
12327 Sun Mar 4 06:09:08 2018
new/usr/src/man/man1m/mount_smbfs.1m
Convert mount_smbfs.1m to mdoc

```

1  \"
1  ' te
2  \" Copyright (c) 2008, Sun Microsystems, Inc. All Rights Reserved.
3  \" Portions Copyright 1994-2008 The FreeBSD Project. All rights reserved.
4  \" Redistribution and use in source and binary forms, with or without modificat
5  \" the following disclaimer. 2. Redistributions in binary form must reproduce t
6  \" BY THE AUTHOR AND CONTRIBUTORS \"AS IS\" AND ANY EXPRESS OR IMPLIED WARRANTIES
7  \" SPECIAL, EXEMPLARY, OR CONSEQUENTIAL DAMAGES (INCLUDING, BUT NOT LIMITED TO,
8  \" OR OTHERWISE) ARISING IN ANY WAY OUT OF THE USE OF THIS SOFTWARE, EVEN IF AD
9  \" Copyright 2012 Nexenta Systems, Inc. All rights reserved.
10 .Dd March 4, 2018
11 .Dt MOUNT_SMBFS 1M
12 .Os
13 .Sh NAME
14 .Nm mount_smbfs ,
15 .Nm umount_smbfs
16 .Nd mount and unmount a shared resource from an SMB file server
17 .Sh SYNOPSIS
18 .Nm mount
19 .Op Fl F Cm smbfs
20 .Op Ar generic-options
21 .Op Fl o Ar name Ns Oo = Ns Ar value Oc
22 .Op Fl O
23 .Ar resource
24 .Nm mount
25 .Op Fl F Cm smbfs
26 .Op Ar generic-options
27 .Op Fl o Ar name Ns Oo = Ns Ar value Oc
28 .Op Fl O
29 .Ar mount-point
30 .Nm mount
31 .Op Fl F Cm smbfs
32 .Op Ar generic-options
33 .Op Fl o Ar name Ns Oo = Ns Ar value Oc
34 .Op Fl O
35 .Ar resource mount-point
36 .Nm umount
37 .Op Fl F Cm smbfs
38 .Op Ar generic-options
39 .Ar mount-point
40 .Sh DESCRIPTION
10 .TH MOUNT_SMBFS 1M "Jan 2, 2012"
11 .SH NAME
12 mount_smbfs, umount_smbfs \- mount and unmount a shared resource from a CIFS
13 file server
14 .SH SYNOPSIS
15 .LP
16 .nf
17 \fB/sbin/mount\fR [\fB-F smbfs\fR] [\fIgeneric-options\fR] [\fB-o\fR \fIname\fR\
18 .fi

20 .LP
21 .nf
22 \fB/sbin/mount\fR [\fB-F smbfs\fR] [\fIgeneric-options\fR] [\fB-o\fR \fIname\fR\
23 .fi

25 .LP
26 .nf
27 \fB/sbin/mount\fR [\fB-F smbfs\fR] [\fIgeneric-options\fR] [\fB-o\fR \fIname\fR\
28 [\fB-O\fR] \fIresource\fR \fImount-point\fR
29 .fi

```

```

31 .LP
32 .nf
33 \fB/sbin/umount\fR [\fB-F smbfs\fR] [\fIgeneric-options\fR] \fImount-point\fR
34 .fi

36 .SH DESCRIPTION
41 .sp
42 .LP
43 The \fBmount\fR utility attaches a named resource, \fIresource\fR, to the file
44 system hierarchy at the path name location, \fImount-point\fR, which must
45 already exist.
46 .sp
47 .LP
48 The
49 .Nm mount
50 utility attaches a named resource,
51 .Ar resource ,
52 to the file system hierarchy at the path name location,
53 .Ar mount-point ,
54 which must already exist.
55 .Pp
56 If
57 .Ar mount-point
58 has any contents prior to the
59 .Nm mount
60 operation, those contents remain hidden until the resource is unmounted.
61 An authorized user with the
62 .Dv SYS_MOUNT
63 privilege can perform a
64 .Nm mount
65 operation.
66 Also, a user can perform SMBFS mount operations on a directory the user owns.
67 .Pp
68 If the resource is listed in the
69 .Pa /etc/vfstab
70 file, you can specify either
71 .Ar resource
72 or
73 .Ar mount-point
74 as the
75 .Nm mount
76 command will consult the
77 .Pa /etc/vfstab
78 file for more information.
79 If the
80 .Fl F
81 option is omitted,
82 .Nm mount
83 takes the file system type from the entry in the
84 .Pa /etc/vfstab
85 file.
86 .Pp
87 If the resource is not listed in the
88 .Pa /etc/vfstab
89 file, the command line must specify both

```

```

90 .Ar resource
91 and
92 .Ar mount-point .
93 .Pp
94 The
95 .Nm mount
96 utility detaches a mounted file system from the file system hierarchy.
97 An authorized user with the
98 .Dv SYS_MOUNT
99 privilege can perform a
100 .Nm mount
101 operation.
102 Also, a user can perform SMBFS unmount operations on a directory the user owns.
103 .Pp
104 The
105 .Em network/smb/client
106 service must be enabled to successfully mount an SMB share.
107 This service is enabled, by default.
108 .Pp
109 To enable the service, enter the following
110 .Xr svcadm 1M
111 command:
112 .Bd -literal
113 # svcadm enable network/smb/client
114 .Ed
115 .Ss Operands
116 The
117 .Nm mount
118 command supports the following operands:
119 .Bl -tag -width Ds
120 .It Xo
121 .Ar resource
122 .No // Ns Oo Ar workgroup Ns \&; Ns Oc Ns
123 .Oo Ar user Ns Oo : Ns Ar password Ns Oc Ns @ Ns Oc Ns
124 .Ar server Ns / Ns Ar share
125 .Xc
126 The name of the resource to be mounted.
127 In addition to its name, you can specify the following information about the
128 resource:
129 .Bl -bullet
130 .It
131 .Ar password
132 is the password associated with
133 .Ar user .
134 If
135 .Ar password
136 is not specified, the mount first attempts to use the password stored by the
137 .Nm smbutil Cm login
138 command (if any).
139 If that password fails to authenticate, the
140 .Nm mount_smbfs
141 prompts you for a password.
142 .It
143 .Ar server
144 is the DNS or NetBIOS name of the remote computer.
145 .It
146 .Ar share
147 is the resource name on the remote server.
148 .It
149 .Ar user
150 is the remote user name.
151 If
152 .Ar user
153 is omitted, the logged in user ID is used.
154 .It
155 .Ar workgroup

```

```

156 is the name of the workgroup or the Windows domain in which the user name is
157 defined.
158 .Pp
159 .sp
160 .LP
161 If the resource is not listed in the \fB/etc/vfstab\fR file, the command line
162 must specify both \fIresource\fR and \fImount-point\fR.
163 .sp
164 .LP
165 The \fBumount\fR utility detaches a mounted file system from the file system
166 hierarchy. An authorized user with the \fBSYS_MOUNT\fR privilege can perform a
167 \fBumount\fR operation. Also, a user can perform SMBFS unmount operations on a
168 directory the user owns.
169 .sp
170 .LP
171 The \fBnetwork/smb/client\fR service must be enabled to successfully mount a
172 CIFS share. This service is enabled, by default.
173 .sp
174 .LP
175 To enable the service, enter the following \fBsvcadm\fR(1M) command:
176 .sp
177 .in +2
178 .nf
179 # \fBsvcadm enable network/smb/client\fR
180 .fi
181 .in -2
182 .sp
183 .SS "Operands"
184 .sp
185 .LP
186 The \fBmount\fR command supports the following operands:
187 .sp
188 .ne 2
189 .na
190 \fB\fIresource\fR
191 //[\fIworkgroup\fR;][\fIuser\fR[:\fIpassword\fR@]\fIserver\fR/\fIshare\fR]\fR
192 .ad
193 .sp .6
194 .RS 4n
195 .sp
196 .LP
197 The name of the resource to be mounted. In addition to its name, you can
198 specify the following information about the resource:
199 .RS +4
200 .TP
201 .ie t \ (bu
202 .el o
203 \fIpassword\fR is the password associated with \fIuser\fR. If \fIpassword\fR is
204 not specified, the mount first attempts to use the password stored by the
205 \fBsmbutil login\fR command (if any). If that password fails to authenticate,
206 the \fBmount_smbfs\fR prompts you for a password.
207 .RE
208 .RS +4
209 .TP
210 .ie t \ (bu
211 .el o
212 \fIserver\fR is the DNS or NetBIOS name of the remote computer.
213 .RE
214 .RS +4
215 .TP
216 .ie t \ (bu
217 .el o
218 \fIshare\fR is the resource name on the remote server.
219 .RE
220 .RS +4

```

```

118 .TP
119 .ie t \(bu
120 .el o
121 \fIuser\fR is the remote user name. If \fIuser\fR is omitted, the logged in
122 user ID is used.
123 .RE
124 .RS +4
125 .TP
126 .ie t \(bu
127 .el o
128 \fIworkgroup\fR is the name of the workgroup or the Windows domain in which the
129 user name is defined.
130 .sp
131 If the resource includes a workgroup, you must escape the semicolon that
132 appears after the workgroup name to prevent it from being interpreted by the
133 command shell.
134 For instance, surround the entire resource name with double quotes:
135 .Bd -literal
136 mount -F smbfs "//SALES;george@RSERVER" /mnt
137 .Ed
138 .El
139 .It Ar mount-point
140 command shell. For instance, surround the entire resource name with double
141 quotes: \fBmount -F smbfs "//SALES;george@RSERVER" /mnt\fR.
142 .RE
143 .RE
144 .sp
145 .ne 2
146 .na
147 \fB\fImount-point\fR\fR
148 .ad
149 .sp .6
150 .RS 4n
151 The path to the location where the file system is to be mounted or unmounted.
152 The
153 .Nm mount
154 command maintains a table of mounted file systems in the
155 .Pa /etc/mnttab
156 file.
157 See the
158 .Xr mnttab 4
159 man page.
160 .El
161 .Sh OPTIONS
162 See the
163 .Xr mount 1M
164 man page for the list of supported
165 .Ar generic-options .
166 .Bl -tag -width Ds
167 .It Fl o Ar name Ns Oo = Ns Ar value Oc
168 Sets the file system-specific properties.
169 You can specify more than one name-value pair as a list of comma-separated
170 pairs.
171 No spaces are permitted in the list.
172 The properties are as follows:
173 .Bl -tag -width Ds
174 .It Cm acl Ns | Ns Cm noacl
175 The \fBmount\fR command maintains a table of mounted file systems in the
176 \fB/etc/mnttab\fR file. See the \fBmnttab\fR(4) man page.
177 .RE
178 .SH OPTIONS
179 .sp
180 .LP
181 See the \fBmount\fR(1M) man page for the list of supported

```

```

154 \fIgeneric-options\fR.
155 .sp
156 .ne 2
157 .na
158 \fB\fB-o\fR \fIname\fR\fB=\fR\fIvalue\fR or\fR
159 .ad
160 .br
161 .na
162 \fB\fB-o\fR \fIname\fR\fR
163 .ad
164 .sp .6
165 .RS 4n
166 Sets the file system-specific properties. You can specify more than one
167 name-value pair as a list of comma-separated pairs. No spaces are permitted in
168 the list. The properties are as follows:
169 .sp
170 .ne 2
171 .na
172 \fB\fBacl\fR|\fBnoacl\fR\fR
173 .ad
174 .sp .6
175 .RS 4n
176 Enable (or disable) presentation of Access Control Lists (ACLs)
177 on files and directories under this
178 .Xr smbfs 7FS
179 mount.
180 The default behavior is
181 .Cm noacl ,
182 which presents files and directories as owned by the owner of the mount point
183 and having permissions based on
184 .Cm fileperms
185 or
186 .Cm dirperms .
187 With the
188 .Cm acl
189 mount option, files are presented with ACLs obtained from the SMB server.
190 Setting the
191 .Cm acl
192 mount option is not advised unless the system is joined to an Active Directory
193 domain and using
194 .Xr ldap 1
195 on files and directories under this \fBsmbfs\fR(7FS) mount.
196 The default behavior is \fBnoacl\fR, which presents files and
197 directories as owned by the owner of the mount point and having
198 permissions based on \fBfileperms\fR or \fBdirperms\fR.
199 With the \fBacl\fR mount option, files are presented with ACLs
200 obtained from the SMB server.
201 Setting the \fBacl\fR mount option is not advised unless the system
202 is joined to an Active Directory domain and using \fBldap\fR(1)
203 so it can correctly present ACL identities from the SMB server.
204 .It Cm dirperms Ns = Ns Ar octaltriplet
205 Specifies the permissions to be assigned to directories.
206 The value must be specified as an octal triplet, such as
207 .Ql 755 .
208 The default value for the directory mode is taken from the
209 .Cm fileperms
210 setting, with execute permission added where
211 .Cm fileperms
212 has read permission.
213 .Pp
214 Note that these permissions have no relation to the rights granted by the SMB
215 .RE
216 .sp
217 .ne 2

```

```

191 .na
192 \fB\fBdirperms=\fR\fIoctaltriplet\fR\fR
193 .ad
194 .sp .6
195 .RS 4n
196 Specifies the permissions to be assigned to directories. The value must be
197 specified as an octal triplet, such as \fB755\fR. The default value for the
198 directory mode is taken from the \fBfileperms\fR setting, with execute
199 permission added where \fBfileperms\fR has read permission.
200 .sp
201 Note that these permissions have no relation to the rights granted by the CIFS
223 server.
224 .It Cm fileperms Ns = Ns Ar octaltriplet
225 Specifies the permissions to be assigned to files.
226 The value must be specified as an octal triplet, such as
227 .Ql 644 .
228 The default value is
229 .Ql 700 .
230 .Pp
231 Note that these permissions have no relation to the rights granted by the SMB
203 .RE

205 .sp
206 .ne 2
207 .na
208 \fB\fBfileperms=\fR\fIoctaltriplet\fR\fR
209 .ad
210 .sp .6
211 .RS 4n
212 Specifies the permissions to be assigned to files. The value must be specified
213 as an octal triplet, such as \fB644\fR. The default value is \fB700\fR.
214 .sp
215 Note that these permissions have no relation to the rights granted by the CIFS
232 server.
233 .It Cm gid Ns = Ns Ar groupid
234 Assigns the specified group ID to files.
235 The default value is the group ID of the directory where the volume is mounted.
236 .It Cm intr Ns | Ns Cm nointr
237 Enable (or disable) cancellation of
238 .Xr smbfs 7FS
239 I/O operations when the user interrupts the calling thread (for example, by
240 hitting Ctrl-C while an operation is underway).
241 The default is
242 .Cm intr
243 (interruption enabled), so cancellation is normally allowed.
244 .It Cm noprompt
245 Suppresses the prompting for a password when mounting a share.
246 This property enables you to permit anonymous access to a share.
247 Anonymous access does not require a password.
248 .Pp
249 The
250 .Nm mount
251 operation fails if a password is required, the
252 .Cm noprompt
253 property is set, and no password is stored by the
254 .Nm smbutil Cm login
255 command.
256 .It Cm retry_count Ns = Ns Ar number
217 .RE

219 .sp
220 .ne 2
221 .na
222 \fB\fBgid=\fR\fIgroupid\fR\fR
223 .ad
224 .sp .6

```

```

225 .RS 4n
226 Assigns the specified group ID to files. The default value is the group ID of
227 the directory where the volume is mounted.
228 .RE

230 .sp
231 .ne 2
232 .na
233 \fB\fBintr\fR|\fBnointr\fR\fR
234 .ad
235 .sp .6
236 .RS 4n
237 Enable (or disable) cancellation of \fBsmbfs\fR(7FS) I/O operations when the
238 user interrupts the calling thread (for example, by hitting Ctrl-C while an
239 operation is underway). The default is \fBintr\fR (interruption enabled), so
240 cancellation is normally allowed.
241 .RE

243 .sp
244 .ne 2
245 .na
246 \fB\fBnoprompt\fR\fR
247 .ad
248 .sp .6
249 .RS 4n
250 Suppresses the prompting for a password when mounting a share. This property
251 enables you to permit anonymous access to a share. Anonymous access does not
252 require a password.
253 .sp
254 The \fBmount\fR operation fails if a password is required, the \fBnoprompt\fR
255 property is set, and no password is stored by the \fBsmbutil login\fR command.
256 .RE

258 .sp
259 .ne 2
260 .na
261 \fB\fBretry_count=\fR\fInumber\fR\fR
262 .ad
263 .sp .6
264 .RS 4n
265 Specifies the number of SMBFS retries to attempt before the connection is
266 marked as broken.
267 By default, 4 attempts are made.
268 .Pp
269 The
270 .Cm retry_count
271 property value set by the
272 .Nm mount
273 command overrides the global value set in SMF or the value set in your
274 .Pa \&.nsmbrc
275 file.
276 .It Cm timeout Ns = Ns Ar seconds
277 Specifies the SMB request timeout.
278 By default, the timeout is 15 seconds.
279 .Pp
280 The
281 .Cm timeout
282 property value set by the
283 .Nm mount
284 command overrides the global value set in SMF or the value set in your
285 .Pa \&.nsmbrc
286 file.
287 .It Cm uid Ns = Ns Ar userid
288 Assigns the specified user ID files.
289 The default value is the owner ID of the directory where the volume is mounted.
290 .It Cm xattr Ns | Ns Cm noxattr

```

```

283 Enable (or disable) Extended Attributes in this mount point.
284 This option defaults to
285 .Cm xattr
286 (enabled Extended Attributes), but note: if the SMB server does not support SMB
287 "named streams",
288 .Xr smbfs 7FS
289 forces this option to
290 .Cm noxattr .
291 When a mount has the
292 .Cm noxattr
293 option, attempts to use Extended attributes fail with
294 .Er EINVAL .
295 .El
296 .It Fl 0
297 Overlays mount.
298 Allow the file system to be mounted over an existing mount point, making the
299 underlying file system inaccessible.
300 If a mount is attempted on a pre-existing mount point without setting this flag,
301 the mount fails, producing the error "device busy."
302 .El
303 .Sh FILES
304 .Bl -tag -width Pa
305 .It Pa /etc/mnttab
306 Table of mounted file systems.
307 .It Pa /etc/dfs/fstypes
308 Default distributed file system type.
309 .It Pa /etc/vfstab
310 Table of automatically mounted resources.
311 .It Pa $HOME/.nsmbrc
312 User-settable mount point configuration file to store the description for each
313 connection.
314 .El
315 .Sh EXAMPLES
316 .Bl -tag -width Ds
317 .It Sy Example 1 No Mounting an SMBFS Share
318 The following example shows how to mount the
319 .Pa /tmp
320 share from the
321 .Em nano
322 server in the
323 .Em SALES
324 workgroup on the local
325 .Pa /mnt
326 mount point.
327 You must supply the password for the root user to successfully perform the mount
328 operation.
329 .Bd -literal
330 # mount -F smbfs "//SALES;root@nano.sfbay/tmp" /mnt
331 marked as broken. By default, 4 attempts are made.
332 .sp
333 The \fBretry_count\fR property value set by the \fBmount\fR command overrides
334 the global value set in SMF or the value set in your \fB\&.nsmbrc\fR file.
335 .RE
336 .sp
337 .ne 2
338 .na
339 \fB\fBtimeout=\fR\fIseconds\fR\fR
340 .ad
341 .sp .6
342 .RS 4n
343 Specifies the CIFS request timeout. By default, the timeout is 15 seconds.
344 .sp
345 The \fBtimeout\fR property value set by the \fBmount\fR command overrides the
346 global value set in SMF or the value set in your \fB\&.nsmbrc\fR file.
347 .RE

```

```

285 .sp
286 .ne 2
287 .na
288 \fB\fBuid=\fR\fIuserid\fR\fR
289 .ad
290 .sp .6
291 .RS 4n
292 Assigns the specified user ID files. The default value is the owner ID of the
293 directory where the volume is mounted.
294 .RE
295 .sp
296 .ne 2
297 .na
298 \fB\fBxattr\fR|\fBnoxattr\fR\fR
299 .ad
300 .sp .6
301 .RS 4n
302 Enable (or disable) Solaris Extended Attributes in this mount point. This
303 option defaults to \fBxattr\fR (enabled Extended Attributes), but note: if the
304 CIFS server does not support CIFS "named streams", \fBsmbfs\fR(7FS) forces this
305 option to \fBnoxattr\fR. When a mount has the \fBnoxattr\fR option, attempts to
306 use Solaris Extended attributes fail with EINVAL.
307 .RE
308 .RE
309 .RE
310 .RE
311 .sp
312 .ne 2
313 .na
314 \fB\fB-O\fR\fR
315 .ad
316 .sp .6
317 .RS 4n
318 Overlays mount. Allow the file system to be mounted over an existing mount
319 point, making the underlying file system inaccessible. If a mount is attempted
320 on a pre-existing mount point without setting this flag, the mount fails,
321 producing the error "device busy."
322 .RE
323 .RE
324 .SH EXAMPLES
325 .LP
326 \fBExample 1 \fRMounting an SMBFS Share
327 .sp
328 .LP
329 The following example shows how to mount the \fB/tmp\fR share from the
330 \fBnano\fR server in the \fBSALES\fR workgroup on the local \fB/mnt\fR mount
331 point. You must supply the password for the \fBroot\fR user to successfully
332 perform the mount operation.
333 .sp
334 .in +2
335 .nf
336 # \fBmount -F smbfs "//SALES;root@nano.sfbay/tmp" /mnt\fR
337 Password:
338 .Ed
339 .It Sy Example 2 No Verifying That an SMBFS File System Is Mounted
340 The following example shows how to mount the
341 .Pa /tmp
342 share from the
343 .Em nano
344 server on the local
345 .Pa /mnt
346 mount point.
347 You must supply the password for the root user to successfully perform the mount

```

```

342 operation.
343 .Bd -literal
344 # mount -F smbfs //root@nano.sfbay/tmp /mnt
345 .fi
346 .in -2
347 .sp
348 .LP
349 \fBExample 2 \fRVerifying That an SMBFS File System Is Mounted
350 .sp
351 .LP
352 The following example shows how to mount the \fB/tmp\fR share from the
353 \fBnano\fR server on the local \fB/mnt\fR mount point. You must supply the
354 password for the \fBroot\fR user to successfully perform the mount operation.
355 .sp
356 .in +2
357 .nf
358 # \fBmount -F smbfs //root@nano.sfbay/tmp /mnt\fR
359 Password:
360 .Ed
361 .Pp
362 .fi
363 .in -2
364 .sp
365 .sp
366 .LP
367 You can verify that the share is mounted in the following ways:
368 .Bb -bullet
369 .It
370 View the file system entry in the
371 \fB/etc/mnttab\fR file.
372 .Bd -literal
373 # grep root /etc/mnttab
374 .RS +4
375 .TP
376 .ie t \fB(bu
377 .el o
378 View the file system entry in the \fB/etc/mnttab\fR file.
379 .sp
380 .in +2
381 .nf
382 # \fBgrep root /etc/mnttab\fR
383 //root@nano.sfbay/tmp /mnt smbfs dev=4900000 1177097833
384 .Ed
385 .It
386 View the output of the
387 \fBmount -l\fR command.
388 .Bd -literal
389 # mount | grep root
390 .fi
391 .in -2
392 .sp
393 .RE
394 .RS +4
395 .TP
396 .ie t \fB(bu
397 .el o
398 View the output of the \fBdf /mnt\fR command.
399 .sp
400 .in +2
401 .nf
402 # \fBdf /mnt\fR
403 /mnt (/root@nano.sfbay/tmp): 3635872 blocks -1 files
404 .Ed
405 .El
406 .Pp
407 Obtain information about the mounted share by viewing the output of the
408 \fBdf -k /mnt\fR command.
409 .Bd -literal
410 # df -k /mnt
411 .fi
412 .in -2
413 .sp
414 .RE
415 .sp
416 .LP
417 Obtain information about the mounted share by viewing the output of the \fBdf
418 -k /mnt\fR command.
419 .sp
420 .in +2
421 .nf
422 # \fBdf -k /mnt\fR
423 Filesystem kbytes used avail capacity Mounted on
424 //root@nano.sfbay/tmp 1882384 64448 1817936 4% /mnt
425 .Ed
426 .It Sy Example 3 No Unmounting an SMB Share
427 This example assumes that an SMB share has been mounted on the
428 \fB/mnt\fR mount point.
429 The following command line unmounts the share from the mount point.
430 .Bd -literal
431 # umount /mnt
432 .Ed
433 .El
434 .Sh INTERFACE STABILITY
435 .Sy Committed
436 .Sh SEE ALSO
437 .Xr ldap 1 ,
438 .Xr smbutil 1 ,
439 .Xr mount 1m ,

```

```

387 # \fBmount | grep root\fR
388 /mnt on //root@nano.sfbay/tmp read/write/setuid/devices/dev=4900000 on
389 Fri Apr 20 13:37:13 2007
390 .Ed
391 .It
392 View the output of the
393 \fBdf /mnt\fR command.
394 .Bd -literal
395 # df /mnt
396 .fi
397 .in -2
398 .sp
399 .RE
400 .RS +4
401 .TP
402 .ie t \fB(bu
403 .el o
404 View the output of the \fBdf /mnt\fR command.
405 .sp
406 .in +2
407 .nf
408 # \fBdf /mnt\fR
409 /mnt (/root@nano.sfbay/tmp): 3635872 blocks -1 files
410 .Ed
411 .El
412 .Pp
413 Obtain information about the mounted share by viewing the output of the
414 \fBdf -k /mnt\fR command.
415 .Bd -literal
416 # df -k /mnt
417 .fi
418 .in -2
419 .sp
420 .RE
421 .sp
422 .LP
423 Obtain information about the mounted share by viewing the output of the \fBdf
424 -k /mnt\fR command.
425 .sp
426 .in +2
427 .nf
428 # \fBdf -k /mnt\fR
429 Filesystem kbytes used avail capacity Mounted on
430 //root@nano.sfbay/tmp 1882384 64448 1817936 4% /mnt
431 .Ed
432 .It Sy Example 3 No Unmounting an SMB Share
433 This example assumes that an SMB share has been mounted on the
434 \fB/mnt\fR mount point.
435 The following command line unmounts the share from the mount point.
436 .Bd -literal
437 # umount /mnt
438 .Ed
439 .El
440 .Sh INTERFACE STABILITY
441 .Sy Committed
442 .Sh SEE ALSO
443 .Xr ldap 1 ,
444 .Xr smbutil 1 ,
445 .Xr mount 1m ,

```

```

401 .Xr mountall 1M ,
402 .Xr svcadm 1M ,
403 .Xr acl 2 ,
404 .Xr fcntl 2 ,
405 .Xr link 2 ,
406 .Xr mknod 2 ,
407 .Xr mount 2 ,
408 .Xr symlink 2 ,
409 .Xr umount 2 ,
410 .Xr mnttab 4 ,
411 .Xr nsmbrc 4 ,
412 .Xr vfstab 4 ,
413 .Xr attributes 5 ,
414 .Xr pcfs 7FS ,
415 .Xr smbfs 7FS
416 .Sh AUTHORS
417 This manual page contains material originally authored by
418 .An Boris Popov
419 .Aq Mt bp@butya.kz ,
420 .Aq Mt bp@FreeBSD.org .
421 .Sh NOTES
422 The SMB client always attempts to use
423 .Xr gethostbyname 3NSL
424 to resolve host names.
425 If the host name cannot be resolved, the SMB client uses NetBIOS name
426 resolution (NBNS).
427 By default, the SMB client permits the use of NBNS to enable SMB clients in
428 Windows environments to work without additional configuration.
429 .Pp
430 Since NBNS has been exploited in the past, you might want to disable it.
431 To disable NBNS, set the
432 .Em nbns-enabled
433 service management facility property to
434 .Cm false .
435 By default,
436 .Em nbns-enabled
437 is set to
438 .Cm true .
439 .Pp
442 .fi
443 .in -2
444 .sp
446 .LP
447 \fBExample 3 \fRUNmounting a CIFS Share
448 .sp
449 .LP
450 This example assumes that a CIFS share has been mounted on the \fB/mnt\fR mount
451 point. The following command line unmounts the share from the mount point.
452 .sp
453 .in +2
454 .nf
455 # \fBumount /mnt\fR
456 .fi
457 .in -2
458 .sp
461 .SH FILES
462 .sp
463 .ne 2
464 .na
465 \fB\fB/etc/mnttab\fR\fR
466 .ad
467 .sp .6
468 .RS 4n

```

```

449 Table of mounted file systems.
450 .RE
452 .sp
453 .ne 2
454 .na
455 \fB\fB/etc/dfs/fstypes\fR\fR
456 .ad
457 .sp .6
458 .RS 4n
459 Default distributed file system type.
460 .RE
462 .sp
463 .ne 2
464 .na
465 \fB\fB/etc/vfstab\fR\fR
466 .ad
467 .sp .6
468 .RS 4n
469 Table of automatically mounted resources.
470 .RE
472 .sp
473 .ne 2
474 .na
475 \fB\fB$HOME/.nsmbrc\fR\fR
476 .ad
477 .sp .6
478 .RS 4n
479 User-settable mount point configuration file to store the description for each
480 connection.
481 .RE
483 .SH ATTRIBUTES
484 .sp
485 .LP
486 See the \fBattributes\fR(5) man page for descriptions of the following
487 attributes:
488 .sp
490 .sp
491 .TS
492 box;
493 c | c
494 l | l .
495 ATTRIBUTE TYPE ATTRIBUTE VALUE
496 _
497 Interface Stability Committed
498 .TE
500 .SH SEE ALSO
501 .sp
502 .LP
503 \fB\fBldap\fR(1), \fB\fBsmbutil\fR(1),
504 \fB\fBmount\fR(1M), \fB\fBmountall\fR(1M), \fB\fBsvcadm\fR(1M),
505 \fB\fBacl\fR(2), \fB\fBfcntl\fR(2), \fB\fBlink\fR(2), \fB\fBmknod\fR(2), \fB\fBmount\fR(2),
506 \fB\fBsymlink\fR(2), \fB\fBumount\fR(2), \fB\fBmnttab\fR(4), \fB\fBnsmbrc\fR(4),
507 \fB\fBvfstab\fR(4), \fB\fBattributes\fR(5), \fB\fBpcfs\fR(7FS), \fB\fBsmbfs\fR(7FS)
508 .SH AUTHORS
509 .sp
510 .LP
511 This manual page contains material originally authored by Boris Popov,
512 \fB\fBbpATbutya.kz\fR, \fB\fBbpATFreeBSD.org\fR.
513 .SH NOTES
514 .sp

