmids — discover all shared memory identifiers, 251
- shmop — shared memory operations, 253
- shutdown, — uadmin, 297
- sigaction — detailed signal management, 255
- sigaltstack — set or get signal alternate stack context, 258
- signal a condition variable
 - _lwp_cond_broadcast, 125
 - _lwp_cond_signal, 125
- signal alternate stack, set or get context — sigaltstack, 258
- signal management, detailed — sigaction, 255
- signal mask
 - change and/or examine — sigprocmask, 263
 - install, and suspend process until signal — sigsuspend, 266
- signals, examine blocked and pending ones — sigpending, 262
- _signotifywait — deliver process signals to specific LWPs, 260
- sigpending — examine signals that are blocked and pending, 262
- sigprocmask — change and/or examine calling process's signal mask, 263
- sigsend — send a signal to a process or a group of processes, 264
- sigsendset — provides an alternate interface to sigsend for sending signals to sets of processes, 264
- sigsuspend — install a signal mask and suspend process until signal, 266
- sigwait() — wait until a signal is posted, 268
- special files, create a new one — mknod, 148
- stat — get file status, 277
- statistics, get for mounted file system — ustat, 306
- statvfs — get file system information, 280
- stime — set system time and date, 282

STREAMS

- get next message off a stream — getmsg, 107
- I/O control functions — ioctl, 116
- send a message on a stream — putmsg, 214
- super block, update — sync, 289
- swap space, manage — swapctl, 283
- swapctl — manage swap space, 283
- symbolic link, make one to a file —
 - symlink, 287
- symlink — make a symbolic link to a file, 287
- sync — update super block, 289
- sysinfo — get and set system information strings, 291
- system administration, administrative control
 - uadmin, 297
- system clock, synchronization — adjtime, 45
- system information, get and set strings — sysinfo, 291
- system operation, update super block — sync, 289
- system resources, control maximum system resource consumption — getrlimit, setrlimit, 111

T

- terminate process
 - _exit, 78
 - exit, 78
- terminate the calling LWP — _lwp_exit, 131
- time — get time, 294
- time
 - correct the time to allow synchronization of the system clock — adjtime, 45
 - set system time and date — stime, 282
- time-accounting, single LWP — _lwp_info, 132
- times — get process and child process times, 295

U

- ulimit — get and set process limits, 299
- umask — set and get file creation mask, 300
- umount — unmount a file system, 301
- umount2 — unmount a file system, 301

- uname — get name of current operating system, 303
- unlink — remove directory entry, 304
- unmount a file system — umount2, 301
- unmount a file system — umount, 301
- user audit identity
 - get user audit identity — getaudit, 99
 - set user audit identity — setaudit, 99
- user ID, set real and effective — setreuid(), 242
- user IDs
 - get — getuid, geteuid, 115
 - set — setuid, 245
- utime — set file access and modification times, 307
- utimes — set file access and modification times, 309

V

- vfork — spawn new process in a virtual memory efficient way, 311
- vhangup — the current controlling terminal, 313

W

- wait — wait for child process to stop or terminate, 314
- wait on a condition variable —
 - _lwp_cond_reltimedwait, 126
- wait on a condition variable —
 - _lwp_cond_timedwait, 126
- wait on a condition variable —
 - _lwp_cond_wait, 126
- wait for a LWP to terminate — _lwp_wait, 140
- waitid — wait for child process to change state, 316
- waitpid — wait for child process to change state, 318
- wracct — get, put, or write extended accounting data, 95
- write on a file
 - write, 320
- write — write on a file, 320

Y

- yield — yield execution to another lightweight process, 326
- yield execution to another lightweight process
 - yield, 326