

Gretl Command Reference



Gnu Regression, Econometrics and Time-series Library

Allin Cottrell
Department of Economics
Wake Forest University

Riccardo "Jack" Lucchetti
Dipartimento di Economia
Università Politecnica delle Marche

September, 2015

Permission is granted to copy, distribute and/or modify this document under the terms of the *GNU Free Documentation License*, Version 1.1 or any later version published by the Free Software Foundation (see <http://www.gnu.org/licenses/fdl.html>).

Contents

1	Gretl commands	1
1.1	Introduction	1
1.2	Commands	1
	add	1
	adf	2
	anova	3
	append	4
	ar	5
	ar1	5
	arbond	6
	arch	7
	arima	7
	biprobit	9
	boxplot	9
	break	10
	catch	10
	chow	10
	clear	11
	coeffsum	11
	coint	11
	coint2	12
	corr	13
	corrgm	13
	cusum	14
	data	14
	dataset	15
	debug	16
	delete	17
	diff	17
	difftest	17
	discrete	18
	dpanel	18
	dummify	19

duration	20
elif	20
else	20
end	20
endif	20
endloop	20
eqnprint	21
equation	21
estimate	21
eval	22
fcast	22
flush	23
foreign	24
fractint	24
freq	25
function	25
garch	26
genr	27
gmm	28
gnuplot	30
graphpg	31
hausman	32
heckit	33
help	33
hsk	33
hurst	34
if	34
include	35
info	35
install	35
intreg	36
join	36
kalman	37
kpss	37
labels	38
lad	38
lags	39
ldiff	39
leverage	39

levinlin	40
logistic	40
logit	41
logs	42
loop	42
mahal	42
makepkg	43
markers	43
meantest	44
mle	44
modeltab	45
modprint	45
modtest	46
mpols	47
negbin	47
nls	48
normtest	49
nulldata	49
ols	49
omit	50
open	51
orthdev	52
outfile	52
panel	53
pca	54
pergm	54
plot	55
poisson	56
print	56
printf	57
probit	58
pvalue	59
qlrtest	59
qqplot	60
quantreg	60
quit	61
rename	61
reset	61
restrict	61

rmplot	63
run	64
runs	64
scatters	64
sdiff	65
set	65
setinfo	69
setmiss	70
setobs	70
setopt	71
shell	72
smpl	73
spearman	74
sprintf	74
square	74
store	75
summary	76
system	76
tabprint	78
textplot	78
tobit	79
tsls	79
var	80
varlist	81
vartest	81
vecm	82
vif	83
wls	83
xcrrgm	83
xtab	84
1.3 Commands by topic	84
Estimation	84
Tests	85
Transformations	85
Statistics	85
Dataset	85
Graphs	86
Printing	86
Prediction	86

Programming	86
Utilities	86
1.4 Short-form command options	87
2 Gretl functions	88
2.1 Introduction	88
2.2 Accessors	88
\$ahat	88
\$aic	88
\$bic	88
\$chisq	88
\$coeff	88
\$command	89
\$compan	89
\$datatype	89
\$depvar	89
\$df	89
\$dpval	89
\$ec	89
\$error	90
\$ess	90
\$evals	90
\$fcast	90
\$fcerr	90
\$fevd	90
\$Fstat	91
\$gmmcrit	91
\$h	91
\$hausman	91
\$hqc	91
\$huge	91
\$jalpha	91
\$jbeta	91
\$jybeta	92
\$lang	92
\$lit	93
\$lnl	93
\$macheps	93
\$mnprobs	93

\$ncoeff	93
\$nobs	93
\$nvars	93
\$obsdate	93
\$obsmajor	94
\$obsmicro	94
\$obsminor	94
\$pd	94
\$pi	94
\$pvalue	94
\$qlrbreak	94
\$rho	95
\$rsq	95
\$sample	95
\$sargan	95
\$sigma	95
\$stderr	95
\$stopwatch	96
\$sysA	96
\$sysB	96
\$sysGamma	96
\$sysinfo	96
\$T	97
\$t1	97
\$t2	97
\$test	97
\$trsq	97
\$uhat	97
\$unit	98
\$vcv	98
\$vecGamma	98
\$version	98
\$vma	98
\$windows	98
\$xlist	98
\$xtxinv	99
\$yhat	99
\$ylist	99
2.3 Functions proper	99

abs	99
acos	99
acosh	99
aggregate	99
argname	100
asin	101
asinh	101
atan	101
atanh	101
atof	101
bessel	101
BFGSmax	102
bkfilt	102
boxcox	103
bread	103
bwfilt	103
bwrite	104
cdemean	104
cdf	104
cdiv	105
ceil	105
cholesky	105
chowlin	105
cmult	106
cnorm	106
colname	106
colnames	106
cols	106
corr	107
corrgm	107
cos	107
cosh	107
cov	107
critical	108
cum	108
curl	108
deseas	109
det	109
diag	110

diagcat	110
diff	110
digamma	110
dnorm	110
dsort	111
dummify	111
easterday	111
eigengen	111
eigensym	112
eigsolve	112
epochday	112
errmsg	112
exp	113
fcstats	113
fdjac	113
fft	114
ffti	114
filter	114
firstobs	115
fixname	115
floor	116
fracdiff	116
gammafun	116
genseries	116
getenv	117
getline	117
ghk	118
gini	118
ginv	119
halton	119
hdprod	119
hpfilt	120
I	120
imaxc	120
imaxr	120
imhof	120
iminc	120
iminr	121
inbundle	121

infnorm	121
inlist	121
int	121
inv	121
invcdf	122
invmills	122
invpd	122
irf	123
irr	123
isconst	123
isdiscrete	123
isdummy	124
isnan	124
isnull	124
isoconv	124
isodate	124
iwishart	125
jsonget	125
kdensity	125
kfilter	126
kpsscrit	126
ksimul	126
ksmooth	126
kurtosis	127
lags	127
lastobs	127
ldet	127
ldiff	127
lincomb	128
linearize	128
ljungbox	128
lngamma	128
log	128
log10	128
log2	129
loess	129
logistic	129
lower	129
lrvar	129

max	130
maxc	130
maxr	130
mcorr	130
mcov	130
mcovg	131
mean	131
meanc	131
meanr	131
median	131
mexp	131
min	132
minc	132
minr	132
missing	132
misszero	132
mlag	133
mnormal	133
mols	133
monthlen	133
movavg	134
mpols	134
mrang	134
mread	135
mreverse	135
mrls	135
mshape	136
msortby	136
muniform	136
mwrite	136
mxtab	137
nadarwat	137
nelem	138
ngetenv	138
nlines	138
nobs	138
normal	138
npv	139
NRmax	139

nullspace	139
obs	139
obslabel	140
obsnum	140
ok	140
onenorm	140
ones	141
orthdev	141
pdf	141
pergm	141
pexpand	142
pmax	142
pmean	142
pmin	142
pnobs	143
polroots	143
polyfit	143
princomp	143
prodc	144
prodr	144
psd	144
psdroot	144
pshrink	144
psum	145
pvalue	145
pxsum	145
qform	146
qlrpval	146
qnorm	146
qrdecomp	146
quadtable	147
quantile	147
randgen	148
randgen1	148
randint	149
rank	149
ranking	149
rcond	149
readfile	150

regsub	150
remove	150
replace	150
resample	151
round	151
rownames	152
rows	152
sd	152
sdc	152
sdiff	152
seasonals	153
selifc	153
selifr	153
seq	153
setnote	154
simann	154
sin	154
sinh	154
skewness	154
sort	155
sortby	155
sprintf	155
square	155
sqrt	155
sscanf	156
sst	156
stringify	157
strlen	157
strncmp	157
strsplit	157
strstr	157
strstrip	158
strsub	158
strvals	158
substr	158
sum	158
sumall	158
sumc	159
sumr	159

svd	159
tan	159
tanh	159
toepsolv	160
tolower	160
toupper	160
tr	160
transp	160
trimr	160
typestr	161
uniform	161
uniq	161
unvech	161
upper	161
urcpval	162
values	162
var	162
varname	162
varnum	163
varsimul	163
vec	163
vech	163
weekday	163
wmean	163
wsd	164
wvar	164
xmax	164
xmin	164
zeromiss	165
zeros	165
3 Operators	166
3.1 Precedence	166
3.2 Assignment	167
3.3 Increment and decrement	167
4 Comments in scripts	169
5 Options, arguments and path-searching	171
5.1 Invoking gretl	171

Contents	xiv
5.2 Preferences dialog	171
5.3 Invoking gretlcli	172
5.4 Path searching	173
MS Windows	173
6 Reserved Words	174
Bibliography	176

Chapter 1

Gretl commands

1.1 Introduction

The commands defined below may be executed interactively in the command-line client program or in the console window of the GUI program. They may also be placed in a “script” or batch file for non-interactive execution.

The following notational conventions are used below:

- A *typewriter font* is used for material that you would type directly, and also for internal names of variables.
- Terms in a *slanted font* are place-holders: you should substitute some specific replacement. For example, you might type *income* in place of the generic *xvar*.
- The construction [*arg*] means that the argument *arg* is optional: you may supply it or not (but in any case don't type the brackets).
- The phrase “estimation command” means a command that generates estimates for a given model, for example *ols*, *ar* or *wls*.

In general, each line of a command script should contain one and only one complete **gretl** command. There are, however, two means of continuing a long command from one line of input to another. First, if the last non-space character on a line is a backslash, this is taken as an indication that the command is continued on the following line. In addition, if the comma is a valid character in a given command (for instance, as a separator between function arguments, or as punctuation in the command `printf`) then a trailing comma also indicates continuation. To emphasize the point: a backslash may be inserted “arbitrarily” to indicate continuation, but a comma works in this capacity only if it is syntactically valid as part of the command.

1.2 Commands

add

Argument: *varlist*

Options: `--lm` (do an LM test, OLS only)
 `--quiet` (print only the basic test result)
 `--silent` (don't print anything)
 `--vcv` (print covariance matrix for augmented model)
 `--both` (IV estimation only, see below)

Examples: `add 5 7 9`
 `add xx yy zz --quiet`

Must be invoked after an estimation command. Performs a joint test for the addition of the specified variables to the last model, the results of which may be retrieved using the accessors `$test` and `$pvalue`.

By default an augmented version of the original model is estimated, including the variables in *varlist*. The test is a Wald test on the augmented model, which replaces the original as the “current model” for the purposes of, for example, retrieving the residuals as `$uhat` or doing further tests.

Alternatively, given the `--lm` option (available only for the models estimated via OLS), an LM test is performed. An auxiliary regression is run in which the dependent variable is the residual from the last model and the independent variables are those from the last model plus *varlist*. Under the null hypothesis that the added variables have no additional explanatory power, the sample size times the unadjusted R-squared from this regression is distributed as chi-square with degrees of freedom equal to the number of added regressors. In this case the original model is not replaced.

The `--both` option is specific to two-stage least squares: it specifies that the new variables should be added both to the list of regressors and the list of instruments, the default in this case being to add to the regressors only.

Menu path: Model window, /Tests/Add variables

adf

Arguments: *order varlist*

Options: `--nc` (test without a constant)
`--c` (with constant only)
`--ct` (with constant and trend)
`--ctt` (with constant, trend and trend squared)
`--seasonals` (include seasonal dummy variables)
`--gls` (de-mean or de-trend using GLS)
`--verbose` (print regression results)
`--quiet` (suppress printing of results)
`--difference` (use first difference of variable)
`--test-down[=criterion]` (automatic lag order)
`--perron-qu` (see below)

Examples: `adf 0 y`
`adf 2 y --nc --c --ct`
`adf 12 y --c --test-down`
 See also `jgm-1996.inp`

The options shown above and the discussion which follows pertain to the use of the `adf` command with regular time series data. For use of this command with panel data please see below.

Computes a set of Dickey-Fuller tests on each of the the listed variables, the null hypothesis being that the variable in question has a unit root. (But if the `--difference` flag is given, the first difference of the variable is taken prior to testing, and the discussion below must be taken as referring to the transformed variable.)

By default, two variants of the test are shown: one based on a regression containing a constant and one using a constant and linear trend. You can control the variants that are presented by specifying one or more of the option flags.

The `--gls` option can be used in conjunction with one or other of the flags `--c` and `--ct` (the model with constant, or model with constant and trend). The effect of this option is that the de-meaning or de-trending of the variable to be tested is done using the GLS procedure suggested by [Elliott et al. \(1996\)](#), which gives a test of greater power than the standard Dickey-Fuller approach. This option is not compatible with `--nc`, `--ctt` or `--seasonals`.

In all cases the dependent variable is the first difference of the specified variable, y , and the key independent variable is the first lag of y . The model is constructed so that the coefficient on lagged

y equals the root in question minus 1. For example, the model with a constant may be written as

$$(1 - L)y_t = \beta_0 + (\alpha - 1)y_{t-1} + \epsilon_t$$

Under the null hypothesis of a unit root the coefficient on lagged y equals zero; under the alternative that y is stationary this coefficient is negative.

If the *order* argument (henceforth, k) is greater than 0, then k lags of the dependent variable are included on the right-hand side of the test regressions. If the order is given as -1 , k is set following the recommendation of [Schwert \(1989\)](#), namely the integer part of $12(T/100)^{0.25}$, where T is the sample size. In either case, however, if the `--test-down` option is given then k is taken as the *maximum* lag and the actual lag order used is obtained by testing down. The criterion for testing down can be selected using the option parameter, which should be one of AIC, BIC or `tstat`; AIC is the default.

When testing down via AIC or BIC is called for, the final lag order for the ADF equation is that which optimizes the chosen information criterion (Akaike or Schwarz Bayesian). The exact procedure depends on whether or not the `--gls` option is given: when GLS detrending is specified, AIC and BIC are the “modified” versions described in [Ng and Perron \(2001\)](#), otherwise they are the standard versions. In the GLS case a refinement is available: if the additional option `--perron-qu` is given, the modified information criteria are computed according to the revised method recommended by [Perron and Qu \(2007\)](#).

When testing down via the t -statistic method is called for, the procedure is as follows:

1. Estimate the Dickey-Fuller regression with k lags of the dependent variable.
2. Is the last lag significant? If so, execute the test with lag order k . Otherwise, let $k = k - 1$; if k equals 0, execute the test with lag order 0, else go to step 1.

In the context of step 2 above, “significant” means that the t -statistic for the last lag has an asymptotic two-sided p -value, against the normal distribution, of 0.10 or less.

P -values for the Dickey-Fuller tests are based on [MacKinnon \(1996\)](#). The relevant code is included by kind permission of the author. In the case of the test with linear trend using GLS these P -values are not applicable; critical values from Table 1 in [Elliott et al. \(1996\)](#) are shown instead.

Panel data

When the `adf` command is used with panel data, to produce a panel unit root test, the applicable options and the results shown are somewhat different.

First, while you may give a list of variables for testing in the regular time-series case, with panel data only one variable may be tested per command. Second, the options governing the inclusion of deterministic terms become mutually exclusive: you must choose between no-constant, constant only, and constant plus trend; the default is constant only. In addition, the `--seasonals` option is not available. Third, the `--verbose` option has a different meaning: it produces a brief account of the test for each individual time series (the default being to show only the overall result).

The overall test (null hypothesis: the series in question has a unit root for all the panel units) is calculated in one or both of two ways: using the method of [Im et al. \(2003\)](#) or that of [Choi \(2001\)](#).

Menu path: /Variable/Unit root tests/Augmented Dickey-Fuller test

anova

Arguments: *response treatment* [*block*]

Option: `--quiet` (don't print results)

Analysis of Variance: *response* is a series measuring some effect of interest and *treatment* must be a discrete variable that codes for two or more types of treatment (or non-treatment). For two-way

ANOVA, the *block* variable (which should also be discrete) codes for the values of some control variable.

Unless the `--quiet` option is given, this command prints a table showing the sums of squares and mean squares along with an *F*-test. The *F*-test and its P-value can be retrieved using the accessors `$test` and `$pvalue` respectively.

The null hypothesis for the *F*-test is that the mean response is invariant with respect to the treatment type, or in words that the treatment has no effect. Strictly speaking, the test is valid only if the variance of the response is the same for all treatment types.

Note that the results shown by this command are in fact a subset of the information given by the following procedure, which is easily implemented in gretl. Create a set of dummy variables coding for all but one of the treatment types. For two-way ANOVA, in addition create a set of dummies coding for all but one of the “blocks”. Then regress *response* on a constant and the dummies using `ols`. For a one-way design the ANOVA table is printed via the `--anova` option to `ols`. In the two-way case the relevant *F*-test is found by using the `omit` command. For example (assuming *y* is the response, *xt* codes for the treatment, and *xb* codes for blocks):

```
# one-way
list dxt = dummify(xt)
ols y 0 dxt --anova
# two-way
list dxb = dummify(xb)
ols y 0 dxt dxb
# test joint significance of dxt
omit dxt --quiet
```

Menu path: /Model/Other linear models/ANOVA

append

Argument: *filename*

Options: `--time-series` (see below)

`--update-overlap` (see below)

See below for additional specialized options

Opens a data file and appends the content to the current dataset, if the new data are compatible. The program will try to detect the format of the data file (native, plain text, CSV, Gnumeric, Excel, etc.).

The appended data may take the form of either additional observations on variables already present in the dataset, or new variables. In the case of adding variables, compatibility requires either (a) that the number of observations for the new data equals that for the current data, or (b) that the new data carries clear observation information so that gretl can work out how to place the values.

A special feature is supported for appending to a panel dataset. Let n denote the number of cross-sectional units in the panel, T denote the number of time periods, and m denote the number of observations for the new data. If $m = n$ the new data are taken to be time-invariant, and are copied into place for each time period. On the other hand, if $m = T$ the data are treated as non-varying across the panel units, and are copied into place for each unit. If the panel is “square”, and m equals both n and T , an ambiguity arises. The default in this case is to treat the new data as time-invariant, but you can force gretl to treat the new data as time series via the `--time-series` option. (This option is ignored in all other cases.)

When a data file is selected for appending, there may be an area of overlap with the existing dataset; that is, one or more series may have one or more observations in common across the two sources. If the option `--update-overlap` is given, the append operation will replace any overlapping obser-

variations with the values from the selected data file, otherwise the values currently in place will be unaffected.

The additional specialized options `--sheet`, `--coloffset`, `--rowoffset` and `--fixed-cols` work in the same way as with `open`; see that command for explanations.

See also `join` for more sophisticated handling of multiple data sources.

Menu path: /File/Append data

ar

Arguments: *lags* ; *depvar indepvars*

Option: `--vcv` (print covariance matrix)

Example: `ar 1 3 4 ; y 0 x1 x2 x3`

Computes parameter estimates using the generalized Cochrane-Orcutt iterative procedure; see Section 9.5 of [Ramanathan \(2002\)](#). Iteration is terminated when successive error sums of squares do not differ by more than 0.005 percent or after 20 iterations.

lags is a list of lags in the residuals, terminated by a semicolon. In the above example, the error term is specified as

$$u_t = \rho_1 u_{t-1} + \rho_3 u_{t-3} + \rho_4 u_{t-4} + e_t$$

Menu path: /Model/Time series/Autoregressive estimation

ar1

Arguments: *depvar indepvars*

Options: `--hilu` (use Hildreth-Lu procedure)

`--pwe` (use Prais-Winsten estimator)

`--vcv` (print covariance matrix)

`--no-corc` (do not fine-tune results with Cochrane-Orcutt)

`--loose` (use looser convergence criterion)

Examples: `ar1 1 0 2 4 6 7`

`ar1 y 0 xlist --pwe`

`ar1 y 0 xlist --hilu --no-corc`

Computes feasible GLS estimates for a model in which the error term is assumed to follow a first-order autoregressive process.

The default method is the Cochrane-Orcutt iterative procedure; see for example section 9.4 of [Ramanathan \(2002\)](#). The criterion for convergence is that successive estimates of the autocorrelation coefficient do not differ by more than 1e-6, or if the `--loose` option is given, by more than 0.001. If this is not achieved within 100 iterations an error is flagged.