```

```
515 .LP
516 The Solaris CIFS client always attempts to use \fBgethostbyname()\fR to resolve
517 host names. If the host name cannot be resolved, the CIFS client uses NetBIOS
518 name resolution (NBNS). By default, the Solaris CIFS client permits the use of
519 NBNS to enable Solaris CIFS clients in Windows environments to work without
520 additional configuration.
521 .sp
522 .LP
523 Since NBNS has been exploited in the past, you might want to disable it. To
524 disable NBNS, set the \fBnbns-enabled\fR service management facility property
525 to \fBfalse\fR. By default, \fBnbns-enabled\fR is set to \fBtrue\fR.
526 .sp
527 .LP
440 If the directory on which a file system is to be mounted is a symbolic link,
441 the file system is mounted on the directory to which the symbolic link refers,
442 rather than being mounted on top of the symbolic link itself.
```

new/usr/src/uts/common/fs/smbcInt/smbfs/smbfs.h

1

6523 Sun Mar 4 06:09:09 2018
new/usr/src/uts/common/fs/smbcInt/smbfs/smbfs.h
5404 smbfs needs mmap support
Portions contributed by: Gordon Ross <gordon.w.ross@gmail.com>

```
1 /*
2  * Copyright (c) 2000-2001, Boris Popov
3  * All rights reserved.
4  *
5  * Redistribution and use in source and binary forms, with or without
6  * modification, are permitted provided that the following conditions
7  * are met:
8  * 1. Redistributions of source code must retain the above copyright
9  *   notice, this list of conditions and the following disclaimer.
10 * 2. Redistributions in binary form must reproduce the above copyright
11 *   notice, this list of conditions and the following disclaimer in the
12 *   documentation and/or other materials provided with the distribution.
13 * 3. All advertising materials mentioning features or use of this software
14 *   must display the following acknowledgement:
15 *   This product includes software developed by Boris Popov.
16 * 4. Neither the name of the author nor the names of any co-contributors
17 *   may be used to endorse or promote products derived from this software
18 *   without specific prior written permission.
19 *
20 * THIS SOFTWARE IS PROVIDED BY THE AUTHOR AND CONTRIBUTORS "AS IS" AND
21 * ANY EXPRESS OR IMPLIED WARRANTIES, INCLUDING, BUT NOT LIMITED TO, THE
22 * IMPLIED WARRANTIES OF MERCHANTABILITY AND FITNESS FOR A PARTICULAR PURPOSE
23 * ARE DISCLAIMED. IN NO EVENT SHALL THE AUTHOR OR CONTRIBUTORS BE LIABLE
24 * FOR ANY DIRECT, INDIRECT, INCIDENTAL, SPECIAL, EXEMPLARY, OR CONSEQUENTIAL
25 * DAMAGES (INCLUDING, BUT NOT LIMITED TO, PROCUREMENT OF SUBSTITUTE GOODS
26 * OR SERVICES; LOSS OF USE, DATA, OR PROFITS; OR BUSINESS INTERRUPTION)
27 * HOWEVER CAUSED AND ON ANY THEORY OF LIABILITY, WHETHER IN CONTRACT, STRICT
28 * LIABILITY, OR TORT (INCLUDING NEGLIGENCE OR OTHERWISE) ARISING IN ANY WAY
29 * OUT OF THE USE OF THIS SOFTWARE, EVEN IF ADVISED OF THE POSSIBILITY OF
30 * SUCH DAMAGE.
31 *
32 * $Id: smbfs.h,v 1.30.100.1 2005/05/27 02:35:28 lindak Exp $
33 */

35 /*
36  * Copyright (c) 2008, 2010, Oracle and/or its affiliates. All rights reserved.
37  */

39 #ifndef _SMBFS_SMBFS_H
40 #define _SMBFS_SMBFS_H

42 /*
43  * FS-specific VFS structures for smbfs.
44  * (per-mount stuff, etc.)
45  *
46  * This file used to have mount args stuff,
47  * but that's now in sys/fs/smbfs_mount.h
48  */

50 #include <sys/param.h>
51 #include <sys/fstyp.h>
52 #include <sys/avl.h>
53 #include <sys/list.h>
54 #include <sys/t_lock.h>
55 #include <sys/vfs.h>
56 #include <sys/vfs_opreg.h>
57 #include <sys/fs/smbfs_mount.h>
58 #include <sys/zone.h>

60 /*
```

new/usr/src/uts/common/fs/smbcInt/smbfs/smbfs.h

2

```
61  * Path component length
62  *
63  * The generic fs code uses MAXNAMELEN to represent
64  * what the largest component length is, but note:
65  * that length DOES include the terminating NULL.
66  * SMB_MAXFNAMELEN does NOT include the NULL.
67  */
68 #define SMB_MAXFNAMELEN      (MAXNAMELEN-1) /* 255 */

70 /*
71  * SM_MAX_STATFSTIME is the maximum time to cache statvfs data. Since this
72  * should be a fast call on the server, the time the data cached is short.
73  * That lets the cache handle bursts of statvfs() requests without generating
74  * lots of network traffic.
75  */
76 #define SM_MAX_STATFSTIME 2

78 /* Mask values for smbmount structure sm_status field */
79 #define SM_STATUS_STATFS_BUSY 0x00000001 /* statvfs is in progress */
80 #define SM_STATUS_STATFS_WANT 0x00000002 /* statvfs wakeup is wanted */
81 #define SM_STATUS_TIMEO 0x00000004 /* this mount is not responding */
82 #define SM_STATUS_DEAD 0x00000010 /* connection gone - unmount this */

84 extern const struct fs_operation_def      smbfs_vnodeops_template[];
85 extern struct vnodeops                    *smbfs_vnodeops;

87 struct smbnode;
88 struct smb_share;

90 /*
91  * The values for smi_flags (from nfs_cInt.h)
92  */
93 #define SMI_INT          0x04          /* interrupts allowed */
94 #define SMI_NOAC         0x10          /* don't cache attributes */
95 #define SMI_LLOCK        0x80          /* local locking only */
96 #define SMI_ACL          0x2000       /* share supports ACLs */
97 #define SMI_DIRECTIO     0x40000      /* do direct I/O */
98 #define SMI_EXTATTR      0x80000     /* share supports ext. attrs */
99 #define SMI_DEAD         0x200000     /* mount has been terminated */

101 /*
102  * Stuff returned by smbfs_smb_qfsattr
103  * See [CIFS] SMB_QUERY_FS_ATTRIBUTE_INFO
104  */
105 typedef struct smb_fs_attr_info {
106     uint32_t      fsa_aflags; /* Attr. flags [CIFS 4.1.6.6] */
107     uint32_t      fsa_maxname; /* max. component length */
108     char          fsa_tname[FSTYPSZ]; /* type name, i.e. "NTFS" */
109 } smb_fs_attr_info_t;
110 unchanged_portion_omitted
```

new/usr/src/uts/common/fs/smbclnt/smbfs/smbfs_client.c

1

20424 Sun Mar 4 06:09:09 2018
new/usr/src/uts/common/fs/smbclnt/smbfs/smbfs_client.c
Kill flags arg in smbfs_purge_caches
Lots of comment cleanup
5404 smbfs needs mmap support
Portions contributed by: Gordon Ross <gordon.w.ross@gmail.com>

```
1 /*
2  * CDDL HEADER START
3  *
4  * The contents of this file are subject to the terms of the
5  * Common Development and Distribution License (the "License").
6  * You may not use this file except in compliance with the License.
7  *
8  * You can obtain a copy of the license at usr/src/OPENSOLARIS.LICENSE
9  * or http://www.opensolaris.org/os/licensing.
10 * See the License for the specific language governing permissions
11 * and limitations under the License.
12 *
13 * When distributing Covered Code, include this CDDL HEADER in each
14 * file and include the License file at usr/src/OPENSOLARIS.LICENSE.
15 * If applicable, add the following below this CDDL HEADER, with the
16 * fields enclosed by brackets "[]" replaced with your own identifying
17 * information: Portions Copyright [yyyy] [name of copyright owner]
18 *
19 * CDDL HEADER END
20 */
21
22 /*
23  * Copyright (c) 2008, 2010, Oracle and/or its affiliates. All rights reserved.
24  *
25  * Copyright (c) 1983,1984,1985,1986,1987,1988,1989 AT&T.
26  * All rights reserved.
27 */
28
29 #include <sys/param.h>
30 #include <sys/systm.h>
31 #include <sys/thread.h>
32 #include <sys/t_lock.h>
33 #include <sys/time.h>
34 #include <sys/vnode.h>
35 #include <sys/vfs.h>
36 #include <sys/errno.h>
37 #include <sys/buf.h>
38 #include <sys/stat.h>
39 #include <sys/cred.h>
40 #include <sys/kmem.h>
41 #include <sys/debug.h>
42 #include <sys/vmsystem.h>
43 #include <sys/flock.h>
44 #include <sys/share.h>
45 #include <sys/cmn_err.h>
46 #include <sys/tiuser.h>
47 #include <sys/sysmacros.h>
48 #include <sys/callb.h>
49 #include <sys/acl.h>
50 #include <sys/kstat.h>
51 #include <sys/signal.h>
52 #include <sys/list.h>
53 #include <sys/zone.h>
54
55 #include <netsmb/smb.h>
56 #include <netsmb/smb_conn.h>
57 #include <netsmb/smb_subr.h>
```

new/usr/src/uts/common/fs/smbclnt/smbfs/smbfs_client.c

2

```
59 #include <smbfs/smbfs.h>
60 #include <smbfs/smbfs_node.h>
61 #include <smbfs/smbfs_subr.h>
62
63 #include <vm/hat.h>
64 #include <vm/as.h>
65 #include <vm/page.h>
66 #include <vm/pvn.h>
67 #include <vm/seg.h>
68 #include <vm/seg_map.h>
69 #include <vm/seg_vn.h>
70
71 #define ATTRCACHE_VALID(vp) (gethrtime() < VTOSMB(vp)->r_attrtime)
72
73 static int smbfs_getattr_cache(vnode_t *, smbfsattr_t *);
74 static void smbfsattr_to_vattr(vnode_t *, smbfsattr_t *, vattr_t *);
75 static void smbfsattr_to_xvattr(smbfsattr_t *, vattr_t *);
76 static int smbfs_getattr_otw(vnode_t *, struct smbfsattr *, cred_t *);
77
78
79 /*
80  * The following code provide zone support in order to perform an action
81  * for each smbfs mount in a zone. This is also where we would add
82  * per-zone globals and kernel threads for the smbfs module (since
83  * they must be terminated by the shutdown callback).
84  */
85
86 struct smi_globals {
87     kmutex_t      smg_lock; /* lock protecting smg_list */
88     list_t        smg_list; /* list of SMBFS mounts in zone */
89     boolean_t      smg_destructor_called;
90 };
91 typedef struct smi_globals smi_globals_t;
92
93 static zone_key_t smi_list_key;
94
95 /*
96  * Attributes caching:
97  *
98  * Attributes are cached in the smbnode in struct vattr form.
99  * There is a time associated with the cached attributes (r_attrtime)
100 * which tells whether the attributes are valid. The time is initialized
101 * to the difference between current time and the modify time of the vnode
102 * when new attributes are cached. This allows the attributes for
103 * files that have changed recently to be timed out sooner than for files
104 * that have not changed for a long time. There are minimum and maximum
105 * timeout values that can be set per mount point.
106 */
107
108 /*
109  * Helper for _validate_caches
110  * Validate caches by checking cached attributes. If they have timed out
111  * get the attributes from the server and compare mtimes. If mtimes are
112  * different purge all caches for this vnode.
113  */
114 int
115 smbfs_waitfor_purge_complete(vnode_t *vp)
116 {
117     smbnode_t *np;
118     k_sigset_t smask;
119
120     np = VTOSMB(vp);
121     if (np->r_serial != NULL && np->r_serial != curthread) {
122         mutex_enter(&np->r_statelock);
123         sigintr(&smask, VTOSMI(vp)->smi_flags & SMI_INT);
124         while (np->r_serial != NULL) {
```

```

122         if (!cv_wait_sig(&np->r_cv, &np->r_statelock)) {
123             sigunintr(&smask);
124             mutex_exit(&np->r_statelock);
125             return (EINTR);
126         }
127     }
128     sigunintr(&smask);
129     mutex_exit(&np->r_statelock);
130 }
131 return (0);
132 }

134 /*
135  * Validate caches by checking cached attributes. If the cached
136  * attributes have timed out, then get new attributes from the server.
137  * As a side affect, this will do cache invalidation if the attributes
138  * have changed.
139  *
140  * If the attributes have not timed out and if there is a cache
141  * invalidation being done by some other thread, then wait until that
142  * thread has completed the cache invalidation.
143  */
144 int
145 smbfs_validate_caches(
146     struct vnode *vp,
147     cred_t *cr)
148 {
149     struct smbfsattr fa;
150     int error;
151     struct vaattr va;
152
153     if (ATTRCACHE_VALID(vp)) {
154         error = smbfs_waitfor_purge_complete(vp);
155         if (error)
156             return (error);
157         return (0);
158     }
159
160     return (smbfs_getattr_otw(vp, &fa, cr));
161     va.va_mask = AT_SIZE;
162     return (smbfs_getattr(vp, &va, cr));
163 }

162 /*
163  * Purge all of the various data caches.
164  *
165  * Here NFS also had a flags arg to control what gets flushed.
166  * We only have the page cache, so no flags arg.
167  */
168 /* ARGSUSED */
169 void
170 smbfs_purge_caches(struct vnode *vp, cred_t *cr)
171 {
172     /* not yet: mmap support */
173     /*
174      * Here NFS has: Purge the DNLC for this vp,
175      * NFS: Purge the DNLC for this vp,
176      * Clear any readdr state bits,
177      * the readlink response cache, ...
178      */
179     smbnode_t *np = VTOSMB(vp);

```

```

180     * Flush the page cache.
181     */
182     if (vn_has_cached_data(vp)) {
183         (void) VOP_PUTPAGE(vp, (u_offset_t)0, 0, B_INVAL, cr, NULL);
184     }
185
186     /*
187      * Here NFS has: Flush the readdr response cache.
188      * No readdr cache in smbfs.
189      */
190 #endif /* not yet */
191 }

192 /*
193  * Here NFS has:
194  * nfs_purge_rddir_cache()
195  * nfs3_cache_post_op_attr()
196  * nfs3_cache_post_op_vattr()
197  * nfs3_cache_wcc_data()
198  */

199 /*
200  * Check the attribute cache to see if the new attributes match
201  * those cached. If they do, the various 'data' caches are
202  * considered to be good. Otherwise, purge the cached data.
203  */
204 static void
205 smbfs_cache_check(
206     struct vnode *vp,
207     struct smbfsattr *fap,
208     cred_t *cr)
209 {
210     struct smbfsattr *fap;
211     smbnode_t *np;
212     int purge_data = 0;
213     int purge_acl = 0;
214
215     np = VTOSMB(vp);
216     mutex_enter(&np->r_statelock);
217
218     /*
219      * Compare with NFS macro: CACHE_VALID
220      * If the mtime or size has changed,
221      * purge cached data.
222      */
223     if (np->r_attr.fa_mtime.tv_sec != fap->fa_mtime.tv_sec ||
224         np->r_attr.fa_mtime.tv_nsec != fap->fa_mtime.tv_nsec)
225         purge_data = 1;
226     if (np->r_attr.fa_size != fap->fa_size)
227         purge_data = 1;
228
229     if (np->r_attr.fa_ctime.tv_sec != fap->fa_ctime.tv_sec ||
230         np->r_attr.fa_ctime.tv_nsec != fap->fa_ctime.tv_nsec)
231         purge_acl = 1;
232
233     if (purge_acl) {
234         /* just invalidate r_secattr (XXX: OK?) */
235         np->r_secattr = gethrtime();
236     }
237
238     mutex_exit(&np->r_statelock);
239
240     if (purge_data)
241         smbfs_purge_caches(vp, cr);
242     smbfs_purge_caches(vp);

```

```

241 }

243 /*
244  * Set attributes cache for given vnode using vnode attributes.
245  * From NFS: nfs_attrcache_va
246  */
247 #if 0 /* not yet (not sure if we need this) */
248 void
249 smbfs_attrcache_va(vnode_t *vp, struct vattr *vap)
250 {
251     smbfsattr_t fa;
252     smbnode_t *np;

253     vattr_to_fattr(vp, vap, &fa);
254     smbfs_attrcache_fa(vp, &fa);
255 }
256 #endif /* not yet */

257 /*
258  * Set attributes cache for given vnode using SMB fattr
259  * and update the attribute cache timeout.
260  */
261 * Based on NFS: nfs_attrcache, nfs_attrcache_va
262 * From NFS: nfs_attrcache, nfs_attrcache_va
263 */
264 void
265 smbfs_attrcache_fa(vnode_t *vp, struct smbfsattr *fap)
266 {
267     smbnode_t *np;
268     smbmntinfo_t *smi;
269     hrttime_t delta, now;
270     u_offset_t newsize;
271     vtype_t vtype, oldvt;
272     mode_t mode;

273     np = VTOSMB(vp);
274     smi = VTOSMI(vp);

275     /*
276      * We allow v_type to change, so set that here
277      * (and the mode, which depends on the type).
278      */
279     if (fap->fa_attr & SMB_FA_DIR) {
280         vtype = VDIR;
281         mode = smi->smi_dmode;
282     } else {
283         vtype = VREG;
284         mode = smi->smi_fmode;
285     }

286     mutex_enter(&np->r_statelock);
287     now = gethrtime();

288     /*
289      * Delta is the number of nanoseconds that we will
290      * cache the attributes of the file. It is based on
291      * the number of nanoseconds since the last time that
292      * we detected a change. The assumption is that files
293      * that changed recently are likely to change again.
294      * There is a minimum and a maximum for regular files
295      * and for directories which is enforced though.
296      *
297      * Using the time since last change was detected
298      * eliminates direct comparison or calculation
299      * using mixed client and server times. SMBFS
300      * does not make any assumptions regarding the

```

```

290     * client and server clocks being synchronized.
291     */
292     if (fap->fa_mtime.tv_sec != np->r_attr.fa_mtime.tv_sec ||
293         fap->fa_mtime.tv_nsec != np->r_attr.fa_mtime.tv_nsec ||
294         fap->fa_size != np->r_attr.fa_size)
295         np->r_mtime = now;

296     if ((smi->smi_flags & SMI_NOAC) || (vp->v_flag & VNOCACHE))
297         delta = 0;
298     else {
299         delta = now - np->r_mtime;
300         if (vtype == VDIR) {
301             if (delta < smi->smi_acdirmin)
302                 delta = smi->smi_acdirmin;
303             else if (delta > smi->smi_acdirmax)
304                 delta = smi->smi_acdirmax;
305         } else {
306             if (delta < smi->smi_acregmin)
307                 delta = smi->smi_acregmin;
308             else if (delta > smi->smi_acregmax)
309                 delta = smi->smi_acregmax;
310         }
311     }
312     np->r_attrtime = now + delta;
313     np->r_attr = *fap;
314     np->n_mode = mode;
315     oldvt = vp->v_type;
316     vp->v_type = vtype;

317     /*
318      * Shall we update r_size? (local notion of size)
319      *
320      * The real criteria for updating r_size should be:
321      * if the file has grown on the server, or if
322      * the client has not modified the file.
323      *
324      * Also deal with the fact that SMB presents
325      * directories as having size=0. Doing that
326      * here and leaving fa_size as returned 0w
327      * avoids fixing the size lots of places.
328      */
329     newsize = fap->fa_size;
330     if (vtype == VDIR && newsize < DEV_BSIZE)
331         newsize = DEV_BSIZE;

332     if (np->r_size != newsize &&
333         (!vn_has_cached_data(vp) ||
334          (!(np->r_flags & RDIRTY) && np->r_count == 0))) {
335         if (np->r_size != newsize) {
336             /* not yet: mmap support */
337             if (!vn_has_cached_data(vp) || ...)
338                 /* XXX: See NFS page cache code. */
339                 /* OK to set the size. */
340                 np->r_size = newsize;
341         }

342     /*
343      * Here NFS has:
344      * nfs_setswaplike(vp, va);
345      * np->r_flags &= ~RWRITEATTR;
346      * (not needed here)
347      */
348     /* NFS: np->r_flags &= ~RWRITEATTR; */

```

```

350     np->n_flag &= ~NATTRCHANGED;

351     mutex_exit(&np->r_statelock);

353     if (oldvt != vtype) {
354         SMBVDEBUG("vtype change %d to %d\n", oldvt, vtype);
355     }
356 }
    unchanged_portion_omitted_

388 /*
389  * Get attributes over-the-wire and update attributes cache
390  * if no error occurred in the over-the-wire operation.
391  * Return 0 if successful, otherwise error.
392  * From NFS: nfs_getattr_otw
393  */
394 static int
395 smbfs_getattr_otw(vnode_t *vp, struct smbfsattr *fap, cred_t *cr)
396 {
397     struct smbnode *np;
398     struct smb_cred scred;
399     int error;

401     np = VTOSMB(vp);

403     /*
404     * Here NFS uses the ACL RPC (if smi_flags & SMI_ACL)
405     * NFS uses the ACL rpc here (if smi_flags & SMI_ACL)
406     * With SMB, getting the ACL is a significantly more
407     * expensive operation, so we do that only when asked
408     * for the uid/gid. See smbfsgetattr().
409     */

410     /* Shared lock for (possible) n_fid use. */
411     if (smbfs_rw_enter_sig(&np->r_lkserlock, RW_READER, SMBINTR(vp)))
412         return (EINTR);
413     smb_credinit(&scred, cr);

415     bzero(fap, sizeof(*fap));
416     error = smbfs_smb_getfattr(np, fap, &scred);

418     smb_credrele(&scred);
419     smbfs_rw_exit(&np->r_lkserlock);

421     if (error) {
422         /* Here NFS has: PURGE_STALE_FH(error, vp, cr) */
423         /* NFS had: PURGE_STALE_FH(error, vp, cr) */
424         smbfs_attrcache_remove(np);
425         if (error == ENOENT || error == ENOTDIR) {
426             /*
427              * Getattr failed because the object was
428              * removed or renamed by another client.
429              * Remove any cached attributes under it.
430              */
431             smbfs_attrcache_prune(np);
432         }
433     }
    return (error);
}

435 /*
436  * Here NFS has: nfs_cache_fattr(vap, fa, vap, t, cr);
437  * NFS: smbfs_cache_fattr(vap, fa, vap, t, cr);
438  * which did: fattr_to_vattr, nfs_attr_cache.
439  * We cache the fattr form, so just do the
440  * cache check and store the attributes.

```

```

440     /*
441     smbfs_cache_check(vp, fap, cr);
442     smbfs_cache_check(vp, fap);
443     smbfs_attrcache_fa(vp, fap);
444     return (0);
445 }

447 /*
448  * Return either cached or remote attributes. If we get remote attrs,
449  * Return either cached or remote attributes. If get remote attr
450  * use them to check and invalidate caches, then cache the new attributes.
451  * From NFS: nfsgetattr()
452  */
453 int
454 smbfsgetattr(vnode_t *vp, struct vattr *vap, cred_t *cr)
455 {
456     struct smbfsattr fa;
457     smbmntinfo_t *smi;
458     uint_t mask;
459     int error;

461     smi = VTOSMI(vp);

463     ASSERT(curproc->p_zone == smi->smi_zone_ref.zref_zone);

465     /*
466     * If asked for UID or GID, update n_uid, n_gid.
467     */
468     mask = AT_ALL;
469     if (vap->va_mask & (AT_UID | AT_GID)) {
470         if (smi->smi_flags & SMI_ACL)
471             (void) smbfs_acl_getids(vp, cr);
472         /* else leave as set in make_smbnode */
473     } else {
474         mask &= ~(AT_UID | AT_GID);
475     }

477     /*
478     * If we've got cached attributes, just use them;
479     * otherwise go to the server to get attributes,
480     * which will update the cache in the process.
481     */
482     error = smbfs_getattr_cache(vp, &fa);
483     if (error)
484         error = smbfs_getattr_otw(vp, &fa, cr);
485     if (error)
486         return (error);
487     vap->va_mask |= mask;

489     /*
490     * Re. client's view of the file size, see:
491     * smbfs_attrcache_fa, smbfs_getattr_otw
492     */
493     smbfsattr_to_vattr(vp, &fa, vap);
494     if (vap->va_mask & AT_XVATTR)
495         smbfsattr_to_xvattr(&fa, vap);

497     return (0);
498 }
    unchanged_portion_omitted_

596 /*
597  * Here NFS has:
598  * nfs_async... stuff

```

```

599 * which we're not using (no async I/O), and:
600 *   writerp(),
601 *   nfs_putpages()
602 *   nfs_invalidate_pages()
603 * which we have in smbfs_vnops.c, and
604 *   nfs_printfhndle()
605 *   nfs_write_error()
606 * not needed here.
607 */

609 /*
610 * Helper function for smbfs_sync
611 *
612 * Walk the per-zone list of smbfs mounts, calling smbfs_rflush
613 * on each one. This is a little tricky because we need to exit
614 * the list mutex before each _rflush call and then try to resume
615 * where we were in the list after re-entering the mutex.
616 */
617 void
618 smbfs_flushall(cred_t *cr)
619 {
620     smi_globals_t *smg;
621     smbmntinfo_t *tmp_smi, *cur_smi, *next_smi;

622     smg = zone_getspecific(smi_list_key, crgetzone(cr));
623     ASSERT(smg != NULL);

624     mutex_enter(&smg->smg_lock);
625     cur_smi = list_head(&smg->smg_list);
626     if (cur_smi == NULL) {
627         mutex_exit(&smg->smg_lock);
628         return;
629     }
630     VFS_HOLD(cur_smi->smi_vfsp);
631     mutex_exit(&smg->smg_lock);

632     flush:
633     smbfs_rflush(cur_smi->smi_vfsp, cr);

634     mutex_enter(&smg->smg_lock);
635     /*
636     * Resume after cur_smi if that's still on the list,
637     * otherwise restart at the head.
638     */
639     for (tmp_smi = list_head(&smg->smg_list);
640          tmp_smi != NULL;
641          tmp_smi = list_next(&smg->smg_list, tmp_smi))
642         if (tmp_smi == cur_smi)
643             break;
644     if (tmp_smi != NULL)
645         next_smi = list_next(&smg->smg_list, tmp_smi);
646     else
647         next_smi = list_head(&smg->smg_list);

648     if (next_smi != NULL)
649         VFS_HOLD(next_smi->smi_vfsp);
650     VFS_RELE(cur_smi->smi_vfsp);

651     mutex_exit(&smg->smg_lock);

652     if (next_smi != NULL) {
653         cur_smi = next_smi;
654         goto flush;
655     }
656 }

```

```

665 /*
666 * SMB Client initialization and cleanup.
667 * Much of it is per-zone now.
668 */

671 /* ARGSUSED */
672 static void *
673 smbfs_zone_init(zoneid_t zoneid)
674 {
675     smi_globals_t *smg;

676     smg = kmem_alloc(sizeof (*smg), KM_SLEEP);
677     mutex_init(&smg->smg_lock, NULL, MUTEX_DEFAULT, NULL);
678     list_create(&smg->smg_list, sizeof (smbmntinfo_t),
679               offsetof(smbmntinfo_t, smi_zone_node));
680     smg->smg_destructor_called = B_FALSE;
681     return (smg);
682 }
683 }

```

unchanged_portion_omitted

new/usr/src/uts/common/fs/smbclnt/smbfs/smbfs_node.c

1

6280 Sun Mar 4 06:09:09 2018

new/usr/src/uts/common/fs/smbclnt/smbfs/smbfs_node.c

Lots of comment cleanup

_____unchanged_portion_omitted_____

```
204 /*
205  * smbfs_attrcache_enter, smbfs_attrcache_lookup replaced by
206  * code more closely resembling NFS. See smbfs_client.c
207  */
```

```
209 /*
210  * Update the local notion of the mtime of some directory.
211  * See comments re. r_mtime in smbfs_node.h
212  */
213 void
214 smbfs_attr_touchdir(struct smbnode *dnp)
215 {
```

```
216     mutex_enter(&dnp->r_statelock);
```

```
217     /*
218     * Now that we keep the client's notion of mtime
219     * separately from the server, this is easy.
220     */
221     dnp->r_mtime = gethrtime();
```

```
222     /*
223     * Invalidate the cache, so that we go to the wire
224     * to check that the server doesn't have a better
225     * timestamp next time we care.
226     */
227     smbfs_attrcache_rm_locked(dnp);
228     mutex_exit(&dnp->r_statelock);
229 }
```

_____unchanged_portion_omitted_____

```

*****
11806 Sun Mar  4 06:09:09 2018
new/usr/src/uts/common/fs/smbclnt/smbfs/smbfs_node.h
5404 smbfs needs mmap support
Portions contributed by: Gordon Ross <gordon.w.ross@gmail.com>
*****
_____unchanged_portion_omitted_____

```

```

125 /*
126  * Below is the SMBFS-specific representation of a "node".
127  * This struct is a mixture of Sun NFS and Darwin code.
128  * Fields starting with "r_" came from NFS struct "rnode"
129  * and fields starting with "n_" came from Darwin, or
130  * were added during the Solaris port. We have avoided
131  * renaming fields so we would not cause excessive
132  * changes in the code using this struct.
133  *
134  * Now using an AVL tree instead of hash lists, but kept the
135  * "hash" in some member names and functions to reduce churn.
136  * One AVL tree per mount replaces the global hash buckets.
137  *
138  * Notes carried over from the NFS code:
139  *
140  * The smbnode is the "inode" for remote files. It contains all the
141  * information necessary to handle remote file on the client side.
142  *
143  * Note on file sizes: we keep two file sizes in the smbnode: the size
144  * according to the client (r_size) and the size according to the server
145  * (r_attr.fa_size). They can differ because we modify r_size during a
146  * write system call (smbfs_rdw), before the write request goes over the
147  * wire (before the file is actually modified on the server). If an OTW
148  * request occurs before the cached data is written to the server the file
149  * size returned from the server (r_attr.fa_size) may not match r_size.
150  * r_size is the one we use, in general. r_attr.fa_size is only used to
151  * determine whether or not our cached data is valid.
152  *
153  * Each smbnode has 3 locks associated with it (not including the smbnode
154  * "hash" AVL tree and free list locks):
155  *
156  *     r_rwlock:      Serializes smbfs_write and smbfs_setattr requests
157  *                     and allows smbfs_read requests to proceed in parallel.
158  *                     Serializes reads/updates to directories.
159  *
160  *     r_lkserlock:   Serializes lock requests with map, write, and
161  *                     readahead operations.
162  *
163  *     r_statelock:   Protects all fields in the smbnode except for
164  *                     those listed below. This lock is intended
165  *                     to be held for relatively short periods of
166  *                     time (not accross entire putpage operations,
167  *                     for example).
168  *
169  * The following members are protected by the mutex smbfreelist_lock:
170  *     r_freef
171  *     r_freeb
172  *
173  * The following members are protected by the AVL tree rwlock:
174  *     r_avl_node      (r_hdr.hdr_avl_node)
175  *
176  * Note: r_modaddr is only accessed when the r_statelock mutex is held.
177  * Its value is also controlled via r_rwlock. It is assumed that
178  * there will be only 1 writer active at a time, so it safe to
179  * set r_modaddr and release r_statelock as long as the r_rwlock
180  * writer lock is held.
181  *
182  * 64-bit offsets: the code formerly assumed that atomic reads of

```

```

183  * r_size were safe and reliable; on 32-bit architectures, this is
184  * not true since an intervening bus cycle from another processor
185  * could update half of the size field. The r_statelock must now
186  * be held whenever any kind of access of r_size is made.
187  *
188  * Lock ordering:
189  *     r_rwlock > r_lkserlock > r_statelock
190  */

192 typedef struct smbnode {
193     /* Our linkage in the node cache AVL tree (see above). */
194     smbfs_node_hdr_t      r_hdr;

196     /* short-hand names for r_hdr members */
197     #define r_avl_node      r_hdr.hdr_avl_node
198     #define n_rpath         r_hdr.hdr_n_rpath
199     #define n_rplen        r_hdr.hdr_n_rplen

201     smbmntinfo_t          *n_mount;      /* VFS data */
202     vnode_t                *r_vnode;     /* associated vnode */

204     /*
205      * Linkage in smbfreelist, for reclaiming nodes.
206      * Lock for the free list is: smbfreelist_lock
207      */
208     struct smbnode        *r_freef;      /* free list forward pointer */
209     struct smbnode        *r_freeb;      /* free list back pointer */

211     smbfs_rwlock_t        r_rwlock;      /* serialize write/setattr requests */
212     smbfs_rwlock_t        r_lkserlock;   /* serialize lock with other ops */
213     kmutex_t              r_statelock;    /* protect (most) smbnode fields */

215     /*
216      * File handle, directory search handle,
217      * and reference counts for them, etc.
218      * Lock for these is: r_lkserlock
219      */
220     int                    n_dirrefs;
221     struct smbfs_ctx       *n_dirseq;     /* ff context */
222     int                    n_dirofs;      /* last ff offset */
223     int                    n_fidrefs;
224     uint16_t              n_fid;          /* file handle */
225     enum vtype             n_ovtype;      /* vnode type opened */
226     uint32_t              n_rights;       /* granted rights */
227     int                    n_vcgenid;     /* generation no. (reconnect) */

229     /*
230      * Misc. bookkeeping
231      */
232     cred_t                 *r_cred;       /* current credentials */
233     u_offset_t             r_nextr;       /* next read offset (read-ahead) */
234     long                   r_mapcnt;      /* count of mmaped pages */
235     uint_t                 r_inmap;       /* to serialize read/write and mmap */
236     uint_t                 r_count;       /* # of refs not reflect in v_count */
237     uint_t                 r_await;       /* # of outstanding async write */
238     uint_t                 r_gcount;      /* getattrs waiting to flush pages */
239     uint_t                 r_flags;       /* flags, see below */
240     uint32_t               n_flag;        /* N--- flags below */
241     uint32_t               n_flag;        /* NXXX flags below */
242     uint_t                 r_error;       /* async write error */
243     kcondvar_t             r_cv;          /* condvar for blocked threads */
244     avl_tree_t             r_dir;         /* cache of readdir responses */
245     rddir_cache            *r_direof;     /* pointer to the EOF entry */
246     u_offset_t             r_modaddr;     /* address for page in writenp */
247     kthread_t              *r_serial;     /* id of purging thread */
248     list_t                 r_indeimap;    /* list of delmap callers */

```

```
249     /*
250     * Attributes: local, and as last seen on the server.
251     * See notes above re: r_size vs r_attr.fa_size, etc.
252     */
253     smbattr_t      r_attr;          /* attributes from the server */
254     hrttime_t      r_attrtime;     /* time attributes become invalid */
255     hrttime_t      r_mtime;        /* client time file last modified */
256     len_t          r_size;         /* client's view of file size */
257
258     /*
259     * Security attributes.
260     */
261     vsecattr_t     r_secattr;
262     hrttime_t      r_sectime;
263
264     /*
265     * Other attributes, not carried in smbattr_t
266     */
267     u_longlong_t   n_ino;
268     uid_t          n_uid;
269     gid_t          n_gid;
270     mode_t         n_mode;
271 } smbnode_t;
272 unchanged_portion_omitted
```

new/usr/src/uts/common/fs/smbcInt/smbfs/smbfs_subr.h

1

11294 Sun Mar 4 06:09:09 2018
new/usr/src/uts/common/fs/smbcInt/smbfs/smbfs_subr.h
Implement ioctl_FIODIRECTIO
Kill flags arg in smbfs_purge_caches
5404 smbfs needs mmap support
Portions contributed by: Gordon Ross <gordon.w.ross@gmail.com>

unchanged portion omitted
149 typedef struct smbfs_fctx smbfs_fctx_t;

151 #define f_rq f_urq.uf_rq
152 #define f_t2 f_urq.uf_t2

154 /*
155 * smb level (smbfs_smb.c)
156 */
157 int smbfs_smb_lock(struct smbnode *np, int op, caddr_t id,
158 offset_t start, uint64_t len, int largelock,
159 struct smb_cred *scrp, uint32_t timeout);
160 int smbfs_smb_qfsattr(struct smb_share *ssp, struct smb_fs_attr_info *,
161 struct smb_cred *scrp);
162 int smbfs_smb_statfs(struct smb_share *ssp, statvfs64_t *sbp,
163 struct smb_cred *scrp);

165 int smbfs_smb_setdisp(struct smbnode *np, uint16_t fid, uint8_t newdisp,
166 struct smb_cred *scrp);
167 int smbfs_smb_setfsz(struct smbnode *np, uint16_t fid, uint64_t newsize,
168 struct smb_cred *scrp);

170 int smbfs_smb_getfattr(struct smbnode *np, struct smb_fattr *fap,
171 struct smb_cred *scrp);

173 int smbfs_smb_setfattr(struct smbnode *np, int fid,
174 uint32_t attr, struct timespec *mtime, struct timespec *atime,
175 struct smb_cred *scrp);

177 int smbfs_smb_open(struct smbnode *np, const char *name, int nmlen,
178 int xattr, uint32_t rights, struct smb_cred *scrp,
179 uint16_t *fidp, uint32_t *rightsp, struct smb_fattr *fap);
180 int smbfs_smb_tmpopen(struct smbnode *np, uint32_t rights,
181 struct smb_cred *scrp, uint16_t *fidp);
182 int smbfs_smb_close(struct smb_share *ssp, uint16_t fid,
183 struct timespec *mtime, struct smb_cred *scrp);
184 int smbfs_smb_tmpclose(struct smbnode *ssp, uint16_t fid,
185 struct smb_cred *scrp);
186 int smbfs_smb_create(struct smbnode *dnp, const char *name, int nmlen,
187 int xattr, uint32_t disp, struct smb_cred *scrp, uint16_t *fidp);
188 int smbfs_smb_delete(struct smbnode *np, struct smb_cred *scrp,
189 const char *name, int len, int xattr);
190 int smbfs_smb_rename(struct smbnode *src, struct smbnode *tdnp,
191 const char *tname, int tnmlen, struct smb_cred *scrp);
192 int smbfs_smb_t2rename(struct smbnode *np, const char *tname, int tnmlen,
193 struct smb_cred *scrp, uint16_t fid, int replace);
194 int smbfs_smb_move(struct smbnode *src, struct smbnode *tdnp,
195 const char *tname, int tnmlen, uint16_t flags, struct smb_cred *scrp);
196 int smbfs_smb_mkdir(struct smbnode *dnp, const char *name, int len,
197 struct smb_cred *scrp);
198 int smbfs_smb_rmdir(struct smbnode *np, struct smb_cred *scrp);
199 int smbfs_smb_findopen(struct smbnode *dnp, const char *wildcard, int wclen,
200 int attr, struct smb_cred *scrp, struct smb_fctx **ctxpp);
201 int smbfs_smb_findnext(struct smbfs_fctx *ctx, int limit,
202 struct smb_cred *scrp);
203 int smbfs_smb_findclose(struct smbfs_fctx *ctx, struct smb_cred *scrp);
204 int smbfs_fullpath(struct mbchain *mbp, struct smb_vc *vcp,
205 struct smbnode *dnp, const char *name, int nmlen, uint8_t sep);

new/usr/src/uts/common/fs/smbcInt/smbfs/smbfs_subr.h

2

206 int smbfs_smb_lookup(struct smbnode *dnp, const char **namep, int *nmlenp,
207 struct smb_fattr *fap, struct smb_cred *scrp);
208 int smbfs_smb_hideit(struct smbnode *np, const char *name, int len,
209 struct smb_cred *scrp);
210 int smbfs_smb_unhideit(struct smbnode *np, const char *name, int len,
211 struct smb_cred *scrp);
212 int smbfs_smb_flush(struct smbnode *np, struct smb_cred *scrp);
213 int smbfs_0extend(vnode_t *vp, uint16_t fid, len_t from, len_t to,
214 struct smb_cred *scrdp, int timo);

216 /* get/set security descriptor */
217 int smbfs_smb_getsec_m(struct smb_share *ssp, uint16_t fid,
218 struct smb_cred *scrp, uint32_t selector,
219 mblk_t **res, uint32_t *reslen);
220 int smbfs_smb_setsec_m(struct smb_share *ssp, uint16_t fid,
221 struct smb_cred *scrp, uint32_t selector, mblk_t **mp);

223 /*
224 * VFS-level init, fini stuff
225 */

227 int smbfs_vfsinit(void);
228 void smbfs_vfsfini(void);
229 int smbfs_subrinit(void);
230 void smbfs_subrfini(void);
231 int smbfs_clntinit(void);
232 void smbfs_clntfini(void);

234 void smbfs_zonelist_add(smbmntinfo_t *smi);
235 void smbfs_zonelist_remove(smbmntinfo_t *smi);

237 int smbfs_check_table(struct vfs *vfsp, struct smbnode *srp);
238 void smbfs_destroy_table(struct vfs *vfsp);
239 void smbfs_rflush(struct vfs *vfsp, cred_t *cr);
240 void smbfs_flushall(cred_t *cr);

242 int smbfs_directio(vnode_t *vp, int cmd, cred_t *cr);

244 uint32_t smbfs_newnum(void);
245 int smbfs_newname(char *buf, size_t buflen);

247 /*
248 * Function definitions - those having to do with
249 * smbfs nodes, vnodes, etc
250 */

252 void smbfs_attrcache_prune(struct smbnode *np);
253 void smbfs_attrcache_remove(struct smbnode *np);
254 void smbfs_attrcache_rm_locked(struct smbnode *np);
255 #ifdef DEBBUG
256 #define smbfs_attrcache_rm_locked(np) (np)->r_attrtime = gethrtime()
257 #endif
258 void smbfs_attr_touchdir(struct smbnode *dnp);
259 void smbfs_attrcache_fa(vnode_t *vp, struct smb_fattr *fap);
257 void smbfs_cache_check(struct vnode *vp, struct smb_fattr *fap);

261 int smbfs_validate_caches(struct vnode *vp, cred_t *cr);
262 void smbfs_purge_caches(struct vnode *vp, cred_t *cr);

264 void smbfs_addfree(struct smbnode *sp);
265 void smbfs_rmdhash(struct smbnode *);

267 /* See avl_create in smbfs_vfsops.c */
268 void smbfs_init_hash_avl(avl_tree_t *);

270 uint32_t smbfs_gethash(const char *rpath, int prlen);

```
271 uint32_t smbfs_getino(struct smbnode *dnp, const char *name, int nmlen);

273 extern struct smbfsattr smbfs_fattr0;
274 smbnode_t *smbfs_node_findcreate(smbmntinfo_t *mi,
275     const char *dir, int dirlen,
276     const char *name, int nmlen,
277     char sep, struct smbfsattr *fap);

279 int smbfs_nget(vnode_t *dvp, const char *name, int nmlen,
280     struct smbfsattr *fap, vnode_t **vpp);

282 void smbfs_fname_tolocal(struct smbfs_fctx *ctx);
283 char *smbfs_name_alloc(const char *name, int nmlen);
284 void smbfs_name_free(const char *name, int nmlen);

286 int smbfs_readvnode(vnode_t *, uio_t *, cred_t *, struct vattr *);
287 int smbfs_writevnode(vnode_t *vp, uio_t *uiop, cred_t *cr,
288     int ioflag, int timo);
289 int smbfsgetattr(vnode_t *vp, struct vattr *vap, cred_t *cr);

291 /* nfs: writerp writenp */
292 /* nfs_putpages? */
293 void smbfs_invalidate_pages(vnode_t *vp, u_offset_t off, cred_t *cr);

295 /* smbfs ACL support */
296 int smbfs_acl_getids(vnode_t *, cred_t *);
297 int smbfs_acl_setids(vnode_t *, vattr_t *, cred_t *);
298 int smbfs_acl_getvsa(vnode_t *, vsecattr_t *, int, cred_t *);
299 int smbfs_acl_setvsa(vnode_t *, vsecattr_t *, int, cred_t *);
300 int smbfs_acl_iocget(vnode_t *, intptr_t, int, cred_t *);
301 int smbfs_acl_iocset(vnode_t *, intptr_t, int, cred_t *);

303 /* smbfs_xattr.c */
304 int smbfs_get_xattrdir(vnode_t *dvp, vnode_t **vpp, cred_t *cr, int);
305 int smbfs_xa_parent(vnode_t *vp, vnode_t **vpp);
306 int smbfs_xa_exists(vnode_t *vp, cred_t *cr);
307 int smbfs_xa_getfattr(struct smbnode *np, struct smbfsattr *fap,
308     struct smb_cred *scrp);
309 int smbfs_xa_findopen(struct smbfs_fctx *ctx, struct smbnode *dnp,
310     const char *name, int nmlen);
311 int smbfs_xa_findnext(struct smbfs_fctx *ctx, uint16_t limit);
312 int smbfs_xa_findclose(struct smbfs_fctx *ctx);

314 /* For Solaris, interruptible rwlock */
315 int smbfs_rw_enter_sig(smbfs_rwlock_t *l, krw_t rw, int intr);
316 int smbfs_rw_tryenter(smbfs_rwlock_t *l, krw_t rw);
317 void smbfs_rw_exit(smbfs_rwlock_t *l);
318 int smbfs_rw_lock_held(smbfs_rwlock_t *l, krw_t rw);
319 void smbfs_rw_init(smbfs_rwlock_t *l, char *name, krw_type_t type, void *arg);
320 void smbfs_rw_destroy(smbfs_rwlock_t *l);

322 #endif /* !_FS_SMBFS_SMBFS_SUBR_H */
```

new/usr/src/uts/common/fs/smbcInt/smbfs/smbfs_subr2.c

1

32428 Sun Mar 4 06:09:09 2018
new/usr/src/uts/common/fs/smbcInt/smbfs/smbfs_subr2.c
cstyle

Implement ioctl_FIODIRECTIO

Lots of comment cleanup

5404 smbfs needs mmap support

Portions contributed by: Gordon Ross <gordon.w.ross@gmail.com>

```
1 /*
2  * CDDL HEADER START
3  *
4  * The contents of this file are subject to the terms of the
5  * Common Development and Distribution License (the "License").
6  * You may not use this file except in compliance with the License.
7  *
8  * You can obtain a copy of the license at usr/src/OPENSOLARIS.LICENSE
9  * or http://www.opensolaris.org/os/licensing.
10 * See the License for the specific language governing permissions
11 * and limitations under the License.
12 *
13 * When distributing Covered Code, include this CDDL HEADER in each
14 * file and include the License file at usr/src/OPENSOLARIS.LICENSE.
15 * If applicable, add the following below this CDDL HEADER, with the
16 * fields enclosed by brackets "[]" replaced with your own identifying
17 * information: Portions Copyright [yyyy] [name of copyright owner]
18 *
19 * CDDL HEADER END
20 */
21 /*
22 * Copyright 2010 Sun Microsystems, Inc. All rights reserved.
23 * Use is subject to license terms.
24 *
25 * Copyright (c) 1983,1984,1985,1986,1987,1988,1989 AT&T.
26 * All rights reserved.
27 */
28 /*
29 * Copyright (c) 2017 by Delphix. All rights reserved.
30 */
31
32 /*
33 * Node hash implementation initially borrowed from NFS (nfs_subr.c)
34 * but then heavily modified. It's no longer an array of hash lists,
35 * but an AVL tree per mount point. More on this below.
36 */
37
38 #include <sys/param.h>
39 #include <sys/systm.h>
40 #include <sys/time.h>
41 #include <sys/vnode.h>
42 #include <sys/bitmap.h>
43 #include <sys/dnld.h>
44 #include <sys/kmem.h>
45 #include <sys/sunddi.h>
46 #include <sys/sysmacros.h>
47 #include <sys/fcntl.h>
48
49 #include <netsmb/smb_osdep.h>
50
51 #include <netsmb/smb.h>
52 #include <netsmb/smb_conn.h>
53 #include <netsmb/smb_subr.h>
54 #include <netsmb/smb_rq.h>
55
56 #include <smbfs/smbfs.h>
57 #include <smbfs/smbfs_node.h>
```

new/usr/src/uts/common/fs/smbcInt/smbfs/smbfs_subr2.c

2

58 #include <smbfs/smbfs_subr.h>

```
60 /*
61 * The AVL trees (now per-mount) allow finding an smbfs node by its
62 * full remote path name. It also allows easy traversal of all nodes
63 * below (path wise) any given node. A reader/writer lock for each
64 * (per mount) AVL tree is used to control access and to synchronize
65 * lookups, additions, and deletions from that AVL tree.
66 *
67 * Previously, this code use a global array of hash chains, each with
68 * its own rwlock. A few struct members, functions, and comments may
69 * still refer to a "hash", and those should all now be considered to
70 * refer to the per-mount AVL tree that replaced the old hash chains.
71 * (i.e. member smi_hash_lk, function sn_hashfind, etc.)
72 *
73 * The smbnode freelist is organized as a doubly linked list with
74 * a head pointer. Additions and deletions are synchronized via
75 * a single mutex.
76 *
77 * In order to add an smbnode to the free list, it must be linked into
78 * the mount's AVL tree and the exclusive lock for the AVL must be held.
79 * If an smbnode is not linked into the AVL tree, then it is destroyed
80 * because it represents no valuable information that can be reused
81 * about the file. The exclusive lock for the AVL tree must be held
82 * in order to prevent a lookup in the AVL tree from finding the
83 * smbnode and using it and assuming that the smbnode is not on the
84 * freelist. The lookup in the AVL tree will have the AVL tree lock
85 * held, either exclusive or shared.
86 *
87 * The vnode reference count for each smbnode is not allowed to drop
88 * below 1. This prevents external entities, such as the VM
89 * subsystem, from acquiring references to vnodes already on the
90 * freelist and then trying to place them back on the freelist
91 * when their reference is released. This means that the when an
92 * smbnode is looked up in the AVL tree, then either the smbnode
93 * is removed from the freelist and that reference is tranfered to
94 * the new reference or the vnode reference count must be incremented
95 * accordingly. The mutex for the freelist must be held in order to
96 * accurately test to see if the smbnode is on the freelist or not.
97 * The AVL tree lock might be held shared and it is possible that
98 * two different threads may race to remove the smbnode from the
99 * freelist. This race can be resolved by holding the mutex for the
100 * freelist. Please note that the mutex for the freelist does not
101 * need to held if the smbnode is not on the freelist. It can not be
102 * placed on the freelist due to the requirement that the thread
103 * putting the smbnode on the freelist must hold the exclusive lock
104 * for the AVL tree and the thread doing the lookup in the AVL tree
105 * is holding either a shared or exclusive lock for the AVL tree.
106 *
107 * The lock ordering is:
108 *
109 *     AVL tree lock -> vnode lock
110 *     AVL tree lock -> freelist lock
111 */
112
113 static kmutex_t smbfreelist_lock;
114 static smbnode_t *smbfreelist = NULL;
115 static ulong_t smbnodenew = 0;
116 long nsmbnode = 0;
117
118 static struct kmem_cache *smbnode_cache;
119
120 static const vsecattr_t smbfs_vsa0 = { 0 };
121
122 /*
123 * Mutex to protect the following variables:
```

```

124 *      smbfs_major
125 *      smbfs_minor
126 */
127 kmutex_t smbfs_minor_lock;
128 int smbfs_major;
129 int smbfs_minor;

131 /* See smbfs_node_findcreate() */
132 struct smbfsattr smbfs_attr0;

134 /*
135  * Local functions.
136  * SN for Smb Node
137 */
138 static void sn_rmfree(smbnode_t *);
139 static void sn_inactive(smbnode_t *);
140 static void sn_addhash_locked(smbnode_t *, avl_index_t);
141 static void sn_rmhash_locked(smbnode_t *);
142 static void sn_destroy_node(smbnode_t *);
143 void smbfs_kmem_reclaim(void *cdrarg);

145 static smbnode_t *
146 sn_hashfind(smbmntinfo_t *, const char *, int, avl_index_t *);

148 static smbnode_t *
149 make_smbnode(smbmntinfo_t *, const char *, int, int *);

151 /*
152  * Free the resources associated with an smbnode.
153  * Note: This is different from smbfs_inactive
154  *
155  * From NFS: nfs_subr.c:rinactive
156  * NFS: nfs_subr.c:rinactive
157 */
157 static void
158 sn_inactive(smbnode_t *np)
159 {
160     vsecattr_t    ovsa;
161     cred_t        *oldcr;
162     char          *orpath;
163     int           orplen;
164     vnode_t       *vp;

166     /*
167      * Here NFS has:
168      * Flush and invalidate all pages (done by caller)
169      * Flush and invalidate all pages (todo)
170      * Free any held credentials and caches...
171      * etc. (See NFS code)
172      */
172     mutex_enter(&np->r_statelock);

174     ovsa = np->r_secattr;
175     np->r_secattr = smbfs_vsa0;
176     np->r_sectime = 0;

178     oldcr = np->r_cred;
179     np->r_cred = NULL;

181     orpath = np->n_rpath;
182     orplen = np->n_rplen;
183     np->n_rpath = NULL;
184     np->n_rplen = 0;

186     mutex_exit(&np->r_statelock);

```

```

188     vp = SMBTOV(np);
189     if (vn_has_cached_data(vp)) {
190         ASSERT3P(vp, ==, NULL);
191     }

193     if (ovsa.vsa_aclentp != NULL)
194         kmem_free(ovsa.vsa_aclentp, ovs.vsa_aclentsz);

196     if (oldcr != NULL)
197         crfree(oldcr);

199     if (orpath != NULL)
200         kmem_free(orpath, orplen + 1);
201 }

203 /*
204  * Find and optionally create an smbnode for the passed
205  * mountinfo, directory, separator, and name. If the
206  * desired smbnode already exists, return a reference.
207  * If the file attributes pointer is non-null, the node
208  * is created if necessary and linked into the AVL tree.
209  *
210  * Callers that need a node created but don't have the
211  * real attributes pass smbfs_attr0 to force creation.
212  *
213  * Note: make_smbnode() may upgrade the "hash" lock to exclusive.
214  *
215  * Based on NFS: nfs_subr.c:makenfsnode
216  * NFS: nfs_subr.c:makenfsnode
217 */
217 smbnode_t *
218 smbfs_node_findcreate(
219     smbmntinfo_t *mi,
220     const char *dirnm,
221     int dirlen,
222     const char *name,
223     int nmlen,
224     char sep,
225     struct smbfsattr *fap)
226 {
227     char tmpbuf[256];
228     size_t rpalloc;
229     char *p, *rpath;
230     int rplen;
231     smbnode_t *np;
232     vnode_t *vp;
233     int newnode;

235     /*
236      * Build the search string, either in tmpbuf or
237      * in allocated memory if larger than tmpbuf.
238      */
239     rplen = dirlen;
240     if (sep != '\0')
241         rplen++;
242     rplen += nmlen;
243     if (rplen < sizeof(tmpbuf)) {
244         /* use tmpbuf */
245         rpalloc = 0;
246         rpath = tmpbuf;
247     } else {
248         rpalloc = rplen + 1;
249         rpath = kmem_alloc(rpalloc, KM_SLEEP);
250     }
251     p = rpath;
252     bcopy(dirnm, p, dirlen);

```

```

253     p += dirlen;
254     if (sep != '\0')
255         *p++ = sep;
256     if (name != NULL) {
257         bcopy(name, p, nmlen);
258         p += nmlen;
259     }
260     ASSERT(p == rpath + rplen);

262     /*
263      * Find or create a node with this path.
264      */
265     rw_enter(&mi->smbhash_lk, RW_READER);
266     if (fap == NULL)
267         np = sn_hashfind(mi, rpath, rplen, NULL);
268     else
269         np = make_smbnode(mi, rpath, rplen, &newnode);
270     rw_exit(&mi->smbhash_lk);

272     if (rpalloc)
273         kmem_free(rpath, rpalloc);

275     if (fap == NULL) {
276         /*
277          * Caller is "just looking" (no create)
278          * so np may or may not be NULL here.
279          * Either way, we're done.
280          */
281         return (np);
282     }

284     /*
285      * We should have a node, possibly created.
286      * Do we have (real) attributes to apply?
287      */
288     ASSERT(np != NULL);
289     if (fap == &smbfs_fattr0)
290         return (np);

292     /*
293      * Apply the given attributes to this node,
294      * dealing with any cache impact, etc.
295      */
296     vp = SMBTOV(np);
297     if (!newnode) {
298         /*
299          * Found an existing node.
300          * Maybe purge caches...
301          */
302         smbfs_cache_check(vp, fap);
303     }
304     smbfs_attrcache_fa(vp, fap);

306     /*
307      * Note NFS sets vp->v_type here, assuming it
308      * can never change for the life of a node.
309      * We allow v_type to change, and set it in
310      * smbfs_attrcache(). Also: mode, uid, gid
311      */
312     return (np);
313 }

314 /*
315  * Here NFS has: nfs_subr.c:rtablehash
316  * NFS: nfs_subr.c:rtablehash
317  * We use smbfs_hash().

```

```

311  */

313  /*
314   * Find or create an smbnode.
315   * From NFS: nfs_subr.c:make_rnode
316   * NFS: nfs_subr.c:make_rnode
317   */
318  static smbnode_t *
319  make_smbnode(
320      smbntinfo_t *mi,
321      const char *rpath,
322      int rplen,
323      int *newnode)
324  {
325      smbnode_t *np;
326      smbnode_t *tnp;
327      vnode_t *vp;
328      vfs_t *vfsp;
329      avl_index_t where;
330      char *new_rpath = NULL;

332      ASSERT(RW_READ_HELD(&mi->smbhash_lk));
333      vfsp = mi->smb_vfsp;

335  start:
336      np = sn_hashfind(mi, rpath, rplen, NULL);
337      if (np != NULL) {
338          *newnode = 0;
339          return (np);
340      }

342      /* Note: will retake this lock below. */
343      rw_exit(&mi->smbhash_lk);

345      /*
346       * see if we can find something on the freelist
347       */
348      mutex_enter(&smbfreelist_lock);
349      if (smbfreelist != NULL && smbnodecount >= nsmbnode) {
350          np = smbfreelist;
351          sn_rmfree(np);
352          mutex_exit(&smbfreelist_lock);

354          vp = SMBTOV(np);

356          if (np->r_flags & RHASHED) {
357              smbntinfo_t *tmp_mi = np->n_mount;
358              ASSERT(tmp_mi != NULL);
359              rw_enter(&tmp_mi->smbhash_lk, RW_WRITER);
360              mutex_enter(&vp->v_lock);
361              if (vp->v_count > 1) {
362                  VN_RELE_LOCKED(vp);
363                  mutex_exit(&vp->v_lock);
364                  rw_exit(&tmp_mi->smbhash_lk);
365                  /* start over */
366                  rw_enter(&mi->smbhash_lk, RW_READER);
367                  goto start;
368              }
369              mutex_exit(&vp->v_lock);
370              sn_rhash_locked(np);
371              rw_exit(&tmp_mi->smbhash_lk);
372          }

374          sn_inactive(np);

376          mutex_enter(&vp->v_lock);

```

```

376     if (vp->v_count > 1) {
377         VN_RELE_LOCKED(vp);
378         mutex_exit(&vp->v_lock);
379         rw_enter(&mi->smb_hash_lk, RW_READER);
380         goto start;
381     }
382     mutex_exit(&vp->v_lock);
383     vn_invalid(vp);
384     /*
385      * destroy old locks before bzero'ing and
386      * recreating the locks below.
387      */
388     smbfs_rw_destroy(&np->r_rwlock);
389     smbfs_rw_destroy(&np->r_lkserlock);
390     mutex_destroy(&np->r_statelock);
391     cv_destroy(&np->r_cv);
392     /*
393      * Make sure that if smbnode is recycled then
394      * VFS count is decremented properly before
395      * reuse.
396      */
397     VFS_RELE(vp->v_vfsp);
398     vn_reinit(vp);
399 } else {
400     /*
401      * allocate and initialize a new smbnode
402      */
403     vnode_t *new_vp;
404
405     mutex_exit(&smbfreelist_lock);
406
407     np = kmem_cache_alloc(smbnode_cache, KM_SLEEP);
408     new_vp = vn_alloc(KM_SLEEP);
409
410     atomic_inc_ulong((ulong_t *)&smbnodenew);
411     vp = new_vp;
412 }
413
414 /*
415  * Allocate and copy the rpath we'll need below.
416  */
417 new_rpath = kmem_alloc(rplen + 1, KM_SLEEP);
418 bcopy(rpath, new_rpath, rplen);
419 new_rpath[rplen] = '\0';
420
421 /* Initialize smbnode_t */
422 bzero(np, sizeof(*np));
423
424 smbfs_rw_init(&np->r_rwlock, NULL, RW_DEFAULT, NULL);
425 smbfs_rw_init(&np->r_lkserlock, NULL, RW_DEFAULT, NULL);
426 mutex_init(&np->r_statelock, NULL, MUTEX_DEFAULT, NULL);
427 cv_init(&np->r_cv, NULL, CV_DEFAULT, NULL);
428 /* cv_init(&np->r_commit.c_cv, NULL, CV_DEFAULT, NULL); */
429
430 np->r_vnode = vp;
431 np->n_mount = mi;
432
433 np->n_fid = SMB_FID_UNUSED;
434 np->n_uid = mi->smb_uid;
435 np->n_gid = mi->smb_gid;
436 /* Leave attributes "stale." */
437 #if 0 /* XXX dircache */
438 /*
439  * Here NFS has avl_create(&np->r_dir, ...)
440  * for the readdir cache (not used here).

```

```

439     * We don't know if it's a directory yet.
440     * Let the caller do this? XXX
441     */
442     avl_create(&np->r_dir, compar, sizeof(rddir_cache),
443               offsetof(rddir_cache, tree));
444 #endif
445
446 /* Now fill in the vnode. */
447 vn_setops(vp, smbfs_vnodeops);
448 vp->v_data = (caddr_t)np;
449 VFS_HOLD(vfsp);
450 vp->v_vfsp = vfsp;
451 vp->v_type = VNON;
452
453 /*
454  * We entered with mi->smb_hash_lk held (reader).
455  * Retake it now, (as the writer).
456  * Will return with it held.
457  */
458 rw_enter(&mi->smb_hash_lk, RW_WRITER);
459
460 /*
461  * There is a race condition where someone else
462  * may alloc the smbnode while no locks are held,
463  * so check again and recover if found.
464  */
465 tnp = sn_hashfind(mi, rpath, rplen, &where);
466 if (tnp != NULL) {
467     /*
468      * Lost the race. Put the node we were building
469      * on the free list and return the one we found.
470      */
471     rw_exit(&mi->smb_hash_lk);
472     kmem_free(new_rpath, rplen + 1);
473     smbfs_addfree(np);
474     rw_enter(&mi->smb_hash_lk, RW_READER);
475     *newnode = 0;
476     return (tnp);
477 }
478
479 /*
480  * Hash search identifies nodes by the remote path
481  * (n_rpath) so fill that in now, before linking
482  * this node into the node cache (AVL tree).
483  */
484 np->n_rpath = new_rpath;
485 np->n_rplen = rplen;
486 np->n_ino = smbfs_gethash(new_rpath, rplen);
487
488 sn_addhash_locked(np, where);
489 *newnode = 1;
490 return (np);
491 }
492
493 /*
494  * smbfs_addfree
495  * Put an smbnode on the free list, or destroy it immediately
496  * if it offers no value were it to be reclaimed later. Also
497  * destroy immediately when we have too many smbnodes, etc.
498  *
499  * Normally called by smbfs_inactive, but also
500  * called in here during cleanup operations.
501  *
502  * From NFS: nfs_subr.c:rp_addfree
503  * NFS: nfs_subr.c:rp_addfree
504  */

```

```

501 void
502 smbfs_addfree(smbnode_t *np)
503 {
504     vnode_t *vp;
505     struct vfs *vfsp;
506     smbmntinfo_t *mi;
507
508     ASSERT(np->r_freeb == NULL && np->r_freeb == NULL);
509
510     vp = SMBTOV(np);
511     ASSERT(vp->v_count >= 1);
512
513     vfsp = vp->v_vfsp;
514     mi = VFTOSMI(vfsp);
515
516     /*
517      * If there are no more references to this smbnode and:
518      * we have too many smbnodes allocated, or if the node
519      * is no longer accessible via the AVL tree (!RHASHED),
520      * or an i/o error occurred while writing to the file,
521      * or it's part of an unmounted FS, then try to destroy
522      * it instead of putting it on the smbnode freelist.
523      */
524     if (np->r_count == 0 && (
525         (np->r_flags & RHASHED) == 0 ||
526         (np->r_error != 0) ||
527         (vfsp->vfs_flag & VFS_UNMOUNTED) ||
528         (smbnodenew > nsmbnode))) {
529
530         /* Try to destroy this node. */
531
532         if (np->r_flags & RHASHED) {
533             rw_enter(&mi->smi_hash_lk, RW_WRITER);
534             mutex_enter(&vp->v_lock);
535             if (vp->v_count > 1) {
536                 VN_RELE_LOCKED(vp);
537                 mutex_exit(&vp->v_lock);
538                 rw_exit(&mi->smi_hash_lk);
539                 return;
540             }
541             /*
542              * Will get another call later,
543              * via smbfs_inactive.
544              */
545             mutex_exit(&vp->v_lock);
546             sn_rmhash_locked(np);
547             rw_exit(&mi->smi_hash_lk);
548         }
549
550         sn_inactive(np);
551
552         /*
553          * Recheck the vnode reference count. We need to
554          * make sure that another reference has not been
555          * acquired while we were not holding v_lock. The
556          * smbnode is not in the smbnode "hash" AVL tree, so
557          * the only way for a reference to have been acquired
558          * is for a VOP_PUTPAGE because the smbnode was marked
559          * with RDIRTY or for a modified page. This vnode
560          * reference may have been acquired before our call
561          * to sn_inactive. The i/o may have been completed,
562          * thus allowing sn_inactive to complete, but the
563          * reference to the vnode may not have been released
564          * yet. In any case, the smbnode can not be destroyed
565          * until the other references to this vnode have been
566          * released. The other references will take care of

```

```

567     * either destroying the smbnode or placing it on the
568     * smbnode freelist. If there are no other references,
569     * then the smbnode may be safely destroyed.
570     */
571     mutex_enter(&vp->v_lock);
572     if (vp->v_count > 1) {
573         VN_RELE_LOCKED(vp);
574         mutex_exit(&vp->v_lock);
575         return;
576     }
577     mutex_exit(&vp->v_lock);
578
579     sn_destroy_node(np);
580     return;
581 }
582
583 /*
584  * Lock the AVL tree and then recheck the reference count
585  * to ensure that no other threads have acquired a reference
586  * to indicate that the smbnode should not be placed on the
587  * freelist. If another reference has been acquired, then
588  * just release this one and let the other thread complete
589  * the processing of adding this smbnode to the freelist.
590  */
591 rw_enter(&mi->smi_hash_lk, RW_WRITER);
592
593 mutex_enter(&vp->v_lock);
594 if (vp->v_count > 1) {
595     VN_RELE_LOCKED(vp);
596     mutex_exit(&vp->v_lock);
597     rw_exit(&mi->smi_hash_lk);
598     return;
599 }
600 mutex_exit(&vp->v_lock);
601
602 /*
603  * Put this node on the free list.
604  */
605 mutex_enter(&smbfreelist_lock);
606 if (smbfreelist == NULL) {
607     np->r_freeb = np;
608     np->r_freeb = np;
609     smbfreelist = np;
610 } else {
611     np->r_freeb = smbfreelist;
612     np->r_freeb = smbfreelist->r_freeb;
613     smbfreelist->r_freeb->r_freeb = np;
614     smbfreelist->r_freeb = np;
615 }
616 mutex_exit(&smbfreelist_lock);
617
618 rw_exit(&mi->smi_hash_lk);
619 }
620
621 /*
622  * Remove an smbnode from the free list.
623  */
624 * The caller must be holding smbfreelist_lock and the smbnode
625 * must be on the freelist.
626
627 * From NFS: nfs_subr.c:rp_rmfree
628 * NFS: nfs_subr.c:rp_rmfree
629 */
630 static void
631 sn_rmfree(smbnode_t *np)
632 {

```

```

633     ASSERT(MUTEX_HELD(&smbfreelist_lock));
634     ASSERT(np->r_freef != NULL && np->r_freeb != NULL);

636     if (np == smbfreelist) {
637         smbfreelist = np->r_freef;
638         if (np == smbfreelist)
639             smbfreelist = NULL;
640     }

642     np->r_freeb->r_freef = np->r_freef;
643     np->r_freef->r_freeb = np->r_freeb;

645     np->r_freef = np->r_freeb = NULL;
646 }

648 /*
649  * Put an smbnode in the "hash" AVL tree.
650  * The caller must be hold the rwlock as writer.
651  *
652  * From NFS: nfs_subr.c:rp_addhash
653  * NFS: nfs_subr.c:rp_addhash
654  */
655 static void
656 sn_addhash_locked(smbnode_t *np, avl_index_t where)
657 {
658     smbmntinfo_t *mi = np->n_mount;

660     ASSERT(RW_WRITE_HELD(&mi->smi_hash_lk));

662     mutex_enter(&np->r_statelock);
663     if ((np->r_flags & RHASHED) == 0) {
664         avl_insert(&mi->smi_hash_avl, np, where);
665         np->r_flags |= RHASHED;
666     }
667     mutex_exit(&np->r_statelock);
668 }

670 /*
671  * Remove an smbnode from the "hash" AVL tree.
672  * The caller must hold the rwlock as writer.
673  *
674  * From NFS: nfs_subr.c:rp_rmlhash_locked
675  * NFS: nfs_subr.c:rp_rmlhash_locked
676  */
677 static void
678 sn_rmlhash_locked(smbnode_t *np)
679 {
680     smbmntinfo_t *mi = np->n_mount;

682     ASSERT(RW_WRITE_HELD(&mi->smi_hash_lk));

684     mutex_enter(&np->r_statelock);
685     if ((np->r_flags & RHASHED) != 0) {
686         np->r_flags &= ~RHASHED;
687         avl_remove(&mi->smi_hash_avl, np);
688     }
689     mutex_exit(&np->r_statelock);
690 }

```

unchanged_portion_omitted

```

707 /*
708  * Lookup an smbnode by remote pathname
709  */

```

```

710  * The caller must be holding the AVL rwlock, either shared or exclusive.
711  *
712  * From NFS: nfs_subr.c:rfind
713  * NFS: nfs_subr.c:rfind
714  */
715 static smbnode_t *
716 sn_hashfind(
717     smbmntinfo_t *mi,
718     const char *rpath,
719     int rplen,
720     avl_index_t *pwhere) /* optional */
721 {
722     smbfs_node_hdr_t nhdr;
723     smbnode_t *np;
724     vnode_t *vp;

725     ASSERT(RW_LOCK_HELD(&mi->smi_hash_lk));

727     bzero(&nhdr, sizeof (nhdr));
728     nhdr.hdr_n_rpath = (char *)rpath;
729     nhdr.hdr_n_rplen = rplen;

731     /* See smbfs_node_cmp below. */
732     np = avl_find(&mi->smi_hash_avl, &nhdr, pwhere);

734     if (np == NULL)
735         return (NULL);

737     /*
738      * Found it in the "hash" AVL tree.
739      * Remove from free list, if necessary.
740      */
741     vp = SMBTOV(np);
742     if (np->r_freef != NULL) {
743         mutex_enter(&smbfreelist_lock);
744         /*
745          * If the smbnode is on the freelist,
746          * then remove it and use that reference
747          * as the new reference. Otherwise,
748          * need to increment the reference count.
749          */
750         if (np->r_freef != NULL) {
751             sn_rmfree(np);
752             mutex_exit(&smbfreelist_lock);
753         } else {
754             mutex_exit(&smbfreelist_lock);
755             VN_HOLD(vp);
756         }
757     } else
758         VN_HOLD(vp);

760     return (np);
761 }

```

unchanged_portion_omitted

```

851 #ifdef SMB_VNODE_DEBUG
852 int smbfs_check_table_debug = 1;
853 #else /* SMB_VNODE_DEBUG */
854 int smbfs_check_table_debug = 0;
855 #endif /* SMB_VNODE_DEBUG */

858 /*
859  * Return 1 if there is a active vnode belonging to this vfs in the
860  * smbnode cache.
861  */

```

```

862 * Several of these checks are done without holding the usual
863 * locks. This is safe because destroy_smbtable(), smbfs_addfree(),
864 * etc. will redo the necessary checks before actually destroying
865 * any smbnodes.
866 *
867 * From NFS: nfs_subr.c:check_rtable
870 * NFS: nfs_subr.c:check_rtable
868 *
869 * Debugging changes here relative to NFS.
870 * Relatively harmless, so left 'em in.
871 */
872 int
873 smbfs_check_table(struct vfs *vfsp, smbnode_t *rtnp)
874 {
875     smbmntinfo_t *mi;
876     smbnode_t *np;
877     vnode_t *vp;
878     int busycnt = 0;
879
880     mi = VFTOSMI(vfsp);
881     rw_enter(&mi->smb_hash_lk, RW_READER);
882     for (np = avl_first(&mi->smb_hash_avl); np != NULL;
883          np = avl_walk(&mi->smb_hash_avl, np, AVL_AFTER)) {
884
885         if (np == rtnp)
886             continue; /* skip the root */
887         vp = SMBTOV(np);
888
889         /* Now the 'busy' checks: */
890         /* Not on the free list? */
891         if (np->r_freef == NULL) {
892             SMBVDEBUG("r_freef: node=0x%p, rpath=%s\n",
893                      (void *)np, np->n_rpath);
894             busycnt++;
895         }
896
897         /* Has dirty pages? */
898         if (vn_has_cached_data(vp) &&
899             (np->r_flags & RDIRTY)) {
900             SMBVDEBUG("is dirty: node=0x%p, rpath=%s\n",
901                      (void *)np, np->n_rpath);
902             busycnt++;
903         }
904
905         /* Other refs? (not reflected in v_count) */
906         if (np->r_count > 0) {
907             SMBVDEBUG("r_count: node=0x%p, rpath=%s\n",
908                      (void *)np, np->n_rpath);
909             busycnt++;
910         }
911
912         if (busycnt && !smbfs_check_table_debug)
913             break;
914     }
915     rw_exit(&mi->smb_hash_lk);
916
917     return (busycnt);
918 }
919
920 /*
921 * Destroy inactive vnodes from the AVL tree which belong to this
922 * vfs. It is essential that we destroy all inactive vnodes during a
923 * forced unmount as well as during a normal unmount.
924 *
925 * Based on NFS: nfs_subr.c:destroy_rtable