If the `--pwe` option is given, the Prais-Winsten estimator is used. This involves an iteration similar to Cochrane-Orcutt; the difference is that while Cochrane-Orcutt discards the first observation, Prais-Winsten makes use of it. See, for example, Chapter 13 of [Greene \(2000\)](#) for details.

If the `--hilu` option is given, the Hildreth-Lu search procedure is used. The results are then fine-tuned using the Cochrane-Orcutt method, unless the `--no-corc` flag is specified. The `--no-corc` option is ignored for estimators other than Hildreth-Lu.

Menu path: /Model/Time series/AR(1)

arbond

Argument: *p* [*q*] ; *depvar indepvars* [; *instruments*]

Options: --quiet (don't show estimated model)
 --vcv (print covariance matrix)
 --two-step (perform 2-step GMM estimation)
 --time-dummies (add time dummy variables)
 --asymptotic (uncorrected asymptotic standard errors)

Examples: arbond 2 ; y Dx1 Dx2
 arbond 2 5 ; y Dx1 Dx2 ; Dx1
 arbond 1 ; y Dx1 Dx2 ; Dx1 GMM(x2,2,3)
 See also arbond91.inp

Carries out estimation of dynamic panel data models (that is, panel models including one or more lags of the dependent variable) using the GMM-DIF method set out by [Arellano and Bond \(1991\)](#). Please see [dpanel](#) for an updated and more flexible version of this command which handles GMM-SYS as well as GMM-DIF.

The parameter *p* represents the order of the autoregression for the dependent variable. The optional parameter *q* indicates the maximum lag of the level of the dependent variable to be used as an instrument. If this argument is omitted, or given as 0, all available lags are used.

The dependent variable should be given in levels form; it will be automatically differenced (since this estimator uses differencing to cancel out the individual effects). The independent variables are not automatically differenced; if you want to use differences (which will generally be the case for ordinary quantitative variables, though perhaps not for, say, time dummy variables) you should create the differences first then specify these as the regressors.

The last (optional) field in the command is for specifying instruments. If no instruments are given, it is assumed that all the independent variables are strictly exogenous. If you specify any instruments, you should include in the list any strictly exogenous independent variables. For predetermined regressors, you can use the GMM function to include a specified range of lags in block-diagonal fashion. This is illustrated in the third example above. The first argument to GMM is the name of the variable in question, the second is the minimum lag to be used as an instrument, and the third is the maximum lag. If the third argument is given as 0, all available lags are used.

By default the results of 1-step estimation are reported (with robust standard errors). You may select 2-step estimation as an option. In both cases tests for autocorrelation of orders 1 and 2 are provided, as well as the Sargan overidentification test and a Wald test for the joint significance of the regressors. Note that in this differenced model first-order autocorrelation is not a threat to the validity of the model, but second-order autocorrelation violates the maintained statistical assumptions.

In the case of 2-step estimation, standard errors are by default computed using the finite-sample correction suggested by [Windmeijer \(2005\)](#). The standard asymptotic standard errors associated with the 2-step estimator are generally reckoned to be an unreliable guide to inference, but if for some reason you want to see them you can use the --asymptotic option to turn off the Windmeijer correction.

If the --time-dummies option is given, a set of time dummy variables is added to the specified regressors. The number of dummies is one less than the maximum number of periods used in estimation, to avoid perfect collinearity with the constant. The dummies are entered in levels; if you wish to use time dummies in first-differenced form, you will have to define and add these variables manually.

arch

Arguments: *order depvar indepvars*

Example: `arch 4 y 0 x1 x2 x3`

This command is retained at present for backward compatibility, but you are better off using the maximum likelihood estimator offered by the [garch](#) command; for a plain ARCH model, set the first GARCH parameter to 0.

Estimates the given model specification allowing for ARCH (Autoregressive Conditional Heteroskedasticity). The model is first estimated via OLS, then an auxiliary regression is run, in which the squared residual from the first stage is regressed on its own lagged values. The final step is weighted least squares estimation, using as weights the reciprocals of the fitted error variances from the auxiliary regression. (If the predicted variance of any observation in the auxiliary regression is not positive, then the corresponding squared residual is used instead).

The α values displayed below the coefficients are the estimated parameters of the ARCH process from the auxiliary regression.

See also [garch](#) and [modtest](#) (the `--arch` option).

Menu path: /Model/Time series/ARCH

arima

Arguments: *p d q [; P D Q] ; depvar [indepvars]*

Options: `--verbose` (print details of iterations)
`--vcv` (print covariance matrix)
`--hessian` (see below)
`--opg` (see below)
`--nc` (do not include a constant)
`--conditional` (use conditional maximum likelihood)
`--x-12-arima` (use X-12-ARIMA for estimation)
`--lbfgs` (use L-BFGS-B maximizer)
`--y-diff-only` (ARIMAX special, see below)
`--save-ehat` (see below)

Examples: `arima 1 0 2 ; y`
`arima 2 0 2 ; y 0 x1 x2 --verbose`
`arima 0 1 1 ; 0 1 1 ; y --nc`

If no *indepvars* list is given, estimates a univariate ARIMA (Autoregressive, Integrated, Moving Average) model. The values *p*, *d* and *q* represent the autoregressive (AR) order, the differencing order, and the moving average (MA) order respectively. These values may be given in numerical form, or as the names of pre-existing scalar variables. A *d* value of 1, for instance, means that the first difference of the dependent variable should be taken before estimating the ARMA parameters.

If you wish to include only specific AR or MA lags in the model (as opposed to all lags up to a given order) you can substitute for *p* and/or *q* either (a) the name of a pre-defined matrix containing a set of integer values or (b) an expression such as {1,4}; that is, a set of lags separated by commas and enclosed in braces.

The optional integer values *P*, *D* and *Q* represent the seasonal AR order, the order for seasonal differencing, and the seasonal MA order, respectively. These are applicable only if the data have a frequency greater than 1 (for example, quarterly or monthly data). These orders may be given in numerical form or as scalar variables.

In the univariate case the default is to include an intercept in the model but this can be suppressed with the `--nc` flag. If *indepvars* are added, the model becomes ARMAX; in this case the constant

should be included explicitly if you want an intercept (as in the second example above).

An alternative form of syntax is available for this command: if you do not want to apply differencing (either seasonal or non-seasonal), you may omit the d and D fields altogether, rather than explicitly entering 0. In addition, `arma` is a synonym or alias for `arima`. Thus for example the following command is a valid way to specify an ARMA(2, 1) model:

```
arma 2 1 ; y
```

The default is to use the “native” gretl ARMA functionality, with estimation by exact ML using the Kalman filter; estimation via conditional ML is available as an option. (If X-12-ARIMA is installed you have the option of using it instead of native code.) For details regarding these options, please see the *Gretl User's Guide*.

When the native exact ML code is used, estimated standard errors are by default based on a numerical approximation to the (negative inverse of) the Hessian, with a fallback to the outer product of the gradient (OPG) if calculation of the numerical Hessian should fail. Two (mutually exclusive) option flags can be used to force the issue: the `--opg` option forces use of the OPG method, with no attempt to compute the Hessian, while the `--hessian` flag disables the fallback to OPG. Note that failure of the numerical Hessian computation is generally an indicator of a misspecified model.

The option `--lbfgs` is specific to estimation using native ARMA code and exact ML: it calls for use of the “limited memory” L-BFGS-B algorithm in place of the regular BFGS maximizer. This may help in some instances where convergence is difficult to achieve.

The option `--y-diff-only` is specific to estimation of ARIMAX models (models with a non-zero order of integration and including exogenous regressors), and applies only when gretl’s native exact ML is used. For such models the default behavior is to difference both the dependent variable and the regressors, but when this option is specified only the dependent variable is differenced, the regressors remaining in level form.

The option `--save-ehat` is applicable only when using native exact ML estimation. The effect is to make available a vector holding the optimal estimate as of period t of the t -dated disturbance or innovation: this can be retrieved via the accessor `$ehat`. These values differ from the residual series (`$uhat`), which holds the one-step ahead forecast errors.

The AIC value given in connection with ARIMA models is calculated according to the definition used in X-12-ARIMA, namely

$$\text{AIC} = -2\ell + 2k$$

where ℓ is the log-likelihood and k is the total number of parameters estimated. Note that X-12-ARIMA does not produce information criteria such as AIC when estimation is by conditional ML.

The AR and MA roots shown in connection with ARMA estimation are based on the following representation of an ARMA(p, q) process:

$$(1 - \phi_1 L - \phi_2 L^2 - \dots - \phi_p L^p)Y = c + (1 + \theta_1 L + \theta_2 L^2 + \dots + \theta_q L^q)\varepsilon_t$$

The AR roots are therefore the solutions to

$$1 - \phi_1 z - \phi_2 z^2 - \dots - \phi_p z^p = 0$$

and stability requires that these roots lie outside the unit circle.

The “frequency” figure printed in connection with the AR and MA roots is the λ value that solves $z = r e^{i2\pi\lambda}$, where z is the root in question and r is its modulus.

Menu path: /Model/Time series/ARIMA

Other access: Main window pop-up menu (single selection)

biprobit

Arguments: *depvar1 depvar2 indepvars1 [; indepvars2]*

Options: --vcv (print covariance matrix)
 --robust (robust standard errors)
 --cluster=*clustvar* (see [logit](#) for explanation)
 --opg (see below)
 --save-xbeta (see below)
 --verbose (print extra information)

Examples: biprobit y1 y2 0 x1 x2
 biprobit y1 y2 0 x11 x12 ; 0 x21 x22
 See also `biprobit.inp`

Estimates a bivariate probit model, using the Newton-Raphson method to maximize the likelihood.

The argument list starts with the two (binary) dependent variables, followed by a list of regressors. If a second list is given, separated by a semicolon, this is interpreted as a set of regressors specific to the second equation, with *indepvars1* being specific to the first equation; otherwise *indepvars1* is taken to represent a common set of regressors.

By default, standard errors are computed using a numerical approximation to the Hessian at convergence. But if the `--opg` option is given the covariance matrix is based on the Outer Product of the Gradient (OPG), or if the `--robust` option is given QML standard errors are calculated, using a “sandwich” of the inverse of the Hessian and the OPG.

After successful estimation, the accessor `$uhat` retrieves a matrix with two columns holding the generalized residuals for the two equations; that is, the expected values of the disturbances conditional on the observed outcomes and covariates. By default `$yhat` retrieves a matrix with four columns, holding the estimated probabilities of the four possible joint outcomes for (y_1, y_2) , in the order (1,1), (1,0), (0,1), (0,0). Alternatively, if the option `--save-xbeta` is given, `$yhat` has two columns and holds the values of the index functions for the respective equations.

The output includes a likelihood ratio test of the null hypothesis that the disturbances in the two equations are uncorrelated.

boxplot

Argument: *varlist*

Options: --notches (show 90 percent interval for median)
 --factorized (see below)
 --panel (see below)
 --matrix=*name* (plot columns of named matrix)
 --output=*filename* (send output to specified file)

These plots display the distribution of a variable. The central box encloses the middle 50 percent of the data, i.e. it is bounded by the first and third quartiles. The “whiskers” extend from each end of the box for a range equal to 1.5 times the interquartile range. Observations outside that range are considered outliers and represented via dots. A line is drawn across the box at the median. A “+” sign is used to indicate the mean. If the option of showing a confidence interval for the median is selected, this is computed via the bootstrap method and shown in the form of dashed horizontal lines above and/or below the median.

The `--factorized` option allows you to examine the distribution of a chosen variable conditional on the value of some discrete factor. For example, if a data set contains wages and a gender dummy variable you can select the wage variable as the target and gender as the factor, to see side-by-side boxplots of male and female wages, as in

```
boxplot wage gender --factorized
```

Note that in this case you must specify exactly two variables, with the factor given second.

If the current data set is a panel, and just one variable is specified, the `--panel` option produces a series of side-by-side boxplots, one for each panel “unit” or group.

Generally, the argument *varlist* is required, and refers to one or more series in the current dataset (given either by name or ID number). But if a named matrix is supplied via the `--matrix` option this argument becomes optional: by default a plot is drawn for each column of the specified matrix.

Gretl’s boxplots are generated using gnuplot, and it is possible to specify the plot more fully by appending additional gnuplot commands, enclosed in braces. For details, please see the help for the [gnuplot](#) command.

In interactive mode the result is displayed immediately. In batch mode the default behavior is that a gnuplot command file is written in the user’s working directory, with a name on the pattern `gpttmpN.plt`, starting with `N = 01`. The actual plots may be generated later using gnuplot (under MS Windows, `wgnuplot`). This behavior can be modified by use of the `--output=filename` option. For details, please see the [gnuplot](#) command.

Menu path: /View/Graph specified vars/Boxplots

break

Break out of a loop. This command can be used only within a loop; it causes command execution to break out of the current (innermost) loop. See also [loop](#).

catch

Syntax: `catch command`

This is not a command in its own right but can be used as a prefix to most regular commands: the effect is to prevent termination of a script if an error occurs in executing the command. If an error does occur, this is registered in an internal error code which can be accessed as `$error` (a zero value indicates success). The value of `$error` should always be checked immediately after using `catch`, and appropriate action taken if the command failed.

The `catch` keyword cannot be used before `if`, `elif` or `endif`.

chow

Variants: `chow obs`

`chow dummyvar --dummy`

Options: `--dummy` (use a pre-existing dummy variable)

`--quiet` (don’t print estimates for augmented model)

`--limit-to=list` (limit test to subset of regressors)

Examples: `chow 25`

`chow 1988:1`

`chow female --dummy`

Must follow an OLS regression. If an observation number or date is given, provides a test for the null hypothesis of no structural break at the given split point. The procedure is to create a dummy variable which equals 1 from the split point specified by *obs* to the end of the sample, 0 otherwise, and also interaction terms between this dummy and the original regressors. If a dummy variable is given, tests the null hypothesis of structural homogeneity with respect to that dummy. Again, interaction terms are added. In either case an augmented regression is run including the additional terms.

By default an F statistic is calculated, taking the augmented regression as the unrestricted model and the original as the restricted. But if the original model used a robust estimator for the covariance matrix, the test statistic is a Wald chi-square value based on a robust estimator of the covariance matrix for the augmented regression.

The `--limit-to` option can be used to limit the set of interactions with the split dummy variable to a subset of the original regressors. The parameter for this option must be a named list, all of whose members are among the original regressors. The list should not include the constant.

Menu path: Model window, /Tests/Chow test

clear

Option: `--dataset` (clear dataset only)

With no options, clears all saved objects, including the current dataset if any, out of memory. Note that opening a new dataset, or using the `nulldata` command to create an empty dataset, also has this effect, so use of `clear` is not usually necessary.

If the `--dataset` option is given, then only the dataset is cleared (plus any named lists of series); other saved objects such as named matrices and scalars are preserved.

coeffsum

Argument: *varlist*

Example: `coeffsum xt xt_1 xr_2
restrict.inp`

Must follow a regression. Calculates the sum of the coefficients on the variables in *varlist*. Prints this sum along with its standard error and the p-value for the null hypothesis that the sum is zero.

Note the difference between this and [omit](#), which tests the null hypothesis that the coefficients on a specified subset of independent variables are *all* equal to zero.

Menu path: Model window, /Tests/Sum of coefficients

coint

Arguments: *order depvar indepvars*

Options: `--nc` (do not include a constant)
`--ct` (include constant and trend)
`--ctt` (include constant and quadratic trend)
`--skip-df` (no DF tests on individual variables)
`--test-down[=criterion]` (automatic lag order)
`--verbose` (print extra details of regressions)
`--silent` (don't print anything)

Examples: `coint 4 y x1 x2`
`coint 0 y x1 x2 --ct --skip-df`

The [Engle and Granger \(1987\)](#) cointegration test. The default procedure is: (1) carry out Dickey-Fuller tests on the null hypothesis that each of the variables listed has a unit root; (2) estimate the cointegrating regression; and (3) run a DF test on the residuals from the cointegrating regression. If the `--skip-df` flag is given, step (1) is omitted.

If the specified lag order is positive all the Dickey-Fuller tests use that order, with this qualification: if the `--test-down` option is given, the given value is taken as the maximum and the actual lag order used in each case is obtained by testing down. See the [adf](#) command for details of this procedure.

By default, the cointegrating regression contains a constant. If you wish to suppress the constant, add the `--nc` flag. If you wish to augment the list of deterministic terms in the cointegrating regression with a linear or quadratic trend, add the `--ct` or `--ctt` flag. These option flags are mutually exclusive.

P-values for this test are based on [MacKinnon \(1996\)](#). The relevant code is included by kind permission of the author.

Menu path: /Model/Time series/Cointegration test/Engle-Granger

coint2

Arguments: `order ylist [ ; xlist ] [ ; rxlist ]`

Options: `--nc` (no constant)
`--rc` (restricted constant)
`--uc` (unrestricted constant)
`--crt` (constant and restricted trend)
`--ct` (constant and unrestricted trend)
`--seasonals` (include centered seasonal dummies)
`--asy` (record asymptotic p-values)
`--quiet` (print just the tests)
`--silent` (don't print anything)
`--verbose` (print details of auxiliary regressions)

Examples: `coint2 2 y x`
`coint2 4 y x1 x2 --verbose`
`coint2 3 y x1 x2 --rc`

Carries out the Johansen test for cointegration among the variables in *ylist* for the given lag order. For details of this test see the *Gretl User's Guide* or [Hamilton \(1994\)](#), Chapter 20. *P*-values are computed via Doornik's gamma approximation ([Doornik, 1998](#)). Two sets of *p*-values are shown for the trace test, straight asymptotic values and values adjusted for the sample size. By default the `$pvalue` accessor gets the adjusted variant, but the `--asy` flag may be used to record the asymptotic values instead.

The inclusion of deterministic terms in the model is controlled by the option flags. The default if no option is specified is to include an "unrestricted constant", which allows for the presence of a non-zero intercept in the cointegrating relations as well as a trend in the levels of the endogenous variables. In the literature stemming from the work of Johansen (see for example his 1995 book) this is often referred to as "case 3". The first four options given above, which are mutually exclusive, produce cases 1, 2, 4 and 5 respectively. The meaning of these cases and the criteria for selecting a case are explained in the *Gretl User's Guide*.

The optional lists *xlist* and *rxlist* allow you to control for specified exogenous variables: these enter the system either unrestrictedly (*xlist*) or restricted to the cointegration space (*rxlist*). These lists are separated from *ylist* and from each other by semicolons.

The `--seasonals` option, which may be combined with any of the other options, specifies the inclusion of a set of centered seasonal dummy variables. This option is available only for quarterly or monthly data.

The following table is offered as a guide to the interpretation of the results shown for the test, for the 3-variable case. H_0 denotes the null hypothesis, H_1 the alternative hypothesis, and *c* the number of cointegrating relations.

Rank	Trace test		λ -max test	
	H_0	H_1	H_0	H_1
0	$c = 0$	$c = 3$	$c = 0$	$c = 1$
1	$c = 1$	$c = 3$	$c = 1$	$c = 2$
2	$c = 2$	$c = 3$	$c = 2$	$c = 3$

See also the [vecm](#) command.

Menu path: /Model/Time series/Cointegration test/Johansen

corr

Argument: [*varlist*]

Options: --uniform (ensure uniform sample)
 --spearman (Spearman's rho)
 --kendall (Kendall's tau)
 --verbose (print rankings)

Examples: corr y x1 x2 x3
 corr ylist --uniform
 corr x y --spearman

By default, prints the pairwise correlation coefficients (Pearson's product-moment correlation) for the variables in *varlist*, or for all variables in the data set if *varlist* is not given. The standard behavior is to use all available observations for computing each pairwise coefficient, but if the --uniform option is given the sample is limited (if necessary) so that the same set of observations is used for all the coefficients. This option has an effect only if there are differing numbers of missing values for the variables used.