```

```

929 * NFS: nfs_subr.c:destroy_rtable
927 *
928 * In here, we're normally destroying all or most of the AVL tree,
929 * so the natural choice is to use avl_destroy_nodes. However,
930 * there may be a few busy nodes that should remain in the AVL
931 * tree when we're done. The solution: use a temporary tree to
932 * hold the busy nodes until we're done destroying the old tree,
933 * then copy the temporary tree over the (now empty) real tree.
934 */
935 void
936 smbfs_destroy_table(struct vfs *vfsp)
937 {
938     avl_tree_t tmp_avl;
939     smbmntinfo_t *mi;
940     smbnode_t *np;
941     smbnode_t *rlist;
942     void *v;
943
944     mi = VFTOSMI(vfsp);
945     rlist = NULL;
946     smbfs_init_hash_avl(&tmp_avl);
947
948     rw_enter(&mi->smb_hash_lk, RW_WRITER);
949     v = NULL;
950     while ((np = avl_destroy_nodes(&mi->smb_hash_avl, &v)) != NULL) {
951
952         mutex_enter(&smbfreelist_lock);
953         if (np->r_freef == NULL) {
954             /*
955              * Busy node (not on the free list).
956              * Will keep in the final AVL tree.
957              */
958             mutex_exit(&smbfreelist_lock);
959             avl_add(&tmp_avl, np);
960         } else {
961             /*
962              * It's on the free list. Remove and
963              * arrange for it to be destroyed.
964              */
965             sn_rmfree(np);
966             mutex_exit(&smbfreelist_lock);
967
968             /*
969              * Last part of sn_rmhask_locked().
970              * NB: avl_destroy_nodes has already
971              * removed this from the "hash" AVL.
972              */
973             mutex_enter(&np->r_statelock);
974             np->r_flags &= ~RHASHED;
975             mutex_exit(&np->r_statelock);
976
977             /*
978              * Add to the list of nodes to destroy.
979              * Borrowing avl_child[0] for this list.
980              */
981             np->r_avl_node.avl_child[0] =
982                 (struct avl_node *)rlist;
983             rlist = np;
984         }
985     }
986     avl_destroy(&mi->smb_hash_avl);
987
988     /*
989     * Replace the (now destroyed) "hash" AVL with the
990     * temporary AVL, which restores the busy nodes.
991     */

```

```

992      mi->smi_hash_avl = tmp_avl;
993      rw_exit(&mi->smi_hash_lk);

995      /*
996       * Now destroy the nodes on our temporary list (rlist).
997       * This call to smbfs_addfree will end up destroying the
998       * smbnode, but in a safe way with the appropriate set
999       * of checks done.
1000      */
1001      while ((np = rlist) != NULL) {
1002          rlist = (smbnode_t *)np->r_avl_node.avl_child[0];
1003          smbfs_addfree(np);
1004      }
1005 }

1007 /*
1008 * This routine destroys all the resources associated with the smbnode
1009 * and then the smbnode itself. Note: sn_inactive has been called.
1010 *
1011 * From NFS: nfs_subr.c:destroy_rnode
1012 * NFS: nfs_subr.c:destroy_rnode
1013 */
1014 static void
1015 sn_destroy_node(smbnode_t *np)
1016 {
1017     vnode_t *vp;
1018     vfs_t *vfsp;

1019     vp = SMBTOV(np);
1020     vfsp = vp->v_vfsp;

1022     ASSERT(vp->v_count == 1);
1023     ASSERT(np->r_count == 0);
1024     ASSERT(np->r_mapcnt == 0);
1025     ASSERT(np->r_secattr.vsa_aclentp == NULL);
1026     ASSERT(np->r_cred == NULL);
1027     ASSERT(np->n_rpath == NULL);
1028     ASSERT(!(np->r_flags & RHASHED));
1029     ASSERT(np->r_freeb == NULL && np->r_freeb == NULL);
1030     atomic_dec_ulong((ulong_t *)&smbnodenew);
1031     vn_invalid(vp);
1032     vn_free(vp);
1033     kmem_cache_free(smbnode_cache, np);
1034     VFS_RELE(vfsp);
1035 }

1037 /*
1038 * From NFS rflush()
1039 * Flush all vnodes in this (or every) vfs.
1040 * Used by smbfs_sync and by smbfs_unmount.
1041 * Used by nfs_sync and by nfs_unmount.
1042 */
1043 /*ARGSUSED*/
1044 void
1045 smbfs_rflush(struct vfs *vfsp, cred_t *cr)
1046 {
1047     smbmntinfo_t *mi;
1048     smbnode_t *np;
1049     vnode_t *vp, **vplist;
1050     long num, cnt;

1051     mi = VFTOSMI(vfsp);

1053     /*
1054      * Check to see whether there is anything to do.
1055      */

```

```

1056     num = avl_numnodes(&mi->smi_hash_avl);
1057     if (num == 0)
1058         return;

1060     /*
1061      * Allocate a slot for all currently active rnodes on the
1062      * supposition that they all may need flushing.
1063      */
1064     vplist = kmem_alloc(num * sizeof (*vplist), KM_SLEEP);
1065     cnt = 0;

1067     /*
1068      * Walk the AVL tree looking for rnodes with page
1069      * lists associated with them. Make a list of these
1070      * files.
1071      */
1072     rw_enter(&mi->smi_hash_lk, RW_READER);
1073     for (np = avl_first(&mi->smi_hash_avl); np != NULL;
1074          np = avl_walk(&mi->smi_hash_avl, np, AVL_AFTER)) {
1075         vp = SMBTOV(np);
1076         /*
1077          * Don't bother sync'ing a vp if it
1078          * is part of virtual swap device or
1079          * if VFS is read-only
1080          */
1081         if (IS_SWAPVP(vp) || vn_is_readonly(vp))
1082             continue;
1083         /*
1084          * If the vnode has pages and is marked as either
1085          * dirty or mmap'd, hold and add this vnode to the
1086          * list of vnodes to flush.
1087          */
1088         if (vn_has_cached_data(vp) &&
1089             ((np->r_flags & RDIRTY) || np->r_mapcnt > 0)) {
1090             VN_HOLD(vp);
1091             vplist[cnt++] = vp;
1092             if (cnt == num)
1093                 break;
1094         }
1095     }
1096     rw_exit(&mi->smi_hash_lk);

1098     /*
1099      * Flush and release all of the files on the list.
1100      */
1101     while (cnt-- > 0) {
1102         vp = vplist[cnt];
1103         (void) VOP_PUTPAGE(vp, (u_offset_t)0, 0, B_ASYNC, cr, NULL);
1104         VN_RELE(vp);
1105     }

1107     kmem_free(vplist, num * sizeof (vnode_t *));
1108     /* Todo: mmap support. */

1110     /* Here NFS has access cache stuff (nfs_subr.c) not used here */
1111     /* access cache (nfs_subr.c) not used here */

1112     /*
1113      * Set or Clear direct I/O flag
1114      * VOP_RWLOCK() is held for write access to prevent a race condition
1115      * which would occur if a process is in the middle of a write when
1116      * directio flag gets set. It is possible that all pages may not get flushed.
1117      * From nfs_common.c
1118      */

```

```

1120 /* ARGSUSED */
1121 int
1122 smbfs_directio(vnode_t *vp, int cmd, cred_t *cr)
1123 {
1124     int     error = 0;
1125     smbnode_t *np;
1126
1127     np = VTOSMB(vp);
1128
1129     if (cmd == DIRECTIO_ON) {
1130
1131         if (np->r_flags & RDIRECTIO)
1132             return (0);
1133
1134         /*
1135          * Flush the page cache.
1136          */
1137
1138         (void) VOP_RWLOCK(vp, V_WRITELOCK_TRUE, NULL);
1139
1140         if (np->r_flags & RDIRECTIO) {
1141             VOP_RWUNLOCK(vp, V_WRITELOCK_TRUE, NULL);
1142             return (0);
1143         }
1144
1145         /* Here NFS also checks ->r_awcount */
1146         if (vn_has_cached_data(vp) &&
1147             (np->r_flags & RDIRTY) != 0) {
1148             error = VOP_PUTPAGE(vp, (offset_t)0, (uint_t)0,
1149                 B_INVAL, cr, NULL);
1150             if (error) {
1151                 if (error == ENOSPC || error == EDQUOT) {
1152                     mutex_enter(&np->r_statelock);
1153                     if (!np->r_error)
1154                         np->r_error = error;
1155                     mutex_exit(&np->r_statelock);
1156                 }
1157                 VOP_RWUNLOCK(vp, V_WRITELOCK_TRUE, NULL);
1158                 return (error);
1159             }
1160         }
1161
1162         mutex_enter(&np->r_statelock);
1163         np->r_flags |= RDIRECTIO;
1164         mutex_exit(&np->r_statelock);
1165         VOP_RWUNLOCK(vp, V_WRITELOCK_TRUE, NULL);
1166         return (0);
1167     }
1168
1169     if (cmd == DIRECTIO_OFF) {
1170         mutex_enter(&np->r_statelock);
1171         np->r_flags &= ~RDIRECTIO; /* disable direct mode */
1172         mutex_exit(&np->r_statelock);
1173         return (0);
1174     }
1175
1176     return (EINVAL);
1177 }

```

```

1179 static kmutex_t smbfs_newnum_lock;
1180 static uint32_t smbfs_newnum_val = 0;

```

```

1182 /*
1183  * Return a number 0..0xffffffff that's different from the last
1184  * 0xffffffff numbers this returned. Used for unlinked files.
1185  * From NFS nfs_subr.c newnum

```

```

1059 /* (This too was copied from nfs_subr.c)
1060 */
1061 uint32_t
1062 smbfs_newnum(void)
1063 {
1064     uint32_t id;
1065
1066     mutex_enter(&smbfs_newnum_lock);
1067     if (smbfs_newnum_val == 0)
1068         smbfs_newnum_val = (uint32_t)gethrtime_sec();
1069     id = smbfs_newnum_val++;
1070     mutex_exit(&smbfs_newnum_lock);
1071     return (id);
1072 }

```

unchanged_portion_omitted

```

1215 /*
1216  * initialize resources that are used by smbfs_subr.c
1217  * this is called from the _init() routine (by the way of smbfs_clntinit())
1218  */
1219 * From NFS: nfs_subr.c:nfs_subrinit
1220 * NFS: nfs_subr.c:nfs_subrinit
1221 */
1222 int
1223 smbfs_subrinit(void)
1224 {
1225     ulong_t nsmbnode_max;
1226
1227     /*
1228      * Allocate and initialize the smbnode cache
1229      */
1230     if (nsmbnode <= 0)
1231         nsmbnode = ncsz; /* dnld.h */
1232     nsmbnode_max = (ulong_t)((kmem_maxavail() >> 2) /
1233         sizeof(struct smbnode));
1234     if (nsmbnode > nsmbnode_max || (nsmbnode == 0 && ncsz == 0)) {
1235         zcmn_err(GLOBAL_ZONEID, CE_NOTE,
1236             "setting nsmbnode to max value of %ld", nsmbnode_max);
1237         nsmbnode = nsmbnode_max;
1238     }
1239
1240     smbnode_cache = kmem_cache_create("smbnode_cache", sizeof(smbnode_t),
1241         0, NULL, NULL, smbfs_kmem_reclaim, NULL, NULL, 0);
1242
1243     /*
1244      * Initialize the various mutexes and reader/writer locks
1245      */
1246     mutex_init(&smbfreelist_lock, NULL, MUTEX_DEFAULT, NULL);
1247     mutex_init(&smbfs_minor_lock, NULL, MUTEX_DEFAULT, NULL);
1248
1249     /*
1250      * Assign unique major number for all smbfs mounts
1251      */
1252     if ((smbfs_major = getudev()) == -1) {
1253         zcmn_err(GLOBAL_ZONEID, CE_WARN,
1254             "smbfs: init: can't get unique device number");
1255         smbfs_major = 0;
1256     }
1257     smbfs_minor = 0;
1258
1259     return (0);
1260 }

```

```

1261 /*
1262  * free smbfs hash table, etc.

```

```
1263  * From NFS: nfs_subr.c:nfs_subrfini
1137  * NFS: nfs_subr.c:nfs_subrfini
1264  */
1265  void
1266  smbfs_subrfini(void)
1267  {
1269      /*
1270       * Destroy the smbnode cache
1271       */
1272      kmem_cache_destroy(smbnode_cache);
1274      /*
1275       * Destroy the various mutexes and reader/writer locks
1276       */
1277      mutex_destroy(&smbfreelist_lock);
1278      mutex_destroy(&smbfs_minor_lock);
1279  }
unchanged_portion_omitted
```

```
1338  /*
1339  * Here NFS has failover stuff and
1340  * nfs_rw_xxx - see smbfs_rwlock.c
1341  */
1212  /* nfs failover stuff */
1213  /* nfs_rw_xxx - see smbfs_rwlock.c */
```

new/usr/src/uts/common/fs/smbclnt/smbfs/smbfs_vfsops.c

1

```
*****
25639 Sun Mar  4 06:09:09 2018
new/usr/src/uts/common/fs/smbclnt/smbfs/smbfs_vfsops.c
5404 smbfs needs mmap support
Portions contributed by: Gordon Ross <gordon.w.ross@gmail.com>
*****
```

```
_____unchanged_portion_omitted_____

884 static kmutex_t smbfs_syncbusy;

884 /*
885  * Flush dirty smbfs files for file system vfsp.
886  * If vfsp == NULL, all smbfs files are flushed.
887  */
888 /*ARGSUSED*/
889 static int
890 smbfs_sync(vfs_t *vfsp, short flag, cred_t *cr)
891 {
893     /*
894     * SYNC_ATTR is used by fsflush() to force old filesystems like UFS
895     * to sync metadata, which they would otherwise cache indefinitely.
896     * Semantically, the only requirement is that the sync be initiated.
897     * Assume the server-side takes care of attribute sync.
898     * Cross-zone calls are OK here, since this translates to a
899     * VOP_PUTPAGE(B_ASYNC), which gets picked up by the right zone.
900     */
901     if (flag & SYNC_ATTR)
902         return (0);
903
904     if (vfsp == NULL) {
905         /*
906         * Flush ALL smbfs mounts in this zone.
907         */
908         smbfs_flushall(cr);
909         return (0);
910     }
911     if (!(flag & SYNC_ATTR) && mutex_tryenter(&smbfs_syncbusy) != 0) {
912         smbfs_rflush(vfsp, cr);
913         mutex_exit(&smbfs_syncbusy);
914     }
915
916     smbfs_rflush(vfsp, cr);
917
918     return (0);
919 }

920 /*
921  * Initialization routine for VFS routines.  Should only be called once
922  */
923 int
924 smbfs_vfsinit(void)
925 {
926     mutex_init(&smbfs_syncbusy, NULL, MUTEX_DEFAULT, NULL);
927     return (0);
928 }

929 /*
930  * Shutdown routine for VFS routines.  Should only be called once
931  */
932 void
933 smbfs_vfsfini(void)
934 {
935     mutex_destroy(&smbfs_syncbusy);
936 }
_____unchanged_portion_omitted_____
```

new/usr/src/uts/common/fs/smbclnt/smbfs/smbfs_vnops.c

1

```
*****
123699 Sun Mar  4 06:09:10 2018
new/usr/src/uts/common/fs/smbclnt/smbfs/smbfs_vnops.c
cstyle
Fix missing logic for ESTALE and NFS_EOF
Implement ioctl_FIODIRECTIO
Kill flags arg in smbfs_purge_caches
Lots of comment cleanup
5404 smbfs needs mmap support
Portions contributed by: Gordon Ross <gordon.w.ross@gmail.com>
*****
```

```
1 /*
2  * Copyright (c) 2000-2001 Boris Popov
3  * All rights reserved.
4  *
5  * Redistribution and use in source and binary forms, with or without
6  * modification, are permitted provided that the following conditions
7  * are met:
8  * 1. Redistributions of source code must retain the above copyright
9  *   notice, this list of conditions and the following disclaimer.
10 * 2. Redistributions in binary form must reproduce the above copyright
11 *   notice, this list of conditions and the following disclaimer in the
12 *   documentation and/or other materials provided with the distribution.
13 * 3. All advertising materials mentioning features or use of this software
14 *   must display the following acknowledgement:
15 *   This product includes software developed by Boris Popov.
16 * 4. Neither the name of the author nor the names of any co-contributors
17 *   may be used to endorse or promote products derived from this software
18 *   without specific prior written permission.
19 *
20 * THIS SOFTWARE IS PROVIDED BY THE AUTHOR AND CONTRIBUTORS ``AS IS'' AND
21 * ANY EXPRESS OR IMPLIED WARRANTIES, INCLUDING, BUT NOT LIMITED TO, THE
22 * IMPLIED WARRANTIES OF MERCHANTABILITY AND FITNESS FOR A PARTICULAR PURPOSE
23 * ARE DISCLAIMED. IN NO EVENT SHALL THE AUTHOR OR CONTRIBUTORS BE LIABLE
24 * FOR ANY DIRECT, INDIRECT, INCIDENTAL, SPECIAL, EXEMPLARY, OR CONSEQUENTIAL
25 * DAMAGES (INCLUDING, BUT NOT LIMITED TO, PROCUREMENT OF SUBSTITUTE GOODS
26 * OR SERVICES; LOSS OF USE, DATA, OR PROFITS; OR BUSINESS INTERRUPTION)
27 * HOWEVER CAUSED AND ON ANY THEORY OF LIABILITY, WHETHER IN CONTRACT, STRICT
28 * LIABILITY, OR TORT (INCLUDING NEGLIGENCE OR OTHERWISE) ARISING IN ANY WAY
29 * OUT OF THE USE OF THIS SOFTWARE, EVEN IF ADVISED OF THE POSSIBILITY OF
30 * SUCH DAMAGE.
31 *
32 * $Id: smbfs_vnops.c,v 1.128.36.1 2005/05/27 02:35:28 lindak Exp $
33 */
34
35 /*
36 * Copyright (c) 2008, 2010, Oracle and/or its affiliates. All rights reserved.
37 */
38
39 /*
40  * Vnode operations
41  *
42  * This file is similar to nfs3_vnops.c
43  */
44
45 #include <sys/param.h>
46 #include <sys/system.h>
47 #include <sys/cred.h>
48 #include <sys/vnode.h>
49 #include <sys/vfs.h>
50 #include <sys/filio.h>
51 #include <sys/uio.h>
52 #include <sys/dirent.h>
53 #include <sys/errno.h>
54 #include <sys/sunddi.h>
55 #include <sys/sysmacros.h>
```

new/usr/src/uts/common/fs/smbclnt/smbfs/smbfs_vnops.c

2

```
56 #include <sys/kmem.h>
57 #include <sys/cmn_err.h>
58 #include <sys/vfs_opreg.h>
59 #include <sys/policy.h>
60 #include <sys/sdt.h>
61 #include <sys/zone.h>
62 #include <sys/vmsysm.h>
63
64 #include <vm/hat.h>
65 #include <vm/as.h>
66 #include <vm/page.h>
67 #include <vm/pvn.h>
68 #include <vm/seg.h>
69 #include <vm/seg_map.h>
70 #include <vm/seg_kpm.h>
71 #include <vm/seg_vn.h>
72
73 #include <netsmb/smb_osdep.h>
74 #include <netsmb/smb.h>
75 #include <netsmb/smb_conn.h>
76 #include <netsmb/smb_subr.h>
77
78 #include <smbfs/smbfs.h>
79 #include <smbfs/smbfs_node.h>
80 #include <smbfs/smbfs_subr.h>
81
82 #include <sys/fs/smbfs_ioctl.h>
83 #include <fs/fs_subr.h>
84
85 /*
86  * We assign directory offsets like the NFS client, where the
87  * offset increments by _one_ after each directory entry.
88  * Further, the entries "." and ".." are always at offsets
89  * zero and one (respectively) and the "real" entries from
90  * the server appear at offsets starting with two. This
91  * macro is used to initialize the n_dirofs field after
92  * setting n_dirseq with a _findopen call.
93  */
94 #define FIRST_DIROFS 2
95
96 /*
97  * These characters are illegal in NTFS file names.
98  * ref: http://support.microsoft.com/kb/147438
99  */
100 * Careful! The check in the XATTR case skips the
101 * first character to allow colon in XATTR names.
102 */
103 static const char illegal_chars[] = {
104     ':', /* colon - keep this first! */
105     '\\', /* back slash */
106     '/', /* slash */
107     '*', /* asterisk */
108     '?', /* question mark */
109     '"', /* double quote */
110     '<', /* less than sign */
111     '>', /* greater than sign */
112     '|', /* vertical bar */
113     0
114 };
115
116 /*
117  * Turning this on causes nodes to be created in the cache
118  * during directory listings, normally avoiding a second
119  * OtW attribute fetch just after a readdir.
120  */
121 int smbfs_fastlookup = 1;
```

```

123 struct vnodeops *smbfs_vnodeops = NULL;
125 /* local static function defines */

127 static int      smbfslookup_cache(vnode_t *, char *, int, vnode_t **,
128                                   cred_t *);
129 static int      smbfslookup(vnode_t *dvp, char *nm, vnode_t **vpp, cred_t *cr,
130                              int cache_ok, caller_context_t *);
131 static int      smbfsremove(vnode_t *dvp, vnode_t *vp, struct smb_cred *scred,
132                              int flags);
133 static int      smbfsrename(vnode_t *odvp, vnode_t *ovp, vnode_t *ndvp,
134                              char *nnm, struct smb_cred *scred, int flags);
135 static int      smbfssetattr(vnode_t *, struct vattr *, int, cred_t *);
136 static int      smbfsaccessx(void *, int, cred_t *);
137 static int      smbfsreadydir(vnode_t *vp, uio_t *uio, cred_t *cr, int *eofp,
138                               caller_context_t *);
139 static void      smbfs_rele_fid(smbnode_t *, struct smb_cred *);
140 static uint32_t  xvattr_to_dosattr(smbnode_t *, struct vattr *);

142 static int      smbfs_rdwrlbn(vnode_t *, page_t *, u_offset_t, size_t, int,
143                               cred_t *);
144 static int      smbfs_bio(struct buf *, int, cred_t *);
145 static int      smbfs_writenp(smbnode_t *np, caddr_t base, int tcount,
146                               struct uio *uiop, int pgcreated);

148 static int      smbfs_fsync(vnode_t *, int, cred_t *, caller_context_t *);
149 static int      smbfs_putpage(vnode_t *, offset_t, size_t, int, cred_t *,
150                               caller_context_t *);
151 static int      smbfs_getpage(vnode_t *, u_offset_t, size_t, uint_t *,
152                               page_t [], size_t, struct seg *, caddr_t,
153                               enum seg_rw, cred_t *);
154 static int      smbfs_putpage(vnode_t *, page_t *, u_offset_t *, size_t *,
155                               int, cred_t *);
156 static void      smbfs_delmap_callback(struct as *, void *, uint_t);

158 /*
159  * Error flags used to pass information about certain special errors
160  * which need to be handled specially.
161  */
162 #define SMBFS_EOF                -98

164 /* When implementing OtW locks, make this a real function. */
165 #define smbfs_lm_has_sleep(vp) 0

167 /*
168  * These are the vnode ops routines which implement the vnode interface to
169  * the networked file system. These routines just take their parameters,
170  * make them look networkish by putting the right info into interface structs,
171  * and then calling the appropriate remote routine(s) to do the work.
172  *
173  * Note on directory name lookup cacheing: If we detect a stale fhandle,
174  * we purge the directory cache relative to that vnode. This way, the
175  * user won't get burned by the cache repeatedly. See <smbfs/smbnode.h> for
176  * more details on smbnode locking.
177  */

133 static int      smbfs_open(vnode_t **, int, cred_t *, caller_context_t *);
134 static int      smbfs_close(vnode_t *, int, int, offset_t, cred_t *,
135                              caller_context_t *);
136 static int      smbfs_read(vnode_t *, struct uio *, int, cred_t *,
137                              caller_context_t *);
138 static int      smbfs_write(vnode_t *, struct uio *, int, cred_t *,
139                              caller_context_t *);
140 static int      smbfs_ioctl(vnode_t *, int, intptr_t, int, cred_t *, int *,
141                              caller_context_t *);

```

```

142 static int      smbfs_getattr(vnode_t *, struct vattr *, int, cred_t *,
143                               caller_context_t *);
144 static int      smbfs_setattr(vnode_t *, struct vattr *, int, cred_t *,
145                               caller_context_t *);
146 static int      smbfs_access(vnode_t *, int, int, cred_t *, caller_context_t *);
147 static int      smbfs_fsync(vnode_t *, int, cred_t *, caller_context_t *);
148 static void      smbfs_inactive(vnode_t *, cred_t *, caller_context_t *);
149 static int      smbfs_lookup(vnode_t *, char *, vnode_t **, struct pathname *,
150                               int, vnode_t *, cred_t *, caller_context_t *,
151                               int *, pathname_t *);
152 static int      smbfs_create(vnode_t *, char *, struct vattr *, enum vcexcl,
153                               int, vnode_t **, cred_t *, int, caller_context_t *,
154                               vsecattr_t *);
155 static int      smbfs_remove(vnode_t *, char *, cred_t *, caller_context_t *,
156                               int);
157 static int      smbfs_rename(vnode_t *, char *, vnode_t *, char *, cred_t *,
158                               caller_context_t *, int);
159 static int      smbfs_mkdir(vnode_t *, char *, struct vattr *, vnode_t **,
160                               cred_t *, caller_context_t *, int, vsecattr_t *);
161 static int      smbfs_rmdir(vnode_t *, char *, vnode_t *, cred_t *,
162                               caller_context_t *, int);
163 static int      smbfs_readdir(vnode_t *, struct uio *, cred_t *, int *,
164                               caller_context_t *, int);
165 static int      smbfs_rwlock(vnode_t *, int, caller_context_t *);
166 static void      smbfs_rwunlock(vnode_t *, int, caller_context_t *);
167 static int      smbfs_seek(vnode_t *, offset_t, offset_t *, caller_context_t *);
168 static int      smbfs_flock(vnode_t *, int, struct flock64 *, int, offset_t,
169                               struct flk_callback *, cred_t *, caller_context_t *);
170 static int      smbfs_space(vnode_t *, int, struct flock64 *, int, offset_t,
171                               cred_t *, caller_context_t *);
172 static int      smbfs_pathconf(vnode_t *, int, ulong_t *, cred_t *,
173                               caller_context_t *);
174 static int      smbfs_setsecattr(vnode_t *, vsecattr_t *, int, cred_t *,
175                               caller_context_t *);
176 static int      smbfs_getsecattr(vnode_t *, vsecattr_t *, int, cred_t *,
177                               caller_context_t *);
178 static int      smbfs_shrlock(vnode_t *, int, struct shrlock *, int, cred_t *,
179                               caller_context_t *);

181 /* Dummy function to use until correct function is ported in */
182 int noop_vnodeop() {
183     return (0);
184 }

186 struct vnodeops *smbfs_vnodeops = NULL;

188 /*
189  * Most unimplemented ops will return ENOSYS because of fs_nosys().
190  * The only ops where that won't work are ACCESS (due to open(2)
191  * failures) and ... (anything else left?)
192  */
193 const fs_operation_def_t smbfs_vnodeops_template[] = {
194     { VOPNAME_OPEN,          { .vop_open = smbfs_open } },
195     { VOPNAME_CLOSE,        { .vop_close = smbfs_close } },
196     { VOPNAME_READ,         { .vop_read = smbfs_read } },
197     { VOPNAME_WRITE,        { .vop_write = smbfs_write } },
198     { VOPNAME_IOCTL,        { .vop_ioctl = smbfs_ioctl } },
199     { VOPNAME_GETATTR,      { .vop_getattr = smbfs_getattr } },
200     { VOPNAME_SETATTR,      { .vop_setattr = smbfs_setattr } },
201     { VOPNAME_ACCESS,       { .vop_access = smbfs_access } },
202     { VOPNAME_LOOKUP,       { .vop_lookup = smbfs_lookup } },
203     { VOPNAME_CREATE,       { .vop_create = smbfs_create } },
204     { VOPNAME_REMOVE,       { .vop_remove = smbfs_remove } },
205     { VOPNAME_LINK,         { .error = fs_nosys } }, /* smbfs_link, */
206     { VOPNAME_RENAME,       { .vop_rename = smbfs_rename } },
207     { VOPNAME_MKDIR,        { .vop_mkdir = smbfs_mkdir } },

```

```

208 { VOPNAME_RMDIR, { .vop_rmdir = smbfs_rmdir } },
209 { VOPNAME_READDIR, { .vop_readdir = smbfs_readdir } },
210 { VOPNAME_SYMLINK, { .error = fs_nosys } }, /* smbfs_symlink, */
211 { VOPNAME_READLINK, { .error = fs_nosys } }, /* smbfs_readlink, */
212 { VOPNAME_FSYNC, { .vop_fsync = smbfs_fsync } },
213 { VOPNAME_INACTIVE, { .vop_inactive = smbfs_inactive } },
214 { VOPNAME_FID, { .error = fs_nosys } }, /* smbfs_fid, */
215 { VOPNAME_RWLOCK, { .vop_rwlock = smbfs_rwlock } },
216 { VOPNAME_RWUNLOCK, { .vop_rwunlock = smbfs_rwunlock } },
217 { VOPNAME_SEEK, { .vop_seek = smbfs_seek } },
218 { VOPNAME_FRLOCK, { .vop_frlock = smbfs_frlock } },
219 { VOPNAME_SPACE, { .vop_space = smbfs_space } },
220 { VOPNAME_REALVP, { .error = fs_nosys } }, /* smbfs_realvp, */
221 { VOPNAME_GETPAGE, { .error = fs_nosys } }, /* smbfs_getpage, */
222 { VOPNAME_PUTPAGE, { .error = fs_nosys } }, /* smbfs_putpage, */
223 { VOPNAME_MAP, { .error = fs_nosys } }, /* smbfs_map, */
224 { VOPNAME_ADDMAP, { .error = fs_nosys } }, /* smbfs_addmap, */
225 { VOPNAME_DELMAP, { .error = fs_nosys } }, /* smbfs_delmmap, */
226 { VOPNAME_DUMP, { .error = fs_nosys } }, /* smbfs_dump, */
227 { VOPNAME_PATHCONF, { .vop_pathconf = smbfs_pathconf } },
228 { VOPNAME_PAGEIO, { .error = fs_nosys } }, /* smbfs_pageio, */
229 { VOPNAME_SETSECATTR, { .vop_setsecattr = smbfs_setsecattr } },
230 { VOPNAME_GETSECATTR, { .vop_getsecattr = smbfs_getsecattr } },
231 { VOPNAME_SHRLOCK, { .vop_shrlock = smbfs_shrlock } },
232 { NULL, NULL }
233 };

235 /*
181 * XXX
182 * When new and relevant functionality is enabled, we should be
183 * calling vfs_set_feature() to inform callers that pieces of
184 * functionality are available, per PSARC 2007/227.
185 */
186 /* ARGSUSED */
187 static int
188 smbfs_open(vnode_t **vpp, int flag, cred_t *cr, caller_context_t *ct)
189 {
190     smbnode_t *np;
191     vnode_t *vp;
192     smbfsattr_t fa;
193     u_int32_t rights, rightsrcvd;
194     u_int16_t fid, oldfid;
195     int oldgenid;
196     struct smb_cred scred;
197     smbmntinfo_t *smi;
198     smb_share_t *ssp;
199     cred_t *oldcr;
200     int tmperror;
201     int error = 0;

203     vp = *vpp;
204     np = VTOSMB(vp);
205     smi = VTOSMI(vp);
206     ssp = smi->smi_share;

208     if (curproc->p_zone != smi->smi_zone_ref.zref_zone)
209         return (EIO);

211     if (smi->smi_flags & SMI_DEAD || vp->v_vfsp->vfs_flag & VFS_UNMOUNTED)
212         return (EIO);

214     if (vp->v_type != VREG && vp->v_type != VDIR) { /* XXX VLNK? */
215         SMBVDEBUG("open eaccess vtype=%d\n", vp->v_type);
216         return (EACCES);
217     }

```

```

219     /*
220     * Get exclusive access to n_fid and related stuff.
221     * No returns after this until out.
222     */
223     if (smbfs_rw_enter_sig(&np->r_lkserlock, RW_WRITER, SMBINTR(vp)))
224         return (EINTR);
225     smb_credinit(&scred, cr);

227     /*
228     * Keep track of the vnode type at first open.
229     * It may change later, and we need close to do
230     * cleanup for the type we opened. Also deny
231     * open of new types until old type is closed.
232     * XXX: Per-open instance nodes should help.
233     */
234     if (np->n_ovtype == VNON) {
235         ASSERT(np->n_dirrefs == 0);
236         ASSERT(np->n_fidrefs == 0);
237     } else if (np->n_ovtype != vp->v_type) {
238         SMBVDEBUG("open n_ovtype=%d v_type=%d\n",
239             np->n_ovtype, vp->v_type);
240         error = EACCES;
241         goto out;

243     /*
244     * Directory open. See smbfs_readvdir()
245     */
246     if (vp->v_type == VDIR) {
247         if (np->n_dirseq == NULL) {
248             /* first open */
249             error = smbfs_smb_findopen(np, "", 1,
250                 SMB_FA_SYSTEM | SMB_FA_HIDDEN | SMB_FA_DIR,
251                 &scred, &np->n_dirseq);
252             if (error != 0)
253                 goto out;
254         }
255         np->n_dirofs = FIRST_DIROFS;
256         np->n_dirrefs++;
257         goto have_fid;
258     }

260     /*
261     * If caller specified O_TRUNC/FTRUNC, then be sure to set
262     * FWRITE (to drive successful setattr(size=0) after open)
263     */
264     if (flag & FTRUNC)
265         flag |= FWRITE;

267     /*
268     * If we already have it open, and the FID is still valid,
269     * check whether the rights are sufficient for FID reuse.
270     */
271     if (np->n_fidrefs > 0 &&
272         np->n_vcgenid == ssp->ss_vcgenid) {
273         int upgrade = 0;

275         if ((flag & FWRITE) &&
276             !(np->n_rights & SA_RIGHT_FILE_WRITE_DATA))
277             upgrade = 1;
278         if ((flag & FREAD) &&
279             !(np->n_rights & SA_RIGHT_FILE_READ_DATA))
280             upgrade = 1;
281         if (!upgrade) {
282             /*
283             * the existing open is good enough

```

```

284         */
285         np->n_fidrefs++;
286         goto have_fid;
287     }
288 }
289 rights = np->n_fidrefs ? np->n_rights : 0;

291 /*
292  * we always ask for READ_CONTROL so we can always get the
293  * owner/group IDs to satisfy a stat. Ditto attributes.
294  */
295 rights |= (STD_RIGHT_READ_CONTROL_ACCESS |
296            SA_RIGHT_FILE_READ_ATTRIBUTES);
297 if ((flag & FREAD))
298     rights |= SA_RIGHT_FILE_READ_DATA;
299 if ((flag & FWRITE))
300     rights |= SA_RIGHT_FILE_WRITE_DATA |
301              SA_RIGHT_FILE_APPEND_DATA |
302              SA_RIGHT_FILE_WRITE_ATTRIBUTES;

304 bzero(&fa, sizeof (fa));
305 error = smbfs_smb_open(np,
306                        NULL, 0, 0, /* name nmlen xattr */
307                        rights, &scred,
308                        &fid, &rightsrcvd, &fa);
309 if (error)
310     goto out;
311 smbfs_attrcache_fa(vp, &fa);

313 /*
314  * We have a new FID and access rights.
315  */
316 oldfid = np->n_fid;
317 oldgenid = np->n_vcgenid;
318 np->n_fid = fid;
319 np->n_vcgenid = ssp->ss_vcgenid;
320 np->n_rights = rightsrcvd;
321 np->n_fidrefs++;
322 if (np->n_fidrefs > 1 &&
323     oldgenid == ssp->ss_vcgenid) {
324     /*
325      * We already had it open (presumably because
326      * it was open with insufficient rights.)
327      * Close old wire-open.
328      */
329     tmperror = smbfs_smb_close(ssp,
330                                oldfid, NULL, &scred);
331     if (tmperror)
332         SMBVDEBUG("error %d closing %s\n",
333                   tmperror, np->n_rpath);
334 }

336 /*
337  * This thread did the open.
338  * Save our credentials too.
339  */
340 mutex_enter(&np->r_statelock);
341 oldcr = np->r_cred;
342 np->r_cred = cr;
343 crhold(cr);
344 if (oldcr)
345     crfree(oldcr);
346 mutex_exit(&np->r_statelock);

348 have_fid:
349 /*

```

```

350     * Keep track of the vnode type at first open.
351     * (see comments above)
352     */
353     if (np->n_ovtype == VNON)
354         np->n_ovtype = vp->v_type;

356 out:
357     smb_credrele(&scred);
358     smbfs_rw_exit(&np->r_lkserlock);
359     return (error);
360 }

362 /*ARGSUSED*/
363 static int
364 smbfs_close(vnode_t *vp, int flag, int count, offset_t offset, cred_t *cr,
365             caller_context_t *ct)
366 {
367     smbnode_t *np;
368     smbmntinfo_t *smi;
369     struct smb_cred scred;
370     int error = 0;

372     np = VTOSMB(vp);
373     smi = VTOSMI(vp);

375     /*
376      * Don't "bail out" for VFS_UNMOUNTED here,
377      * as we want to do cleanup, etc.
378      */

380     /*
381      * zone_enter(2) prevents processes from changing zones with SMBFS files
382      * open; if we happen to get here from the wrong zone we can't do
383      * anything over the wire.
384      */
385     if (smi->smi_zone_ref.zref_zone != curproc->p_zone) {
386         /*
387          * We could attempt to clean up locks, except we're sure
388          * that the current process didn't acquire any locks on
389          * the file: any attempt to lock a file belong to another zone
390          * will fail, and one can't lock an SMBFS file and then change
391          * zones, as that fails too.
392          *
393          * Returning an error here is the sane thing to do. A
394          * subsequent call to VN_RELE() which translates to a
395          * smbfs_inactive() will clean up state: if the zone of the
396          * vnode's origin is still alive and kicking, an async worker
397          * thread will handle the request (from the correct zone), and
398          * everything (minus the final smbfs_getattr_otw() call) should
399          * be OK. If the zone is going away smbfs_async_inactive() will
400          * throw away cached pages inline.
401          */
402         return (EIO);
403     }

405     /*
406      * If we are using local locking for this filesystem, then
407      * release all of the SYSV style record locks. Otherwise,
408      * we are doing network locking and we need to release all
409      * of the network locks. All of the locks held by this
410      * process on this file are released no matter what the
411      * incoming reference count is.
412      */
413     if (smi->smi_flags & SMI_LLOCK) {
414         pid_t pid = ddi_get_pid();
415         cleanlocks(vp, pid, 0);

```

```

416         cleanshares(vp, pid);
417     }
418     /*
419     * else doing OtW locking. SMB servers drop all locks
420     * on the file ID we close here, so no _lockrelease()
421     */
422
423     /*
424     * This (passed in) count is the ref. count from the
425     * user's file_t before the closef call (fio.c).
426     * The rest happens only on last close.
427     * We only care when the reference goes away.
428     */
429     if (count > 1)
430         return (0);
431
432     /* NFS has DNLC purge here. */
433
434     /*
435     * If the file was open for write and there are pages,
436     * then make sure dirty pages written back.
437     *
438     * NFS does this async when "close-to-open" is off
439     * (MI_NOCTO flag is set) to avoid blocking the caller.
440     * For now, always do this synchronously (no B_ASYNC).
441     */
442     if ((flag & FWRITE) && vn_has_cached_data(vp)) {
443         error = smbfs_putpage(vp, (offset_t)0, 0, 0, cr, ct);
444         if (error == EAGAIN)
445             error = 0;
446     }
447     if (error == 0) {
448         mutex_enter(&np->r_statelock);
449         np->r_flags &= ~RSTALE;
450         np->r_error = 0;
451         mutex_exit(&np->r_statelock);
452     }
453
454     /*
455     * Decrement the reference count for the FID
456     * and possibly do the OtW close.
457     * Exclusive lock for modifying n_fid stuff.
458     * Don't want this one ever interruptible.
459     */
460     (void) smbfs_rw_enter_sig(&np->r_lkserlock, RW_WRITER, 0);
461     smb_credinit(&scred, cr);
462
463     smbfs_rele_fid(np, &scred);
464
465     smb_credrele(&scred);
466     smbfs_rw_exit(&np->r_lkserlock);
467
468     return (0);
469 }
470
471 unchanged_portion_omitted
472
473 /* ARGSUSED */
474 static int
475 smbfs_read(vnode_t *vp, struct uio *uiop, int ioflag, cred_t *cr,
476 caller_context_t *ct)
477 {
478     struct smb_cred scred;
479     struct vattr va;
480     smbnode_t *np;
481     smbmntinfo_t *smi;

```

```

559     smb_share_t *ssp;
560     offset_t endoff;
561     ssize_t past_eof;
562     int error;
563
564     caddr_t base;
565     u_offset_t off;
566     size_t n;
567     int on;
568     uint_t flags;
569
570     np = VTOSMB(vp);
571     smi = VTOSMI(vp);
572     ssp = smi->smi_share;
573
574     if (curproc->p_zone != smi->smi_zone_ref.zref_zone)
575         return (EIO);
576
577     if (smi->smi_flags & SMI_DEAD || vp->v_vfsp->vfs_flag & VFS_UNMOUNTED)
578         return (EIO);
579
580     ASSERT(smbfs_rw_lock_held(&np->r_rwlock, RW_READER));
581
582     if (vp->v_type != VREG)
583         return (EISDIR);
584
585     if (uiop->uio_resid == 0)
586         return (0);
587
588     /*
589     * Like NFS3, just check for 63-bit overflow.
590     * Our SMB layer takes care to return EFBIG
591     * when it has to fallback to a 32-bit call.
592     */
593     endoff = uiop->uio_loffset + uiop->uio_resid;
594     if (uiop->uio_loffset < 0 || endoff < 0)
595         return (EINVAL);
596
597     /* get vnode attributes from server */
598     va.va_mask = AT_SIZE | AT_MTIME;
599     if (error = smbfsgetattr(vp, &va, cr))
600         return (error);
601
602     /* Update mtime with mtime from server here? */
603
604     /* if offset is beyond EOF, read nothing */
605     if (uiop->uio_loffset >= va.va_size)
606         return (0);
607
608     /*
609     * Limit the read to the remaining file size.
610     * Do this by temporarily reducing uio_resid
611     * by the amount the lies beyond the EOF.
612     */
613     if (endoff > va.va_size) {
614         past_eof = (ssize_t)(endoff - va.va_size);
615         uiop->uio_resid -= past_eof;
616     } else
617         past_eof = 0;
618
619     /*
620     * Bypass VM if caching has been disabled (e.g., locking) or if
621     * using client-side direct I/O and the file is not mmap'd and
622     * there are no cached pages.
623     */
624     if ((vp->v_flag & VNOCACHE) ||

```

```

625 (((np->r_flags & RDIRECTIO) || (smi->smi_flags & SMI_DIRECTIO)) &&
626 np->r_mapcnt == 0 && np->r_inmap == 0 &&
627 !vn_has_cached_data(vp))) {
628
629     /* Shared lock for n_fid use in smb_rwuio */
630     if (smbfs_rw_enter_sig(&np->r_lkserlock, RW_READER, SMBINTR(vp))
631         return (EINTR);
632     smb_credinit(&scred, cr);
633
634     /* After reconnect, n_fid is invalid */
635     if (np->n_vcgenid != ssp->ss_vcgenid)
636         error = ESTALE;
637     else
638         error = smb_rwuio(ssp, np->n_fid, UIO_READ,
639             uiop, &scred, smb_timo_read);
640
641     smb_credrele(&scred);
642     smbfs_rw_exit(&np->r_lkserlock);
643
644     /* undo adjustment of resid */
645     uiop->uio_resid += past_eof;
646
647     return (error);
648 }
649
650 /* (else) Do I/O through segmap. */
651 do {
652     off = uiop->uio_loffset & MAXBMASK; /* mapping offset */
653     on = uiop->uio_loffset & MAXBOFFSET; /* Relative offset */
654     n = MIN(MAXBSIZE - on, uiop->uio_resid);
655
656     error = smbfs_validate_caches(vp, cr);
657     if (error)
658         break;
659
660     /* NFS waits for RINCCACHEPURGE here. */
661
662     if (vpm_enable) {
663         /*
664          * Copy data.
665          */
666         error = vpm_data_copy(vp, off + on, n, uiop,
667             1, NULL, 0, S_READ);
668     } else {
669         base = segmap_getmapflt(segkmap, vp, off + on, n, 1,
670             S_READ);
671
672         error = uiomove(base + on, n, UIO_READ, uiop);
673     }
674
675     if (!error) {
676         /*
677          * If read a whole block or read to eof,
678          * won't need this buffer again soon.
679          */
680         mutex_enter(&np->r_statelock);
681         if (n + on == MAXBSIZE ||
682             uiop->uio_loffset == np->r_size)
683             flags = SM_DONTNEED;
684         else
685             flags = 0;
686         mutex_exit(&np->r_statelock);
687         if (vpm_enable) {
688             error = vpm_sync_pages(vp, off, n, flags);
689         } else {
690             error = segmap_release(segkmap, base, flags);

```

```

691     }
692     } else {
693         if (vpm_enable) {
694             (void) vpm_sync_pages(vp, off, n, 0);
695         } else {
696             (void) segmap_release(segkmap, base, 0);
697         }
698     }
699     } while (!error && uiop->uio_resid > 0);
700
701     /* undo adjustment of resid */
702     uiop->uio_resid += past_eof;
703
704     return (error);
705 }
706
707 /* ARGSUSED */
708 static int
709 smbfs_write(vnode_t *vp, struct uio *uiop, int ioflag, cred_t *cr,
710     caller_context_t *ct)
711 {
712     struct smb_cred scred;
713     struct vattr va;
714     smbnode_t *np;
715     smbmntinfo_t *smi;
716     smb_share_t *ssp;
717     off_t endoff, limit;
718     ssize_t past_limit;
719     int error, timo;
720     caddr_t base;
721     u_offset_t off;
722     size_t n;
723     int on;
724     uint_t flags;
725     u_offset_t last_off;
726     size_t last_resid;
727     uint_t bsize;
728
729     np = VTOSMB(vp);
730     smi = VTOSMI(vp);
731     ssp = smi->smi_share;
732
733     if (curproc->p_zone != smi->smi_zone_ref.zref_zone)
734         return (EIO);
735
736     if (smi->smi_flags & SMI_DEAD || vp->v_vfsp->vfs_flag & VFS_UNMOUNTED)
737         return (EIO);
738
739     ASSERT(smbfs_rw_lock_held(&np->r_rwlock, RW_WRITER));
740
741     if (vp->v_type != VREG)
742         return (EISDIR);
743
744     if (uiop->uio_resid == 0)
745         return (0);
746
747     /*
748      * Handle ioflag bits: (FAPPEND|FSYNC|FDSYNC)
749      */
750     if (ioflag & (FAPPEND | FSYNC)) {
751         if (np->n_flag & NMODIFIED) {
752             smbfs_attrcache_remove(np);
753             /* XXX: smbfs_vinvalbuf? */
754         }
755     }

```

```

756     if (ioflag & FAPPEND) {
757         /*
758          * File size can be changed by another client
759          *
760          * Todo: Consider redesigning this to use a
761          * handle opened for append instead.
762          */
763         va.va_mask = AT_SIZE;
764         if (error = smbfsgetattr(vp, &va, cr))
765             return (error);
766         uiop->uio_loffset = va.va_size;
767     }

769     /*
770      * Like NFS3, just check for 63-bit overflow.
771      */
772     endoff = uiop->uio_loffset + uiop->uio_resid;
773     if (uiop->uio_loffset < 0 || endoff < 0)
774         return (EINVAL);

776     /*
777      * Check to make sure that the process will not exceed
778      * its limit on file size. It is okay to write up to
779      * the limit, but not beyond. Thus, the write which
780      * reaches the limit will be short and the next write
781      * will return an error.
782      *
783      * So if we're starting at or beyond the limit, EFBIG.
784      * Otherwise, temporarily reduce resid to the amount
785      * that is after the limit.
786      * the falls after the limit.
787      */
788     limit = uiop->uio_llimit;
789     if (limit == RLIM64_INFINITY || limit > MAXOFFSET_T)
790         limit = MAXOFFSET_T;
791     if (uiop->uio_loffset >= limit) {
792         proc_t *p = ttoproc(curthread);

793         mutex_enter(&p->p_lock);
794         (void) rctl_action(rctlproc_legacy[RLIMIT_FSIZE],
795             p->p_rctls, p, RCA_UNSAFE_SIGINFO);
796         mutex_exit(&p->p_lock);
797         if (uiop->uio_loffset >= limit)
798             return (EFBIG);
799     }
800     if (endoff > limit) {
801         past_limit = (ssize_t)(endoff - limit);
802         uiop->uio_resid -= past_limit;
803     } else
804         past_limit = 0;

805     /*
806      * Bypass VM if caching has been disabled (e.g., locking) or if
807      * using client-side direct I/O and the file is not mmap'd and
808      * there are no cached pages.
809      */
810     if ((vp->v_flag & VNOCACHE) ||
811         (((np->r_flags & RDIRECTIO) || (smi->smi_flags & SMI_DIRECTIO)) &&
812         np->r_mapcnt == 0 && np->r_inmap == 0 &&
813         !vn_has_cached_data(vp))) {

815     smbfs_fwrite:
816         if (np->r_flags & RSTALE) {
817             last_resid = uiop->uio_resid;
818             last_off = uiop->uio_loffset;
819             error = np->r_error;

```

```

820         /*
821          * A close may have cleared r_error, if so,
822          * propagate ESTALE error return properly
823          */
824         if (error == 0)
825             error = ESTALE;
826         goto bottom;
827     }

829     /* Timeout: longer for append. */
830     tmo = smb_tmo_write;
831     if (endoff > np->r_size)
832         tmo = smb_tmo_append;

834     /* Shared lock for n_fid use in smb_rwuio */
835     if (smbfs_rw_enter_sig(&np->r_lkserlock, RW_READER, SMBINTR(vp))
836         return (EINTR);
837     smb_credinit(&scrd, cr);

839     /* After reconnect, n_fid is invalid */
840     if (np->n_vcgenid != ssp->ss_vcgenid)
841         error = ESTALE;
842     else
843         error = smb_rwuio(ssp, np->n_fid, UIO_WRITE,
844             uiop, &scrd, tmo);

846     if (error == 0) {
847         mutex_enter(&np->r_statelock);
848         np->n_flag |= (NFLUSHWIRE | NATTRCHANGED);
849         if (uiop->uio_loffset > (offset_t)np->r_size)
850             np->r_size = (len_t)uiop->uio_loffset;
851         mutex_exit(&np->r_statelock);
852         if (ioflag & (FSYNC | FDSYNC)) {
853             /* Don't error the I/O if this fails. */
854             (void) smbfs_smb_flush(np, &scrd);
855         }
856     }

858     smb_credrele(&scrd);
859     smbfs_rw_exit(&np->r_lkserlock);

861     /* undo adjustment of resid */
862     uiop->uio_resid += past_limit;

864     return (error);
865 }

867 /* (else) Do I/O through segmap. */
868 bsize = vp->v_vfsp->vfs_bsize;

870 do {
871     off = uiop->uio_loffset & MAXBMASK; /* mapping offset */
872     on = uiop->uio_loffset & MAXBOFFSET; /* Relative offset */
873     n = MIN(MAXBSIZE - on, uiop->uio_resid);

875     last_resid = uiop->uio_resid;
876     last_off = uiop->uio_loffset;

878     if (np->r_flags & RSTALE) {
879         error = np->r_error;
880         /*
881          * A close may have cleared r_error, if so,
882          * propagate ESTALE error return properly
883          */
884         if (error == 0)

```

```

885         error = ESTALE;
886         break;
887     }
888
889     /*
890     * From NFS: Don't create dirty pages faster than they
891     * can be cleaned.
892     *
893     * Here NFS also checks for async writes (np->r_awaitcount)
894     */
895     mutex_enter(&np->r_statelock);
896     while (np->r_gcount > 0) {
897         if (SMBINTR(vp)) {
898             klp_t *lwp = ttolwp(curthread);
899
900             if (lwp != NULL)
901                 lwp->lwp_nostop++;
902             if (!cv_wait_sig(&np->r_cv, &np->r_statelock)) {
903                 mutex_exit(&np->r_statelock);
904                 if (lwp != NULL)
905                     lwp->lwp_nostop--;
906                 error = EINTR;
907                 goto bottom;
908             }
909             if (lwp != NULL)
910                 lwp->lwp_nostop--;
911         } else
912             cv_wait(&np->r_cv, &np->r_statelock);
913     }
914     mutex_exit(&np->r_statelock);
915
916     /*
917     * Touch the page and fault it in if it is not in core
918     * before segmap_getmapflt or vpm_data_copy can lock it.
919     * This is to avoid the deadlock if the buffer is mapped
920     * to the same file through mmap which we want to write.
921     */
922     uio_prefaultpages((long)n, uiop);
923
924     if (vpm_enable) {
925         /*
926         * It will use kpm mappings, so no need to
927         * pass an address.
928         */
929         error = smbfs_writenp(np, NULL, n, uiop, 0);
930     } else {
931         if (segmap_kpm) {
932             int pon = uiop->uio_loffset & PAGEOFFSET;
933             size_t pn = MIN(PAGESIZE - pon,
934                             uiop->uio_resid);
935             int pagecreate;
936
937             mutex_enter(&np->r_statelock);
938             pagecreate = (pon == 0) && (pn == PAGESIZE ||
939                                     uiop->uio_loffset + pn >= np->r_size);
940             mutex_exit(&np->r_statelock);
941
942             base = segmap_getmapflt(segkmap, vp, off + on,
943                                     pn, !pagecreate, S_WRITE);
944
945             error = smbfs_writenp(np, base + pon, n, uiop,
946                                   pagecreate);
947
948         } else {
949             base = segmap_getmapflt(segkmap, vp, off + on,
950                                     n, 0, S_READ);
951         }
952     }

```

```

951         error = smbfs_writenp(np, base + on, n, uiop, 0)
952     }
953 }
954
955 if (!error) {
956     if (smi->smi_flags & SMI_NOAC)
957         flags = SM_WRITE;
958     else if ((uiop->uio_loffset % bsize) == 0 ||
959             IS_SWAPVP(vp)) {
960         /*
961         * Have written a whole block.
962         * Start an asynchronous write
963         * and mark the buffer to
964         * indicate that it won't be
965         * needed again soon.
966         */
967         flags = SM_WRITE | SM_ASYNC | SM_DONTNEED;
968     } else
969         flags = 0;
970     if ((ioflag & (FSYNC|FDSYNC)) ||
971         (np->r_flags & ROUTOFSPACE)) {
972         flags &= ~SM_ASYNC;
973         flags |= SM_WRITE;
974     }
975     if (vpm_enable) {
976         error = vpm_sync_pages(vp, off, n, flags);
977     } else {
978         error = segmap_release(segkmap, base, flags);
979     }
980 } else {
981     if (vpm_enable) {
982         (void) vpm_sync_pages(vp, off, n, 0);
983     } else {
984         (void) segmap_release(segkmap, base, 0);
985     }
986     /*
987     * In the event that we got an access error while
988     * faulting in a page for a write-only file just
989     * force a write.
990     */
991     if (error == EACCES)
992         goto smbfs_fwrite;
993 }
994 } while (!error && uiop->uio_resid > 0);
995
996 bottom:
997     /* undo adjustment of resid */
998     if (error) {
999         uiop->uio_resid = last_resid + past_limit;
1000         uiop->uio_loffset = last_off;
1001     } else {
1002         uiop->uio_resid += past_limit;
1003     }
1004
1005     return (error);
1006 }
1007
1008 /*
1009 * Like nfs_client.c: writerp()
1010 *
1011 * Write by creating pages and uio move data onto them.
1012 */
1013
1014 int
1015 smbfs_writenp(smbnode_t *np, caddr_t base, int tcount, struct uio *uiop,
1016               int pgcreated)

```

```

1017 {
1018     int            pagecreate;
1019     int            n;
1020     int            saved_n;
1021     caddr_t        saved_base;
1022     u_offset_t     offset;
1023     int            error;
1024     int            sm_error;
1025     vnode_t        *vp = SMBTOV(np);

1027     ASSERT(tcoun <= MAXBSIZE && tcoun <= uio->uio_resid);
1028     ASSERT(smbfs_rw_lock_held(&np->r_rwlock, RW_WRITER));
1029     if (!vpm_enable) {
1030         ASSERT(((uintptr_t)base & MAXBOFFSET) + tcoun <= MAXBSIZE);
1031     }

1033     /*
1034      * Move bytes in at most PAGE_SIZE chunks. We must avoid
1035      * spanning pages in uiomove() because page faults may cause
1036      * the cache to be invalidated out from under us. The r_size is not
1037      * updated until after the uiomove. If we push the last page of a
1038      * file before r_size is correct, we will lose the data written past
1039      * the current (and invalid) r_size.
1040      */
1041     do {
1042         offset = uio->uio_loffset;
1043         pagecreate = 0;

1045         /*
1046          * n is the number of bytes required to satisfy the request
1047          * or the number of bytes to fill out the page.
1048          */
1049         n = (int)MIN((PAGE_SIZE - (offset & PAGEOFFSET)), tcoun);

1051         /*
1052          * Check to see if we can skip reading in the page
1053          * and just allocate the memory. We can do this
1054          * if we are going to rewrite the entire mapping
1055          * or if we are going to write to or beyond the current
1056          * end of file from the beginning of the mapping.
1057          *
1058          * The read of r_size is now protected by r_statelock.
1059          */
1060         mutex_enter(&np->r_statelock);
1061         /*
1062          * When pgcreated is nonzero the caller has already done
1063          * a segmap_getmapflt with forcefault 0 and S_WRITE. With
1064          * segkpm this means we already have at least one page
1065          * created and mapped at base.
1066          */
1067         pagecreate = pgcreated ||
1068             ((offset & PAGEOFFSET) == 0 &&
1069              (n == PAGE_SIZE || ((offset + n) >= np->r_size)));

1071         mutex_exit(&np->r_statelock);
1072         if (!vpm_enable && pagecreate) {
1073             /*
1074              * The last argument tells segmap_pagecreate() to
1075              * always lock the page, as opposed to sometimes
1076              * returning with the page locked. This way we avoid a
1077              * fault on the ensuing uiomove(), but also
1078              * more importantly (to fix bug 1094402) we can
1079              * call segmap_fault() to unlock the page in all
1080              * cases. An alternative would be to modify
1081              * segmap_pagecreate() to tell us when it is
1082              * locking a page, but that's a fairly major

```

```

1083         * interface change.
1084         */
1085         if (pgcreated == 0)
1086             (void) segmap_pagecreate(segkmap, base,
1087                                     (uint_t)n, 1);
1088         saved_base = base;
1089         saved_n = n;
1090     }

1092     /*
1093      * The number of bytes of data in the last page can not
1094      * be accurately determined while page is being
1095      * uiomove'd to and the size of the file being updated.
1096      * Thus, inform threads which need to know accurately
1097      * how much data is in the last page of the file. They
1098      * will not do the i/o immediately, but will arrange for
1099      * the i/o to happen later when this modify operation
1100      * will have finished.
1101      */
1102     ASSERT(!(np->r_flags & RMODINPROGRESS));
1103     mutex_enter(&np->r_statelock);
1104     np->r_flags |= RMODINPROGRESS;
1105     np->r_modaddr = (offset & MAXBMASK);
1106     mutex_exit(&np->r_statelock);

1108     if (vpm_enable) {
1109         /*
1110          * Copy data. If new pages are created, part of
1111          * the page that is not written will be initialized
1112          * with zeros.
1113          */
1114         error = vpm_data_copy(vp, offset, n, uio,
1115                               !pagecreate, NULL, 0, S_WRITE);
1116     } else {
1117         error = uiomove(base, n, UIO_WRITE, uio);
1118     }

1120     /*
1121      * r_size is the maximum number of
1122      * bytes known to be in the file.
1123      * Make sure it is at least as high as the
1124      * first unwritten byte pointed to by uio_loffset.
1125      */
1126     mutex_enter(&np->r_statelock);
1127     if (np->r_size < uio->uio_loffset)
1128         np->r_size = uio->uio_loffset;
1129     np->r_flags &= ~RMODINPROGRESS;
1130     np->r_flags |= RDIRTY;
1131     mutex_exit(&np->r_statelock);

1133     /* n = # of bytes written */
1134     n = (int)(uio->uio_loffset - offset);

1136     if (!vpm_enable) {
1137         base += n;
1138     }
1139     tcoun -= n;
1140     /*
1141      * If we created pages w/o initializing them completely,
1142      * we need to zero the part that wasn't set up.
1143      * This happens on a most EOF write cases and if
1144      * we had some sort of error during the uiomove.
1145      */
1146     if (!vpm_enable && pagecreate) {
1147         if ((uio->uio_loffset & PAGEOFFSET) || n == 0)
1148             (void) kzero(base, PAGE_SIZE - n);

```

```

1150         if (pgcreated) {
1151             /*
1152              * Caller is responsible for this page,
1153              * it was not created in this loop.
1154              */
1155             pgcreated = 0;
1156         } else {
1157             /*
1158              * For bug 1094402: segmap_pagecreate locks
1159              * page. Unlock it. This also unlocks the
1160              * pages allocated by page_create_va() in
1161              * segmap_pagecreate().
1162              */
1163             sm_error = segmap_fault(kas.a_hat, segkmap,
1164                 saved_base, saved_n,
1165                 F_SOFTUNLOCK, S_WRITE);
1166             if (error == 0)
1167                 error = sm_error;
1168         }
1169     } while (tcount > 0 && error == 0);
1170
1172     return (error);
1173 }

```

```

1175 /*
1176  * Flags are composed of {B_ASYNC, B_INVALID, B_FREE, B_DONTNEED}
1177  * Like nfs3_rdwrlbn()
1178  */
1179 static int
1180 smbfs_rdwrlbn(vnode_t *vp, page_t *pp, u_offset_t off, size_t len,
1181     int flags, cred_t *cr)
1182 {
1183     smbmntinfo_t *smi = VTOSMI(vp);
1184     struct buf *bp;
1185     int error;
1186     int sync;

```

```

1188     if (curproc->p_zone != smi->smi_zone_ref.zref_zone)
1189         return (EIO);

```

```

1191     if (smi->smi_flags & SMI_DEAD || vp->v_vfsp->vfs_flag & VFS_UNMOUNTED)
1192         return (EIO);

```

```

1194     bp = pageio_setup(pp, len, vp, flags);
1195     ASSERT(bp != NULL);

```

```

1197     /*
1198      * pageio_setup should have set b_addr to 0. This
1199      * is correct since we want to do I/O on a page
1200      * boundary. bp_mapin will use this addr to calculate
1201      * an offset, and then set b_addr to the kernel virtual
1202      * address it allocated for us.
1203      */
1204     ASSERT(bp->b_un.b_addr == 0);

```

```

1206     bp->b_edev = 0;
1207     bp->b_dev = 0;
1208     bp->b_lblkno = lbtodb(off);
1209     bp->b_file = vp;
1210     bp->b_offset = (offset_t)off;
1211     bp_mapin(bp);

```

```

1213     /*
1214      * Calculate the desired level of stability to write data.

```

```

1215     /*
1216     if ((flags & (B_WRITE|B_ASYNC)) == (B_WRITE|B_ASYNC) &&
1217         freemem > desfree) {
1218         sync = 0;
1219     } else {
1220         sync = 1;
1221     }

```

```

1223     error = smbfs_bio(bp, sync, cr);

```

```

1225     bp_mapout(bp);
1226     pageio_done(bp);

```

```

1228     return (error);
1229 }

```

```

1232 /*
1233  * Corresponds to nfs3_vnops.c : nfs3_bio(), though the NFS code
1234  * uses nfs3read()/nfs3write() where we use smb_rwuio(). Also,
1235  * NFS has this later in the file. Move it up here closer to
1236  * the one call site just above.
1237  */

```

```

1239 static int
1240 smbfs_bio(struct buf *bp, int sync, cred_t *cr)
1241 {
1242     struct iovec aiov[1];
1243     struct uio auio;
1244     struct smb_cred scred;
1245     smbnode_t *np = VTOSMB(bp->b_vp);
1246     smbmntinfo_t *smi = np->n_mount;
1247     smb_share_t *ssp = smi->smi_share;
1248     offset_t offset;
1249     offset_t endoff;
1250     size_t count;
1251     size_t past_eof;
1252     int error;

```

```

1254     ASSERT(curproc->p_zone == smi->smi_zone_ref.zref_zone);

```

```

1256     offset = ldbtob(bp->b_lblkno);
1257     count = bp->b_bcount;
1258     endoff = offset + count;
1259     if (offset < 0 || endoff < 0)
1260         return (EINVAL);

```

```

1262     /*
1263      * Limit file I/O to the remaining file size, but see
1264      * the notes in smbfs_getpage about SMBFS_EOF.
1265      */
1266     mutex_enter(&np->r_statelock);
1267     if (offset >= np->r_size) {
1268         mutex_exit(&np->r_statelock);
1269         if (bp->b_flags & B_READ) {
1270             return (SMBFS_EOF);
1271         } else {
1272             return (EINVAL);
1273         }
1274     }
1275     if (endoff > np->r_size) {
1276         past_eof = (size_t)(endoff - np->r_size);
1277         count -= past_eof;
1278     } else
1279         past_eof = 0;
1280     mutex_exit(&np->r_statelock);

```

```

1281     ASSERT(count > 0);

1283     /* Caller did bmapin(). Mapped address is... */
1284     aiov[0].iov_base = bp->b_un.b_addr;
1285     aiov[0].iov_len = count;
1286     auio.uio_iov = aiov;
1287     auio.uio_iovcnt = 1;
1288     auio.uio_loffset = offset;
1289     auio.uio_segflg = UIO_SYSSPACE;
1290     auio.uio_fmode = 0;
1291     auio.uio_resid = count;

1293     /* Shared lock for n_fid use in smb_rwuio */
1294     if (smbfs_rw_enter_sig(&np->r_lkserlock, RW_READER,
1295         smi->smi_flags & SMI_INT))
1296         return (EINTR);
1297     smb_credinit(&scred, cr);

1299     DTRACE_IO1(start, struct buf *, bp);

1301     if (bp->b_flags & B_READ) {

1303         /* After reconnect, n_fid is invalid */
1304         if (np->n_vcgenid != ssp->ss_vcgenid)
1305             error = ESTALE;
1306         else
1307             error = smb_rwuio(ssp, np->n_fid, UIO_READ,
1308                 &auio, &scred, smb_timo_read);

1310         /* Like NFS, only set b_error here. */
1311         bp->b_error = error;
1312         bp->b_resid = auio.uio_resid;

1314         if (!error && auio.uio_resid != 0)
1315             error = EIO;
1316         if (!error && past_eof != 0) {
1317             /* Zero the memory beyond EOF. */
1318             bzero(bp->b_un.b_addr + count, past_eof);
1319         }
1320     } else {

1322         /* After reconnect, n_fid is invalid */
1323         if (np->n_vcgenid != ssp->ss_vcgenid)
1324             error = ESTALE;
1325         else
1326             error = smb_rwuio(ssp, np->n_fid, UIO_WRITE,
1327                 &auio, &scred, smb_timo_write);

1329         /* Like NFS, only set b_error here. */
1330         bp->b_error = error;
1331         bp->b_resid = auio.uio_resid;

1333         if (!error && auio.uio_resid != 0)
1334             error = EIO;
1335         if (!error && sync) {
1336             (void) smbfs_smb_flush(np, &scred);
1337         }
1338     }

1340     /*
1341     * This comes from nfs3_commit()
1342     */
1343     if (error != 0) {
1344         mutex_enter(&np->r_statelock);
1345         if (error == ESTALE)
1346             np->r_flags |= RSTALE;

```

```

1347         if (!np->r_error)
1348             np->r_error = error;
1349         mutex_exit(&np->r_statelock);
1350         bp->b_flags |= B_ERROR;
1351     }

1353     DTRACE_IO1(done, struct buf *, bp);

1355     smb_credrele(&scred);
1356     smbfs_rw_exit(&np->r_lkserlock);

1358     if (error == ESTALE)
1359         smbfs_attrcache_remove(np);

1361     return (error);
1362 }

1364 /*
1365  * Here NFS has: nfs3write, nfs3read
1366  * We use smb_rwuio instead.
1367  */

1369 /* ARGSUSED */
1370 static int
1371 smbfs_ioctl(vnode_t *vp, int cmd, intptr_t arg, int flag,
1372     cred_t *cr, int *rvalp, caller_context_t *ct)
1373 {
1374     int error;
1375     smbmntinfo_t *smi;

1377     smi = VTOSMI(vp);

1379     if (curproc->p_zone != smi->smi_zone_ref.zref_zone)
1380         return (EIO);

1382     if (smi->smi_flags & SMI_DEAD || vp->v_vfsp->vfs_flag & VFS_UNMOUNTED)
1383         return (EIO);

1385     switch (cmd) {
1386     /* First three from ZFS. XXX - need these? */

1387     case _FIOFFS:
1388         error = smbfs_fsync(vp, 0, cr, ct);
1389         break;

1391         /*
1392         * The following two ioctls are used by bfu.
1393         * Silently ignore to avoid bfu errors.
1394         */
1395     case _FIOGDI0:
1396     case _FIOSDI0:
1397         error = 0;
1398         break;

1400 #if 0 /* Todo - SMB ioctl query regions */
1401     case _FIO_SEEK_DATA:
1402     case _FIO_SEEK_HOLE:
1403 #endif

1405 #ifdef NOT_YET /* XXX - from the NFS code. */
1406     case _FIODIRECTIO:
1407         error = smbfs_directio(vp, (int)arg, cr);
1408         break;
1409 #endif

1409         /*