The (mutually exclusive) options --spearman and --kendall produce, respectively, Spearman's rank correlation rho and Kendall's rank correlation tau in place of the default Pearson coefficient. When either of these options is given, *varlist* should contain just two variables.

When a rank correlation is computed, the --verbose option can be used to print the original and ranked data (otherwise this option is ignored).

Menu path: /View/Correlation matrix

Other access: Main window pop-up menu (multiple selection)

corrgm

Arguments: *series* [*order*]

Options: --bartlett (use Bartlett standard errors)
 --plot=*mode-or-filename* (see below)

Example: corrgm x 12

Prints the values of the autocorrelation function (ACF) for *series*, which may be specified by name or number. The values are defined as $\hat{\rho}(u_t, u_{t-s})$, where u_t is the t^{th} observation of the variable u and s denotes the number of lags.

The partial autocorrelations (PACF, calculated using the Durbin-Levinson algorithm) are also shown: these are net of the effects of intervening lags. In addition the Ljung-Box Q statistic is printed. This may be used to test the null hypothesis that the series is "white noise"; it is asymptotically distributed as chi-square with degrees of freedom equal to the number of lags used.

Asterisks are used to indicate statistical significance of the individual autocorrelations. By default this is assessed using a standard error of one over the square root of the sample size, but if the

`--bartlett` option is given then Bartlett standard errors are used for the ACF. This option also governs the confidence band drawn in the ACF plot, if applicable.

If an *order* value is specified the length of the correlogram is limited to at most that number of lags, otherwise the length is determined automatically, as a function of the frequency of the data and the number of observations.

By default, a plot of the correlogram is produced: a gnuplot graph in interactive mode or an ASCII graphic in batch mode. This can be adjusted via the `--plot` option. The acceptable parameters to this option are `none` (to suppress the plot); `ascii` (to produce a text graphic even when in interactive mode); `display` (to produce a gnuplot graph even when in batch mode); or a file name. The effect of providing a file name is as described for the `--output` option of the [gnuplot](#) command.

Upon successful completion, the accessors `$test` and `$pvalue` contain the corresponding figures of the Ljung-Box test for the maximum order displayed. Note that if you just want to compute the *Q* statistic, you'll probably want to use the [ljungbox](#) function instead.

Menu path: /Variable/Correlogram

Other access: Main window pop-up menu (single selection)

cusum

Options: `--squares` (perform the CUSUMSQ test)
`--quiet` (just print the Harvey-Collier test)

Must follow the estimation of a model via OLS. Performs the CUSUM test—or if the `--squares` option is given, the CUSUMSQ test—for parameter stability. A series of one-step ahead forecast errors is obtained by running a series of regressions: the first regression uses the first *k* observations and is used to generate a prediction of the dependent variable at observation *k* + 1; the second uses the first *k* + 1 observations and generates a prediction for observation *k* + 2, and so on (where *k* is the number of parameters in the original model).

The cumulated sum of the scaled forecast errors, or the squares of these errors, is printed and graphed. The null hypothesis of parameter stability is rejected at the 5 percent significance level if the cumulated sum strays outside of the 95 percent confidence band.

In the case of the CUSUM test, the Harvey-Collier *t*-statistic for testing the null hypothesis of parameter stability is also printed. See Greene's *Econometric Analysis* for details. For the CUSUMSQ test, the 95 percent confidence band is calculated using the algorithm given in [Edgerton and Wells \(1994\)](#).

Menu path: Model window, /Tests/CUSUM(SQ)

data

Argument: *varlist*
Options: `--compact=method` (specify compaction method)
`--interpolate` (do interpolation for low-frequency data)
`--quiet` (don't report results except on error)

Reads the variables in *varlist* from a database (gretl, RATS 4.0 or PcGive), which must have been opened previously using the [open](#) command. The data frequency and sample range may be established via the [setobs](#) and [smp1](#) commands prior to using this command. Here is a full example:

```
open macrodat.rat
setobs 4 1959:1
smp1 ; 1999:4
data GDP_JP GDP_UK
```

The commands above open a database named `macrodat.rat`, establish a quarterly data set starting in the first quarter of 1959 and ending in the fourth quarter of 1999, and then import the series named `GDP_JP` and `GDP_UK`.

If `setobs` and `smp1` are not specified in this way, the data frequency and sample range are set using the first variable read from the database.

If the series to be read are of higher frequency than the working dataset, you may specify a compaction method as below:

```
data LHUR PUNEW --compact=average
```

The four available compaction methods are “average” (takes the mean of the high frequency observations), “last” (uses the last observation), “first” and “sum”. If no method is specified, the default is to use the average.

If the series to be read are of lower frequency than the working dataset, the default is to repeat the values of the added data as required, but the `--interpolate` option can be used to request interpolation using the method of [Chow and Lin \(1971\)](#): the regressors are a constant and quadratic trend and an AR(1) disturbance process is assumed. Note, however, that this option is available only for conversion from quarterly data to monthly or annual data to quarterly.

In the case of native gretl databases (only), the “glob” characters `*` and `?` can be used in *varlist* to import series that match the given pattern. For example, the following will import all series in the database whose names begin with `cpi`:

```
data cpi*
```

Menu path: /File/Databases

dataset

Arguments: *keyword parameters*

Examples: `dataset addobs 24`
`dataset insobs 10`
`dataset compact 1`
`dataset compact 4 last`
`dataset expand interp`
`dataset transpose`
`dataset sortby x1`
`dataset resample 500`
`dataset renumber x 4`
`dataset clear`

Performs various operations on the data set as a whole, depending on the given *keyword*, which must be `addobs`, `insobs`, `clear`, `compact`, `expand`, `transpose`, `sortby`, `dsortby`, `resample` or `renumber`. Note: with the exception of `clear`, these actions are not available when the dataset is currently subsampled by selection of cases on some Boolean criterion.

addobs: Must be followed by a positive integer. Adds the specified number of extra observations to the end of the working dataset. This is primarily intended for forecasting purposes. The values of most variables over the additional range will be set to missing, but certain deterministic variables are recognized and extended, namely, a simple linear trend and periodic dummy variables.

insobs: Must be followed by a positive integer no greater than the current number of observations. Inserts a single observation at the specified position. All subsequent data are shifted by one place

and the dataset is extended by one observation. All variables apart from the constant are given missing values at the new observation. This action is not available for panel datasets.

clear: No parameter required. Clears out the current data, returning gretl to its initial “empty” state.

compact: Must be followed by a positive integer representing a new data frequency, which should be lower than the current frequency (for example, a value of 4 when the current frequency is 12 indicates compaction from monthly to quarterly). This command is available for time series data only; it compacts all the series in the data set to the new frequency. A second parameter may be given, namely one of `sum`, `first` or `last`, to specify, respectively, compaction using the sum of the higher-frequency values, start-of-period values or end-of-period values. The default is to compact by averaging.

expand: This command is only available for annual or quarterly time series data: annual data can be expanded to quarterly, and quarterly data to monthly frequency. By default all the series in the data set are padded out to the new frequency by repeating the existing values, but if the modifier `interp` is appended then the series are expanded using Chow-Lin interpolation (see [Chow and Lin \(1971\)](#)): the regressors are a constant and quadratic trend and an AR(1) disturbance process is assumed.

transpose: No additional parameter required. Transposes the current data set. That is, each observation (row) in the current data set will be treated as a variable (column), and each variable as an observation. This command may be useful if data have been read from some external source in which the rows of the data table represent variables.

sortby: The name of a single series or list is required. If one series is given, the observations on all variables in the dataset are re-ordered by increasing value of the specified series. If a list is given, the sort proceeds hierarchically: if the observations are tied in sort order with respect to the first key variable then the second key is used to break the tie, and so on until the tie is broken or the keys are exhausted. Note that this command is available only for undated data.

dsortby: Works as `sortby` except that the re-ordering is by decreasing value of the key series.

resample: Constructs a new dataset by random sampling, with replacement, of the rows of the current dataset. One argument is required, namely the number of rows to include. This may be less than, equal to, or greater than the number of observations in the original data. The original dataset can be retrieved via the command `smpl full`.

renumber: Requires the name of an existing series followed by an integer between 1 and the number of series in the dataset minus one. Moves the specified series to the specified position in the dataset, renumbering the other series accordingly. (Position 0 is occupied by the constant, which cannot be moved.)

Menu path: /Data

debug

Argument: *function*

Experimental debugger for user-defined functions, available in the command-line program, `gretlcli`, and in the GUI console. The `debug` command should be invoked after the function in question is defined but before it is called. The effect is that execution pauses when the function is called and a special prompt is shown.

At the debugging prompt you can type `next` to execute the next command in the function, or `continue` to allow execution of the function to continue unimpeded. These commands can be abbreviated as `n` and `c` respectively. You can also interpolate an instruction at this prompt, for example a `print` command to reveal the current value of some variable of interest.

delete

Variants: `delete varlist`
`delete varname`
`delete --type=type-name`

Option: `--db` (delete series from database)

This command is an all-purpose destructor for named variables (whether series, scalars, matrices, strings or bundles). It should be used with caution; no confirmation is asked.

In the first form above, *varlist* is a list of series, given by name or ID number. Note that when you delete series any series with higher ID numbers than those on the deletion list will be re-numbered. If the `--db` option is given, this command deletes the listed series not from the current dataset but from a gretl database, assuming that a database has been opened, and the user has write permission for file in question. See also the [open](#) command.

In the second form, the name of a scalar, matrix, string or bundle may be given for deletion. The `--db` option is not applicable in this case. Note that series and variables of other types should not be mixed in a given call to `delete`.

In the third form, the `--type` option must be accompanied by one of the following type-names: `matrix`, `bundle`, `string`, `list`, `scalar` or `array`. The effect is to delete all variables of the given type. In this case no argument other than the option should be given.

Menu path: Main window pop-up (single selection)

diff

Argument: *varlist*

The first difference of each variable in *varlist* is obtained and the result stored in a new variable with the prefix `d_`. Thus `diff x y` creates the new variables

$$\begin{aligned}d_x &= x(t) - x(t-1) \\ d_y &= y(t) - y(t-1)\end{aligned}$$

Menu path: /Add/First differences of selected variables

difftest

Arguments: *series1 series2*

Options: `--sign` (Sign test, the default)
`--rank-sum` (Wilcoxon rank-sum test)
`--signed-rank` (Wilcoxon signed-rank test)
`--verbose` (print extra output)

Carries out a nonparametric test for a difference between two populations or groups, the specific test depending on the option selected.

With the `--sign` option, the Sign test is performed. This test is based on the fact that if two samples, x and y , are drawn randomly from the same distribution, the probability that $x_i > y_i$, for each observation i , should equal 0.5. The test statistic is w , the number of observations for which $x_i > y_i$. Under the null hypothesis this follows the Binomial distribution with parameters $(n, 0.5)$, where n is the number of observations.

With the `--rank-sum` option, the Wilcoxon rank-sum test is performed. This test proceeds by ranking the observations from both samples jointly, from smallest to largest, then finding the sum of the ranks of the observations from one of the samples. The two samples do not have to be of the same size, and if they differ the smaller sample is used in calculating the rank-sum. Under the null

hypothesis that the samples are drawn from populations with the same median, the probability distribution of the rank-sum can be computed for any given sample sizes; and for reasonably large samples a close Normal approximation exists.

With the `--signed-rank` option, the Wilcoxon signed-rank test is performed. This is designed for matched data pairs such as, for example, the values of a variable for a sample of individuals before and after some treatment. The test proceeds by finding the differences between the paired observations, $x_i - y_i$, ranking these differences by absolute value, then assigning to each pair a signed rank, the sign agreeing with the sign of the difference. One then calculates W_+ , the sum of the positive signed ranks. As with the rank-sum test, this statistic has a well-defined distribution under the null that the median difference is zero, which converges to the Normal for samples of reasonable size.

For the Wilcoxon tests, if the `--verbose` option is given then the ranking is printed. (This option has no effect if the Sign test is selected.)

discrete

Argument: *varlist*

Option: `--reverse` (mark variables as continuous)

Marks each variable in *varlist* as being discrete. By default all variables are treated as continuous; marking a variable as discrete affects the way the variable is handled in frequency plots, and also allows you to select the variable for the command [dummify](#).

If the `--reverse` flag is given, the operation is reversed; that is, the variables in *varlist* are marked as being continuous.

Menu path: /Variable/Edit attributes

dpanel

Argument: *p* ; *depvar indepvars* [; *instruments*]

Options: `--quiet` (don't show estimated model)

`--vcv` (print covariance matrix)

`--two-step` (perform 2-step GMM estimation)

`--system` (add equations in levels)

`--time-dummies` (add time dummy variables)

`--dpdstyle` (emulate DPD package for Ox)

`--asymptotic` (uncorrected asymptotic standard errors)

Examples: `dpanel 2 ; y x1 x2`

`dpanel 2 ; y x1 x2 --system`

`dpanel {2 3} ; y x1 x2 ; x1`

`dpanel 1 ; y x1 x2 ; x1 GMM(x2,2,3)`

See also `bbond98.inp`

Carries out estimation of dynamic panel data models (that is, panel models including one or more lags of the dependent variable) using either the GMM-DIF or GMM-SYS method.

The parameter *p* represents the order of the autoregression for the dependent variable. In the simplest case this is a scalar value, but a pre-defined matrix may be given for this argument, to specify a set of (possibly non-contiguous) lags to be used.

The dependent variable and regressors should be given in levels form; they will be differenced automatically (since this estimator uses differencing to cancel out the individual effects).

The last (optional) field in the command is for specifying instruments. If no instruments are given, it

is assumed that all the independent variables are strictly exogenous. If you specify any instruments, you should include in the list any strictly exogenous independent variables. For predetermined regressors, you can use the GMM function to include a specified range of lags in block-diagonal fashion. This is illustrated in the third example above. The first argument to GMM is the name of the variable in question, the second is the minimum lag to be used as an instrument, and the third is the maximum lag. The same syntax can be used with the GMMlevel function to specify GMM-type instruments for the equations in levels.

By default the results of 1-step estimation are reported (with robust standard errors). You may select 2-step estimation as an option. In both cases tests for autocorrelation of orders 1 and 2 are provided, as well as the Sargan overidentification test and a Wald test for the joint significance of the regressors. Note that in this differenced model first-order autocorrelation is not a threat to the validity of the model, but second-order autocorrelation violates the maintained statistical assumptions.

In the case of 2-step estimation, standard errors are by default computed using the finite-sample correction suggested by Windmeijer (2005). The standard asymptotic standard errors associated with the 2-step estimator are generally reckoned to be an unreliable guide to inference, but if for some reason you want to see them you can use the `--asymptotic` option to turn off the Windmeijer correction.

If the `--time-dummies` option is given, a set of time dummy variables is added to the specified regressors. The number of dummies is one less than the maximum number of periods used in estimation, to avoid perfect collinearity with the constant. The dummies are entered in differenced form unless the `--dpdstyle` option is given, in which case they are entered in levels.

For further details and examples, please see the *Gretl User's Guide*.

Menu path: /Model/Panel/Dynamic panel model

dummify

Argument: *varlist*

Options: `--drop-first` (omit lowest value from encoding)
`--drop-last` (omit highest value from encoding)

For any suitable variables in *varlist*, creates a set of dummy variables coding for the distinct values of that variable. Suitable variables are those that have been explicitly marked as discrete, or those that take on a fairly small number of values all of which are “fairly round” (multiples of 0.25).

By default a dummy variable is added for each distinct value of the variable in question. For example if a discrete variable *x* has 5 distinct values, 5 dummy variables will be added to the data set, with names *Dx_1*, *Dx_2* and so on. The first dummy variable will have value 1 for observations where *x* takes on its smallest value, 0 otherwise; the next dummy will have value 1 when *x* takes on its second-smallest value, and so on. If one of the option flags `--drop-first` or `--drop-last` is added, then either the lowest or the highest value of each variable is omitted from the encoding (which may be useful for avoiding the “dummy variable trap”).

This command can also be embedded in the context of a regression specification. For example, the following line specifies a model where *y* is regressed on the set of dummy variables coding for *x*. (Option flags cannot be passed to `dummify` in this context.)

```
ols y dummify(x)
```

Other access: Main window pop-up menu (single selection)

duration

Arguments: *depvar indepvars* [; *censvar*]

Options: --exponential (use exponential distribution)
 --loglogistic (use log-logistic distribution)
 --lognormal (use log-normal distribution)
 --medians (fitted values are medians)
 --robust (robust (QML) standard errors)
 --cluster=*clustvar* (see [logit](#) for explanation)
 --vcv (print covariance matrix)
 --verbose (print details of iterations)

Examples: duration y 0 x1 x2
 duration y 0 x1 x2 ; cens

Estimates a duration model: the dependent variable (which must be positive) represents the duration of some state of affairs, for example the length of spells of unemployment for a cross-section of respondents. By default the Weibull distribution is used but the exponential, log-logistic and log-normal distributions are also available.

If some of the duration measurements are right-censored (e.g. an individual's spell of unemployment has not come to an end within the period of observation) then you should supply the trailing argument *censvar*, a series in which non-zero values indicate right-censored cases.

By default the fitted values obtained via the accessor *\$yhat* are the conditional means of the durations, but if the --medians option is given then *\$yhat* provides the conditional medians instead.

Please see the *Gretl User's Guide* for details.

Menu path: /Model/Limited dependent variable/Duration data...

elif

See [if](#).

else

See [if](#). Note that **else** requires a line to itself, before the following conditional command. You can append a comment, as in

```
else # OK, do something different
```

But you cannot append a command, as in

```
else x = 5 # wrong!
```

end

Ends a block of commands of some sort. For example, **end system** terminates an equation [system](#).

endif

See [if](#).

endloop

Marks the end of a command loop. See [loop](#).

eqnprint

Argument: [*-f filename*]

Option: --complete (Create a complete document)

Must follow the estimation of a model. Prints the estimated model in the form of a $\text{\LaTeX}$ equation. If a filename is specified using the *-f* flag output goes to that file, otherwise it goes to a file with a name of the form *equation_N.tex*, where N is the number of models estimated to date in the current session. See also [tabprint](#).

If the --complete flag is given, the $\text{\LaTeX}$ file is a complete document, ready for processing; otherwise it must be included in a document.

Menu path: Model window, / $\text{\LaTeX}$

equation

Arguments: *depvar indepvars*

Example: `equation y x1 x2 x3 const`

Specifies an equation within a system of equations (see [system](#)). The syntax for specifying an equation within an SUR system is the same as that for, e.g., [ols](#). For an equation within a Three-Stage Least Squares system you may either (a) give an OLS-type equation specification and provide a common list of instruments using the *instr* keyword (again, see [system](#)), or (b) use the same equation syntax as for [tsls](#).

estimate

Arguments: [*systemname*] [*estimator*]

Options: --iterate (iterate to convergence)

--no-df-corr (no degrees of freedom correction)

--geomean (see below)

--quiet (don't print results)

--verbose (print details of iterations)

Examples: `estimate "Klein Model 1" method=fiml`

`estimate Sys1 method=sur`

`estimate Sys1 method=sur --iterate`

Calls for estimation of a system of equations, which must have been previously defined using the [system](#) command. The name of the system should be given first, surrounded by double quotes if the name contains spaces. The estimator, which must be one of *ols*, *tsls*, *sur*, *3sls*, *fiml* or *liml*, is preceded by the string *method=*. These arguments are optional if the system in question has already been estimated and occupies the place of the “last model”; in that case the estimator defaults to the previously used value.

If the system in question has had a set of restrictions applied (see the [restrict](#) command), estimation will be subject to the specified restrictions.

If the estimation method is *sur* or *3sls* and the --iterate flag is given, the estimator will be iterated. In the case of SUR, if the procedure converges the results are maximum likelihood estimates. Iteration of three-stage least squares, however, does not in general converge on the full-information maximum likelihood results. The --iterate flag is ignored for other methods of estimation.