```

```

1410         * Allow get/set with "raw" security descriptor (SD) data.
1411         * Useful for testing, diagnosing idmap problems, etc.
1412         */
1413     case SMBFSIO_GETSD:
1414         error = smbfs_acl_iocget(vp, arg, flag, cr);
1415         break;
1417     case SMBFSIO_SETSD:
1418         error = smbfs_acl_iocset(vp, arg, flag, cr);
1419         break;
1421     default:
1422         error = ENOTTY;
1423         break;
1424 }
1426 return (error);
1427 }

1430 /*
1431  * Return either cached or remote attributes. If get remote attr
1432  * use them to check and invalidate caches, then cache the new attributes.
1433  */
1434 /* ARGSUSED */
1435 static int
1436 smbfs_getattr(vnode_t *vp, struct vattr *vap, int flags, cred_t *cr,
1437             caller_context_t *ct)
1438 {
1439     smbnode_t *np;
1440     smbmntinfo_t *smi;
1441     int error;

1443     smi = VTOSMI(vp);

1445     if (curproc->p_zone != smi->smi_zone_ref.zref_zone)
1446         return (EIO);

1448     if (smi->smi_flags & SMI_DEAD || vp->v_vfsp->vfs_flag & VFS_UNMOUNTED)
1449         return (EIO);

1451     /*
1452      * If it has been specified that the return value will
1453      * just be used as a hint, and we are only being asked
1454      * for size, fsid or rdev, then return the client's
1455      * notion of these values without checking to make sure
1456      * that the attribute cache is up to date.
1457      * The whole point is to avoid an over the wire GETATTR
1458      * call.
1459      */
1460     np = VTOSMB(vp);
1461     if (flags & ATTR_HINT) {
1462         if (vap->va_mask ==
1463             (vap->va_mask & (AT_SIZE | AT_FSID | AT_RDEV))) {
1464             mutex_enter(&np->r_statelock);
1465             if (vap->va_mask | AT_SIZE)
1466                 vap->va_size = np->r_size;
1467             if (vap->va_mask | AT_FSID)
1468                 vap->va_fsid = vp->v_vfsp->vfs_dev;
1469             if (vap->va_mask | AT_RDEV)
1470                 vap->va_rdev = vp->v_rdev;
1471             mutex_exit(&np->r_statelock);
1472             return (0);
1473         }
1474     }

```

```

1476     /*
1477      * Only need to flush pages if asking for the mtime
1478      * and if there any dirty pages.
1479      * Here NFS also checks for async writes (np->r_awcount)
1480      */
1481     if (vap->va_mask & AT_MTIME) {
1482         if (vn_has_cached_data(vp) &&
1483             ((np->r_flags & RDIRTY) != 0)) {
1484             mutex_enter(&np->r_statelock);
1485             np->r_gcount++;
1486             mutex_exit(&np->r_statelock);
1487             error = smbfs_putpage(vp, (offset_t)0, 0, 0, cr, ct);
1488             mutex_enter(&np->r_statelock);
1489             if (error && (error == ENOSPC || error == EDQUOT)) {
1490                 if (!np->r_error)
1491                     np->r_error = error;
1492             }
1493             if (--np->r_gcount == 0)
1494                 cv_broadcast(&np->r_cv);
1495             mutex_exit(&np->r_statelock);
1496         }
1497     }
1498 }

1500 return (smbfsgetattr(vp, vap, cr));
1501 }

1503 /* smbfsgetattr() in smbfs_client.c */

1505 /* ARGSUSED4 */
1506 static int
1507 smbfs_setattr(vnode_t *vp, struct vattr *vap, int flags, cred_t *cr,
1508             caller_context_t *ct)
1509 {
1510     vfs_t *vfsp;
1511     smbmntinfo_t *smi;
1512     int error;
1513     uint_t mask;
1514     struct vattr oldva;

1516     vfsp = vp->v_vfsp;
1517     smi = VFTOSMI(vfsp);

1519     if (curproc->p_zone != smi->smi_zone_ref.zref_zone)
1520         return (EIO);

1522     if (smi->smi_flags & SMI_DEAD || vfsp->vfs_flag & VFS_UNMOUNTED)
1523         return (EIO);

1525     mask = vap->va_mask;
1526     if (mask & AT_NOSET)
1527         return (EINVAL);

1529     if (vfsp->vfs_flag & VFS_RDONLY)
1530         return (EROFS);

1532     /*
1533      * This is a _local_ access check so that only the owner of
1534      * this mount can set attributes. With ACLs enabled, the
1535      * file owner can be different from the mount owner, and we
1536      * need to check the _mount_ owner here. See _access_rwx
1537      */
1538     bzero(&oldva, sizeof (oldva));
1539     oldva.va_mask = AT_TYPE | AT_MODE;
1540     error = smbfsgetattr(vp, &oldva, cr);
1541     if (error)

```

```

1542         return (error);
1543     oldva.va_mask |= AT_UID | AT_GID;
1544     oldva.va_uid = smi->smi_uid;
1545     oldva.va_gid = smi->smi_gid;

1547     error = secpolicy_vnode_setattr(cr, vp, vap, &oldva, flags,
1548         smbfs_accessx, vp);
1549     if (error)
1550         return (error);

1552     if (mask & (AT_UID | AT_GID)) {
1553         if (smi->smi_flags & SMI_ACL)
1554             error = smbfs_acl_setids(vp, vap, cr);
1555         else
1556             error = ENOSYS;
1557         if (error != 0) {
1558             SMBVDEBUG("error %d setting UID/GID on %s",
1559                 error, VTOSMB(vp)->n_rpath);
1560             /*
1561              * It might be more correct to return the
1562              * error here, but that causes complaints
1563              * when root extracts a cpio archive, etc.
1564              * So ignore this error, and go ahead with
1565              * the rest of the setattr work.
1566              */
1567         }
1568     }

1570     error = smbfssetattr(vp, vap, flags, cr);

1572 #ifdef SMBFS_VNEVENT
1573     if (error == 0 && (vap->va_mask & AT_SIZE) && vap->va_size == 0)
1574         vnevent_truncate(vp, ct);
1575 #endif

1577     return (error);
1578     return (smbfssetattr(vp, vap, flags, cr));
1579 }

1580 /*
1581  * Mostly from Darwin smbfs_setattr()
1582  * but then modified a lot.
1583  */
1584 /* ARGSUSED */
1585 static int
1586 smbfssetattr(vnode_t *vp, struct vattr *vap, int flags, cred_t *cr)
1587 {
1588     int error = 0;
1589     smbnode_t *np = VTOSMB(vp);
1590     uint_t mask = vap->va_mask;
1591     struct timespec *mtime, *atime;
1592     struct smb_cred scred;
1593     int cerror, modified = 0;
1594     unsigned short fid;
1595     int have_fid = 0;
1596     uint32_t rights = 0;
1597     uint32_t dosattr = 0;

1599     ASSERT(curproc->p_zone == VTOSMI(vp)->smi_zone_ref.zref_zone);

1601     /*
1602      * There are no settable attributes on the XATTR dir,
1603      * so just silently ignore these. On XATTR files,
1604      * you can set the size but nothing else.
1605      */
1606     if (vp->v_flag & V_XATTRDIR)

```

```

1607         return (0);
1608     if (np->n_flag & N_XATTR) {
1609         if (mask & AT_TIMES)
1610             SMBVDEBUG("ignore set time on xattr\n");
1611         mask &= AT_SIZE;
1612     }

1614     /*
1615      * Only need to flush pages if there are any pages and
1616      * if the file is marked as dirty in some fashion. The
1617      * file must be flushed so that we can accurately
1618      * determine the size of the file and the cached data
1619      * after the SETATTR returns. A file is considered to
1620      * be dirty if it is either marked with RDIRTY, has
1621      * outstanding i/o's active, or is mmap'd. In this
1622      * last case, we can't tell whether there are dirty
1623      * pages, so we flush just to be sure.
1624      */
1625     if (vn_has_cached_data(vp) &&
1626         ((np->r_flags & RDIRTY) ||
1627         np->r_count > 0 ||
1628         np->r_mapcnt > 0)) {
1629         ASSERT(vp->v_type != VCHR);
1630         error = smbfs_putpage(vp, (offset_t)0, 0, 0, cr, NULL);
1631         if (error && (error == ENOSPC || error == EDQUOT)) {
1632             mutex_enter(&np->r_statelock);
1633             if (!np->r_error)
1634                 np->r_error = error;
1635             mutex_exit(&np->r_statelock);
1636         }
1637     }

1639     /*
1640      * If our caller is trying to set multiple attributes, they
1641      * can make no assumption about what order they are done in.
1642      * Here we try to do them in order of decreasing likelihood
1643      * of failure, just to minimize the chance we'll wind up
1644      * with a partially complete request.
1645      */

1647     /* Shared lock for (possible) n_fid use. */
1648     if (smbfs_rw_enter_sig(&np->r_lkserlock, RW_READER, SMBINTR(vp)))
1649         return (EINTR);
1650     smb_credinit(&scred, cr);

1652     /*
1653      * If the caller has provided extensible attributes,
1654      * map those into DOS attributes supported by SMB.
1655      * Note: zero means "no change".
1656      */
1657     if (mask & AT_XVATTR)
1658         dosattr = xvattr_to_dosattr(np, vap);

1660     /*
1661      * Will we need an open handle for this setattr?
1662      * If so, what rights will we need?
1663      */
1664     if (dosattr || (mask & (AT_ATIME | AT_MTIME))) {
1665         rights |=
1666             SA_RIGHT_FILE_WRITE_ATTRIBUTES;
1667     }
1668     if (mask & AT_SIZE) {
1669         rights |=
1670             SA_RIGHT_FILE_WRITE_DATA |
1671             SA_RIGHT_FILE_APPEND_DATA;
1672     }

```

```

1674 /*
1675  * Only SIZE really requires a handle, but it's
1676  * simpler and more reliable to set via a handle.
1677  * Some servers like NT4 won't set times by path.
1678  * Also, we're usually setting everything anyway.
1679  */
1680 if (rights != 0) {
1681     error = smbfs_smb_tmpopen(np, rights, &scred, &fid);
1682     if (error) {
1683         SMBVDEBUG("error %d opening %s\n",
1684             error, np->n_rpath);
1685         goto out;
1686     }
1687     have_fid = 1;
1688 }
1689
1690 /*
1691  * If the server supports the UNIX extensions, right here is where
1692  * we'd support changes to uid, gid, mode, and possibly va_flags.
1693  * For now we claim to have made any such changes.
1694  */
1695
1696 if (mask & AT_SIZE) {
1697     /*
1698      * If the new file size is less than what the client sees as
1699      * the file size, then just change the size and invalidate
1700      * the pages.
1701      * I am commenting this code at present because the function
1702      * smbfs_putapage() is not yet implemented.
1703      */
1704     /*
1705      * Set the file size to vap->va_size.
1706      */
1707     ASSERT(have_fid);
1708     error = smbfs_smb_setfsize(np, fid, vap->va_size, &scred);
1709     if (error) {
1710         SMBVDEBUG("setsize error %d file %s\n",
1711             error, np->n_rpath);
1712     } else {
1713         /*
1714          * Darwin had code here to zero-extend.
1715          * Tests indicate the server will zero-fill,
1716          * so looks like we don't need to do that.
1717          * so looks like we don't need to do this.
1718          * Good thing, as this could take forever.
1719          * XXX: Reportedly, writing one byte of zero
1720          * at the end offset avoids problems here.
1721          */
1722         mutex_enter(&np->r_statelock);
1723         np->r_size = vap->va_size;
1724         mutex_exit(&np->r_statelock);
1725         modified = 1;
1726     }
1727 }
1728
1729 /*
1730  * Todo: Implement setting create_time (which is
1731  * different from ctime).
1732  * XXX: When Solaris has create_time, set that too.
1733  * Note: create_time is different from ctime.
1734  */
1735 mtime = ((mask & AT_MTIME) ? &vap->va_mtime : 0);
1736 atime = ((mask & AT_ETIME) ? &vap->va_atime : 0);

```

```

1731 if (dosattr || mtime || atime) {
1732     /*
1733      * Always use the handle-based set attr call now.
1734      */
1735     ASSERT(have_fid);
1736     error = smbfs_smb_setfattr(np, fid,
1737         dosattr, mtime, atime, &scred);
1738     if (error) {
1739         SMBVDEBUG("set times error %d file %s\n",
1740             error, np->n_rpath);
1741     } else {
1742         modified = 1;
1743     }
1744 }
1745
1746 out:
1747 if (modified) {
1748     /*
1749      * Invalidate attribute cache in case the server
1750      * doesn't set exactly the attributes we asked.
1751      */
1752     smbfs_attrcache_remove(np);
1753 }
1754
1755 if (have_fid) {
1756     error = smbfs_smb_tmpclose(np, fid, &scred);
1757     if (error)
1758         SMBVDEBUG("error %d closing %s\n",
1759             error, np->n_rpath);
1760 }
1761
1762 smb_credrele(&scred);
1763 smbfs_rw_exit(&np->r_lkserlock);
1764
1765 if (modified) {
1766     /*
1767      * Invalidate attribute cache in case the server
1768      * doesn't set exactly the attributes we asked.
1769      */
1770     smbfs_attrcache_remove(np);
1771 }
1772
1773 /*
1774  * If changing the size of the file, invalidate
1775  * any local cached data which is no longer part
1776  * of the file. We also possibly invalidate the
1777  * last page in the file. We could use
1778  * pvn_vpzero(), but this would mark the page as
1779  * modified and require it to be written back to
1780  * the server for no particularly good reason.
1781  * This way, if we access it, then we bring it
1782  * back in. A read should be cheaper than a
1783  * write.
1784  */
1785 if (mask & AT_SIZE) {
1786     smbfs_invalidate_pages(vp,
1787         (vap->va_size & PAGEMASK), cr);
1788 }
1789
1790 return (error);
1791 }
1792
1793 unchanged_portion_omitted
1794
1795 /*
1796  * smbfs_access_rwx()

```

```

1845 * Common function for smbfs_access, etc.
1846 *
1847 * The security model implemented by the FS is unusual
1848 * due to the current "single user mounts" restriction:
1849 * All access under a given mount point uses the CIFS
1850 * credentials established by the owner of the mount.
1851 *
1852 * Most access checking is handled by the CIFS server,
1853 * but we need sufficient Unix access checks here to
1854 * prevent other local Unix users from having access
1855 * to objects under this mount that the uid/gid/mode
1856 * settings in the mount would not allow.
1857 *
1858 * With this model, there is a case where we need the
1859 * ability to do an access check before we have the
1860 * vnode for an object. This function takes advantage
1861 * of the fact that the uid/gid/mode is per mount, and
1862 * avoids the need for a vnode.
1863 *
1864 * We still (sort of) need a vnode when we call
1865 * secpolicy_vnode_access, but that only uses
1866 * the vtype field, so we can use a pair of fake
1867 * vnodes that have only v_type filled in.
1211 *
1212 * XXX: Later, add a new secpolicy_vtype_access()
1213 * that takes the vtype instead of a vnode, and
1214 * get rid of the tmpl_vxxx fake vnodes below.
1868 */
1869 static int
1870 smbfs_access_rwx(vfs_t *vfsp, int vtype, int mode, cred_t *cr)
1871 {
1872     /* See the secpolicy call below. */
1873     static const vnode_t tmpl_vdir = { .v_type = VDIR };
1874     static const vnode_t tmpl_vreg = { .v_type = VREG };
1875     va;
1876     vattr_t      *vattr;
1877     struct smbmntinfo *smi = VFTOSMI(vfsp);
1878     int shift = 0;
1880     /*
1881     * Build our (fabricated) vnode attributes.
1882     * XXX: Could make these templates in the
1883     * per-mount struct and use them here.
1884     */
1885     bzero(&va, sizeof(va));
1886     va.va_mask = AT_TYPE | AT_MODE | AT_UID | AT_GID;
1887     va.va_type = vtype;
1888     va.va_mode = (vtype == VDIR) ?
1889         smi->smi_dmode : smi->smi_fmode;
1890     va.va_uid = smi->smi_uid;
1891     va.va_gid = smi->smi_gid;
1892     /*
1893     * Disallow write attempts on read-only file systems,
1894     * unless the file is a device or fifo node. Note:
1895     * Inline vn_is_readonly and IS_DEVVP here because
1896     * we may not have a vnode ptr. Original expr. was:
1897     * (mode & VWRITE) && vn_is_readonly(vp) && !IS_DEVVP(vp))
1898     */
1899     if ((mode & VWRITE) &&
1900         (vfsp->vfs_flag & VFS_RDONLY) &&
1901         !(vtype == VCHR || vtype == VBLK || vtype == VFIFO))
1902         return (EROFS);
1903     /*
1904     * Disallow attempts to access mandatory lock files.

```

```

1905     * Similarly, expand MANDLOCK here.
1255     * XXX: not sure we need this.
1906     */
1907     if ((mode & (VWRITE | VREAD | VEXEC)) &&
1908         va.va_type == VREG && MANDMODE(va.va_mode))
1909         return (EACCES);
1911     /*
1912     * Access check is based on only
1913     * one of owner, group, public.
1914     * If not owner, then check group.
1915     * If not a member of the group,
1916     * then check public access.
1917     */
1918     if (crgetuid(cr) != va.va_uid) {
1919         shift += 3;
1920         if (!groupmember(va.va_gid, cr))
1921             shift += 3;
1922     }
1924     /*
1925     * We need a vnode for secpolicy_vnode_access,
1926     * but the only thing it looks at is v_type,
1927     * so pass one of the templates above.
1928     */
1929     tvp = (va.va_type == VDIR) ?
1930         (vnode_t *)&tmpl_vdir :
1931         (vnode_t *)&tmpl_vreg;
1933     return (secpolicy_vnode_access2(cr, tvp, va.va_uid,
1934         va.va_mode << shift, mode));
1935 }
1936
1937 _____unchanged_portion_omitted_____
1938
1975 /* ARGSUSED */
1976 static int
1977 smbfs_readlink(vnode_t *vp, struct uio *uiop, cred_t *cr, caller_context_t *ct)
1978 {
1979     /* Not yet... */
1980     return (ENOSYS);
1981 }
1982
1984 /*
1985 * Flush local dirty pages to stable storage on the server.
1986 *
1987 * If FNODESYNC is specified, then there is nothing to do because
1988 * metadata changes are not cached on the client before being
1989 * sent to the server.
1990 */
1991 /* ARGSUSED */
1992 static int
1993 smbfs_fsync(vnode_t *vp, int syncflag, cred_t *cr, caller_context_t *ct)
1994 {
1995     int error = 0;
1996     smbmntinfo_t *smi;
1997     smbnode_t *np;
1998     struct smb_cred scred;
1999
2000     np = VTOSMB(vp);
2001     smi = VTOSMI(vp);
2002
2003     if (curproc->p_zone != smi->smi_zone_ref.zref_zone)
2004         return (EIO);

```

```

2006     if (smi->smi_flags & SMI_DEAD || vp->v_vfsp->vfs_flag & VFS_UNMOUNTED)
2007         return (EIO);

2009     if ((syncflag & FNODSYNC) || IS_SWAPVP(vp))
2010         return (0);

2012     if ((syncflag & (FSYNC|FDSYNC)) == 0)
2013         return (0);

2015     error = smbfs_putpage(vp, (offset_t)0, 0, 0, cr, ct);
2016     if (error)
2017         return (error);

2019     /* Shared lock for n_fid use in _flush */
2020     if (smbfs_rw_enter_sig(&np->r_lkserlock, RW_READER, SMBINTR(vp)))
2021         return (EINTR);
2022     smb_credinit(&scred, cr);

2024     error = smbfs_smb_flush(np, &scred);

2026     smb_credrele(&scred);
2027     smbfs_rw_exit(&np->r_lkserlock);

2029     return (error);
2030 }

2032 /*
2033  * Last reference to vnode went away.
2034  */
2035 /* ARGSUSED */
2036 static void
2037 smbfs_inactive(vnode_t *vp, cred_t *cr, caller_context_t *ct)
2038 {
2039     struct smbnode *np;
2040     struct smb_cred scred;
2041     smbnode_t *np = VTOSMB(vp);
2042     int error;

2043     /*
2044      * Don't "bail out" for VFS_UNMOUNTED here,
2045      * as we want to do cleanup, etc. This call may
2046      * See also pcfs_inactive
2047      */

1385     np = VTOSMB(vp);

2049     /*
2050      * If this is coming from the wrong zone, we let someone in the right
2051      * zone take care of it asynchronously. We can get here due to
2052      * VN_RELE() being called from pageout() or fsflush(). This call may
2053      * potentially turn into an expensive no-op if, for instance, v_count
2054      * gets incremented in the meantime, but it's still correct.
2055      */

2057     /*
2058      * From NFS:r_inactive()
2059      *
2060      * Before freeing anything, wait until all asynchronous
2061      * activity is done on this rnode. This will allow all
2062      * asynchronous read ahead and write behind i/o's to
2063      * finish.
2064      */
2065     mutex_enter(&np->r_statelock);
2066     while (np->r_count > 0)
2067         cv_wait(&np->r_cv, &np->r_statelock);
2068     mutex_exit(&np->r_statelock);

```

```

2070     /*
2071      * Flush and invalidate all pages associated with the vnode.
2072      */
2073     if (vn_has_cached_data(vp)) {
2074         if ((np->r_flags & RDIRTY) && !np->r_error) {
2075             error = smbfs_putpage(vp, (u_offset_t)0, 0, 0, cr, ct);
2076             if (error && (error == ENOSPC || error == EDQUOT)) {
2077                 mutex_enter(&np->r_statelock);
2078                 if (!np->r_error)
2079                     np->r_error = error;
2080                 mutex_exit(&np->r_statelock);
2081             }
2082         }
2083         smbfs_invalidate_pages(vp, (u_offset_t)0, cr);
2084     }
2085     /*
2086      * This vnode should have lost all cached data.
2087      */
2088     ASSERT(vn_has_cached_data(vp) == 0);

2090     /*
2091      * Defend against the possibility that higher-level callers
2092      * might not correctly balance open and close calls. If we
2093      * get here with open references remaining, it means there
2094      * was a missing VOP_CLOSE somewhere. If that happens, do
2095      * the close here so we don't "leak" FIDs on the server.
2096      */
2097     /* Exclusive lock for modifying n_fid stuff.
2098      * Don't want this one ever interruptible.
2099      */
2100     (void) smbfs_rw_enter_sig(&np->r_lkserlock, RW_WRITER, 0);
2101     smb_credinit(&scred, cr);

2103     switch (np->n_ovtype) {
2104     case VNON:
2105         /* not open (OK) */
2106         break;

2108     case VDIR:
2109         if (np->n_dirrefs == 0)
2110             break;
2111         SMBVDEBDEBUG("open dir: refs %d path %s\n",
2112             np->n_dirrefs, np->n_rpath);
2113         /* Force last close. */
2114         np->n_dirrefs = 1;
2115         smbfs_rele_fid(np, &scred);
2116         break;

2118     case VREG:
2119         if (np->n_fidrefs == 0)
2120             break;
2121         SMBVDEBDEBUG("open file: refs %d id 0x%x path %s\n",
2122             np->n_fidrefs, np->n_fid, np->n_rpath);
2123         /* Force last close. */
2124         np->n_fidrefs = 1;
2125         smbfs_rele_fid(np, &scred);
2126         break;

2128     default:
2129         SMBVDEBDEBUG("bad n_ovtype %d\n", np->n_ovtype);
2130         np->n_ovtype = VNON;
2131         break;
2132     }

2134     smb_credrele(&scred);

```

```

2135     smbfs_rw_exit(&np->r_lkserlock);

2137     /*
2138      * XATTR directories (and the files under them) have
2139      * little value for reclaim, so just remove them from
2140      * the "hash" (AVL) as soon as they go inactive.
2141      * Note that the node may already have been removed
2142      * from the hash by smbfsremove.
2143      */
2144     if ((np->n_flag & N_XATTR) != 0 &&
2145         (np->r_flags & RHASHED) != 0)
2146         smbfs_rhash(np);

2148     smbfs_addfree(np);
2149 }
_____unchanged_portion_omitted_____

2203 /* ARGSUSED */
2204 static int
2205 smbfslookup(vnode_t *dvp, char *nm, vnode_t **vpp, cred_t *cr,
2206            int cache_ok, caller_context_t *ct)
2207 {
2208     int            error;
2209     int            supplen; /* supported length */
2210     vnode_t        *vp;
2211     smbnode_t      *np;
2212     smbnode_t      *dnp;
2213     smbmntinfo_t   *smi;
2214     /* struct smb_vc *vcp; */
2215     const char     *ill;
2216     const char     *name = (const char *)nm;
2217     int            nmlen = strlen(nm);
2218     int            rplen;
2219     struct smb_cred scred;
2220     struct smbattr fa;

2222     smi = VTOSMI(dvp);
2223     dnp = VTOSMB(dvp);

2225     ASSERT(curproc->p_zone == smi->smi_zone_ref.zref_zone);

2227 #ifdef NOT_YET
2228     vcp = SSTOVC(smi->smi_share);

2230     /* XXX: Should compute this once and store it in smbmntinfo_t */
2231     supplen = (SMB_DIALECT(vcp) >= SMB_DIALECT_LANMAN2_0) ? 255 : 12;
2232 #else
2233     supplen = 255;
2234 #endif

2236     /*
2237      * Rlock must be held, either reader or writer.
2238      * XXX: Can we check without looking directly
2239      * inside the struct smbfs_rwlock_t?
2240      */
2241     ASSERT(dnp->r_rwlock.count != 0);

2243     /*
2244      * If lookup is for "", just return dvp.
2245      * No need to perform any access checks.
2246      */
2247     if (nmlen == 0) {
2248         VN_HOLD(dvp);
2249         *vpp = dvp;
2250         return (0);
2251     }

```

```

2252     /*
2253      * Can't do lookups in non-directories.
2254      */
2255     if (dvp->v_type != VDIR)
2256         return (ENOTDIR);

2257     /*
2258      * Need search permission in the directory.
2259      */
2260     error = smbfs_access(dvp, VEXEC, 0, cr, ct);
2261     if (error)
2262         return (error);

2264     /*
2265      * If lookup is for ".", just return dvp.
2266      * Access check was done above.
2267      */
2268     if (nmlen == 1 && name[0] == '.') {
2269         VN_HOLD(dvp);
2270         *vpp = dvp;
2271         return (0);
2272     }

2274     /*
2275      * Now some sanity checks on the name.
2276      * First check the length.
2277      */
2278     if (nmlen > supplen)
2279         return (ENAMETOOLONG);

2281     /*
2282      * Avoid surprises with characters that are
2283      * illegal in Windows file names.
2284      * Todo: CATIA mappings?
2285      * Todo: CATIA mappings XXX
2286      */
2287     ill = illegal_chars;
2288     if (dnp->n_flag & N_XATTR)
2289         ill++; /* allow colon */
2290     if (strpbrk(nm, ill))
2291         return (EINVAL);

2292     /*
2293      * Special handling for lookup of ".."
2294      *
2295      * We keep full pathnames (as seen on the server)
2296      * so we can just trim off the last component to
2297      * get the full pathname of the parent. Note:
2298      * We don't actually copy and modify, but just
2299      * compute the trimmed length and pass that with
2300      * the current dir path (not null terminated).
2301      *
2302      * We don't go over-the-wire to get attributes
2303      * for ".." because we know it's a directory,
2304      * and we can just leave the rest "stale"
2305      * until someone does a getattr.
2306      */
2307     if (nmlen == 2 && name[0] == '.' && name[1] == '.') {
2308         if (dvp->v_flag & VROOT) {
2309             /*
2310              * Already at the root. This can happen
2311              * with directory listings at the root,
2312              * which lookup "." and ".." to get the
2313              * inode numbers. Let "." be the same
2314              * as "." in the FS root.

```

```

2315         */
2316         VN_HOLD(dvp);
2317         *vpp = dvp;
2318         return (0);
2319     }
2320
2321     /*
2322     * Special case for XATTR directory
2323     */
2324     if (dvp->v_flag & V_XATTRDIR) {
2325         error = smbfs_xa_parent(dvp, vpp);
2326         return (error);
2327     }
2328
2329     /*
2330     * Find the parent path length.
2331     */
2332     rplen = dnp->n_rplen;
2333     ASSERT(rplen > 0);
2334     while (--rplen >= 0) {
2335         if (dnp->n_rpath[rplen] == '\\')
2336             break;
2337     }
2338     if (rplen <= 0) {
2339         /* Found our way to the root. */
2340         vp = SMBTOV(smi->smi_root);
2341         VN_HOLD(vp);
2342         *vpp = vp;
2343         return (0);
2344     }
2345     np = smbfs_node_findcreate(smi,
2346         dnp->n_rpath, rplen, NULL, 0, 0,
2347         &smbfs_fattr0); /* force create */
2348     ASSERT(np != NULL);
2349     vp = SMBTOV(np);
2350     vp->v_type = VDIR;
2351
2352     /* Success! */
2353     *vpp = vp;
2354     return (0);
2355 }
2356
2357 /*
2358 * Normal lookup of a name under this directory.
2359 * Note we handled "", ".", "." above.
2360 */
2361 if (cache_ok) {
2362     /*
2363     * The caller indicated that it's OK to use a
2364     * cached result for this lookup, so try to
2365     * reclaim a node from the smbfs node cache.
2366     */
2367     error = smbfslookup_cache(dvp, nm, nmlen, &vp, cr);
2368     if (error)
2369         return (error);
2370     if (vp != NULL) {
2371         /* hold taken in lookup_cache */
2372         *vpp = vp;
2373         return (0);
2374     }
2375 }
2376
2377 /*
2378 * OK, go over-the-wire to get the attributes,
2379 * then create the node.
2380 */

```

```

2381     smb_credinit(&scred, cr);
2382     /* Note: this can allocate a new "name" */
2383     error = smbfs_smb_lookup(dnp, &name, &nmlen, &fa, &scred);
2384     smb_credrele(&scred);
2385     if (error == ENOTDIR) {
2386         /*
2387         * Lookup failed because this directory was
2388         * removed or renamed by another client.
2389         * Remove any cached attributes under it.
2390         */
2391         smbfs_attrcache_remove(dnp);
2392         smbfs_attrcache_prune(dnp);
2393     }
2394     if (error)
2395         goto out;
2396
2397     error = smbfs_nget(dvp, name, nmlen, &fa, &vp);
2398     if (error)
2399         goto out;
2400
2401     /* Success! */
2402     *vpp = vp;
2403
2404 out:
2405     /* smbfs_smb_lookup may have allocated name. */
2406     if (name != nm)
2407         smbfs_name_free(name, nmlen);
2408
2409     return (error);
2410 }
2411 unchanged_portion_omitted
2412
2413 /*
2414 * XXX
2415 * vsecattr_t is new to build 77, and we need to eventually support
2416 * it in order to create an ACL when an object is created.
2417 */
2418 * This op should support the new IGNORECASE flag for case-insensitive
2419 * lookups, per PSARC 2007/244.
2420 */
2421 /* ARGSUSED */
2422 static int
2423 smbfs_create(vnode_t *dvp, char *nm, struct vattr *va, enum vcexcl exclusive,
2424     int mode, vnode_t **vpp, cred_t *cr, int lfaware, caller_context_t *ct,
2425     vsecattr_t *vsecp)
2426 {
2427     int error;
2428     int cerror;
2429     vfs_t *vfsp;
2430     vnode_t *vp;
2431 #ifdef NOT_YET
2432     smbnode_t *np;
2433 #endif
2434     smbnode_t *dnp;
2435     smbmntinfo_t *smi;
2436     struct vattr vattr;
2437     struct smb_fattr fattr;
2438     struct smb_cred scred;
2439     const char *name = (const char *)nm;
2440     int nmlen = strlen(nm);
2441     uint32_t disp;
2442     uint16_t fid;
2443     int xattr;
2444
2445     vfsp = dvp->v_vfsp;

```

```

2537     smi = VFTOSMI(vfsp);
2538     dnp = VTOSMB(dvp);
2539     vp = NULL;

2541     if (curproc->p_zone != smi->smi_zone_ref.zref_zone)
2542         return (EPERM);

2544     if (smi->smi_flags & SMI_DEAD || vfs->vfs_flag & VFS_UNMOUNTED)
2545         return (EIO);

2547     /*
2548      * Note: this may break mknod(2) calls to create a directory,
2549      * but that's obscure use. Some other filesystems do this.
2550      * Todo: redirect VDIR type here to _mkdir.
2551      * XXX: Later, redirect VDIR type here to _mkdir.
2552      */
2552     if (va->va_type != VREG)
2553         return (EINVAL);

2555     /*
2556      * If the pathname is "", just use dvp, no checks.
2557      * Do this outside of the rwlock (like zfs).
2558      */
2559     if (nmlen == 0) {
2560         VN_HOLD(dvp);
2561         *vpp = dvp;
2562         return (0);
2563     }

2565     /* Don't allow "." or ".." through here. */
2566     if ((nmlen == 1 && name[0] == '.') ||
2567         (nmlen == 2 && name[0] == '.' && name[1] == '.'))
2568         return (EISDIR);

2570     /*
2571      * We make a copy of the attributes because the caller does not
2572      * expect us to change what va points to.
2573      */
2574     vattr = *va;

2576     if (smbfs_rw_enter_sig(&dnp->r_rwlock, RW_WRITER, SMBINTR(dvp)))
2577         return (EINTR);
2578     smb_credinit(&scred, cr);

2580     /*
2581      * NFS needs to go over the wire, just to be sure whether the
2582      * file exists or not. Using a cached result is dangerous in
2583      * this case when making a decision regarding existence.
2584      *
2585      * The SMB protocol does NOT really need to go OTW here
2586      * thanks to the expressive NTCREATE disposition values.
2587      * Unfortunately, to do Unix access checks correctly,
2588      * we need to know if the object already exists.
2589      * When the object does not exist, we need VWRITE on
2590      * the directory. Note: smbfslookup() checks VEXEC.
2591      */
2592     error = smbfslookup(dvp, nm, &vp, cr, 0, ct);
2593     if (error == 0) {
2594         /*
2595          * The file already exists. Error?
2596          * NB: have a hold from smbfslookup
2597          */
2598         if (exclusive == EXCL) {
2599             error = EEXIST;
2600             VN_RELE(vp);
2601             goto out;

```

```

2602     }
2603     /*
2604      * Verify requested access.
2605      */
2606     error = smbfs_access(vp, mode, 0, cr, ct);
2607     if (error) {
2608         VN_RELE(vp);
2609         goto out;
2610     }

2612     /*
2613      * Truncate (if requested).
2614      */
2615     if ((vattr.va_mask & AT_SIZE) && vp->v_type == VREG) {
2616         np = VTOSMB(vp);
2617         /*
2618          * Check here for large file truncation by
2619          * LF-unaware process, like ufs_create().
2620          */
2621         if (!(lfaware & FOFFMAX)) {
2622             mutex_enter(&np->r_statelock);
2623             if (np->r_size > MAXOFF32_T)
2624                 error = EOVERFLOW;
2625             mutex_exit(&np->r_statelock);
2626         }
2627         if (error) {
2628             VN_RELE(vp);
2629             goto out;
2630         }
2631         if ((vattr.va_mask & AT_SIZE) && vattr.va_size == 0) {
2632             vattr.va_mask = AT_SIZE;
2633             error = smbfssetattr(vp, &vattr, 0, cr);
2634             if (error) {
2635                 VN_RELE(vp);
2636                 goto out;
2637             }
2638             /* Existing file was truncated */
2639             vnevent_create(vp, ct);
2640         }
2641         /* invalidate pages done in smbfssetattr() */
2642     }
2643     /* Success! */
2644     #ifdef NOT_YET
2645     vnevent_create(vp, ct);
2646     #endif

2648     /*
2649      * The file did not exist. Need VWRITE in the directory.
2650      */
2651     error = smbfs_access(dvp, VWRITE, 0, cr, ct);
2652     if (error)
2653         goto out;

2655     /*
2656      * Now things get tricky. We also need to check the
2657      * requested open mode against the file we may create.
2658      * See comments at smbfs_access_rwx
2659      */
2660     error = smbfs_access_rwx(vfsp, VREG, mode, cr);
2661     if (error)
2662         goto out;

```

```

2664  /*
2665  * Now the code derived from Darwin,
2666  * but with greater use of NT_CREATE
2667  * disposition options.  Much changed.
2668  *
2669  * Create (or open) a new child node.
2670  * Note we handled "." and ".." above.
2671  */
2673  if (exclusive == EXCL)
2674      disp = NTCREATEX_DISP_CREATE;
2675  else {
2676      /* Truncate regular files if requested. */
2677      if ((va->va_type == VREG) &&
2678          (va->va_mask & AT_SIZE) &&
2679          (va->va_size == 0))
2680          disp = NTCREATEX_DISP_OVERWRITE_IF;
2681      else
2682          disp = NTCREATEX_DISP_OPEN_IF;
2683  }
2684  xattr = (dnp->n_flag & N_XATTR) ? 1 : 0;
2685  error = smbfs_smb_create(dnp,
2686                          name, nmlen, xattr,
2687                          disp, &scred, &fid);
2688  if (error)
2689      goto out;
2691  /*
2692  * XXX: Missing some code here to deal with
2693  * the case where we opened an existing file,
2694  * it's size is larger than 32-bits, and we're
2695  * setting the size from a process that's not
2696  * aware of large file offsets.  i.e.
2697  * from the NFS3 code:
2698  */
2699  #if NOT_YET /* XXX */
2700  if ((vattn.va_mask & AT_SIZE) &&
2701      vp->v_type == VREG) {
2702      np = VTOSMB(vp);
2703      /*
2704       * Check here for large file handled
2705       * by LF-unaware process (as
2706       * ufs_create() does)
2707       */
2708      if (!(lfaware & FOFFMAX)) {
2709          mutex_enter(&np->r_statelock);
2710          if (np->r_size > MAXOFF32_T)
2711              error = EOVERFLOW;
2712          mutex_exit(&np->r_statelock);
2713      }
2714      if (!error) {
2715          vattn.va_mask = AT_SIZE;
2716          error = smbfssetattr(vp,
2717                              &vattn, 0, cr);
2718      }
2719  }
2720  #endif /* XXX */
2721  /*
2722  * Should use the fid to get/set the size
2723  * while we have it opened here.  See above.
2724  */
2726  cerror = smbfs_smb_close(smi->smi_share, fid, NULL, &scred);
2727  if (cerror)
2728      SMBVDEBUG("error %d closing %s\\%s\n",
2729                cerror, dnp->n_rpath, name);

```

```

2701  /*
2702  * In the open case, the name may differ a little
2703  * from what we passed to create (case, etc.)
2704  * so call lookup to get the (opened) name.
2705  *
2706  * XXX: Could avoid this extra lookup if the
2707  * "createact" result from NT_CREATE says we
2708  * created the object.
2709  */
2710  error = smbfs_smb_lookup(dnp, &name, &nmlen, &fattn, &scred);
2711  if (error)
2712      goto out;
2714  /* update attr and directory cache */
2715  smbfs_attr_touchdir(dnp);
2717  error = smbfs_nget(dvp, name, nmlen, &fattn, &vp);
2718  if (error)
2719      goto out;
2721  /* XXX invalidate pages if we truncated? */
2722  /* Success! */
2723  *vpp = vp;
2724  error = 0;
2725 out:
2726  smb_credrele(&scred);
2727  smbfs_rw_exit(&dnp->r_rwlock);
2728  if (name != nm)
2729      smbfs_name_free(name, nmlen);
2730  return (error);
2731 }
2732 unchanged_portion_omitted
2733 /*
2734 * smbfsremove does the real work of removing in SMBFS
2735 * Caller has done dir access checks etc.
2736 */
2737 /* The normal way to delete a file over SMB is open it (with DELETE access),
2738 * set the "delete-on-close" flag, and close the file.  The problem for Unix
2739 * applications is that they expect the file name to be gone once the unlink
2740 * completes, and the SMB server does not actually delete the file until ALL
2741 * opens of that file are closed.  We can't assume our open handles are the
2742 * only open handles on a file we're deleting, so to be safe we'll try to
2743 * rename the file to a temporary name and then set delete-on-close.  If we
2744 * fail to set delete-on-close (i.e. because other opens prevent it) then
2745 * undo the changes we made and give up with EBUSY.  Note that we might have
2746 * permission to delete a file but lack permission to rename, so we want to
2747 * continue in cases where rename fails.  As an optimization, only do the
2748 * rename when we have the file open.
2749 */
2750 /* This is similar to what NFS does when deleting a file that has local opens,
2751 * but thanks to SMB delete-on-close, we don't need to keep track of when the
2752 * last local open goes away and send a delete.  The server does that for us.
2753 */
2754 /* ARGSUSED */
2755 static int
2756 smbfsremove(vnode_t *dvp, vnode_t *vp, struct smb_cred *scred,
2757             int flags)
2758 {
2759     smbnode_t *dnp = VTOSMB(dvp);
2760     smbnode_t *np = VTOSMB(vp);
2761     char *tmpname = NULL;
2762     int tnlen;

```

```

2812         int          error;
2813         unsigned short fid;
2814         boolean_t     have_fid = B_FALSE;
2815         boolean_t     renamed = B_FALSE;

2817         /*
2818          * The dvp RWlock must be held as writer.
2819          */
2820         ASSERT(dnp->r_rwlock.owner == curthread);

2822         /* Never allow link/unlink directories on SMB. */
2823         if (vp->v_type == VDIR)
2824             return (EPERM);

2826         /*
2827          * We need to flush any dirty pages which happen to
2828          * be hanging around before removing the file. This
2829          * shouldn't happen very often and mostly on file
2830          * systems mounted "nocto".
2831          */
2832         if (vn_has_cached_data(vp) &&
2833             ((np->r_flags & RDIRTY) || np->r_count > 0)) {
2834             error = smbfs_putpage(vp, (offset_t)0, 0, 0,
2835                 scred->scr_cred, NULL);
2836             if (error && (error == ENOSPC || error == EDQUOT)) {
2837                 mutex_enter(&np->r_statelock);
2838                 if (!np->r_error)
2839                     np->r_error = error;
2840                 mutex_exit(&np->r_statelock);
2841             }
2842         }

2844         /* Shared lock for n_fid use in smbfs_smb_setdisp etc. */
2845         if (smbfs_rw_enter_sig(&np->r_lkserlock, RW_READER, SMBINTR(vp)))
2846             return (EINTR);

2153         /* Force lookup to go 0tw */
2154         smbfs_attrcache_remove(np);

2848         /*
2849          * Get a file handle with delete access.
2850          * Close this FID before return.
2851          */
2852         error = smbfs_smb_tmpopen(np, STD_RIGHT_DELETE_ACCESS,
2853             scred, &fid);
2854         if (error) {
2855             SMBVDEBUG("error %d opening %s\n",
2856                 error, np->n_rpath);
2857             goto out;
2858         }
2859         have_fid = B_TRUE;

2861         /*
2862          * If we have the file open, try to rename it to a temporary name.
2863          * If we can't rename, continue on and try setting DoC anyway.
2864          */
2865         if ((vp->v_count > 1) && (np->n_fidrefs > 0)) {
2866             tmpname = kmem_alloc(MAXNAMELEN, KM_SLEEP);
2867             tnlen = smbfs_newname(tmpname, MAXNAMELEN);
2868             error = smbfs_smb_t2rename(np, tmpname, tnlen, scred, fid, 0);
2869             if (error != 0) {
2870                 SMBVDEBUG("error %d renaming %s -> %s\n",
2871                     error, np->n_rpath, tmpname);
2872                 /* Keep going without the rename. */
2873             } else {
2874                 renamed = B_TRUE;

```

```

2875         }
2876     }

2878     /*
2879     * Mark the file as delete-on-close. If we can't,
2880     * undo what we did and err out.
2881     */
2882     error = smbfs_smb_setdisp(np, fid, 1, scred);
2883     if (error != 0) {
2884         SMBVDEBUG("error %d setting DoC on %s\n",
2885             error, np->n_rpath);
2886         /*
2887          * Failed to set DoC. If we renamed, undo that.
2888          * Need np->n_rpath relative to parent (dnp).
2889          * Use parent path name length plus one for
2890          * the separator ('/' or ':')
2891          */
2892         if (renamed) {
2893             char *oldname;
2894             int oldnlen;
2895             int err2;

2897             oldname = np->n_rpath + (dnp->n_rplen + 1);
2898             oldnlen = np->n_rplen - (dnp->n_rplen + 1);
2899             err2 = smbfs_smb_t2rename(np, oldname, oldnlen,
2900                 scred, fid, 0);
2901             SMBVDEBUG("error %d un-renaming %s -> %s\n",
2902                 err2, tmpname, np->n_rpath);
2903         }
2904         error = EBUSY;
2905         goto out;
2906     }
2907     /* Done! */
2908     smbfs_attrcache_prune(np);

2910 #ifdef SMBFS_VNEVENT
2911     vnevent_remove(vp, dvp, nm, ct);
2912 #endif

2914 out:
2915     if (tmpname != NULL)
2916         kmem_free(tmpname, MAXNAMELEN);

2918     if (have_fid)
2919         (void) smbfs_smb_tmpclose(np, fid, scred);
2920     smbfs_rw_exit(&np->r_lkserlock);

2922     if (error == 0) {
2923         /* Keep lookup from finding this node anymore. */
2924         smbfs_rmhash(np);
2925     }

2927     return (error);
2928 }

2931 /* ARGSUSED */
2932 static int
2933 smbfs_link(vnode_t *tdvp, vnode_t *svp, char *tnm, cred_t *cr,
2934     caller_context_t *ct, int flags)
2935 {
2936     /* Not yet... */
2937     return (ENOSYS);
2938 }

```

```

2941 /*
2942  * XXX
2943  * This op should support the new FIGNORECASE flag for case-insensitive
2944  * lookups, per PSARC 2007/244.
2945  */
2946 /* ARGSUSED */
2947 static int
2948 smbfs_rename(vnode_t *odvp, char *onm, vnode_t *ndvp, char *nnm, cred_t *cr,
2949 caller_context_t *ct, int flags)
2950 {
2951     struct smb_cred scred;
2952     smbnode_t *odnp = VTOSMB(odvp);
2953     smbnode_t *ndnp = VTOSMB(ndvp);
2954     vnode_t *ovp;
2955     int error;

2957     if (curproc->p_zone != VTOSMI(odvp)->smi_zone_ref.zref_zone ||
2958         curproc->p_zone != VTOSMI(ndvp)->smi_zone_ref.zref_zone)
2959         return (EPERM);

2961     if (VTOSMI(odvp)->smi_flags & SMI_DEAD ||
2962         VTOSMI(ndvp)->smi_flags & SMI_DEAD ||
2963         odvp->v_vfsp->vfs_flag & VFS_UNMOUNTED ||
2964         ndvp->v_vfsp->vfs_flag & VFS_UNMOUNTED)
2965         return (EIO);

2967     if (strcmp(onm, ".") == 0 || strcmp(onm, "..") == 0 ||
2968         strcmp(nnm, ".") == 0 || strcmp(nnm, "..") == 0)
2969         return (EINVAL);

2971     /*
2972      * Check that everything is on the same filesystem.
2973      * vn_rename checks the fsid's, but in case we don't
2974      * fill those in correctly, check here too.
2975      */
2976     if (odvp->v_vfsp != ndvp->v_vfsp)
2977         return (EXDEV);

2979     /*
2980      * Need write access on source and target.
2981      * Server takes care of most checks.
2982      */
2983     error = smbfs_access(odvp, VWRITE|VEXEC, 0, cr, ct);
2984     if (error)
2985         return (error);
2986     if (odvp != ndvp) {
2987         error = smbfs_access(ndvp, VWRITE, 0, cr, ct);
2988         if (error)
2989             return (error);
2990     }

2992     /*
2993      * Need to lock both old/new dirs as writer.
2994      *
2995      * Avoid deadlock here on old vs new directory nodes
2996      * by always taking the locks in order of address.
2997      * The order is arbitrary, but must be consistent.
2998      */
2999     if (odnp < ndnp) {
3000         if (smbfs_rw_enter_sig(&odnp->r_rwlock, RW_WRITER,
3001             SMBINTR(odvp)))
3002             return (EINTR);
3003         if (smbfs_rw_enter_sig(&ndnp->r_rwlock, RW_WRITER,
3004             SMBINTR(ndvp))) {
3005             smbfs_rw_exit(&odnp->r_rwlock);
3006             return (EINTR);

```

```

3007     } else {
3008         if (smbfs_rw_enter_sig(&ndnp->r_rwlock, RW_WRITER,
3009             SMBINTR(ndvp)))
3010             return (EINTR);
3011         if (smbfs_rw_enter_sig(&odnp->r_rwlock, RW_WRITER,
3012             SMBINTR(odvp))) {
3013             smbfs_rw_exit(&ndnp->r_rwlock);
3014             return (EINTR);
3015         }
3016     }
3017     smb_credinit(&scred, cr);

3020     /* Lookup the "old" name */
3021     error = smbfslookup(odvp, onm, &ovp, cr, 0, ct);
3022     if (error == 0) {
3023         /*
3024          * Do the real rename work
3025          */
3026         error = smbfsrename(odvp, ovp, ndvp, nnm, &scred, flags);
3027         VN_RELE(ovp);
3028     }

3030     smb_credrele(&scred);
3031     smbfs_rw_exit(&odnp->r_rwlock);
3032     smbfs_rw_exit(&ndnp->r_rwlock);

3034     return (error);
3035 }

3037 /*
3038  * smbfsrename does the real work of renaming in SMBFS
3039  * Caller has done dir access checks etc.
3040  */
3041 /* ARGSUSED */
3042 static int
3043 smbfsrename(vnode_t *odvp, vnode_t *ovp, vnode_t *ndvp, char *nnm,
3044 struct smb_cred *scred, int flags)
3045 {
3046     smbnode_t *odnp = VTOSMB(odvp);
3047     smbnode_t *onp = VTOSMB(ovp);
3048     smbnode_t *ndnp = VTOSMB(ndvp);
3049     vnode_t *nvp = NULL;
3050     int error;
3051     int nvp_locked = 0;

3053     /* Things our caller should have checked. */
3054     ASSERT(curproc->p_zone == VTOSMI(odvp)->smi_zone_ref.zref_zone);
3055     ASSERT(odvp->v_vfsp == ndvp->v_vfsp);
3056     ASSERT(odnp->r_rwlock.owner == curthread);
3057     ASSERT(ndnp->r_rwlock.owner == curthread);

3059     /*
3060      * Lookup the target file. If it exists, it needs to be
3061      * checked to see whether it is a mount point and whether
3062      * it is active (open).
3063      */
3064     error = smbfslookup(ndvp, nnm, &nvp, scred->scr_cred, 0, NULL);
3065     if (!error) {
3066         /*
3067          * Target (nvp) already exists. Check that it
3068          * has the same type as the source. The server
3069          * will check this also, (and more reliably) but
3070          * this lets us return the correct error codes.
3071          */
3072         if (ovp->v_type == VDIR) {

```

```

3073         if (nvp->v_type != VDIR) {
3074             error = ENOTDIR;
3075             goto out;
3076         }
3077     } else {
3078         if (nvp->v_type == VDIR) {
3079             error = EISDIR;
3080             goto out;
3081         }
3082     }
3083
3084     /*
3085     * POSIX dictates that when the source and target
3086     * entries refer to the same file object, rename
3087     * must do nothing and exit without error.
3088     */
3089     if (ovp == nvp) {
3090         error = 0;
3091         goto out;
3092     }
3093
3094     /*
3095     * Also must ensure the target is not a mount point,
3096     * and keep mount/umount away until we're done.
3097     */
3098     if (vn_vfsrlock(nvp)) {
3099         error = EBUSY;
3100         goto out;
3101     }
3102     nvp_locked = 1;
3103     if (vn_mountedvfs(nvp) != NULL) {
3104         error = EBUSY;
3105         goto out;
3106     }
3107
3108     /*
3109     * CIFS may give a SHARING_VIOLATION error when
3110     * trying to rename onto an existing object,
3111     * so try to remove the target first.
3112     * (Only for files, not directories.)
3113     */
3114     if (nvp->v_type == VDIR) {
3115         error = EEXIST;
3116         goto out;
3117     }
3118     error = smbfsremove(ndvp, nvp, scred, flags);
3119     if (error != 0)
3120         goto out;
3121
3122     /*
3123     * OK, removed the target file. Continue as if
3124     * lookup target had failed (nvp == NULL).
3125     */
3126     vn_vfsunlock(nvp);
3127     nvp_locked = 0;
3128     VN_RELE(nvp);
3129     nvp = NULL;
3130 } /* nvp */
3131
3132 smbfs_attrcache_remove(onp);
3133
3134 error = smbfs_smb_rename(onp, ndnp, nnm, strlen(nnm), scred);
3135
3136 /*
3137 * If the old name should no longer exist,
3138 * discard any cached attributes under it.

```

```

3138     */
3139     if (error == 0) {
2434         if (error == 0)
3140             smbfs_attrcache_prune(onp);
3141         /* SMBFS_VNEVENT... */
3142     }
3143
3144 out:
3145     if (nvp) {
3146         if (nvp_locked)
3147             vn_vfsunlock(nvp);
3148         VN_RELE(nvp);
3149     }
3150
3151     return (error);
3152 }
3153
3154 /*
3155 * XXX
3156 * vsecattr_t is new to build 77, and we need to eventually support
3157 * it in order to create an ACL when an object is created.
3158 *
3159 * This op should support the new IGNORECASE flag for case-insensitive
3160 * lookups, per PSARC 2007/244.
3161 */
3162 /* ARGSUSED */
3163 static int
3164 smbfs_mkdir(vnode_t *dvp, char *nm, struct vattr *va, vnode_t **vpp,
3165             cred_t *cr, caller_context_t *ct, int flags, vsecattr_t *vsecp)
3166 {
3167     vnode_t *vp;
3168     struct smbnode *dnp = VTOSMB(dvp);
3169     struct smbntinfo *smi = VTOSMI(dvp);
3170     struct smb_cred scred;
3171     struct smbattr fatr;
3172     const char *name = (const char *) nm;
3173     int nmlen = strlen(name);
3174     int error, hiderr;
3175
3176     if (curproc->p_zone != smi->smi_zone_ref.zref_zone)
3177         return (EPERM);
3178
3179     if (smi->smi_flags & SMI_DEAD || dvp->v_vfsp->vfs_flag & VFS_UNMOUNTED)
3180         return (EIO);
3181
3182     if ((nmlen == 1 && name[0] == '.') ||
3183         (nmlen == 2 && name[0] == '.' && name[1] == '.'))
3184         return (EEXIST);
3185
3186     /* Only plain files are allowed in V_XATTRDIR. */
3187     if (dvp->v_flag & V_XATTRDIR)
3188         return (EINVAL);
3189
3190     if (smbfs_rw_enter_sig(&dnp->r_rwlock, RW_WRITER, SMBINTR(dvp)))
3191         return (EINTR);
3192     smb_credinit(&scred, cr);
3193
3194     /*
2488     * XXX: Do we need r_lkserlock too?
2489     * No use of any shared fid or fctx...
2490     */
2491
2492     /*
3195     * Require write access in the containing directory.
3196     */
3197     error = smbfs_access(dvp, VWRITE, 0, cr, ct);

```

```

3198         if (error)
3199             goto out;

3201     error = smbfs_smb_mkdir(dnp, name, nmlen, &scred);
3202     if (error)
3203         goto out;

3205     error = smbfs_smb_lookup(dnp, &name, &nmlen, &fattr, &scred);
3206     if (error)
3207         goto out;

3209     smbfs_attr_touchdir(dnp);

3211     error = smbfs_nget(dvp, name, nmlen, &fattr, &vp);
3212     if (error)
3213         goto out;

3215     if (name[0] == '.')
3216         if ((hiderr = smbfs_smb_hideit(VTOSMB(vp), NULL, 0, &scred)))
3217             SMBVDEBUG("hide failure %d\n", hiderr);

3219     /* Success! */
3220     *vpp = vp;
3221     error = 0;
3222 out:
3223     smb_credrele(&scred);
3224     smbfs_rw_exit(&dnp->r_rwlock);

3226     if (name != nm)
3227         smbfs_name_free(name, nmlen);

3229     return (error);
3230 }
_____unchanged_portion_omitted_____

3335 /* ARGSUSED */
3336 static int
3337 smbfs_symlink(vnode_t *dvp, char *lnm, struct vattr *tva, char *tnm, cred_t *cr,
3338 caller_context_t *ct, int flags)
3339 {
3340     /* Not yet... */
3341     return (ENOSYS);
3342 }

3345 /* ARGSUSED */
3346 static int
3347 smbfs_readdir(vnode_t *vp, struct uio *uiop, cred_t *cr, int *eofp,
3348 caller_context_t *ct, int flags)
3349 {
3350     struct smbnode *np = VTOSMB(vp);
3351     int error = 0;
3352     smbmntinfo_t *smi;

3354     smi = VTOSMI(vp);

3356     if (curproc->p_zone != smi->smi_zone_ref.zref_zone)
3357         return (EIO);

3359     if (smi->smi_flags & SMI_DEAD || vp->v_vfsp->vfs_flag & VFS_UNMOUNTED)
3360         return (EIO);

3362     /*
3363      * Require read access in the directory.
3364      */

```

```

3365     error = smbfs_access(vp, VREAD, 0, cr, ct);
3366     if (error)
3367         return (error);

3369     ASSERT(smbfs_rw_lock_held(&np->r_rwlock, RW_READER));

3371     /*
3372      * Todo readdir cache here
3373      * XXX: Todo readdir cache here
3374      * Note: NFS code is just below this.
3375      * I am serializing the entire readdir operation
3376      * now since we have not yet implemented readdir
3377      * cache. This fix needs to be revisited once
3378      * we implement readdir cache.
3379      */
3379     if (smbfs_rw_enter_sig(&np->r_lkserlock, RW_WRITER, SMBINTR(vp)))
3380         return (EINTR);

3382     error = smbfs_readdir(vp, uiop, cr, eofp, ct);

3384     smbfs_rw_exit(&np->r_lkserlock);

3386     return (error);
3387 }
_____unchanged_portion_omitted_____

3606 /*
3607  * Here NFS has: nfs3_bio
3608  * See smbfs_bio above.
3609  */

3611 /* ARGSUSED */
3612 static int
3613 smbfs_fid(vnode_t *vp, fid_t *fidp, caller_context_t *ct)
3614 {
3615     return (ENOSYS);
3616 }

3619 /*
3620  * The pair of functions VOP_RWLOCK, VOP_RWUNLOCK
3621  * are optional functions that are called by:
3622  * getdents, before/after VOP_READDIR
3623  * pread, before/after ... VOP_READ
3624  * pwrite, before/after ... VOP_WRITE
3625  * (other places)
3626  *
3627  * Careful here: None of the above check for any
3628  * error returns from VOP_RWLOCK / VOP_RWUNLOCK!
3629  * In fact, the return value from _rwlock is NOT
3630  * an error code, but V_WRITELOCK_TRUE / _FALSE.
3631  *
3632  * Therefore, it's up to _this_ code to make sure
3633  * the lock state remains balanced, which means
3634  * we can't "bail out" on interrupts, etc.
3635  */

3637 /* ARGSUSED2 */
3638 static int
3639 smbfs_rwlock(vnode_t *vp, int write_lock, caller_context_t *ctp)
3640 {
3641     smbnode_t *np = VTOSMB(vp);

3643     if (!write_lock) {
3644         (void) smbfs_rw_enter_sig(&np->r_rwlock, RW_READER, FALSE);

```

```

3645         return (V_WRITELOCK_FALSE);
3646     }

3649     (void) smbfs_rw_enter_sig(&np->r_rwlock, RW_WRITER, FALSE);
3650     return (V_WRITELOCK_TRUE);
3651 }
    unchanged_portion_omitted

3692 /* mmap support ***** */

3694 #ifdef DEBUG
3695 static int smbfs_lostpage = 0; /* number of times we lost original page */
3696 #endif

3698 /*
3699  * Return all the pages from [off..off+len) in file
3700  * Like nfs3_getpage
3701  */
3702 /* ARGSUSED */
3703 static int
3704 smbfs_getpage(vnode_t *vp, offset_t off, size_t len, uint_t *protp,
3705              page_t *pl[], size_t plsz, struct seg *seg, caddr_t addr,
3706              enum seg_rw rw, cred_t *cr, caller_context_t *ct)
3707 {
3708     smbnode_t      *np;
3709     smbmntinfo_t   *smi;
3710     int             error;

3712     np = VTOSMB(vp);
3713     smi = VTOSMI(vp);

3715     if (curproc->p_zone != smi->smi_zone_ref.zref_zone)
3716         return (EIO);

3718     if (smi->smi_flags & SMI_DEAD || vp->v_vfsp->vfs_flag & VFS_UNMOUNTED)
3719         return (EIO);

3721     if (vp->v_flag & VNOMAP)
3722         return (ENOSYS);

3724     if (protp != NULL)
3725         *protp = PROT_ALL;

3727     /*
3728      * Now validate that the caches are up to date.
3729      */
3730     error = smbfs_validate_caches(vp, cr);
3731     if (error)
3732         return (error);

3734 retry:
3735     mutex_enter(&np->r_statelock);

3737     /*
3738      * Don't create dirty pages faster than they
3739      * can be cleaned ... (etc. see nfs)
3740      *
3741      * Here NFS also tests:
3742      * (mi->mi_max_threads != 0 &&
3743      *  rp->r_awcount > 2 * mi->mi_max_threads)
3744      */
3745     if (rw == S_CREATE) {
3746         while (np->r_gcount > 0)
3747             cv_wait(&np->r_cv, &np->r_statelock);
3748     }

```

```

3750     /*
3751      * If we are getting called as a side effect of a write
3752      * operation the local file size might not be extended yet.
3753      * In this case we want to be able to return pages of zeroes.
3754      */
3755     if (off + len > np->r_size + PAGEOFFSET && seg != segkmap) {
3756         mutex_exit(&np->r_statelock);
3757         return (EFAULT); /* beyond EOF */
3758     }

3760     mutex_exit(&np->r_statelock);

3762     error = pvn_getpages(smbfs_getapage, vp, off, len, protp,
3763                        pl, plsz, seg, addr, rw, cr);

3765     switch (error) {
3766     case SMBFS_EOF:
3767         smbfs_purge_caches(vp, cr);
3768         goto retry;
3769     case ESTALE:
3770         /*
3771          * Here NFS has: PURGE_STALE_FH(error, vp, cr);
3772          * In-line here as we only use it once.
3773          */
3774         mutex_enter(&np->r_statelock);
3775         np->r_flags |= RSTALE;
3776         if (!np->r_error)
3777             np->r_error = (error);
3778         mutex_exit(&np->r_statelock);
3779         if (vn_has_cached_data(vp))
3780             smbfs_invalidate_pages(vp, (u_offset_t)0, cr);
3781         smbfs_purge_caches(vp, cr);
3782         break;
3783     default:
3784         break;
3785     }

3787     return (error);
3788 }

3790 /*
3791  * Called from pvn_getpages to get a particular page.
3792  * Like nfs3_getapage
3793  */
3794 /* ARGSUSED */
3795 static int
3796 smbfs_getapage(vnode_t *vp, u_offset_t off, size_t len, uint_t *protp,
3797               page_t *pl[], size_t plsz, struct seg *seg, caddr_t addr,
3798               enum seg_rw rw, cred_t *cr)
3799 {
3800     smbnode_t      *np;
3801     smbmntinfo_t   *smi;

3803     uint_t          bsize;
3804     struct buf      *bp;
3805     page_t          *pp;
3806     u_offset_t      lbn;
3807     u_offset_t      io_off;
3808     u_offset_t      blkoff;
3809     size_t          io_len;
3810     uint_t          blksize;
3811     int error;
3812     /* int readahead; */
3813     int readahead_issued = 0;
3814     /* int ra_window; * readahead window */

```

```

3815     page_t *pagefound;
3817     np = VTOSMB(vp);
3818     smi = VTOSMI(vp);
3820     if (curproc->p_zone != smi->smi_zone_ref.zref_zone)
3821         return (EIO);
3823     if (smi->smi_flags & SMI_DEAD || vp->v_vfsp->vfs_flag & VFS_UNMOUNTED)
3824         return (EIO);
3826     bsize = MAX(vp->v_vfsp->vfs_bsize, PAGESIZE);
3828 reread:
3829     bp = NULL;
3830     pp = NULL;
3831     pagefound = NULL;
3833     if (pl != NULL)
3834         pl[0] = NULL;
3836     error = 0;
3837     lbn = off / bsize;
3838     blkoff = lbn * bsize;
3840     /*
3841      * NFS queues up readahead work here.
3842      */
3844 again:
3845     if ((pagefound = page_exists(vp, off)) == NULL) {
3846         if (pl == NULL) {
3847             (void) 0; /* Todo: smbfs_async_readahead(); */
3848         } else if (rw == S_CREATE) {
3849             /*
3850              * Block for this page is not allocated, or the offset
3851              * is beyond the current allocation size, or we're
3852              * allocating a swap slot and the page was not found,
3853              * so allocate it and return a zero page.
3854              */
3855             if ((pp = page_create_va(vp, off,
3856                                     PAGESIZE, PG_WAIT, seg, addr)) == NULL)
3857                 cmn_err(CE_PANIC, "smbfs_getapage: page_create")
3858             io_len = PAGESIZE;
3859             mutex_enter(&np->r_statelock);
3860             np->r_nextr = off + PAGESIZE;
3861             mutex_exit(&np->r_statelock);
3862         } else {
3863             /*
3864              * Need to go to server to get a BLOCK, exception to
3865              * that being while reading at offset = 0 or doing
3866              * random i/o, in that case read only a PAGE.
3867              */
3868             mutex_enter(&np->r_statelock);
3869             if (blkoff < np->r_size &&
3870                 blkoff + bsize >= np->r_size) {
3871                 /*
3872                  * If only a block or less is left in
3873                  * the file, read all that is remaining.
3874                  */
3875                 if (np->r_size <= off) {
3876                     /*
3877                      * Trying to access beyond EOF,
3878                      * set up to get at least one page.
3879                      */
3880                     blksize = off + PAGESIZE - blkoff;

```

```

3881         } else
3882             blksize = np->r_size - blkoff;
3883     } else if ((off == 0) ||
3884               (off != np->r_nextr && !readahead_issued)) {
3885         blksize = PAGESIZE;
3886         blkoff = off; /* block = page here */
3887     } else
3888         blksize = bsize;
3889     mutex_exit(&np->r_statelock);
3891     pp = pvn_read_kluster(vp, off, seg, addr, &io_off,
3892                          &io_len, blkoff, blksize, 0);
3894     /*
3895      * Some other thread has entered the page,
3896      * so just use it.
3897      */
3898     if (pp == NULL)
3899         goto again;
3901     /*
3902      * Now round the request size up to page boundaries.
3903      * This ensures that the entire page will be
3904      * initialized to zeroes if EOF is encountered.
3905      */
3906     io_len = ptob(btopr(io_len));
3908     bp = pageio_setup(pp, io_len, vp, B_READ);
3909     ASSERT(bp != NULL);
3911     /*
3912      * pageio_setup should have set b_addr to 0. This
3913      * is correct since we want to do I/O on a page
3914      * boundary. bp_mapin will use this addr to calculate
3915      * an offset, and then set b_addr to the kernel virtual
3916      * address it allocated for us.
3917      */
3918     ASSERT(bp->b_un.b_addr == 0);
3920     bp->b_edev = 0;
3921     bp->b_dev = 0;
3922     bp->b_lblkno = lbtodb(io_off);
3923     bp->b_file = vp;
3924     bp->b_offset = (offset_t)off;
3925     bp_mapin(bp);
3927     /*
3928      * If doing a write beyond what we believe is EOF,
3929      * don't bother trying to read the pages from the
3930      * server, we'll just zero the pages here. We
3931      * don't check that the rw flag is S_WRITE here
3932      * because some implementations may attempt a
3933      * read access to the buffer before copying data.
3934      */
3935     mutex_enter(&np->r_statelock);
3936     if (io_off >= np->r_size && seg == segkmap) {
3937         mutex_exit(&np->r_statelock);
3938         bzero(bp->b_un.b_addr, io_len);
3939     } else {
3940         mutex_exit(&np->r_statelock);
3941         error = smbfs_bio(bp, 0, cr);
3942     }
3944     /*
3945      * Unmap the buffer before freeing it.
3946      */

```

```

3947         bp_mapout(bp);
3948         pageio_done(bp);
3950         /* Here NFS3 updates all pp->p_fsdata */
3952         if (error == SMBFS_EOF) {
3953             /*
3954              * If doing a write system call just return
3955              * zeroed pages, else user tried to get pages
3956              * beyond EOF, return error. We don't check
3957              * that the rw flag is S_WRITE here because
3958              * some implementations may attempt a read
3959              * access to the buffer before copying data.
3960              */
3961             if (seg == segkmap)
3962                 error = 0;
3963             else
3964                 error = EFAULT;
3965         }
3967         if (!readahead_issued && !error) {
3968             mutex_enter(&np->r_statelock);
3969             np->r_nextr = io_off + io_len;
3970             mutex_exit(&np->r_statelock);
3971         }
3972     }
3973 }
3975 if (pl == NULL)
3976     return (error);
3978 if (error) {
3979     if (pp != NULL)
3980         pvn_read_done(pp, B_ERROR);
3981     return (error);
3982 }
3984 if (pagefound) {
3985     se_t se = (rw == S_CREATE ? SE_EXCL : SE_SHARED);
3987     /*
3988      * Page exists in the cache, acquire the appropriate lock.
3989      * If this fails, start all over again.
3990      */
3991     if ((pp = page_lookup(vp, off, se)) == NULL) {
3992 #ifdef DEBUG
3993         smbfs_lostpage++;
3994 #endif
3995         goto reread;
3996     }
3997     pl[0] = pp;
3998     pl[1] = NULL;
3999     return (0);
4000 }
4002 if (pp != NULL)
4003     pvn_plist_init(pp, pl, plsz, off, io_len, rw);
4005 return (error);
4006 }
4008 /*
4009  * Here NFS has: nfs3_readahead
4010  * No read-ahead in smbfs yet.
4011  */

```

```

4013 /*
4014  * Flags are composed of {B_INVAL, B_FREE, B_DONTNEED, B_FORCE}
4015  * If len == 0, do from off to EOF.
4016  *
4017  * The normal cases should be len == 0 && off == 0 (entire vp list),
4018  * len == MAXBSIZE (from segmap_release actions), and len == PAGE_SIZE
4019  * (from pageout).
4020  *
4021  * Like nfs3_putpage + nfs_putpages
4022  */
4023 /* ARGSUSED */
4024 static int
4025 smbfs_putpage(vnode_t *vp, offset_t off, size_t len, int flags, cred_t *cr,
4026               caller_context_t *ct)
4027 {
4028     smbnode_t *np;
4029     smbmntinfo_t *smi;
4030     page_t *pp;
4031     u_offset_t eoff;
4032     u_offset_t io_off;
4033     size_t io_len;
4034     int error;
4035     int rdirty;
4036     int err;
4038     np = VTOSMB(vp);
4039     smi = VTOSMI(vp);
4041     if (curproc->p_zone != smi->smi_zone_ref.zref_zone)
4042         return (EIO);
4044     if (smi->smi_flags & SMI_DEAD || vp->v_vfsp->vfs_flag & VFS_UNMOUNTED)
4045         return (EIO);
4047     if (vp->v_flag & VNOMAP)
4048         return (ENOSYS);
4050     /* Here NFS does rp->r_count (++/--) stuff. */
4052     /* Beginning of code from nfs_putpages. */
4054     if (!vn_has_cached_data(vp))
4055         return (0);
4057     /*
4058      * If ROUTOFSpace is set, then all writes turn into B_INVAL
4059      * writes. B_FORCE is set to force the VM system to actually
4060      * invalidate the pages, even if the i/o failed. The pages
4061      * need to get invalidated because they can't be written out
4062      * because there isn't any space left on either the server's
4063      * file system or in the user's disk quota. The B_FREE bit
4064      * is cleared to avoid confusion as to whether this is a
4065      * request to place the page on the freelist or to destroy
4066      * it.
4067      */
4068     if ((np->r_flags & ROUTOFSpace) ||
4069         (vp->v_vfsp->vfs_flag & VFS_UNMOUNTED))
4070         flags = (flags & ~B_FREE) | B_INVAL | B_FORCE;
4072     if (len == 0) {
4073         /*
4074          * If doing a full file synchronous operation, then clear
4075          * the RDIRTY bit. If a page gets dirtied while the flush
4076          * is happening, then RDIRTY will get set again. The
4077          * RDIRTY bit must get cleared before the flush so that
4078          * we don't lose this information.