If the equation-by-equation estimators *ols* or *tsls* are chosen, the default is to apply a degrees of freedom correction when calculating standard errors. This can be suppressed using the --no-df-corr flag. This flag has no effect with the other estimators; no degrees of freedom correction is applied in any case.

By default, the formula used in calculating the elements of the cross-equation covariance matrix is

$$\hat{\sigma}_{i,j} = \frac{\hat{u}_i' \hat{u}_j}{T}$$

If the `--geomean` flag is given, a degrees of freedom correction is applied: the formula is

$$\hat{\sigma}_{i,j} = \frac{\hat{u}_i' \hat{u}_j}{\sqrt{(T - k_i)(T - k_j)}}$$

where the *ks* denote the number of independent parameters in each equation.

If the `--verbose` option is given and an iterative method is specified, details of the iterations are printed.

eval

Argument: *expression*

Examples: `eval x`
`eval inv(X'X)`
`eval sqrt($pi)`

This command makes gretl act like a glorified calculator. The program evaluates *expression* and prints its value. The argument may be the name of a variable, or something more complicated. In any case, it should be an expression which could stand as the right-hand side of an assignment statement.

fcast

Arguments: [*startobs endobs*] [*steps-ahead*] [*varname*]

Options: `--dynamic` (create dynamic forecast)
`--static` (create static forecast)
`--out-of-sample` (generate post-sample forecast)
`--no-stats` (don't print forecast statistics)
`--quiet` (don't print anything)
`--rolling` (see below)
`--plot=filename` (see below)

Examples: `fcast 1997:1 2001:4 f1`
`fcast fit2`
`fcast 2004:1 2008:3 4 rfcast --rolling`

Must follow an estimation command. Forecasts are generated for a certain range of observations: if *startobs* and *endobs* are given, for that range (if possible); otherwise if the `--out-of-sample` option is given, for observations following the range over which the model was estimated; otherwise over the currently defined sample range. If an out-of-sample forecast is requested but no relevant observations are available, an error is flagged. Depending on the nature of the model, standard errors may also be generated; see below. Also see below for the special effect of the `--rolling` option.

If the last model estimated is a single equation, then the optional *varname* argument has the following effect: the forecast values are not printed, but are saved to the dataset under the given name. If the last model is a system of equations, *varname* has a different effect, namely selecting a particular endogenous variable for forecasting (the default being to produce forecasts for all the endogenous variables). In the system case, or if *varname* is not given, the forecast values can be retrieved using the accessor `$fcast`, and the standard errors, if available, via `$fcerr`.

The choice between a static and a dynamic forecast applies only in the case of dynamic models, with an autoregressive error process or including one or more lagged values of the dependent variable as regressors. Static forecasts are one step ahead, based on realized values from the previous period, while dynamic forecasts employ the chain rule of forecasting. For example, if a forecast for y in 2008 requires as input a value of y for 2007, a static forecast is impossible without actual data for 2007. A dynamic forecast for 2008 is possible if a prior forecast can be substituted for y in 2007.

The default is to give a static forecast for any portion of the forecast range that lies within the sample range over which the model was estimated, and a dynamic forecast (if relevant) out of sample. The `--dynamic` option requests a dynamic forecast from the earliest possible date, and the `--static` option requests a static forecast even out of sample.

The `--rolling` option is presently available only for single-equation models estimated via OLS. When this option is given the forecasts are recursive. That is, each forecast is generated from an estimate of the given model using data from a fixed starting point (namely, the start of the sample range for the original estimation) up to the forecast date minus k , where k is the number of steps ahead, which must be given in the *steps-ahead* argument. The forecasts are always dynamic if this is applicable. Note that the *steps-ahead* argument should be given only in conjunction with the `--rolling` option.

The `--plot` option (available only in the case of single-equation estimation) calls for a plot file to be produced, containing a graphical representation of the forecast. The suffix of the *filename* argument to this option controls the format of the plot: `.eps` for EPS, `.pdf` for PDF, `.png` for PNG, `.plt` for a gnuplot command file. The dummy filename `display` can be used to force display of the plot in a window. For example,

```
fcast --plot=fc.pdf
```

will generate a graphic in PDF format. Absolute pathnames are respected, otherwise files are written to the gretl working directory.

The nature of the forecast standard errors (if available) depends on the nature of the model and the forecast. For static linear models standard errors are computed using the method outlined by Davidson and MacKinnon (2004); they incorporate both uncertainty due to the error process and parameter uncertainty (summarized in the covariance matrix of the parameter estimates). For dynamic models, forecast standard errors are computed only in the case of a dynamic forecast, and they do not incorporate parameter uncertainty. For nonlinear models, forecast standard errors are not presently available.

Menu path: Model window, /Analysis/Forecasts

flush

This simple command (no arguments, no options) is intended for use in time-consuming scripts that may be executed via the gretl GUI (it is ignored by the command-line program), to give the user a visual indication that things are moving along and gretl is not “frozen”.

Ordinarily if you launch a script in the GUI no output is shown until its execution is completed, but the effect of invoking `flush` is as follows:

- On the first invocation, gretl opens a window, displays the output so far, and appends the message “Processing...”.
- On subsequent invocations the text shown in the output window is updated, and a new “processing” message is appended.

When execution of the script is completed any remaining output is automatically flushed to the text window.

Please note, there is no point in using `flush` in scripts that take less than (say) 5 seconds to execute. Also note that this command should not be used at a point in the script where there is no further output to be printed, as the “processing” message will then be misleading to the user.

The following illustrates the intended use of `flush`:

```
set echo off
scalar n = 10
loop i=1..n
  # do some time-consuming operation
  loop 100 --quiet
    a = mnormal(200,200)
    b = inv(a)
  endloop
  # print some results
  printf "Iteration %2d done\n", i
  if i < n
    flush
  endif
endloop
```

foreign

Syntax: `foreign language=lang`

Options: `--send-data [=list]` (pre-load data; see below)
`--quiet` (suppress output from foreign program)

This command opens a special mode in which commands to be executed by another program are accepted. You exit this mode with `end foreign`; at this point the stacked commands are executed.

At present the “foreign” programs supported in this way are GNU R (`language=R`), Jurgen Doornik’s Ox (`language=Ox`), GNU Octave (`language=Octave`), Python and, to a lesser extent, Stata. Language names are recognized on a case-insensitive basis.

In connection with R, Octave and Stata the `--send-data` option has the effect of making data from gretl’s workspace available within the target program. By default the entire dataset is sent, but you can limit the data to be sent by giving the name of a predefined list of series. For example:

```
list Rlist = x1 x2 x3
foreign language=R --send-data=Rlist
```

See the *Gretl User’s Guide* for details and examples.

fractint

Arguments: `series [ order ]`

Options: `--gph` (do Geweke and Porter-Hudak test)
`--all` (do both tests)
`--quiet` (don’t print results)

Tests the specified series for fractional integration (“long memory”). The null hypothesis is that the integration order of the series is zero. By default the local Whittle estimator ([Robinson, 1995](#)) is used but if the `--gph` option is given the GPH test ([Geweke and Porter-Hudak, 1983](#)) is performed instead. If the `--all` flag is given then the results of both tests are printed.

For details on this sort of test, see [Phillips and Shimotsu \(2004\)](#).

If the optional *order* argument is not given the order for the test(s) is set automatically as the lesser of $T/2$ and $T^{0.6}$.

The results can be retrieved using the accessors `$test` and `$pvalue`. These values are based on the Local Whittle Estimator unless the `--gph` option is given.

Menu path: /Variable/Unit root tests/Fractional integration

freq

Argument: *var*

Options: `--nbins=n` (specify number of bins)
`--min=minval` (specify minimum, see below)
`--binwidth=width` (specify bin width, see below)
`--quiet` (suppress printing of graph)
`--normal` (test for the normal distribution)
`--gamma` (test for gamma distribution)
`--silent` (don't print anything)
`--show-plot` (see below)
`--matrix=name` (use column of named matrix)

Examples: `freq x`
`freq x --normal`
`freq x --nbins=5`
`freq x --min=0 --binwidth=0.10`

With no options given, displays the frequency distribution for the series *var* (given by name or number), with the number of bins and their size chosen automatically.

If the `--matrix` option is given, *var* (which must be an integer) is instead interpreted as a 1-based index that selects a column from the named matrix. If the matrix in question is in fact a column vector, the *var* argument may be omitted.

To control the presentation of the distribution you may specify *either* the number of bins or the minimum value plus the width of the bins, as shown in the last two examples above. The `--min` option sets the lower limit of the left-most bin.

If the `--normal` option is given, the Doornik-Hansen chi-square test for normality is computed. If the `--gamma` option is given, the test for normality is replaced by Locke's nonparametric test for the null hypothesis that the variable follows the gamma distribution; see [Locke \(1976\)](#), [Shapiro and Chen \(2001\)](#). Note that the parameterization of the gamma distribution used in gretl is (shape, scale).

In interactive mode a graph of the distribution is displayed by default. The `--quiet` flag can be used to suppress this. Conversely, the graph is not usually shown when the `freq` is used in a script, but you can force its display by giving the `--show-plot` option. (This does not apply when using the command-line program, `gretlcli`.)

The `--silent` flag suppresses the usual output entirely. This makes sense only in conjunction with one or other of the distribution test options: the test statistic and its p-value are recorded, and can be retrieved using the accessors `$test` and `$pvalue`.

Menu path: /Variable/Frequency distribution

function

Argument: *fname*

Opens a block of statements in which a function is defined. This block must be closed with `end function`. Please see the *Gretl User's Guide* for details.

garch

Arguments: `p q ; depvar [ indepvars ]`

Options: `--robust` (robust standard errors)
`--verbose` (print details of iterations)
`--vcv` (print covariance matrix)
`--nc` (do not include a constant)
`--stdresid` (standardize the residuals)
`--fcp` (use Fiorentini, Calzolari, Panattoni algorithm)
`--arma-init` (initial variance parameters from ARMA)

Examples: `garch 1 1 ; y`
`garch 1 1 ; y 0 x1 x2 --robust`

Estimates a GARCH model (GARCH = Generalized Autoregressive Conditional Heteroskedasticity), either a univariate model or, if *indepvars* are specified, including the given exogenous variables. The integer values *p* and *q* (which may be given in numerical form or as the names of pre-existing scalar variables) represent the lag orders in the conditional variance equation:

$$h_t = \alpha_0 + \sum_{i=1}^q \alpha_i \varepsilon_{t-i}^2 + \sum_{j=1}^p \beta_j h_{t-j}$$

The parameter *p* therefore represents the Generalized (or “AR”) order, while *q* represents the regular ARCH (or “MA”) order. If *p* is non-zero, *q* must also be non-zero otherwise the model is unidentified. However, you can estimate a regular ARCH model by setting *q* to a positive value and *p* to zero. The sum of *p* and *q* must be no greater than 5. Note that a constant is automatically included in the mean equation unless the `--nc` option is given.

By default native gretl code is used in estimation of GARCH models, but you also have the option of using the algorithm of [Fiorentini et al. \(1996\)](#). The former uses the BFGS maximizer while the latter uses the information matrix to maximize the likelihood, with fine-tuning via the Hessian.

Several variant estimators of the covariance matrix are available with this command. By default, the Hessian is used unless the `--robust` option is given, in which case the QML (White) covariance matrix is used. Other possibilities (e.g. the information matrix, or the Bollerslev-Wooldridge estimator) can be specified using the `set` command.

By default, the estimates of the variance parameters are initialized using the unconditional error variance from initial OLS estimation for the constant, and small positive values for the coefficients on the past values of the squared error and the error variance. The flag `--arma-init` calls for the starting values of these parameters to be set using an initial ARMA model, exploiting the relationship between GARCH and ARMA set out in Chapter 21 of Hamilton's *Time Series Analysis*. In some cases this may improve the chances of convergence.

The GARCH residuals and estimated conditional variance can be retrieved as `$uhat` and `$h` respectively. For example, to get the conditional variance:

```
series ht = $h
```

If the `--stdresid` option is given, the `$uhat` values are divided by the square root of h_t .

Menu path: /Model/Time series/GARCH

genr

Arguments: *newvar* = *formula*

NOTE: this command has undergone numerous changes and enhancements since the following help text was written, so for comprehensive and updated info on this command you'll want to refer to the *Gretl User's Guide*. On the other hand, this help does not contain anything actually erroneous, so take the following as "you have this, plus more".

In the appropriate context, *series*, *scalar*, *matrix*, *string* and *bundle* are synonyms for this command.

Creates new variables, often via transformations of existing variables. See also [diff](#), [logs](#), [lags](#), [ldiff](#), [sdiff](#) and [square](#) for shortcuts. In the context of a *genr* formula, existing variables must be referenced by name, not ID number. The formula should be a well-formed combination of variable names, constants, operators and functions (described below). Note that further details on some aspects of this command can be found in the *Gretl User's Guide*.

```
series c = 10
```

A *genr* command may yield either a series or a scalar result. For example, the formula $x^2 = x * 2$ naturally yields a series if the variable *x* is a series and a scalar if *x* is a scalar. The formulae $x = 0$ and $mx = \text{mean}(x)$ naturally return scalars. Under some circumstances you may want to have a scalar result expanded into a series or vector. You can do this by using *series* as an "alias" for the *genr* command. For example, *series x = 0* produces a series all of whose values are set to 0. You can also use *scalar* as an alias for *genr*. It is not possible to coerce a vector result into a scalar, but use of this keyword indicates that the result *should be* a scalar: if it is not, an error occurs.

When a formula yields a series result, the range over which the result is written to the target variable depends on the current sample setting. It is possible, therefore, to define a series piecewise using the *smp1* command in conjunction with *genr*.

Supported *arithmetical operators* are, in order of precedence: $\wedge$ (exponentiation); $*$, $/$ and $\%$ (modulus or remainder); $+$ and $-$.

The available *Boolean operators* are (again, in order of precedence): $!$ (negation), $\&\&$ (logical AND), $||$ (logical OR), $>$, $<$, $=$, $>=$ (greater than or equal), $<=$ (less than or equal) and $!=$ (not equal). The Boolean operators can be used in constructing dummy variables: for instance $(x > 10)$ returns 1 if $x > 10$, 0 otherwise.

Built-in constants are *pi* and *NA*. The latter is the missing value code: you can initialize a variable to the missing value with *scalar x = NA*.

The *genr* command supports a wide range of mathematical and statistical functions, including all the common ones plus several that are special to econometrics. In addition it offers access to numerous internal variables that are defined in the course of running regressions, doing hypothesis tests, and so on.

For a listing of functions and accessors, see Chapter 2.

Besides the operators and functions noted above there are some special uses of *genr*:

- *genr time* creates a time trend variable (1,2,3,...) called *time*. *genr index* does the same thing except that the variable is called *index*.
- *genr dummy* creates dummy variables up to the periodicity of the data. In the case of quarterly data (periodicity 4), the program creates $dq1 = 1$ for first quarter and 0 in other quarters, $dq2 = 1$ for the second quarter and 0 in other quarters, and so on. With monthly data the dummies are named *dm1*, *dm2*, and so on. With other frequencies the names are *dummy_1*, *dummy_2*, etc.

- `genr unitdum` and `genr timedum` create sets of special dummy variables for use with panel data. The first codes for the cross-sectional units and the second for the time period of the observations.

Note: In the command-line program, `genr` commands that retrieve model-related data always reference the model that was estimated most recently. This is also true in the GUI program, if one uses `genr` in the “gretl console” or enters a formula using the “Define new variable” option under the Add menu in the main window. With the GUI, however, you have the option of retrieving data from any model currently displayed in a window (whether or not it’s the most recent model). You do this under the “Save” menu in the model’s window.

The special variable `obs` serves as an index of the observations. For instance `genr dum = (obs=15)` will generate a dummy variable that has value 1 for observation 15, 0 otherwise. You can also use this variable to pick out particular observations by date or name. For example, `genr d = (obs>1986:4)`, `genr d = (obs>"2008-04-01")`, or `genr d = (obs="CA")`. If daily dates or observation labels are used in this context, they should be enclosed in double quotes. Quarterly and monthly dates (with a colon) may be used unquoted. Note that in the case of annual time series data, the year is not distinguishable syntactically from a plain integer; therefore if you wish to compare observations against `obs` by year you must use the function `obsnum` to convert the year to a 1-based index value, as in `genr d = (obs>obsnum(1986))`.

Scalar values can be pulled from a series in the context of a `genr` formula, using the syntax `varname[obs]`. The `obs` value can be given by number or date. Examples: `x[5]`, `CPI[1996:01]`. For daily data, the form `YYYY-MM-DD` should be used, e.g. `ibm[1970-01-23]`.

An individual observation in a series can be modified via `genr`. To do this, a valid observation number or date, in square brackets, must be appended to the name of the variable on the left-hand side of the formula. For example, `genr x[3] = 30` or `genr x[1950:04] = 303.7`.

Menu path: /Add/Define new variable

Other access: Main window pop-up menu

gmm

Options: `--two-step` (two step estimation)
 `--iterate` (iterated GMM)
 `--vcv` (print covariance matrix)
 `--verbose` (print details of iterations)
 `--lbfgs` (use L-BFGS-B instead of regular BFGS)

Performs Generalized Method of Moments (GMM) estimation using the BFGS (Broyden, Fletcher, Goldfarb, Shanno) algorithm. You must specify one or more commands for updating the relevant quantities (typically GMM residuals), one or more sets of orthogonality conditions, an initial matrix of weights, and a listing of the parameters to be estimated, all enclosed between the tags `gmm` and `end gmm`. Any options should be appended to the `end gmm` line.

Please see the *Gretl User’s Guide* for details on this command. Here we just illustrate with a simple example.

```
gmm e = y - X*b
    orthog e ; W
    weights V
    params b
end gmm
```

In the example above we assume that `y` and `X` are data matrices, `b` is an appropriately sized vector of parameter values, `W` is a matrix of instruments, and `V` is a suitable matrix of weights. The statement

<i>Formula</i>	<i>Comment</i>
<code>y = x1^3</code>	x1 cubed
<code>y = ln((x1+x2)/x3)</code>	
<code>z = x>y</code>	$z(t) = 1$ if $x(t) > y(t)$, otherwise 0
<code>y = x(-2)</code>	x lagged 2 periods
<code>y = x(+2)</code>	x led 2 periods
<code>y = diff(x)</code>	$y(t) = x(t) - x(t-1)$
<code>y = ldiff(x)</code>	$y(t) = \log x(t) - \log x(t-1)$, the instantaneous rate of growth of x
<code>y = sort(x)</code>	sorts x in increasing order and stores in y
<code>y = dsort(x)</code>	sort x in decreasing order
<code>y = int(x)</code>	truncate x and store its integer value as y
<code>y = abs(x)</code>	store the absolute values of x
<code>y = sum(x)</code>	sum x values excluding missing NA entries
<code>y = cum(x)</code>	cumulation: $y_t = \sum_{\tau=1}^t x_\tau$
<code>aa = \$ess</code>	set aa equal to the Error Sum of Squares from last regression
<code>x = \$coeff(sqft)</code>	grab the estimated coefficient on the variable sqft from the last regression
<code>rho4 = \$rho(4)</code>	grab the 4th-order autoregressive coefficient from the last model (presumes an ar model)
<code>cvx1x2 = \$vcv(x1, x2)</code>	grab the estimated coefficient covariance of vars x1 and x2 from the last model
<code>foo = uniform()</code>	uniform pseudo-random variable in range 0-1
<code>bar = 3 * normal()</code>	normal pseudo-random variable, $\mu = 0$, $\sigma = 3$
<code>samp = ok(x)</code>	= 1 for observations where x is not missing.

```
orthog e ; W
```

indicates that the residual vector e is in principle orthogonal to each of the instruments composing the columns of W .

Menu path: /Model/GMM

gnuplot

Arguments: *yvars* *xvar* [*dumvar*]

Options: --with-lines [=varspec] (use lines, not points)
 --with-lp [=varspec] (use lines and points)
 --with-impulses [=varspec] (use vertical lines)
 --time-series (plot against time)
 --single-yaxis (force use of just one y-axis)
 --dummy (see below)
 --fit=*fitspec* (see below)
 --matrix=*name* (plot columns of named matrix)
 --output=*filename* (send output to specified file)
 --input=*filename* (take input from specified file)

Examples: gnuplot y1 y2 x
 gnuplot x --time-series --with-lines
 gnuplot wages educ gender --dummy
 gnuplot y x --fit=quadratic
 gnuplot y1 y2 x --with-lines=y2

The variables in the list *yvars* are graphed against *xvar*. For a time series plot you may either give time as *xvar* or use the option flag --time-series.

By default, data-points are shown as points; this can be overridden by giving one of the options --with-lines, --with-lp or --with-impulses. If more than one variable is to be plotted on the y axis, the effect of these options may be confined to a subset of the variables by using the *varspec* parameter. This should take the form of a comma-separated listing of the names or numbers of the variables to be plotted with lines or impulses respectively. For instance, the final example above shows how to plot y_1 and y_2 against x , such that y_2 is represented by a line but y_1 by points.

If the --dummy option is selected, exactly three variables should be given: a single y variable, an x variable, and *dvar*, a discrete variable. The effect is to plot *yvar* against *xvar* with the points shown in different colors depending on the value of *dvar* at the given observation.

Generally, the arguments *yvars* and *xvar* are required, and refer to series in the current dataset (given either by name or ID number). But if a named matrix is supplied via the --matrix option these arguments become optional: if the specified matrix has k columns, by default the first $k - 1$ columns are treated as the *yvars* and the last column as *xvar*. If the --time-series option is given, however, all k columns are plotted against time. If you wish to plot selected columns of the matrix, you should specify *yvars* and *xvar* in the form of 1-based column numbers. For example if you want a scatterplot of column 2 of matrix *M* against column 1, you can do:

```
gnuplot 2 1 --matrix=M
```

The --fit option is applicable only for bivariate scatterplots and single time-series plots. The default behavior for a scatterplot is to show the OLS fit if the slope coefficient is significant at the 10 percent level, while the default behavior for time-series is not to show any fitted line. You can call for different behavior by using this option along with one of the following *fitspec* parameter values. Note that if the plot is a single time series the place of x is taken by time.

- **linear**: show the OLS fit regardless of its level of statistical significance.
- **none**: don't show any fitted line.
- **inverse**, **quadratic**, **cubic**, **semilog** or **linlog**: show a fitted line based on a regression of the specified type. By **semilog**, we mean a regression of $\log y$ on x ; the fitted line represents the conditional expectation of y , obtained by exponentiation. By **linlog** we mean a regression of y on the log of x .
- **loess**: show the fit from a robust locally weighted regression (also is sometimes known as "lowess").

In interactive mode the plot is displayed immediately. In batch mode the default behavior is that a gnuplot command file is written in the user's working directory, with a name on the pattern `gpttmpN.plt`, starting with `N = 01`. The actual plots may be generated later using gnuplot (under MS Windows, `wgnuplot`). This behavior can be modified by use of the `--output=filename` option. This option controls the filename used, and at the same time allows you to specify a particular output format via the three-letter extension of the file name, as follows: `.eps` results in the production of an Encapsulated PostScript (EPS) file; `.pdf` produces PDF; `.png` produces PNG format, `.emf` calls for EMF (Enhanced MetaFile), `.fig` calls for an Xfig file, and `.svg` for SVG (Scalable Vector Graphics). If the dummy filename "display" is given then the plot is shown on screen as in interactive mode. If a filename with any extension other than those just mentioned is given, a gnuplot command file is written.

A further option to this command is available: following the specification of the variables to be plotted and the option flag (if any), you may add literal gnuplot commands to control the appearance of the plot (for example, setting the plot title and/or the axis ranges). These commands should be enclosed in braces, and each gnuplot command must be terminated with a semi-colon. A backslash may be used to continue a set of gnuplot commands over more than one line. Here is an example of the syntax:

```
{ set title 'My Title'; set yrange [0:1000]; }
```

Menu path: /View/Graph specified vars

Other access: Main window pop-up menu, graph button on toolbar

graphpg

Variants: `graphpg add`
`graphpg fontscale value`
`graphpg show`
`graphpg free`
`graphpg --output=filename`

The session "graph page" will work only if you have the $\text{\LaTeX}$ typesetting system installed, and are able to generate and view PDF or PostScript output.

In the session icon window, you can drag up to eight graphs onto the graph page icon. When you double-click on the graph page (or right-click and select "Display"), a page containing the selected graphs will be composed and opened in a suitable viewer. From there you should be able to print the page.

To clear the graph page, right-click on its icon and select "Clear".

Note that on systems other than MS Windows, you may have to adjust the setting for the program used to view PDF or PostScript files. Find that under the "Programs" tab in the gretl Preferences dialog box (under the Tools menu in the main window).

It's also possible to operate on the graph page via script, or using the console (in the GUI program). The following commands and options are supported:

To add a graph to the graph page, issue the command `graphpg add` after saving a named graph, as in

```
grf1 <- gnuplot Y X
graphpg add
```

To display the graph page: `graphpg show`.

To clear the graph page: `graphpg free`.

To adjust the scale of the font used in the graph page, use `graphpg fontscale scale`, where *scale* is a multiplier (with a default of 1.0). Thus to make the font size 50 percent bigger than the default you can do

```
graphpg fontscale 1.5
```

To call for printing of the graph page to file, use the flag `--output=` plus a filename; the filename should have the suffix `".pdf"`, `".ps"` or `".eps"`. For example:

```
graphpg --output="myfile.pdf"
```

In this context the output uses colored lines by default; to use dot/dash patterns instead of colors you can append the `--monochrome` flag.

hausman

This test is available only after estimating an OLS model using panel data (see also `setobs`). It tests the simple pooled model against the principal alternatives, the fixed effects and random effects models.

The fixed effects model allows the intercept of the regression to vary across the cross-sectional units. An *F*-test is reported for the null hypotheses that the intercepts do not differ. The random effects model decomposes the residual variance into two parts, one part specific to the cross-sectional unit and the other specific to the particular observation. (This estimator can be computed only if the number of cross-sectional units in the data set exceeds the number of parameters to be estimated.) The Breusch-Pagan LM statistic tests the null hypothesis that the pooled OLS estimator is adequate against the random effects alternative.

The pooled OLS model may be rejected against both of the alternatives, fixed effects and random effects. Provided the unit- or group-specific error is uncorrelated with the independent variables, the random effects estimator is more efficient than the fixed effects estimator; otherwise the random effects estimator is inconsistent and the fixed effects estimator is to be preferred. The null hypothesis for the Hausman test is that the group-specific error is not so correlated (and therefore the random effects model is preferable). A low p-value for this test counts against the random effects model and in favor of fixed effects.

Menu path: Model window, /Tests/Panel diagnostics

heckit

Arguments: *depvar indepvars ; selection equation*
Options: --quiet (suppress printing of results)
 --robust (QML standard errors)
 --two-step (perform two-step estimation)
 --vcv (print covariance matrix)
 --verbose (print extra output)
Example: heckit y 0 x1 x2 ; ys 0 x3 x4
 heckit.inp

Heckman-type selection model. In the specification, the list before the semicolon represents the outcome equation, and the second list represents the selection equation. The dependent variable in the selection equation (ys in the example above) must be a binary variable.

By default, the parameters are estimated by maximum likelihood. The covariance matrix of the parameters is computed using the negative inverse of the Hessian. If two-step estimation is desired, use the `--two-step` option. In this case, the covariance matrix of the parameters of the outcome equation is appropriately adjusted as per [Heckman \(1979\)](#).

Please note that in ML estimation a numerical approximation of the Hessian is used; this may lead to inaccuracies in the estimated covariance matrix if the scale of the explanatory variables is such that some of the estimated coefficients are very small in absolute value. This problem will be addressed in future versions; in the meantime, rescaling the offending explanatory variable(s) can be used as a workaround.

Menu path: /Model/Limited dependent variable/Heckit

help

Variants: help
 help functions
 help *command*
 help *function*
Option: --func (select functions help)

If no arguments are given, prints a list of available commands. If the single argument `functions` is given, prints a list of available functions (see [genr](#)).

`help command` describes *command* (e.g. `help smpl`). `help function` describes *function* (e.g. `help ldet`). Some functions have the same names as related commands (e.g. `diff`): in that case the default is to print help for the command, but you can get help on the function by using the `--func` option.

Menu path: /Help

hsk

Arguments: *depvar indepvars*
Options: --no-squares (see below)
 --vcv (print covariance matrix)

This command is applicable where heteroskedasticity is present in the form of an unknown function of the regressors which can be approximated by a quadratic relationship. In that context it offers the possibility of consistent standard errors and more efficient parameter estimates as compared with OLS.

The procedure involves (a) OLS estimation of the model of interest, followed by (b) an auxiliary

regression to generate an estimate of the error variance, then finally (c) weighted least squares, using as weight the reciprocal of the estimated variance.

In the auxiliary regression (b) we regress the log of the squared residuals from the first OLS on the original regressors and their squares (by default), or just on the original regressors (if the `--no-squares` option is given). The log transformation is performed to ensure that the estimated variances are all non-negative. Call the fitted values from this regression u^* . The weight series for the final WLS is then formed as $1/\exp(u^*)$.

Menu path: /Model/Other linear models/Heteroskedasticity corrected

hurst

Argument: *series*

Calculates the Hurst exponent (a measure of persistence or long memory) for a time-series variable having at least 128 observations.

The Hurst exponent is discussed by Mandelbrot. In theoretical terms it is the exponent, H , in the relationship

$$RS(x) = an^H$$

where RS is the “rescaled range” of the variable x in samples of size n and a is a constant. The rescaled range is the range (maximum minus minimum) of the cumulated value or partial sum of x over the sample period (after subtraction of the sample mean), divided by the sample standard deviation.

As a reference point, if x is white noise (zero mean, zero persistence) then the range of its cumulated “wandering” (which forms a random walk), scaled by the standard deviation, grows as the square root of the sample size, giving an expected Hurst exponent of 0.5. Values of the exponent significantly in excess of 0.5 indicate persistence, and values less than 0.5 indicate anti-persistence (negative autocorrelation). In principle the exponent is bounded by 0 and 1, although in finite samples it is possible to get an estimated exponent greater than 1.

In gretl, the exponent is estimated using binary sub-sampling: we start with the entire data range, then the two halves of the range, then the four quarters, and so on. For sample sizes smaller than the data range, the RS value is the mean across the available samples. The exponent is then estimated as the slope coefficient in a regression of the log of RS on the log of sample size.

Menu path: /Variable/Hurst exponent

if

Flow control for command execution. Three sorts of construction are supported, as follows.

```
# simple form
if condition
    commands
endif

# two branches
if condition
    commands1
else
    commands2
endif

# three or more branches
if condition1
    commands1
elif condition2
```

```

        commands2
    else
        commands3
    endif

```

condition must be a Boolean expression, for the syntax of which see [genr](#). More than one `elif` block may be included. In addition, `if ... endif` blocks may be nested.

include

Argument: *filename*
 Examples: `include myfile.inp`
 `include sols.gfn`

Intended for use in a command script, primarily for including definitions of functions. Executes the commands in *filename* then returns control to the main script. To include a packaged function, be sure to include the filename extension.

See also [run](#).

info

Prints out any supplementary information stored with the current datafile.

Menu path: /Data/Dataset info

Other access: Data browser windows

install

Argument: *pkgname*
 Options: `--local` (install from local file)
 `--remove` (see below)
 `--purge` (see below)
 Examples: `install armax`
 `install felogit.gfn`
 `install /path/to/myfile.gfn --local`
 `install http://foo.bar.net/gretl/myfile.gfn`

Installer for gretl function packages (gfn or zip files).

If this command is given the “plain” name of a gretl function package (as in the first two examples) the action is to download the specified package from the gretl server and install it on the local machine. In this case it is not necessary to supply a filename extension.

If the `--local` option is given, the *pkgname* argument should be the path to an uninstalled package file on the local machine, with the correct extension. The action is to copy the file into place (gfn), or unzip it into place (zip), “into place” meaning where the [include](#) command will find it.

When no option is given, if *pkgname* begins with `http://`, the effect is to download a package file from a specified server and install it locally.

With the `--remove` or `--purge` option the inverse operation is performed; that is, an installed package is uninstalled. If just `--remove` is given, the specified package is unloaded from memory and is removed from the GUI menu to which it is attached, if any. If the `--purge` option is given then in addition to the actions just mentioned the package file is deleted. (If the package is installed in its own subdirectory, the whole subdirectory is deleted.)

Menu path: /Tools/Function packages/On server

intreg

Arguments: *minvar maxvar indepvars*

Options: `--quiet` (suppress printing of results)
`--verbose` (print details of iterations)
`--robust` (robust standard errors)
`--cluster=clustvar` (see [logit](#) for explanation)

Example: `intreg lo hi const x1 x2`
`wtp.inp`

Estimates an interval regression model. This model arises when the dependent variable is imperfectly observed for some (possibly all) observations. In other words, the data generating process is assumed to be

$$y_t^* = x_t\beta + \epsilon_t$$

but we only observe

$$m_t \leq y_t \leq M_t$$

(the interval may be left- or right-unbounded). Note that for some observations m may equal M . The variables *minvar* and *maxvar* must contain NAs for left- and right-unbounded observations, respectively.

The model is estimated by maximum likelihood, assuming normality of the disturbance term.

By default, standard errors are computed using the negative inverse of the Hessian. If the `--robust` flag is given, then QML or Huber-White standard errors are calculated instead. In this case the estimated covariance matrix is a “sandwich” of the inverse of the estimated Hessian and the outer product of the gradient.

Menu path: /Model/Limited dependent variable/Interval regression

join

Arguments: *filename varname*

Options: `--data=column-name` (see below)
`--filter=expression` (see below)
`--ikey=inner-key` (see below)
`--okey=outer-key` (see below)
`--aggr=method` (see below)
`--tkey=column-name,format-string` (see below)
`--verbose` (report on progress)

This command imports a data series from the source *filename* (which must be either a delimited text data file or a “native” gretl data file) under the name *varname*. For details please see the *Gretl User's Guide*; here we just give a brief summary of the available options.

The `--data` option can be used to specify the column heading of the data in the source file, if this differs from the name by which the data should be known in gretl.

The `--filter` option can be used to specify a criterion for filtering the source data (that is, selecting a subset of observations).

The `--ikey` and `--okey` options can be used to specify a mapping between observations in the current dataset and observations in the source data (for example, individuals can be matched against the household to which they belong).

The `--aggr` option is used when the mapping between observations in the current dataset and the source is not one-to-one.

The `--tkey` option is applicable only when the current dataset has a time-series structure. It can be used to specify the name of a column containing dates to be matched to the dataset and/or the format in which dates are represented in that column.

See also [append](#) for simpler joining operations.

kalman

Options: `--cross` (allow for cross-correlated disturbances)
`--diffuse` (use diffuse initialization)

Opens a block of statements to set up a Kalman filter. This block should end with the line `end kalman`, to which the options shown above may be appended. The intervening lines specify the matrices that compose the filter. For example,

```
kalman
  obsy y
  obsymat H
  statemat F
  statevar Q
end kalman
```

would set up a linear, time invariant state-space model with observation equation

$$y_t = H' \alpha_t$$

and state transition equation

$$\alpha_{t+1} = F \alpha_t + u_t$$

where $V(u) = Q$.

More complex state-space models are handled via the additional keywords `obsx`, `obsxmat`, `obsvar`, `inistate`, `inivar` and `stconst`. Please see the *Gretl User's Guide* for details.

See also [kfilter](#), [ksimul](#), [ksmooth](#).

kpss

Arguments: *order varlist*

Options: `--trend` (include a trend)
`--seasonals` (include seasonal dummies)
`--verbose` (print regression results)
`--quiet` (suppress printing of results)
`--difference` (use first difference of variable)

Examples: `kpss 8 y`
`kpss 4 x1 --trend`

For use of this command with panel data please see the final section in this entry.

Computes the KPSS test ([Kwiatkowski et al., 1992](#)) for stationarity, for each of the specified variables (or their first difference, if the `--difference` option is selected). The null hypothesis is that the variable in question is stationary, either around a level or, if the `--trend` option is given, around a deterministic linear trend.

The *order* argument determines the size of the window used for Bartlett smoothing. If a negative value is given this is taken as a signal to use an automatic window size of $4(T/100)^{0.25}$, where T is the sample size.

If the `--verbose` option is chosen the results of the auxiliary regression are printed, along with the estimated variance of the random walk component of the variable.

The critical values shown for the test statistic are based on response surfaces estimated in the manner set out by [Sephton \(1995\)](#), which are more accurate for small samples than the values given in the original KPSS article. When the test statistic lies between the 10 percent and 1 percent critical values a p-value is shown; this is obtained by linear interpolation and should not be taken too literally. See the [kpsscrit](#) function for a means of obtaining these critical values programmatically.

Panel data

When the `kpss` command is used with panel data, to produce a panel unit root test, the applicable options and the results shown are somewhat different. While you may give a list of variables for testing in the regular time-series case, with panel data only one variable may be tested per command. And the `--verbose` option has a different meaning: it produces a brief account of the test for each individual time series (the default being to show only the overall result).

When possible, the overall test (null hypothesis: the series in question is stationary for all the panel units) is calculated using the method of [Choi \(2001\)](#). This is not always straightforward, the difficulty being that while the Choi test is based on the p-values of the tests on the individual series, we do not currently have a means of calculating p-values for the KPSS test statistic; we must rely on a few critical values.

If the test statistic for a given series falls between the 10 percent and 1 percent critical values, we are able to interpolate a p-value. But if the test falls short of the 10 percent value, or exceeds the 1 percent value, we cannot interpolate and can at best place a bound on the global Choi test. If the individual test statistic falls short of the 10 percent value for some units but exceeds the 1 percent value for others, we cannot even compute a bound for the global test.

Menu path: /Variable/Unit root tests/KPSS test

labels

```

Variants:  labels [ varlist ]
           labels --to-file=filename
           labels --from-file=filename
           labels --delete

```

In the first form, prints out the informative labels (if present) for the series in *varlist*, or for all series in the dataset if *varlist* is not specified.

With the option `--to-file`, writes to the named file the labels for all series in the dataset, one per line. If no labels are present an error is flagged; if some series have labels and others do not, a blank line is printed for series with no label.

With the option `--from-file`, reads the specified file (which should be plain text) and assigns labels to the series in the dataset, reading one label per line and taking blank lines to indicate blank labels.

The `--delete` option does what you'd expect: it removes all the series labels from the dataset.

Menu path: /Data/Variable labels

lad

```

Arguments:  depvar indepvars
Option:     --vcv (print covariance matrix)

```

Calculates a regression that minimizes the sum of the absolute deviations of the observed from the fitted values of the dependent variable. Coefficient estimates are derived using the Barrodale-Roberts simplex algorithm; a warning is printed if the solution is not unique.

Standard errors are derived using the bootstrap procedure with 500 drawings. The covariance matrix for the parameter estimates, printed when the `--vcv` flag is given, is based on the same bootstrap.

Menu path: /Model/Robust estimation/Least Absolute Deviation

lags

Arguments: [*order* ;] *laglist*

Examples: `lags x y`
`lags 12 ; x y`

Creates new series which are lagged values of each of the series in *varlist*. By default the number of lags created equals the periodicity of the data. For example, if the periodicity is 4 (quarterly), the command `lags x` creates

```
x_1 = x(t-1)
x_2 = x(t-2)
x_3 = x(t-3)
x_4 = x(t-4)
```

The number of lags created can be controlled by the optional first parameter (which, if present, must be followed by a semicolon).