```

```

4079      *
4080      * NFS has B_ASYNC vs sync stuff here.
4081      */
4082      if (off == (u_offset_t)0 &&
4083          (np->r_flags & RDIRTY)) {
4084          mutex_enter(&np->r_statelock);
4085          rdirty = (np->r_flags & RDIRTY);
4086          np->r_flags &= ~RDIRTY;
4087          mutex_exit(&np->r_statelock);
4088      } else
4089          rdirty = 0;
4091
4092      /*
4093       * Search the entire vp list for pages >= off, and flush
4094       * the dirty pages.
4095       */
4096      error = pvn_vplist_dirty(vp, off, smbfs_putapage,
4097                              flags, cr);
4098
4099      /*
4100       * If an error occurred and the file was marked as dirty
4101       * before and we aren't forcibly invalidating pages, then
4102       * reset the RDIRTY flag.
4103       */
4104      if (error && rdirty &&
4105          (flags & (B_INVAL | B_FORCE)) != (B_INVAL | B_FORCE)) {
4106          mutex_enter(&np->r_statelock);
4107          np->r_flags |= RDIRTY;
4108          mutex_exit(&np->r_statelock);
4109      } else {
4110          /*
4111           * Do a range from [off...off + len) looking for pages
4112           * to deal with.
4113           */
4114          error = 0;
4115          io_len = 1; /* quiet warnings */
4116          eoff = off + len;
4117
4118          for (io_off = off; io_off < eoff; io_off += io_len) {
4119              mutex_enter(&np->r_statelock);
4120              if (io_off >= np->r_size) {
4121                  mutex_exit(&np->r_statelock);
4122                  break;
4123              }
4124              mutex_exit(&np->r_statelock);
4125              /*
4126               * If we are not invalidating, synchronously
4127               * freeing or writing pages use the routine
4128               * page_lookup_nowait() to prevent reclaiming
4129               * them from the free list.
4130               */
4131              if ((flags & B_INVAL) || !(flags & B_ASYNC)) {
4132                  pp = page_lookup(vp, io_off,
4133                                  (flags & (B_INVAL | B_FREE)) ?
4134                                  SE_EXCL : SE_SHARED);
4135              } else {
4136                  pp = page_lookup_nowait(vp, io_off,
4137                                          (flags & B_FREE) ? SE_EXCL : SE_SHARED);
4138              }
4139
4140              if (pp == NULL || !pvn_getdirty(pp, flags))
4141                  io_len = PAGE_SIZE;
4142              else {
4143                  err = smbfs_putapage(vp, pp, &io_off,
4144                                      &io_len, flags, cr);

```

```

4145          if (!error)
4146              error = err;
4147          /*
4148           * "io_off" and "io_len" are returned as
4149           * the range of pages we actually wrote.
4150           * This allows us to skip ahead more quickly
4151           * since several pages may've been dealt
4152           * with by this iteration of the loop.
4153           */
4154      }
4155  }
4156  }
4157
4158  return (error);
4159 }
4160
4161 /*
4162  * Write out a single page, possibly klustering adjacent dirty pages.
4163  * Like nfs3_putapage / nfs3_sync_putapage
4164  */
4165 static int
4166 smbfs_putapage(vnode_t *vp, page_t *pp, u_offset_t *offp, size_t *lenp,
4167               int flags, cred_t *cr)
4168 {
4169     smbnode_t *np;
4170     u_offset_t io_off;
4171     u_offset_t lbn_off;
4172     u_offset_t lbn;
4173     size_t io_len;
4174     uint_t bsize;
4175     int error;
4176
4177     np = VTOSMB(vp);
4178
4179     ASSERT(!vn_is_readonly(vp));
4180
4181     bsize = MAX(vp->v_vfsp->vfs_bsize, PAGE_SIZE);
4182     lbn = pp->p_offset / bsize;
4183     lbn_off = lbn * bsize;
4184
4185     /*
4186      * Find a kluster that fits in one block, or in
4187      * one page if pages are bigger than blocks. If
4188      * there is less file space allocated than a whole
4189      * page, we'll shorten the i/o request below.
4190      */
4191     pp = pvn_write_kluster(vp, pp, &io_off, &io_len, lbn_off,
4192                           roundup(bsize, PAGE_SIZE), flags);
4193
4194     /*
4195      * pvn_write_kluster shouldn't have returned a page with offset
4196      * behind the original page we were given. Verify that.
4197      */
4198     ASSERT((pp->p_offset / bsize) >= lbn);
4199
4200     /*
4201      * Now pp will have the list of kept dirty pages marked for
4202      * write back. It will also handle invalidation and freeing
4203      * of pages that are not dirty. Check for page length rounding
4204      * problems.
4205      */
4206     if (io_off + io_len > lbn_off + bsize) {
4207         ASSERT((io_off + io_len) - (lbn_off + bsize) < PAGE_SIZE);
4208         io_len = lbn_off + bsize - io_off;
4209     }
4210 }

```

```

4211 /*
4212  * The RMODINPROGRESS flag makes sure that smbfs_bio() sees a
4213  * consistent value of r_size. RMODINPROGRESS is set in writerp().
4214  * When RMODINPROGRESS is set it indicates that a uiomove() is in
4215  * progress and the r_size has not been made consistent with the
4216  * new size of the file. When the uiomove() completes the r_size is
4217  * updated and the RMODINPROGRESS flag is cleared.
4218  *
4219  * The RMODINPROGRESS flag makes sure that smbfs_bio() sees a
4220  * consistent value of r_size. Without this handshaking, it is
4221  * possible that smbfs_bio() picks up the old value of r_size
4222  * before the uiomove() in writerp() completes. This will result
4223  * in the write through smbfs_bio() being dropped.
4224  *
4225  * More precisely, there is a window between the time the uiomove()
4226  * completes and the time the r_size is updated. If a VOP_PUTPAGE()
4227  * operation intervenes in this window, the page will be picked up,
4228  * because it is dirty (it will be unlocked, unless it was
4229  * pagecreate'd). When the page is picked up as dirty, the dirty
4230  * bit is reset (pv_n_getdirty()). In smbfs_write(), r_size is
4231  * checked. This will still be the old size. Therefore the page will
4232  * not be written out. When segmap_release() calls VOP_PUTPAGE(),
4233  * the page will be found to be clean and the write will be dropped.
4234  */
4235 if (np->r_flags & RMODINPROGRESS) {
4236     mutex_enter(&np->r_statelock);
4237     if ((np->r_flags & RMODINPROGRESS) &&
4238         np->r_modaddr + MAXBSIZE > io_off &&
4239         np->r_modaddr < io_off + io_len) {
4240         page_t *plist;
4241         /*
4242          * A write is in progress for this region of the file.
4243          * If we did not detect RMODINPROGRESS here then this
4244          * path through smbfs_putpage() would eventually go to
4245          * smbfs_bio() and may not write out all of the data
4246          * in the pages. We end up losing data. So we decide
4247          * to set the modified bit on each page in the page
4248          * list and mark the rnode with RDIRTY. This write
4249          * will be restarted at some later time.
4250          */
4251         plist = pp;
4252         while (plist != NULL) {
4253             pp = plist;
4254             page_sub(&plist, pp);
4255             hat_setmod(pp);
4256             page_io_unlock(pp);
4257             page_unlock(pp);
4258         }
4259         np->r_flags |= RDIRTY;
4260         mutex_exit(&np->r_statelock);
4261         if (offp)
4262             *offp = io_off;
4263         if (lenp)
4264             *lenp = io_len;
4265         return (0);
4266     }
4267     mutex_exit(&np->r_statelock);
4268 }
4269
4270 /*
4271  * NFS handles (flags & B_ASYNC) here...
4272  * (See nfs_async_putpage())
4273  *
4274  * This code section from: nfs3_sync_putpage()
4275  */

```

```

4277     flags |= B_WRITE;
4278
4279     error = smbfs_rdwrlbn(vp, pp, io_off, io_len, flags, cr);
4280
4281     if ((error == ENOSPC || error == EDQUOT || error == EFBIG ||
4282         error == EACCES) &&
4283         (flags & (B_INVAL|B_FORCE)) != (B_INVAL|B_FORCE)) {
4284         if (!(np->r_flags & ROUTOFSPACE)) {
4285             mutex_enter(&np->r_statelock);
4286             np->r_flags |= ROUTOFSPACE;
4287             mutex_exit(&np->r_statelock);
4288         }
4289         flags |= B_ERROR;
4290         pv_n_write_done(pp, flags);
4291         /*
4292          * If this was not an async thread, then try again to
4293          * write out the pages, but this time, also destroy
4294          * them whether or not the write is successful. This
4295          * will prevent memory from filling up with these
4296          * pages and destroying them is the only alternative
4297          * if they can't be written out.
4298          *
4299          * Don't do this if this is an async thread because
4300          * when the pages are unlocked in pv_n_write_done,
4301          * some other thread could have come along, locked
4302          * them, and queued for an async thread. It would be
4303          * possible for all of the async threads to be tied
4304          * up waiting to lock the pages again and they would
4305          * all already be locked and waiting for an async
4306          * thread to handle them. Deadlock.
4307          */
4308         if (!(flags & B_ASYNC)) {
4309             error = smbfs_putpage(vp, io_off, io_len,
4310                                   B_INVAL | B_FORCE, cr, NULL);
4311         }
4312     } else {
4313         if (error)
4314             flags |= B_ERROR;
4315         else if (np->r_flags & ROUTOFSPACE) {
4316             mutex_enter(&np->r_statelock);
4317             np->r_flags &= ~ROUTOFSPACE;
4318             mutex_exit(&np->r_statelock);
4319         }
4320         pv_n_write_done(pp, flags);
4321     }
4322
4323     /* Now more code from: nfs3_putpage */
4324
4325     if (offp)
4326         *offp = io_off;
4327     if (lenp)
4328         *lenp = io_len;
4329
4330     return (error);
4331 }
4332
4333 /*
4334  * NFS has this in nfs_client.c (shared by v2,v3,...)
4335  * We have it here so smbfs_putpage can be file scope.
4336  */
4337 void
4338 smbfs_invalidate_pages(vnode_t *vp, u_offset_t off, cred_t *cr)
4339 {
4340     smbnode_t *np;
4341
4342     np = VTOSMB(vp);

```

```

4344     mutex_enter(&np->r_statelock);
4345     while (np->r_flags & RTRUNCATE)
4346         cv_wait(&np->r_cv, &np->r_statelock);
4347     np->r_flags |= RTRUNCATE;

4349     if (off == (u_offset_t)0) {
4350         np->r_flags &= ~RDIRTY;
4351         if (!(np->r_flags & RSTALE))
4352             np->r_error = 0;
4353     }
4354     /* Here NFSv3 has np->r_truncaddr = off; */
4355     mutex_exit(&np->r_statelock);

4357     (void) pvn_vplist_dirty(vp, off, smbfs_putapage,
4358         B_INVAL | B_TRUNC, cr);

4360     mutex_enter(&np->r_statelock);
4361     np->r_flags &= ~RTRUNCATE;
4362     cv_broadcast(&np->r_cv);
4363     mutex_exit(&np->r_statelock);
4364 }

4366 /* Like nfs3_map */

4368 /* ARGSUSED */
4369 static int
4370 smbfs_map(vnode_t *vp, offset_t off, struct as *as, caddr_t *addrp,
4371     size_t len, uchar_t prot, uchar_t maxprot, uint_t flags,
4372     cred_t *cr, caller_context_t *ct)
4373 {
4374     segvn_crargs_t vn_a;
4375     struct vattr va;
4376     smbnode_t *np;
4377     smbmntinfo_t *smi;
4378     int error;

4380     np = VTOSMB(vp);
4381     smi = VTOSMI(vp);

4383     if (curproc->p_zone != smi->smi_zone_ref.zref_zone)
4384         return (EIO);

4386     if (smi->smi_flags & SMI_DEAD || vp->v_vfsp->vfs_flag & VFS_UNMOUNTED)
4387         return (EIO);

4389     if (vp->v_flag & VNOMAP)
4390         return (ENOSYS);

4392     if (off < 0 || off + (ssize_t)len < 0)
4393         return (ENXIO);

4395     if (vp->v_type != VREG)
4396         return (ENODEV);

4398     /*
4399      * NFS does close-to-open consistency stuff here.
4400      * Just get (possibly cached) attributes.
4401      */
4402     va.va_mask = AT_ALL;
4403     if ((error = smbfsgetattr(vp, &va, cr)) != 0)
4404         return (error);

4406     /*
4407      * Check to see if the vnode is currently marked as not cachable.
4408      * This means portions of the file are locked (through VOP_FRLock).

```

```

4409     /* In this case the map request must be refused. We use
4410      * rp->r_lkserlock to avoid a race with concurrent lock requests.
4411      */
4412     /*
4413      * Atomically increment r_inmap after acquiring r_rwlock. The
4414      * idea here is to acquire r_rwlock to block read/write and
4415      * not to protect r_inmap. r_inmap will inform smbfs_read/write()
4416      * that we are in smbfs_map(). Now, r_rwlock is acquired in order
4417      * and we can prevent the deadlock that would have occurred
4418      * when smbfs_addmap() would have acquired it out of order.
4419      *
4420      * Since we are not protecting r_inmap by any lock, we do not
4421      * hold any lock when we decrement it. We atomically decrement
4422      * r_inmap after we release r_lkserlock. Note that rlock is
4423      * re-entered as writer in smbfs_addmap (called via as_map).
4424      */

4426     if (smbfs_rw_enter_sig(&np->r_rwlock, RW_WRITER, SMBINTR(vp)))
4427         return (EINTR);
4428     atomic_inc_uint(&np->r_inmap);
4429     smbfs_rw_exit(&np->r_rwlock);

4431     if (smbfs_rw_enter_sig(&np->r_lkserlock, RW_WRITER, SMBINTR(vp))) {
4432         atomic_dec_uint(&np->r_inmap);
4433         return (EINTR);
4434     }

4436     if (vp->v_flag & VNOCACHE) {
4437         error = EAGAIN;
4438         goto done;
4439     }

4441     /*
4442      * Don't allow concurrent locks and mapping if mandatory locking is
4443      * enabled.
4444      */
4445     if ((flk_has_remote_locks(vp) || smbfs_lm_has_sleep(vp)) &&
4446         MANDLOCK(vp, va.va_mode)) {
4447         error = EAGAIN;
4448         goto done;
4449     }

4451     as_rangelock(as);
4452     error = choose_addr(as, addrp, len, off, ADDR_VACALIGN, flags);
4453     if (error != 0) {
4454         as_rangeunlock(as);
4455         goto done;
4456     }

4458     vn_a.vp = vp;
4459     vn_a.offset = off;
4460     vn_a.type = (flags & MAP_TYPE);
4461     vn_a.prot = (uchar_t)prot;
4462     vn_a.maxprot = (uchar_t)maxprot;
4463     vn_a.flags = (flags & ~MAP_TYPE);
4464     vn_a.cred = cr;
4465     vn_a.amp = NULL;
4466     vn_a.szc = 0;
4467     vn_a.lgrp_mem_policy_flags = 0;

4469     error = as_map(as, *addrp, len, segvn_create, &vn_a);
4470     as_rangeunlock(as);

4472 done:
4473     smbfs_rw_exit(&np->r_lkserlock);
4474     atomic_dec_uint(&np->r_inmap);

```

```

4475     return (error);
4476 }

4478 /* ARGSUSED */
4479 static int
4480 smbfs_addmap(vnode_t *vp, offset_t off, struct as *as, caddr_t addr,
4481             size_t len, uchar_t prot, uchar_t maxprot, uint_t flags,
4482             cred_t *cr, caller_context_t *ct)
4483 {
4484     smbnode_t *np = VTOSMB(vp);
4485     boolean_t inc_fidrefs = B_FALSE;

4487     /*
4488      * When r_mapcnt goes from zero to non-zero,
4489      * increment n_fidrefs
4490      */
4491     mutex_enter(&np->r_statelock);
4492     if (np->r_mapcnt == 0)
4493         inc_fidrefs = B_TRUE;
4494     np->r_mapcnt += btopr(len);
4495     mutex_exit(&np->r_statelock);

4497     if (inc_fidrefs) {
4498         (void) smbfs_rw_enter_sig(&np->r_lkserlock, RW_WRITER, 0);
4499         np->n_fidrefs++;
4500         smbfs_rw_exit(&np->r_lkserlock);
4501     }

4503     return (0);
4504 }

4506 /*
4507  * Use an address space callback to flush pages dirty pages after unmap,
4508  * which we can't do directly in smbfs_delmap due to locking issues.
4509  */
4510 typedef struct smbfs_delmap_args {
4511     vnode_t      *vp;
4512     cred_t       *cr;
4513     offset_t     off;
4514     caddr_t      addr;
4515     size_t       len;
4516     uint_t       prot;
4517     uint_t       maxprot;
4518     uint_t       flags;
4519     boolean_t    dec_fidrefs;
4520 } smbfs_delmap_args_t;

4522 /* ARGSUSED */
4523 static int
4524 smbfs_delmap(vnode_t *vp, offset_t off, struct as *as, caddr_t addr,
4525             size_t len, uint_t prot, uint_t maxprot, uint_t flags,
4526             cred_t *cr, caller_context_t *ct)
4527 {
4528     smbnode_t *np = VTOSMB(vp);
4529     smbfs_delmap_args_t *dmapp;
4530     int error;

4532     dmapp = kmem_zalloc(sizeof (*dmapp), KM_SLEEP);

4534     dmapp->vp = vp;
4535     dmapp->off = off;
4536     dmapp->addr = addr;
4537     dmapp->len = len;
4538     dmapp->prot = prot;
4539     dmapp->maxprot = maxprot;
4540     dmapp->flags = flags;

```

```

4541     dmapp->cr = cr;
4542     dmapp->dec_fidrefs = B_FALSE;

4544     /*
4545      * When r_mapcnt returns to zero, arrange for the
4546      * callback to decrement n_fidrefs
4547      */
4548     mutex_enter(&np->r_statelock);
4549     np->r_mapcnt -= btopr(len);
4550     ASSERT(np->r_mapcnt >= 0);
4551     if (np->r_mapcnt == 0)
4552         dmapp->dec_fidrefs = B_TRUE;
4553     mutex_exit(&np->r_statelock);

4555     error = as_add_callback(as, smbfs_delmap_callback, dmapp,
4556                           AS_UNMAP_EVENT, addr, len, KM_SLEEP);
4557     if (error != 0) {
4558         /*
4559          * So sad, no callback is coming. Can't flush pages
4560          * in delmap (as locks). Just handle n_fidrefs.
4561          */
4562         cmn_err(CE_NOTE, "smbfs_delmap(%p) "
4563                "as_add_callback err=%d",
4564                (void *)vp, error);

4566         if (dmapp->dec_fidrefs) {
4567             struct smb_cred scred;

4569             (void) smbfs_rw_enter_sig(&np->r_lkserlock,
4570                                     RW_WRITER, 0);
4571             smb_credinit(&scred, dmapp->cr);

4573             smbfs_rele_fid(np, &scred);

4575             smb_credrele(&scred);
4576             smbfs_rw_exit(&np->r_lkserlock);
4577         }
4578         kmem_free(dmapp, sizeof (*dmapp));
4579     }

4581     return (0);
4582 }

4584 /*
4585  * Remove some pages from an mmap'd vnode. Flush any
4586  * dirty pages in the unmapped range.
4587  */
4588 /* ARGSUSED */
4589 static void
4590 smbfs_delmap_callback(struct as *as, void *arg, uint_t event)
4591 {
4592     vnode_t      *vp;
4593     smbnode_t    *np;
4594     smbmntinfo_t *smi;
4595     smbfs_delmap_args_t *dmapp = arg;

4597     vp = dmapp->vp;
4598     np = VTOSMB(vp);
4599     smi = VTOSMI(vp);

4601     /* Decrement r_mapcnt in smbfs_delmap */

4603     /*
4604      * Initiate a page flush and potential commit if there are
4605      * pages, the file system was not mounted readonly, the segment
4606      * was mapped shared, and the pages themselves were writeable.

```

```

4607      *
4608      * mark RDIRTY here, will be used to check if a file is dirty when
4609      * unmount smbfs
4610      */
4611      if (vn_has_cached_data(vp) && !vn_is_readonly(vp) &&
4612          dmapp->flags == MAP_SHARED && (dmapp->maxprot & PROT_WRITE)) {
4613          mutex_enter(&np->r_statelock);
4614          np->r_flags |= RDIRTY;
4615          mutex_exit(&np->r_statelock);
4616
4617          /*
4618           * Need to finish the putpage before we
4619           * close the 0tW FID needed for I/O.
4620           */
4621          (void) smbfs_putpage(vp, dmapp->off, dmapp->len, 0,
4622                              dmapp->cr, NULL);
4623      }
4624
4625      if ((np->r_flags & RDIRECTIO) || (smi->smi_flags & SMI_DIRECTIO))
4626          (void) smbfs_putpage(vp, dmapp->off, dmapp->len,
4627                              B_INVALID, dmapp->cr, NULL);
4628
4629      /*
4630       * If r_mapcnt went to zero, drop our FID ref now.
4631       * On the last fidref, this does an 0tW close.
4632       */
4633      if (dmapp->dec_fidrefs) {
4634          struct smb_cred scred;
4635
4636          (void) smbfs_rw_enter_sig(&np->r_lkserlock, RW_WRITER, 0);
4637          smb_credinit(&scred, dmapp->cr);
4638
4639          smbfs_rele_fid(np, &scred);
4640
4641          smb_credrele(&scred);
4642          smbfs_rw_exit(&np->r_lkserlock);
4643      }
4644
4645      (void) as_delete_callback(as, arg);
4646      kmem_free(dmapp, sizeof (*dmapp));
4647 }
4648
4649 /* No smbfs_pageio() or smbfs_dispose() ops. */
4650
4651 /* misc. **** */
4652
4653 /*
4654  * XXX
4655  * This op may need to support PSARC 2007/440, nbmand changes for CIFS Service.
4656  */
4657 static int
4658 smbfs_frlock(vnode_t *vp, int cmd, struct flock64 *bfp, int flag,
4659             offset_t offset, struct flk_callback *flk_cbp, cred_t *cr,
4660             caller_context_t *ct)
4661 {
4662     if (curproc->p_zone != VTOSMI(vp)->smi_zone_ref.zref_zone)
4663         return (EIO);
4664
4665     if (VTOSMI(vp)->smi_flags & SMI_LLOCK)
4666         return (fs_frlock(vp, cmd, bfp, flag, offset, flk_cbp, cr, ct));
4667     else
4668         return (ENOSYS);
4669 }
4670
4671 /*

```

```

4673      * Free storage space associated with the specified vnode. The portion
4674      * to be freed is specified by bfp->l_start and bfp->l_len (already
4675      * normalized to a "whence" of 0).
4676      *
4677      * Called by fcntl(fd, F_FREESP, lkp) for libc:ftruncate, etc.
4678      */
4679      /* ARGSUSED */
4680      static int
4681      smbfs_space(vnode_t *vp, int cmd, struct flock64 *bfp, int flag,
4682                  offset_t offset, cred_t *cr, caller_context_t *ct)
4683      {
4684          int error;
4685          smbmntinfo_t *smi;
4686
4687          smi = VTOSMI(vp);
4688
4689          if (curproc->p_zone != smi->smi_zone_ref.zref_zone)
4690              return (EIO);
4691
4692          if (smi->smi_flags & SMI_DEAD || vp->v_vfsp->vfs_flag & VFS_UNMOUNTED)
4693              return (EIO);
4694
4695          /* Caller (fcntl) has checked v_type */
4696          ASSERT(vp->v_type == VREG);
4697          if (cmd != F_FREESP)
4698              return (EINVAL);
4699
4700          /*
4701           * Like NFS3, no 32-bit offset checks here.
4702           * Our SMB layer takes care to return EFBIG
4703           * when it has to fallback to a 32-bit call.
4704           */
4705
4706          error = convoff(vp, bfp, 0, offset);
4707          if (!error) {
4708              ASSERT(bfp->l_start >= 0);
4709              if (bfp->l_len == 0) {
4710                  struct vattr va;
4711
4712                  /*
4713                   * ftruncate should not change the ctime and
4714                   * mtime if we truncate the file to its
4715                   * previous size.
4716                   */
4717                  va.va_mask = AT_SIZE;
4718                  error = smbfs_getattr(vp, &va, cr);
4719                  if (error || va.va_size == bfp->l_start)
4720                      return (error);
4721                  va.va_mask = AT_SIZE;
4722                  va.va_size = bfp->l_start;
4723                  error = smbfs_setattr(vp, &va, 0, cr);
4724                  /* SMBFS_VNEVENT... */
4725              } else
4726                  error = EINVAL;
4727          }
4728
4729          return (error);
4730 }
4731
4732 /*
4733  * XXX
4734  * This op may need to support PSARC 2007/440, nbmand changes for CIFS Service.
4735  */
4736 static int
4737 smbfs_realvp(vnode_t *vp, vnode_t **vpp, caller_context_t *ct)
4738 {
4739     return (ENOSYS);

```

```

4739 }

4742 /* ARGSUSED */
4743 static int
4744 smbfs_pathconf(vnode_t *vp, int cmd, ulong_t *valp, cred_t *cr,
4745 caller_context_t *ct)
4746 {
4747     vfs_t *vfs;
4748     smbmntinfo_t *smi;
4749     struct smb_share *ssp;

4751     vfs = vp->v_vfsp;
4752     smi = VFTOSMI(vfs);

4754     if (curproc->p_zone != smi->smi_zone_ref.zref_zone)
4755         return (EIO);

4757     if (smi->smi_flags & SMI_DEAD || vp->v_vfsp->vfs_flag & VFS_UNMOUNTED)
4758         return (EIO);

4760     switch (cmd) {
4761     case _PC_FILESIZEBITS:
4762         ssp = smi->smi_share;
4763         if (SSTOVC(ssp)->vc_sopt.sv_caps & SMB_CAP_LARGE_FILES)
4764             *valp = 64;
4765         else
4766             *valp = 32;
4767         break;

4769     case _PC_LINK_MAX:
4770         /* We only ever report one link to an object */
4771         *valp = 1;
4772         break;

4774     case _PC_ACL_ENABLED:
4775         /*
4776          * Always indicate that ACLs are enabled and
4777          * that we support ACE_T format, otherwise
4778          * libsec will ask for ACLNT_T format data
4779          * which we don't support.
4780          */
4781         *valp = _ACL_ACE_ENABLED;
4782         break;

4784     case _PC_SYMLINK_MAX: /* No symlinks until we do Unix extensions */
4785         *valp = 0;
4786         break;

4788     case _PC_XATTR_EXISTS:
4789         if (vfs->vfs_flag & VFS_XATTR) {
4790             *valp = smbfs_xa_exists(vp, cr);
4791             break;
4792         }
4793         return (EINVAL);

4795     case _PC_SATTR_ENABLED:
4796     case _PC_SATTR_EXISTS:
4797         *valp = 1;
4798         break;

4800     case _PC_TIMESTAMP_RESOLUTION:
4801         /*
4802          * Windows times are tenths of microseconds
4803          * (multiples of 100 nanoseconds).
4804          */

```

```

4805         *valp = 100L;
4806         break;

4808     default:
4809         return (fs_pathconf(vp, cmd, valp, cr, ct));
4810     }
4811     return (0);
4812 }
    unchanged_portion_omitted

4920 /*
4921  * Most unimplemented ops will return ENOSYS because of fs_nosys().
4922  * The only ops where that won't work are ACCESS (due to open(2)
4923  * failures) and ... (anything else left?)
4924  */
4925 const fs_operation_def_t smbfs_vnodeops_template[] = {
4926     VOPNAME_OPEN, { .vop_open = smbfs_open },
4927     VOPNAME_CLOSE, { .vop_close = smbfs_close },
4928     VOPNAME_READ, { .vop_read = smbfs_read },
4929     VOPNAME_WRITE, { .vop_write = smbfs_write },
4930     VOPNAME_IOCTL, { .vop_ioctl = smbfs_ioctl },
4931     VOPNAME_GETATTR, { .vop_getattr = smbfs_getattr },
4932     VOPNAME_SETATTR, { .vop_setattr = smbfs_setattr },
4933     VOPNAME_ACCESS, { .vop_access = smbfs_access },
4934     VOPNAME_LOOKUP, { .vop_lookup = smbfs_lookup },
4935     VOPNAME_CREATE, { .vop_create = smbfs_create },
4936     VOPNAME_REMOVE, { .vop_remove = smbfs_remove },
4937     VOPNAME_LINK, { .vop_link = smbfs_link },
4938     VOPNAME_RENAME, { .vop_rename = smbfs_rename },
4939     VOPNAME_MKDIR, { .vop_mkdir = smbfs_mkdir },
4940     VOPNAME_RMDIR, { .vop_rmdir = smbfs_rmdir },
4941     VOPNAME_READDIR, { .vop_readdir = smbfs_readdir },
4942     VOPNAME_SYMLINK, { .vop_symlink = smbfs_symlink },
4943     VOPNAME_READLINK, { .vop_readlink = smbfs_readlink },
4944     VOPNAME_FSYNC, { .vop_fsync = smbfs_fsync },
4945     VOPNAME_INACTIVE, { .vop_inactive = smbfs_inactive },
4946     VOPNAME_FID, { .vop_fid = smbfs_fid },
4947     VOPNAME_RWLOCK, { .vop_rwlock = smbfs_rwlock },
4948     VOPNAME_RWUNLOCK, { .vop_rwunlock = smbfs_rwunlock },
4949     VOPNAME_SEEK, { .vop_seek = smbfs_seek },
4950     VOPNAME_FRLOCK, { .vop_frlock = smbfs_frlock },
4951     VOPNAME_SPACE, { .vop_space = smbfs_space },
4952     VOPNAME_REALVP, { .vop_realvp = smbfs_realvp },
4953     VOPNAME_GETPAGE, { .vop_getpage = smbfs_getpage },
4954     VOPNAME_PUTPAGE, { .vop_putpage = smbfs_putpage },
4955     VOPNAME_MAP, { .vop_map = smbfs_map },
4956     VOPNAME_ADDMAP, { .vop_addmap = smbfs_addmap },
4957     VOPNAME_DELMAP, { .vop_delmap = smbfs_delmap },
4958     VOPNAME_DUMP, { .error = fs_nosys }, /* smbfs_dump, */
4959     VOPNAME_PATHCONF, { .vop_pathconf = smbfs_pathconf },
4960     VOPNAME_PAGEIO, { .error = fs_nosys }, /* smbfs_pageio, */
4961     VOPNAME_SETSECATTR, { .vop_setsecattr = smbfs_setsecattr },
4962     VOPNAME_GETSECATTR, { .vop_getsecattr = smbfs_getsecattr },
4963     VOPNAME_SHRLOCK, { .vop_shrlock = smbfs_shrlock },
4964     #ifdef SMBFS_VNEVENT
4965     VOPNAME_VNEVENT, { .vop_vnevent = fs_vnevent_support },
4966     #endif
4967     { NULL, NULL }
4968 };

```