Menu path: /Add/Lags of selected variables

ldiff

Argument: *varlist*

The first difference of the natural log of each series in *varlist* is obtained and the result stored in a new series with the prefix `ld_`. Thus `ldiff x y` creates the new variables

```
ld_x = log(x) - log(x(-1))
ld_y = log(y) - log(y(-1))
```

Menu path: /Add/Log differences of selected variables

leverage

Options: `--save` (save variables)
`--quiet` (don't print results)

Must follow an `ols` command. Calculates the leverage (h , which must lie in the range 0 to 1) for each data point in the sample on which the previous model was estimated. Displays the residual (u) for each observation along with its leverage and a measure of its influence on the estimates, $uh/(1 - h)$. “Leverage points” for which the value of h exceeds $2k/n$ (where k is the number of parameters being estimated and n is the sample size) are flagged with an asterisk. For details on the concepts of leverage and influence see [Davidson and MacKinnon \(1993\)](#), Chapter 2.

DFFITS values are also computed: these are “studentized residuals” (predicted residuals divided by their standard errors) multiplied by $\sqrt{h/(1 - h)}$. For discussions of studentized residuals and DFFITS see chapter 12 of [Maddala \(1992\)](#) or [Belsley et al. \(1980\)](#).

Briefly, a “predicted residual” is the difference between the observed value of the dependent variable at observation t , and the fitted value for observation t obtained from a regression in which that observation is omitted (or a dummy variable with value 1 for observation t alone has been added); the studentized residual is obtained by dividing the predicted residual by its standard error.

If the `--save` flag is given with this command, then the leverage, influence and DFFITS values are added to the current data set. In that context the `--quiet` flag may be used to suppress the printing of results.

After execution, the `$test` accessor returns the cross-validation criterion, which is defined as

$$\sum_{i=1}^n (y_i - \hat{y}_{-i})^2$$

where $\hat{y}_{-i}$ is the forecast error for the i -th observation, after it has been excluded from the sample. The criterion is, hence, the sum of the squared forecasting errors where all n observations but the i -th one are used to predict it (the so-called *leave-one-out* estimator). For a broader discussion of the cross-validation criterion, see Davidson and MacKinnon's *Econometric Theory and Methods*, pages 685–686, and the references therein.

Menu path: Model window, /Tests/Influential observations

levinlin

Arguments: *order series*

Options: `--nc` (test without a constant)
`--ct` (with constant and trend)
`--quiet` (suppress printing of results)

Examples: `levinlin 0 y`
`levinlin 2 y --ct`
`levinlin {2,2,3,3,4,4} y`

Carries out the panel unit-root test described by [Levin et al. \(2002\)](#). The null hypothesis is that all of the individual time series exhibit a unit root, and the alternative is that none of the series has a unit root. (That is, a common AR(1) coefficient is assumed, although in other respects the statistical properties of the series are allowed to vary across individuals.)

By default the test ADF regressions include a constant; to suppress the constant use the `--nc` option, or to add a linear trend use the `--ct` option. (See the [adf](#) command for explanation of ADF regressions.)

The (non-negative) *order* for the test (governing the number of lags of the dependent variable to include in the ADF regressions) may be given in either of two forms. If a scalar value is given, this is applied to all the individuals in the panel. The alternative is to provide a matrix containing a specific lag order for each individual; this must be a vector with as many elements as there are individuals in the current sample range. Such a matrix can be specified by name, or constructed using braces as illustrated in the last example above.

Menu path: /Variable/Unit root tests/Levin-Lin-Chu test

logistic

Arguments: *depvar indepvars*

Options: `--ymax=value` (specify maximum of dependent variable)
`--vcv` (print covariance matrix)

Examples: `logistic y const x`
`logistic y const x --ymax=50`

Logistic regression: carries out an OLS regression using the logistic transformation of the dependent variable,

$$\log \left(\frac{y}{y^* - y} \right)$$

The dependent variable must be strictly positive. If all its values lie between 0 and 1, the default is to use a y^* value (the asymptotic maximum of the dependent variable) of 1; if its values lie between 0 and 100, the default y^* is 100.

If you wish to set a different maximum, use the `--ymax` option. Note that the supplied value must be greater than all of the observed values of the dependent variable.

The fitted values and residuals from the regression are automatically transformed using

$$y = \frac{y^*}{1 + e^{-x}}$$

where x represents either a fitted value or a residual from the OLS regression using the transformed dependent variable. The reported values are therefore comparable with the original dependent variable.

Note that if the dependent variable is binary, you should use the `logit` command instead.

Menu path: /Model/Limited dependent variable/Logistic

logit

Arguments: *depvar indepvars*

Options: `--robust` (robust standard errors)
`--cluster=clustvar` (clustered standard errors)
`--multinomial` (estimate multinomial logit)
`--vcv` (print covariance matrix)
`--verbose` (print details of iterations)
`--p-values` (show p-values instead of slopes)

If the dependent variable is a binary variable (all values are 0 or 1) maximum likelihood estimates of the coefficients on *indepvars* are obtained via the Newton-Raphson method. As the model is nonlinear the slopes depend on the values of the independent variables. By default the slopes with respect to each of the independent variables are calculated (at the means of those variables) and these slopes replace the usual p-values in the regression output. This behavior can be suppressed by giving the `--p-values` option. The chi-square statistic tests the null hypothesis that all coefficients are zero apart from the constant.

By default, standard errors are computed using the negative inverse of the Hessian. If the `--robust` flag is given, then QML or Huber-White standard errors are calculated instead. In this case the estimated covariance matrix is a “sandwich” of the inverse of the estimated Hessian and the outer product of the gradient; see chapter 10 of [Davidson and MacKinnon \(2004\)](#). But if the `--cluster` option is given, then “cluster-robust” standard errors are produced; see the *Gretl User's Guide* for details.

If the dependent variable is not binary but is discrete, then by default it is interpreted as an ordinal response, and Ordered Logit estimates are obtained. However, if the `--multinomial` option is given, the dependent variable is interpreted as an unordered response, and Multinomial Logit estimates are produced. (In either case, if the variable selected as dependent is not discrete an error is flagged.) In the multinomial case, the accessor `$mn1probs` is available after estimation, to get a matrix containing the estimated probabilities of the outcomes at each observation (observations in rows, outcomes in columns).

If you want to use logit for analysis of proportions (where the dependent variable is the proportion of cases having a certain characteristic, at each observation, rather than a 1 or 0 variable indicating whether the characteristic is present or not) you should not use the `logit` command, but rather construct the logit variable, as in

```
series lgt_p = log(p/(1 - p))
```

and use this as the dependent variable in an OLS regression. See chapter 12 of [Ramanathan \(2002\)](#).

Menu path: /Model/Limited dependent variable/Logit

logs

Argument: *varlist*

The natural log of each of the series in *varlist* is obtained and the result stored in a new series with the prefix `l_` (“el” underscore). For example, `logs x y` creates the new variables `l_x = ln(x)` and `l_y = ln(y)`.

Menu path: /Add/Logs of selected variables

loop

Argument: *control*

Options: `--progressive` (enable special forms of certain commands)
`--verbose` (report details of `genr` commands)
`--quiet` (do not report number of iterations performed)

Examples: `loop 1000`
`loop 1000 --progressive`
`loop while essdiff > .00001`
`loop i=1991..2000`
`loop for (r=-.99; r<=.99; r+=.01)`
`loop foreach i xlist`

This command opens a special mode in which the program accepts commands to be executed repeatedly. You exit the mode of entering loop commands with `endloop`: at this point the stacked commands are executed.

The parameter *control* may take any of five forms, as shown in the examples: an integer number of times to repeat the commands within the loop; “while” plus a boolean condition; a range of integer values for index variable; “for” plus three expressions in parentheses, separated by semicolons (which emulates the `for` statement in the C programming language); or “foreach” plus an index variable and a list.

See the *Gretl User's Guide* for further details and examples. The effect of the `--progressive` option (which is designed for use in Monte Carlo simulations) is explained there. Not all `gretl` commands may be used within a loop; the commands available in this context are also set out there.

mahal

Argument: *varlist*

Options: `--quiet` (don't print anything)
`--save` (add distances to the dataset)
`--vcv` (print covariance matrix)

Computes the Mahalanobis distances between the series in *varlist*. The Mahalanobis distance is the distance between two points in a k -dimensional space, scaled by the statistical variation in each dimension of the space. For example, if p and q are two observations on a set of k variables with covariance matrix C , then the Mahalanobis distance between the observations is given by

$$\sqrt{(p - q)'C^{-1}(p - q)}$$

where $(p - q)$ is a k -vector. This reduces to Euclidean distance if the covariance matrix is the identity matrix.

The space for which distances are computed is defined by the selected variables. For each observation in the current sample range, the distance is computed between the observation and the centroid of the selected variables. This distance is the multidimensional counterpart of a standard z-score, and can be used to judge whether a given observation “belongs” with a group of other observations.

If the `--vcv` option is given, the covariance matrix and its inverse are printed. If the `--save` option is given, the distances are saved to the dataset under the name `mdist` (or `mdist1`, `mdist2` and so on if there is already a variable of that name).

Menu path: /View/Mahalanobis distances

makepkg

Argument: *filename*

Options: `--index` (write auxiliary index file)
`--translations` (write auxiliary strings file)

Supports creation of a gretl function package via the command line. The mode of operation of this command depends on the extension of *filename*, which must be either `.gfn` or `.zip`.

Gfn mode

Writes a gfn file. It is assumed that a package specification file, with the same basename as *filename* but with the extension `.spec`, is accessible, along with any auxiliary files that it references. It is also assumed that all the functions to be packaged have been read into memory.

Zip mode

Writes a zip package file (gfn plus other materials). If a gfn file of the same basename as *filename* is found, it forms the basis of the zip package. If no gfn file is found, the program first attempts to build the gfn, as described above.

Gfn options

The option flags support the writing of auxiliary files, intended for use with gretl “addons”. The index file is a short XML document containing basic information about the package; it has the same basename as the package and the extension `.xml`. The translations file contains strings from the package that may be suitable for translation, in C format; for package `foo` this file is named `foo-i18n.c`. These files are not produced if the command is operating in zip mode and a pre-existing gfn file is used.

For details on all of this, see the the *Gretl Function Package Guide*.

Menu path: /Tools/Function packages/New package

markers

Variants: `markers --to-file=filename`
`markers --from-file=filename`
`markers --delete`

With the option `--to-file`, writes to the named file the observation marker strings from the current dataset, one per line. If no such strings are present an error is flagged.

With the option `--from-file`, reads the specified file (which should be plain text) and assigns observation markers to the rows in the dataset, reading one marker per line. In general there should be at least as many markers in the file as observations in the dataset, but if the dataset is a panel it

is also acceptable if the number of markers in the file matches the number of cross-sectional units (in which case the markers are repeated for each time period.)

The `--delete` option does what you'd expect: it removes the observation marker strings from the dataset.

Menu path: /Data/Observation markers

meantest

Arguments: *series1 series2*

Option: `--unequal-vars` (assume variances are unequal)

Calculates the *t* statistic for the null hypothesis that the population means are equal for the variables *series1* and *series2*, and shows its p-value.

By default the test statistic is calculated on the assumption that the variances are equal for the two variables. With the `--unequal-vars` option the variances are assumed to be different; in this case the degrees of freedom for the test statistic are approximated as per [Satterthwaite \(1946\)](#).

Menu path: /Tools/Test statistic calculator

mle

Arguments: *log-likelihood function* [*derivatives*]

Options: `--quiet` (don't show estimated model)

`--vcv` (print covariance matrix)

`--hessian` (base covariance matrix on the Hessian)

`--robust` (QML covariance matrix)

`--verbose` (print details of iterations)

`--no-gradient-check` (see below)

`--lbfgs` (use L-BFGS-B instead of regular BFGS)

Example: `weibull.inp`

Performs Maximum Likelihood (ML) estimation using either the BFGS (Broyden, Fletcher, Goldfarb, Shanno) algorithm or Newton's method. The user must specify the log-likelihood function. The parameters of this function must be declared and given starting values (using the `genr` command) prior to estimation. Optionally, the user may specify the derivatives of the log-likelihood function with respect to each of the parameters; if analytical derivatives are not supplied, a numerical approximation is computed.

Simple example: Suppose we have a series *X* with values 0 or 1 and we wish to obtain the maximum likelihood estimate of the probability, *p*, that *X* = 1. (In this simple case we can guess in advance that the ML estimate of *p* will simply equal the proportion of *X*s equal to 1 in the sample.)

The parameter *p* must first be added to the dataset and given an initial value. For example, `scalar p = 0.5`.

We then construct the MLE command block:

```
mle loglik = X*log(p) + (1-X)*log(1-p)
    deriv p = X/p - (1-X)/(1-p)
end mle
```

The first line above specifies the log-likelihood function. It starts with the keyword `mle`, then a dependent variable is specified and an expression for the log-likelihood is given (using the same syntax as in the `genr` command). The next line (which is optional) starts with the keyword `deriv` and supplies the derivative of the log-likelihood function with respect to the parameter *p*. If no

derivatives are given, you should include a statement using the keyword `params` which identifies the free parameters: these are listed on one line, separated by spaces and can be either scalars, or vectors, or any combination of the two. For example, the above could be changed to:

```
mle loglik = X*log(p) + (1-X)*log(1-p)
  params p
end mle
```

in which case numerical derivatives would be used.

Note that any option flags should be appended to the ending line of the MLE block.

By default, estimated standard errors are based on the Outer Product of the Gradient. If the `--hessian` option is given, they are instead based on the negative inverse of the Hessian (which is approximated numerically). If the `--robust` option is given, a QML estimator is used (namely, a sandwich of the negative inverse of the Hessian and the covariance matrix of the gradient).

If you supply analytical derivatives, by default gretl runs a numerical check on their plausibility. Occasionally this may produce false positives, instances where correct derivatives appear to be wrong and estimation is refused. To counter this, or to achieve a little extra speed, you can give the option `--no-gradient-check`. Obviously, you should do this only if you are quite confident that the gradient you have specified is right.

For a much more in-depth description of `mle`, please refer to the *Gretl User's Guide*.

Menu path: /Model/Maximum likelihood

modeltab

Variants: `modeltab add`
`modeltab show`
`modeltab free`
`modeltab --output=filename`

Manipulates the gretl “model table”. See the *Gretl User's Guide* for details. The sub-commands have the following effects: `add` adds the last model estimated to the model table, if possible; `show` displays the model table in a window; and `free` clears the table.

To call for printing of the model table, use the flag `--output=` plus a filename. If the filename has the suffix `“.tex”`, the output will be in $\text{\TeX}$ format; if the suffix is `“.rtf”` the output will be RTF; otherwise it will be plain text. In the case of $\text{\TeX}$ output the default is to produce a “fragment”, suitable for inclusion in a document; if you want a stand-alone document instead, use the `--complete` option, for example

```
modeltab --output="myfile.tex" --complete
```

Menu path: Session icon window, Model table icon

modprint

Arguments: `coeffmat names [ addstats ]`

Prints the coefficient table and optional additional statistics for a model estimated “by hand”. Mainly useful for user-written functions.

The argument `coeffmat` should be a k by 2 matrix containing k coefficients and k associated standard errors, and `names` should be a string containing at least k names for the coefficients, separated by commas or spaces. (The `names` argument may be either the name of a string variable or a literal string, enclosed in double quotes.)

The optional argument *addstats* is a vector containing p additional statistics to be printed under the coefficient table. If this argument is given, then *names* should contain $k + p$ comma-separated strings, the additional p strings to be associated with the additional statistics.

modtest

Argument: [*order*]

Options: --normality (normality of residual)
 --logs (non-linearity, logs)
 --autocorr (serial correlation)
 --arch (ARCH)
 --squares (non-linearity, squares)
 --white (heteroskedasticity, White's test)
 --white-nocross (White's test, squares only)
 --breusch-pagan (heteroskedasticity, Breusch-Pagan)
 --robust (robust variance estimate for Breusch-Pagan)
 --panel (heteroskedasticity, groupwise)
 --comfac (common factor restriction, AR1 models only)
 --quiet (don't print details)
 --silent (don't print anything)

Must immediately follow an estimation command. Depending on the option given, this command carries out one of the following: the Doornik-Hansen test for the normality of the error term; a Lagrange Multiplier test for nonlinearity (logs or squares); White's test (with or without cross-products) or the Breusch-Pagan test ([Breusch and Pagan \(1979\)](#)) for heteroskedasticity; the LMF test for serial correlation ([Kiviet, 1986](#)); a test for ARCH (Autoregressive Conditional Heteroskedasticity; see also the *arch* command); or a test of the common factor restriction implied by AR(1) estimation. With the exception of the normality and common factor test most of the options are only available for models estimated via OLS, but see below for details regarding two-stage least squares.

The optional *order* argument is relevant only in case the --autocorr or --arch options are selected. The default is to run these tests using a lag order equal to the periodicity of the data, but this can be adjusted by supplying a specific lag order.

The --robust option applies only when the Breusch-Pagan test is selected; its effect is to use the robust variance estimator proposed by [Koenker \(1981\)](#), making the test less sensitive to the assumption of normality.

The --panel option is available only when the model is estimated on panel data: in this case a test for groupwise heteroskedasticity is performed (that is, for a differing error variance across the cross-sectional units).

The --comfac option is available only when the model is estimated via an AR(1) method such as Hildreth-Lu. The auxiliary regression takes the form of a relatively unrestricted dynamic model, which is used to test the common factor restriction implicit in the AR(1) specification.

By default, the program prints the auxiliary regression on which the test statistic is based, where applicable. This may be suppressed by using the --quiet flag (minimal printed output) or the --silent flag (no printed output). The test statistic and its p-value may be retrieved using the accessors \$test and \$pvalue respectively.

When a model has been estimated by two-stage least squares (see [tsls](#)), the LM principle breaks down and gretl offers some equivalents: the --autocorr option computes Godfrey's test for autocorrelation ([Godfrey, 1994](#)) while the --white option yields the HET1 heteroskedasticity test ([Pesaran and Taylor, 1999](#)).

Menu path: Model window, /Tests

mpols

Arguments: *depvar indepvars*

Options: --vcv (print covariance matrix)
 --simple-print (do not print auxiliary statistics)
 --quiet (suppress printing of results)

Computes OLS estimates for the specified model using multiple precision floating-point arithmetic, with the help of the Gnu Multiple Precision (GMP) library. By default 256 bits of precision are used for the calculations, but this can be increased via the environment variable GRETL_MP_BITS. For example, when using the bash shell one could issue the following command, before starting gretl, to set a precision of 1024 bits.

```
export GRETL_MP_BITS=1024
```

A rather arcane option is available for this command (primarily for testing purposes): if the *indepvars* list is followed by a semicolon and a further list of numbers, those numbers are taken as powers of *x* to be added to the regression, where *x* is the last variable in *indepvars*. These additional terms are computed and stored in multiple precision. In the following example *y* is regressed on *x* and the second, third and fourth powers of *x*:

```
mpols y 0 x ; 2 3 4
```

Menu path: /Model/Other linear models/High precision OLS

negbin

Arguments: *depvar indepvars* [; *offset*]

Options: --model1 (use NegBin 1 model)
 --robust (QML covariance matrix)
 --cluster=*clustvar* (see [logit](#) for explanation)
 --opg (see below)
 --vcv (print covariance matrix)
 --verbose (print details of iterations)

Estimates a Negative Binomial model. The dependent variable is taken to represent a count of the occurrence of events of some sort, and must have only non-negative integer values. By default the model NegBin 2 is used, in which the conditional variance of the count is given by $\mu(1 + \alpha\mu)$, where μ denotes the conditional mean. But if the --model1 option is given the conditional variance is $\mu(1 + \alpha)$.

The optional *offset* series works in the same way as for the [poisson](#) command. The Poisson model is a restricted form of the Negative Binomial in which $\alpha = 0$ by construction.

By default, standard errors are computed using a numerical approximation to the Hessian at convergence. But if the --opg option is given the covariance matrix is based on the Outer Product of the Gradient (OPG), or if the --robust option is given QML standard errors are calculated, using a “sandwich” of the inverse of the Hessian and the OPG.

Menu path: /Model/Limited dependent variable/Count data...

nls

Arguments: *function* [*derivatives*]
Options: --quiet (don't show estimated model)
 --robust (robust standard errors)
 --vcv (print covariance matrix)
 --verbose (print details of iterations)
Example: wg_nls.inp

Performs Nonlinear Least Squares (NLS) estimation using a modified version of the Levenberg-Marquardt algorithm. You must supply a function specification. The parameters of this function must be declared and given starting values (using the `genr` command) prior to estimation. Optionally, you may specify the derivatives of the regression function with respect to each of the parameters. If you do not supply derivatives you should instead give a list of the parameters to be estimated (separated by spaces or commas), preceded by the keyword `params`. In the latter case a numerical approximation to the Jacobian is computed.

It is easiest to show what is required by example. The following is a complete script to estimate the nonlinear consumption function set out in William Greene's *Econometric Analysis* (Chapter 11 of the 4th edition, or Chapter 9 of the 5th). The numbers to the left of the lines are for reference and are not part of the commands. Note that any option flags, such as `--vcv` for printing the covariance matrix of the parameter estimates, should be appended to the final command, `end nls`.

```

1  open greene11_3.gdt
2  ols C 0 Y
3  scalar a = $coeff(0)
4  scalar b = $coeff(Y)
5  scalar g = 1.0
6  nls C = a + b * Y^g
7  deriv a = 1
8  deriv b = Y^g
9  deriv g = b * Y^g * log(Y)
10 end nls --vcv
```

It is often convenient to initialize the parameters by reference to a related linear model; that is accomplished here on lines 2 to 5. The parameters alpha, beta and gamma could be set to any initial values (not necessarily based on a model estimated with OLS), although convergence of the NLS procedure is not guaranteed for an arbitrary starting point.

The actual NLS commands occupy lines 6 to 10. On line 6 the `nls` command is given: a dependent variable is specified, followed by an equals sign, followed by a function specification. The syntax for the expression on the right is the same as that for the `genr` command. The next three lines specify the derivatives of the regression function with respect to each of the parameters in turn. Each line begins with the keyword `deriv`, gives the name of a parameter, an equals sign, and an expression whereby the derivative can be calculated (again, the syntax here is the same as for `genr`). As an alternative to supplying numerical derivatives, you could substitute the following for lines 7 to 9:

```
params a b g
```

Line 10, `end nls`, completes the command and calls for estimation. Any options should be appended to this line.

For further details on NLS estimation please see the *Gretl User's Guide*.

Menu path: /Model/Nonlinear Least Squares

normtest

Argument: *series*
 Options: --dhansen (Doornik–Hansen test, the default)
 --swilk (Shapiro–Wilk test)
 --lillie (Lilliefors test)
 --jbera (Jarque–Bera test)
 --all (do all tests)
 --quiet (suppress printed output)

Carries out a test for normality for the given *series*. The specific test is controlled by the option flags (but if no flag is given, the Doornik–Hansen test is performed). Note: the Doornik–Hansen and Shapiro–Wilk tests are recommended over the others, on account of their superior small-sample properties.

The test statistic and its p-value may be retrieved using the accessors `$test` and `$pvalue`. Please note that if the `--all` option is given, the result recorded is that from the Doornik–Hansen test.

Menu path: /Variable/Normality test

nulldata

Argument: *series-length*
 Option: --preserve (preserve matrices)
 Example: nulldata 500

Establishes a “blank” data set, containing only a constant and an index variable, with periodicity 1 and the specified number of observations. This may be used for simulation purposes: some of the `genr` commands (e.g. `genr uniform()`, `genr normal()`) will generate dummy data from scratch to fill out the data set. This command may be useful in conjunction with `loop`. See also the “seed” option to the `set` command.

By default, this command cleans out all data in gretl’s current workspace. If you give the `--preserve` option, however, any currently defined matrices are retained.

Menu path: /File/New data set

ols

Arguments: *depvar indepvars*
 Options: --vcv (print covariance matrix)
 --robust (robust standard errors)
 --cluster=*clustvar* (clustered standard errors)
 --jackknife (see below)
 --simple-print (do not print auxiliary statistics)
 --quiet (suppress printing of results)
 --anova (print an ANOVA table)
 --no-df-corr (suppress degrees of freedom correction)
 --print-final (see below)
 Examples: ols 1 0 2 4 6 7
 ols y 0 x1 x2 x3 --vcv
 ols y 0 x1 x2 x3 --quiet

Computes ordinary least squares (OLS) estimates with *depvar* as the dependent variable and *indepvars* as the list of independent variables. Variables may be specified by name or number; use the number zero for a constant term.

Besides coefficient estimates and standard errors, the program also prints p-values for t (two-tailed) and F -statistics. A p-value below 0.01 indicates statistical significance at the 1 percent level and is marked with ***. ** indicates significance between 1 and 5 percent and * indicates significance between the 5 and 10 percent levels. Model selection statistics (the Akaike Information Criterion or AIC and Schwarz's Bayesian Information Criterion) are also printed. The formula used for the AIC is that given by Akaike (1974), namely minus two times the maximized log-likelihood plus two times the number of parameters estimated.

If the option `--no-df-corr` is given, the usual degrees of freedom correction is not applied when calculating the estimated error variance (and hence also the standard errors of the parameter estimates).

The option `--print-final` is applicable only in the context of a [loop](#). It arranges for the regression to be run silently on all but the final iteration of the loop. See the *Gretl User's Guide* for details.

Various internal variables may be retrieved following estimation. For example

```
series uh = $uhat
```

saves the residuals under the name `uh`. See the “accessors” section of the *gretl* function reference for details.

The specific formula (“HC” version) used for generating robust standard errors when the `--robust` option is given can be adjusted via the [set](#) command. The `--jackknife` option has the effect of selecting an `hc_version` of 3a. The `--cluster` overrides the selection of HC version, and produces robust standard errors by grouping the observations by the distinct values of *clustvar*; see the *Gretl User's Guide* for details.

Menu path: /Model/Ordinary Least Squares

Other access: Beta-hat button on toolbar

omit

Argument: *varlist*

Options: `--test-only` (don't replace the current model)
`--chi-square` (give chi-square form of Wald test)
`--quiet` (print only the basic test result)
`--silent` (don't print anything)
`--vcv` (print covariance matrix for reduced model)
`--auto[=alpha]` (sequential elimination, see below)

Examples: `omit 5 7 9`
`omit seasonals --quiet`
`omit --auto`
`omit --auto=0.05`

This command must follow an estimation command. It calculates a Wald test for the joint significance of the variables in *varlist*, which should be a subset of the independent variables in the model last estimated. The results of the test may be retrieved using the accessors `$test` and `$pvalue`.

By default the restricted model is estimated and it replaces the original as the “current model” for the purposes of, for example, retrieving the residuals as `$uhat` or doing further tests. This behavior may be suppressed via the `--test-only` option.

By default the F -form of the Wald test is recorded; the `--chi-square` option may be used to record the chi-square form instead.

If the restricted model is both estimated and printed, the `--vcv` option has the effect of printing its covariance matrix, otherwise this option is ignored.

Alternatively, if the `--auto` flag is given, sequential elimination is performed: at each step the variable with the highest p-value is omitted, until all remaining variables have a p-value no greater than some cutoff. The default cutoff is 10 percent (two-sided); this can be adjusted by appending “=” and a value between 0 and 1 (with no spaces), as in the fourth example above. If *varlist* is given this process is confined to the listed variables, otherwise all variables are treated as candidates for omission. Note that the `--auto` and `--test-only` options cannot be combined.

Menu path: Model window, /Tests/Omit variables

open

Argument: *filename*

Options: `--quiet` (don't print list of series)
`--preserve` (preserve variables other than series)
`--frompkg=pkgname` (see below)
`--www` (use a database on the gretl server)
 See below for additional specialized options

Examples: `open data4-1`
`open voter.dta`
`open fedbog --www`

Opens a data file or database. If a data file is already open, it is replaced by the newly opened one. To add data to the current dataset, see [append](#) and (for greater flexibility) [join](#).

If a full path is not given, the program will search some relevant paths to try to find the file. If no filename suffix is given (as in the first example above), gretl assumes a native datafile with suffix `.gdt`. Based on the name of the file and various heuristics, gretl will try to detect the format of the data file (native, plain text, CSV, MS Excel, Stata, SPSS, etc.).

If the `--frompkg` option is used, gretl will look for the specified data file in the subdirectory associated with the function package specified by *pkgname*.

If the *filename* argument takes the form of a URI starting with `http://`, then gretl will attempt to download the indicated data file before opening it.

By default, opening a new data file clears the current gretl session, which includes deletion of all named variables, including matrices, scalars and strings. If you wish to keep your currently defined variables (other than series, which are necessarily cleared out), use the `--preserve` option.

The `open` command can also be used to open a database (gretl, RATS 4.0 or PcGive) for reading. In that case it should be followed by the [data](#) command to extract particular series from the database. If the `www` option is given, the program will try to access a database of the given name on the gretl server — for instance the Federal Reserve interest rates database in the third example above.

When opening a spreadsheet file (Gnumeric, Open Document or MS Excel), you may give up to three additional parameters following the filename. First, you can select a particular worksheet within the file. This is done either by giving its (1-based) number, using the syntax, e.g., `--sheet=2`, or, if you know the name of the sheet, by giving the name in double quotes, as in `--sheet="MacroData"`. The default is to read the first worksheet. You can also specify a column and/or row offset into the worksheet via, e.g.,

```
--coloffset=3 --rowoffset=2
```

which would cause gretl to ignore the first 3 columns and the first 2 rows. The default is an offset of 0 in both dimensions, that is, to start reading at the top-left cell.

With plain text files, gretl generally expects to find the data columns delimited in some standard manner. But there is also a special facility for reading “fixed format” files, in which there are no

delimiters but there is a known specification of the form, e.g., “variable k occupies 8 columns starting at column 24”. To read such files, you should append a string `--fixed-cols=colspec`, where *colspec* is composed of comma-separated integers. These integers are interpreted as a set of pairs. The first element of each pair denotes a starting column, measured in bytes from the beginning of the line with 1 indicating the first byte; and the second element indicates how many bytes should be read for the given field. So, for example, if you say

```
open fixed.txt --fixed-cols=1,6,20,3
```

then for variable 1 gretl will read 6 bytes starting at column 1; and for variable 2, 3 bytes starting at column 20. Lines that are blank, or that begin with #, are ignored, but otherwise the column-reading template is applied, and if anything other than a valid numerical value is found an error is flagged. If the data are read successfully, the variables will be named `v1`, `v2`, etc. It's up to the user to provide meaningful names and/or descriptions using the commands [rename](#) and/or [setinfo](#).

Menu path: /File/Open data

Other access: Drag a data file onto gretl's main window

orthdev

Argument: *varlist*

Applicable with panel data only. A series of forward orthogonal deviations is obtained for each variable in *varlist* and stored in a new variable with the prefix `o_`. Thus `orthdev x y` creates the new variables `o_x` and `o_y`.

The values are stored one step ahead of their true temporal location (that is, `o_x` at observation t holds the deviation that, strictly speaking, belongs at $t - 1$). This is for compatibility with first differences: one loses the first observation in each time series, not the last.

outfile

Variants: `outfile filename option`
`outfile --close`

Options: `--append` (append to file)
`--write` (overwrite file)
`--quiet` (see below)

Examples: `outfile regress.txt --write`
`outfile --close`

Diverts output to *filename*, until further notice. Use the flag `--append` to append output to an existing file or `--write` to start a new file (or overwrite an existing one). Only one file can be opened in this way at any given time.

The `--close` flag is used to close an output file that was previously opened as above. Output will then revert to the default stream. Note that since only one file can be opened via `outfile` at any given time, no filename argument need (nor should) be supplied with this variant of the command.

In the first example command above, the file `regress.txt` is opened for writing, and in the second it is closed. This would make sense as a sequence only if some commands were issued before the `--close`. For example if an `ols` command intervened, its output would go to `regress.txt` rather than the screen.

Three special variants on the above are available. If you give the keyword `null` in place of a real filename along with the `--write` option, the effect is to suppress all printed output until redirection is ended. If either of the keywords `stdout` or `stderr` are given in place of a regular filename the effect is to redirect output to standard output or standard error output respectively.

The `--quiet` option is for use with `--write` or `--append`: its effect is to turn off the echoing of commands and the printing of auxiliary messages while output is redirected. It is equivalent to doing

```
set echo off
set messages off
```

except that when redirection is ended the original values of the `echo` and `messages` variables are restored.

panel

Arguments: *depvar indepvars*

Options: `--vcv` (print covariance matrix)
`--fixed-effects` (estimate with group fixed effects)
`--random-effects` (random effects or GLS model)
`--nerlove` (use the Nerlove transformation)
`--between` (estimate the between-groups model)
`--robust` (robust standard errors; see below)
`--time-dummies` (include time dummy variables)
`--unit-weights` (weighted least squares)
`--iterate` (iterative estimation)
`--matrix-diff` (use matrix-difference method for Hausman test)
`--quiet` (less verbose output)
`--verbose` (more verbose output)

Estimates a panel model. By default the fixed effects estimator is used; this is implemented by subtracting the group or unit means from the original data.

If the `--random-effects` flag is given, random effects estimates are computed, by default using the method of [Swamy and Arora \(1972\)](#). In this case (only) the option `--matrix-diff` forces use of the matrix-difference method (as opposed to the regression method) for carrying out the Hausman test for the consistency of the random effects estimator. Also specific to the random effects estimator is the `--nerlove` flag, which selects the method of [Nerlove \(1971\)](#) as opposed to Swamy and Arora.

Alternatively, if the `--unit-weights` flag is given, the model is estimated via weighted least squares, with the weights based on the residual variance for the respective cross-sectional units in the sample. In this case (only) the `--iterate` flag may be added to produce iterative estimates: if the iteration converges, the resulting estimates are Maximum Likelihood.

As a further alternative, if the `--between` flag is given, the between-groups model is estimated (that is, an OLS regression using the group means).

The `--robust` option is available only for fixed effects models. The default variant is the Arellano HAC estimator, but Beck-Katz “Panel Corrected Standard Errors” can be selected via the command `set pcse on`. When the robust option is specified the joint F test on the fixed effects is performed using the robust method of [Welch \(1951\)](#).

For more details on panel estimation, please see the *Gretl User's Guide*.

Menu path: /Model/Panel

pca

Argument: *varlist*
 Options: --covariance (use the covariance matrix)
 --save[=*n*] (save major components)
 --save-all (save all components)
 --quiet (don't print results)

Principal Components Analysis. Unless the --quiet option is given, prints the eigenvalues of the correlation matrix (or the covariance matrix if the --covariance option is given) for the variables in *varlist*, along with the proportion of the joint variance accounted for by each component. Also prints the corresponding eigenvectors (or “component loadings”).

If you give the --save-all option then all components are saved to the dataset as series, with names PC1, PC2 and so on. These artificial variables are formed as the sum of (component loading) times (standardized X_i), where X_i denotes the i th variable in *varlist*.

If you give the --save option without a parameter value, components with eigenvalues greater than the mean (which means greater than 1.0 if the analysis is based on the correlation matrix) are saved to the dataset as described above. If you provide a value for *n* with this option then the most important *n* components are saved.

See also the [princomp](#) function.

Menu path: /View/Principal components

Other access: Main window pop-up (multiple selection)

pergm

Arguments: *series* [*bandwidth*]
 Options: --bartlett (use Bartlett lag window)
 --log (use log scale)
 --radians (show frequency in radians)
 --degrees (show frequency in degrees)
 --plot=*mode-or-filename* (see below)

Computes and displays the spectrum of the specified series. By default the sample periodogram is given, but optionally a Bartlett lag window is used in estimating the spectrum (see, for example, Greene's *Econometric Analysis* for a discussion of this). The default width of the Bartlett window is twice the square root of the sample size but this can be set manually using the *bandwidth* parameter, up to a maximum of half the sample size.

If the --log option is given the spectrum is represented on a logarithmic scale.

The (mutually exclusive) options --radians and --degrees influence the appearance of the frequency axis when the periodogram is graphed. By default the frequency is scaled by the number of periods in the sample, but these options cause the axis to be labeled from 0 to π radians or from 0 to 180°, respectively.

By default, if the program is not in batch mode a plot of the periodogram is shown. This can be adjusted via the --plot option. The acceptable parameters to this option are none (to suppress the plot); display (to display a plot even when in batch mode); or a file name. The effect of providing a file name is as described for the --output option of the [gnuplot](#) command.

Menu path: /Variable/Periodogram

Other access: Main window pop-up menu (single selection)

plot

Argument: *data*

Options: --with-lines[=*varspec*] (use lines, not points)
 --with-lp[=*varspec*] (use lines and points)
 --with-impulses[=*varspec*] (use vertical lines)
 --time-series (plot against time)
 --single-yaxis (force use of just one y-axis)
 --dummy (see below)
 --fit=*fitspec* (see below)
 --output=*filename* (send output to specified file)

The `plot` block provides an alternative to the [gnuplot](#) command which may be more convenient when you are producing an elaborate plot (with several options and/or `gnuplot` commands to be inserted into the plot file).

A `plot` block starts with the command-word `plot` followed by the required argument, *data*, which specifies the data to be plotted: this should be the name of a list, a matrix, or a single series.

If a list or matrix is given, the last element (list) or column (matrix) is assumed to be the *x*-axis variable and the other(s) the *y*-axis variable(s), unless the `--time-series` option is given in which case all the specified data go on the *y* axis.

The option of supplying a single series name is restricted to time-series data, in which case it is assumed that a time-series plot is wanted; otherwise an error is flagged.

The starting line may be prefixed with the “*savename* <-” apparatus to save a plot as an icon in the GUI program. The block ends with `end plot`.

Inside the block you have zero or more lines of these types, identified by an initial keyword:

- **option**: specify a single option.
- **options**: specify multiple options on a single line, separated by spaces.
- **literal**: a command to be passed to `gnuplot` literally.
- **printf**: a `printf` statement whose result will be passed to `gnuplot` literally.

Note that when you specify an option using the `option` or `options` keywords, it is not necessary to supply the customary double-dash before the option specifier. For details on the effects of the various options please see [gnuplot](#).

The intended use of the `plot` block is best illustrated by example:

```
string title = "My title"
string xname = "My x-variable"
plot plotmat
    options with-lines fit=none
    literal set linetype 3 lc rgb "#0000ff"
    literal set nokey
    printf "set title \"%s\"", title
    printf "set xlabel \"%s\"", xname
end plot --output=display
```

This example assumes that `plotmat` is the name of a matrix with at least 2 columns (or a list with at least two members). Note that it is considered good practice to place the `--output` option (only) on the last line of the block.

poisson

Arguments: *depvar indepvars* [; *offset*]

Options: --robust (robust standard errors)
 --cluster=*clustvar* (see [logit](#) for explanation)
 --vcv (print covariance matrix)
 --verbose (print details of iterations)

Examples: poisson y 0 x1 x2
 poisson y 0 x1 x2 ; S

Estimates a poisson regression. The dependent variable is taken to represent the occurrence of events of some sort, and must take on only non-negative integer values.

If a discrete random variable Y follows the Poisson distribution, then

$$\Pr(Y = y) = \frac{e^{-\nu} \nu^y}{y!}$$

for $y = 0, 1, 2, \dots$. The mean and variance of the distribution are both equal to ν . In the Poisson regression model, the parameter ν is represented as a function of one or more independent variables. The most common version (and the only one supported by gretl) has

$$\nu = \exp(\beta_0 + \beta_1 x_1 + \beta_2 x_2 + \dots)$$

or in other words the log of ν is a linear function of the independent variables.

Optionally, you may add an “offset” variable to the specification. This is a scale variable, the log of which is added to the linear regression function (implicitly, with a coefficient of 1.0). This makes sense if you expect the number of occurrences of the event in question to be proportional, other things equal, to some known factor. For example, the number of traffic accidents might be supposed to be proportional to traffic volume, other things equal, and in that case traffic volume could be specified as an “offset” in a Poisson model of the accident rate. The offset variable must be strictly positive.

By default, standard errors are computed using the negative inverse of the Hessian. If the --robust flag is given, then QML or Huber-White standard errors are calculated instead. In this case the estimated covariance matrix is a “sandwich” of the inverse of the estimated Hessian and the outer product of the gradient.

See also [negbin](#).

Menu path: /Model/Limited dependent variable/Count data...

print

Variants: print *varlist*
 print
 print *object-name*
 print *string-literal*

Options: --byobs (by observations)
 --no-dates (use simple observation numbers)

Examples: print x1 x2 --byobs
 print my_matrix
 print "This is a string"

If *varlist* is given, prints the values of the specified series, or if no argument is given, prints the values of all series in the current dataset. If the --byobs flag is added the data are printed by observation, otherwise they are printed by variable. When printing by observation, the default is to

show the date (with time-series data) or the observation marker string (if any) at the start of each line. The `--no-dates` option suppresses the printing of dates or markers; a simple observation number is shown instead.

Besides printing series, you may give the name of a (single) matrix or scalar variable for printing. Or you may give a literal string argument, enclosed in double quotes, to be printed as is. In these case the option flags are not applicable.

Note that you can gain greater control over the printing format (and so, for example, expose a greater number of digits than are shown by default) by using [printf](#).

Menu path: /Data/Display values

printf

Arguments: *format* , *args*

Prints scalar values, series, matrices, or strings under the control of a format string (providing a subset of the `printf()` statement in the C programming language). Recognized numeric formats are `%e`, `%E`, `%f`, `%g`, `%G` and `%d`, in each case with the various modifiers available in C. Examples: the format `%.10g` prints a value to 10 significant figures; `%12.6f` prints a value to 6 decimal places, with a width of 12 characters. The format `%s` should be used for strings.

The format string itself must be enclosed in double quotes. The values to be printed must follow the format string, separated by commas. These values should take the form of either (a) the names of variables, (b) expressions that are valid for the `genr` command, or (c) the special functions `varname()` or `date()`. The following example prints the values of two variables plus that of a calculated expression:

```
ols 1 0 2 3
scalar b = $coeff[2]
scalar se_b = $stderr[2]
printf "b = %.8g, standard error %.8g, t = %.4f\n",
      b, se_b, b/se_b
```

The next lines illustrate the use of the `varname` and `date` functions, which respectively print the name of a variable, given its ID number, and a date string, given a 1-based observation number.

```
printf "The name of variable %d is %s\n", i, varname(i)
printf "The date of observation %d is %s\n", j, date(j)
```

If a matrix argument is given in association with a numeric format, the entire matrix is printed using the specified format for each element. The same applies to series, except that the range of values printed is governed by the current sample setting.

The maximum length of a format string is 127 characters. The escape sequences `\n` (newline), `\t` (tab), `\v` (vertical tab) and `\\` (literal backslash) are recognized. To print a literal percent sign, use `%%`.

As in C, numerical values that form part of the format (width and or precision) may be given directly as numbers, as in `%10.4f`, or they may be given as variables. In the latter case, one puts asterisks into the format string and supplies corresponding arguments in order. For example,

```
scalar width = 12
scalar precision = 6
printf "x = %*.*f\n", width, precision, x
```

probit

Arguments: *depvar indepvars*

Options: `--robust` (robust standard errors)
`--cluster=clustvar` (see [logit](#) for explanation)
`--vcv` (print covariance matrix)
`--verbose` (print details of iterations)
`--p-values` (show p-values instead of slopes)
`--random-effects` (estimates a random effects panel probit model)
`--quadpoints=k` (number of quadrature points for RE estimation)

Examples: `ooballot.inp`, `oprobit.inp`, `reprobit.inp`

If the dependent variable is a binary variable (all values are 0 or 1) maximum likelihood estimates of the coefficients on *indepvars* are obtained via the Newton-Raphson method. As the model is nonlinear the slopes depend on the values of the independent variables. By default the slopes with respect to each of the independent variables are calculated (at the means of those variables) and these slopes replace the usual p-values in the regression output. This behavior can be suppressed by giving the `--p-values` option. The chi-square statistic tests the null hypothesis that all coefficients are zero apart from the constant.

By default, standard errors are computed using the negative inverse of the Hessian. If the `--robust` flag is given, then QML or Huber-White standard errors are calculated instead. In this case the estimated covariance matrix is a “sandwich” of the inverse of the estimated Hessian and the outer product of the gradient. See chapter 10 of Davidson and MacKinnon for details.

If the dependent variable is not binary but is discrete, then Ordered Probit estimates are obtained. (If the variable selected as dependent is not discrete, an error is flagged.)

Probit for panel data

With the `--random-effects` option, the error term is assumed to be composed of two normally distributed components: one time-invariant term that is specific to the cross-sectional unit or “individual” (and is known as the individual effect); and one term that is specific to the particular observation.

Evaluation of the likelihood for this model involves the use of Gauss-Hermite quadrature for approximating the value of expectations of functions of normal variates. The number of quadrature points used can be chosen through the `--quadpoints` option (the default is 32). Using more points will increase the accuracy of the results, but at the cost of longer compute time; with many quadrature points and a large dataset estimation may be quite time consuming.

Besides the usual parameter estimates (and associated statistics) relating to the included regressors, certain additional information is presented on estimation of this sort of model:

- `lnsigma2`: the maximum likelihood estimate of the log of the variance of the individual effect;
- `sigma_u`: the estimated standard deviation of the individual effect; and
- `rho`: the estimated share of the individual effect in the composite error variance (also known as the intra-class correlation).

The Likelihood Ratio test of the null hypothesis that `rho` equals zero provides a means of assessing whether the random effects specification is needed. If the null is not rejected that suggests that a simple pooled probit specification is adequate.

Menu path: /Model/Limited dependent variable/Probit

pvalue

Arguments: *dist* [*params*] *xval*

Examples: `pvalue z zscore`
`pvalue t 25 3.0`
`pvalue X 3 5.6`
`pvalue F 4 58 fval`
`pvalue G shape scale x`
`pvalue B bprob 10 6`
`pvalue P lambda x`
`pvalue W shape scale x`

Computes the area to the right of *xval* in the specified distribution (z for Gaussian, t for Student's *t*, X for chi-square, F for *F*, G for gamma, B for binomial, P for Poisson, or W for Weibull).

Depending on the distribution, the following information must be given, before the *xval*: for the *t* and chi-square distributions, the degrees of freedom; for *F*, the numerator and denominator degrees of freedom; for gamma, the shape and scale parameters; for the binomial distribution, the “success” probability and the number of trials; for the Poisson distribution, the parameter λ (which is both the mean and the variance); and for the Weibull distribution, shape and scale parameters. As shown in the examples above, the numerical parameters may be given in numeric form or as the names of variables.

The parameters for the gamma distribution are sometimes given as mean and variance rather than shape and scale. The mean is the product of the shape and the scale; the variance is the product of the shape and the square of the scale. So the scale may be found as the variance divided by the mean, and the shape as the mean divided by the scale.

Menu path: /Tools/P-value finder

qlrtest

Options: `--limit-to=list` (limit test to subset of regressors)
`--plot=mode-or-filename` (see below)

For a model estimated on time-series data via OLS, performs the Quandt likelihood ratio (QLR) test for a structural break at an unknown point in time, with 15 percent trimming at the beginning and end of the sample period.

For each potential break point within the central 70 percent of the observations, a Chow test is performed. See [chow](#) for details; as with the regular Chow test, this is a robust Wald test if the original model was estimated with the `--robust` option, an F-test otherwise. The QLR statistic is then the maximum of the individual test statistics.

An asymptotic p-value is obtained using the method of [Hansen \(1997\)](#).

Besides the standard hypothesis test accessors `$test` and `$pvalue`, `$qlrbreak` can be used to retrieve the index of the observation at which the test statistic is maximized.

The `--limit-to` option can be used to limit the set of interactions with the split dummy variable in the Chow tests to a subset of the original regressors. The parameter for this option must be a named list, all of whose members are among the original regressors. The list should not include the constant.

When this command is run interactively (only), a plot of the Chow test statistic is displayed by default. This can be adjusted via the `--plot` option. The acceptable parameters to this option are `none` (to suppress the plot); `display` (to display a plot even when not in interactive mode); or a file name. The effect of providing a file name is as described for the `--output` option of the [gnuplot](#) command.

Menu path: Model window, /Tests/QLR test

qqplot

Variants: `qqplot y`
`qqplot y x`

Options: `--z-scores` (see below)
`--raw` (see below)
`--output=filename` (send output to specified file)

Given just one series argument, displays a plot of the empirical quantiles of the selected series (given by name or ID number) against the quantiles of the normal distribution. The series must include at least 20 valid observations in the current sample range. By default the empirical quantiles are plotted against quantiles of the normal distribution having the same mean and variance as the sample data, but two alternatives are available: if the `--z-scores` option is given the data are standardized, while if the `--raw` option is given the “raw” empirical quantiles are plotted against the quantiles of the standard normal distribution.

The option `--output` has the effect to send the output to the desired filename; use “display” to force output to the screen, for example during a loop.

Given two series arguments, *y* and *x*, displays a plot of the empirical quantiles of *y* against those of *x*. The data values are not standardized.

Menu path: /Variable/Normal Q-Q plot

Menu path: /View/Graph specified vars/Q-Q plot

quantreg

Arguments: *tau depvar indepvars*

Options: `--robust` (robust standard errors)
`--intervals[=level]` (compute confidence intervals)
`--vcv` (print covariance matrix)
`--quiet` (suppress printing of results)

Examples: `quantreg 0.25 y 0 xlist`
`quantreg 0.5 y 0 xlist --intervals`
`quantreg 0.5 y 0 xlist --intervals=.95`
`quantreg tauvec y 0 xlist --robust`
 See also `mrw_qr.inp`

Quantile regression. The first argument, *tau*, is the conditional quantile for which estimates are wanted. It may be given either as a numerical value or as the name of a pre-defined scalar variable; the value must be in the range 0.01 to 0.99. (Alternatively, a vector of values may be given for *tau*; see below for details.) The second and subsequent arguments compose a regression list on the same pattern as [ols](#).

Without the `--intervals` option, standard errors are printed for the quantile estimates. By default, these are computed according to the asymptotic formula given by [Koenker and Bassett \(1978\)](#), but if the `--robust` option is given, standard errors that are robust with respect to heteroskedasticity are calculated using the method of [Koenker and Zhao \(1994\)](#).

When the `--intervals` option is chosen, confidence intervals are given for the parameter estimates instead of standard errors. These intervals are computed using the rank inversion method, and in general they are asymmetrical about the point estimates. The specifics of the calculation are inflected by the `--robust` option: without this, the intervals are computed on the assumption

of IID errors (Koenker, 1994); with it, they use the robust estimator developed by Koenker and Machado (1999).

By default, 90 percent confidence intervals are produced. You can change this by appending a confidence level (expressed as a decimal fraction) to the intervals option, as in `--intervals=0.95`.

Vector-valued *tau*: instead of supplying a scalar, you may give the name of a pre-defined matrix. In this case estimates are computed for all the given *tau* values and the results are printed in a special format, showing the sequence of quantile estimates for each regressor in turn.

Menu path: /Model/Robust estimation/Quantile regression

quit

Exits from the program, giving you the option of saving the output from the session on the way out.

Menu path: /File/Exit

rename

Arguments: *series newname*

Changes the name of *series* (identified by name or ID number) to *newname*. The new name must be of 31 characters maximum, must start with a letter, and must be composed of only letters, digits, and the underscore character.

Menu path: /Variable/Edit attributes

Other access: Main window pop-up menu (single selection)

reset

Options: `--quiet` (don't print the auxiliary regression)
`--squares-only` (compute the test using only the squares)
`--cubes-only` (compute the test using only the cubes)

Must follow the estimation of a model via OLS. Carries out Ramsey's RESET test for model specification (non-linearity) by adding the square and/or the cube of the fitted values to the regression and calculating the *F* statistic for the null hypothesis that the parameters on the added terms are zero.

Both the square and the cube are added, unless one of the options `--squares-only` or `--cubes-only` is given.

Menu path: Model window, /Tests/Ramsey's RESET

restrict

Options: `--quiet` (don't print restricted estimates)
`--silent` (don't print anything)
`--wald` (system estimators only – see below)
`--bootstrap` (bootstrap the test if possible)
`--full` (OLS and VECMs only, see below)

Imposes a set of (usually linear) restrictions on either (a) the model last estimated or (b) a system of equations previously defined and named. In all cases the set of restrictions should be started with the keyword "restrict" and terminated with "end restrict".

In the single equation case the restrictions are always implicitly to be applied to the last model, and they are evaluated as soon as the `restrict` block is closed.

In the case of a system of equations (defined via the `system` command), the initial “restrict” may be followed by the name of a previously defined system of equations. If this is omitted and the last model was a system then the restrictions are applied to the last model. By default the restrictions are evaluated when the system is next estimated, using the `estimate` command. But if the `--wald` option is given the restriction is tested right away, via a Wald chi-square test on the covariance matrix. Note that this option will produce an error if a system has been defined but not yet estimated.

Depending on the context, the restrictions to be tested may be expressed in various ways. The simplest form is as follows: each restriction is given as an equation, with a linear combination of parameters on the left and a scalar value to the right of the equals sign (either a numerical constant or the name of a scalar variable).

In the single-equation case, parameters may be referenced in the form $b[i]$, where i represents the position in the list of regressors (starting at 1), or $b[varname]$, where *varname* is the name of the regressor in question. In the system case, parameters are referenced using b plus two numbers in square brackets. The leading number represents the position of the equation within the system and the second number indicates position in the list of regressors. For example $b[2,1]$ denotes the first parameter in the second equation, and $b[3,2]$ the second parameter in the third equation. The b terms in the equation representing a restriction may be prefixed with a numeric multiplier, for example $3.5*b[4]$.

Here is an example of a set of restrictions for a previously estimated model:

```
restrict
b[1] = 0
b[2] - b[3] = 0
b[4] + 2*b[5] = 1
end restrict
```

And here is an example of a set of restrictions to be applied to a named system. (If the name of the system does not contain spaces, the surrounding quotes are not required.)

```
restrict "System 1"
b[1,1] = 0
b[1,2] - b[2,2] = 0
b[3,4] + 2*b[3,5] = 1
end restrict
```

In the single-equation case the restrictions are by default evaluated via a Wald test, using the covariance matrix of the model in question. If the original model was estimated via OLS then the restricted coefficient estimates are printed; to suppress this, append the `--quiet` option flag to the initial `restrict` command. As an alternative to the Wald test, for models estimated via OLS or WLS only, you can give the `--bootstrap` option to perform a bootstrapped test of the restriction.

In the system case, the test statistic depends on the estimator chosen: a Likelihood Ratio test if the system is estimated using a Maximum Likelihood method, or an asymptotic F -test otherwise.

There are two alternatives to the method of expressing restrictions discussed above. First, a set of g linear restrictions on a k -vector of parameters, β , may be written compactly as $R\beta - q = 0$, where R is an $g \times k$ matrix and q is a g -vector. You can specify a restriction by giving the names of pre-defined, conformable matrices to be used as R and q , as in

```
restrict
R = Rmat
q = qvec
end restrict
```

Secondly, if you wish to test a nonlinear restriction (this is currently available for single-equation models only) you should give the restriction as the name of a function, preceded by “`rfunc =`”, as in

```
restrict
  rfunc = myfunction
end restrict
```

The constraint function should take a single `const matrix` argument; this will be automatically filled out with the parameter vector. And it should return a vector which is zero under the null hypothesis, non-zero otherwise. The length of the vector is the number of restrictions. This function is used as a “callback” by gretl’s numerical Jacobian routine, which calculates a Wald test statistic via the delta method.

Here is a simple example of a function suitable for testing one nonlinear restriction, namely that two pairs of parameter values have a common ratio.

```
function matrix restr (const matrix b)
  matrix v = b[1]/b[2] - b[4]/b[5]
  return v
end function
```

On successful completion of the `restrict` command the accessors `$test` and `$pvalue` give the test statistic and its p-value.

When testing restrictions on a single-equation model estimated via OLS, or on a VECM, the `--full` option can be used to set the restricted estimates as the “last model” for the purposes of further testing or the use of accessors such as `$coeff` and `$vcv`. Note that some special considerations apply in the case of testing restrictions on Vector Error Correction Models. Please see the *Gretl User’s Guide* for details.

Menu path: Model window, /Tests/Linear restrictions

rmplot

Argument: *series*

Options: `--trim` (see below)

`--quiet` (suppress printed output)

Range-mean plot: this command creates a simple graph to help in deciding whether a time series, $y(t)$, has constant variance or not. We take the full sample $t=1,...,T$ and divide it into small subsamples of arbitrary size k . The first subsample is formed by $y(1),...,y(k)$, the second is $y(k+1), ..., y(2k)$, and so on. For each subsample we calculate the sample mean and range (= maximum minus minimum), and we construct a graph with the means on the horizontal axis and the ranges on the vertical. So each subsample is represented by a point in this plane. If the variance of the series is constant we would expect the subsample range to be independent of the subsample mean; if we see the points approximate an upward-sloping line this suggests the variance of the series is increasing in its mean; and if the points approximate a downward sloping line this suggests the variance is decreasing in the mean.

Besides the graph, gretl displays the means and ranges for each subsample, along with the slope coefficient for an OLS regression of the range on the mean and the p-value for the null hypothesis that this slope is zero. If the slope coefficient is significant at the 10 percent significance level then the fitted line from the regression of range on mean is shown on the graph. The t -statistic for the null, and the corresponding p-value, are recorded and may be retrieved using the accessors `$test` and `$pvalue` respectively.

If the `--trim` option is given, the minimum and maximum values in each sub-sample are discarded before calculating the mean and range. This makes it less likely that outliers will distort the analysis.

If the `--quiet` option is given, no graph is shown and no output is printed; only the t -statistic and p -value are recorded.

Menu path: /Variable/Range-mean graph

run

Argument: *filename*

Executes the commands in *filename* then returns control to the interactive prompt. This command is intended for use with the command-line program `gretlcli`, or at the “gretl console” in the GUI program.

See also [include](#).

Menu path: Run icon in script window

runs

Argument: *series*

Options: `--difference` (use first difference of variable)
`--equal` (positive and negative values are equiprobable)

Carries out the nonparametric “runs” test for randomness of the specified *series*, where runs are defined as sequences of consecutive positive or negative values. If you want to test for randomness of deviations from the median, for a variable named `x1` with a non-zero median, you can do the following:

```
series signx1 = x1 - median(x1)
runs signx1
```

If the `--difference` option is given, the variable is differenced prior to the analysis, hence the runs are interpreted as sequences of consecutive increases or decreases in the value of the variable.

If the `--equal` option is given, the null hypothesis incorporates the assumption that positive and negative values are equiprobable, otherwise the test statistic is invariant with respect to the “fairness” of the process generating the sequence, and the test focuses on independence alone.

Menu path: /Tools/Nonparametric tests

scatters

Arguments: *yvar* ; *xvars* or *yvars* ; *xvar*

Options: `--with-lines` (create line graphs)
`--matrix=name` (plot columns of named matrix)
`--output=filename` (send output to specified file)
`--output=filename` (send output to specified file)

Examples: `scatters 1 ; 2 3 4 5`
`scatters 1 2 3 4 5 6 ; 7`
`scatters y1 y2 y3 ; x --with-lines`

Generates pairwise graphs of *yvar* against all the variables in *xvars*, or of all the variables in *yvars* against *xvar*. The first example above puts variable 1 on the y -axis and draws four graphs, the first having variable 2 on the x -axis, the second variable 3 on the x -axis, and so on. The second example

plots each of variables 1 through 6 against variable 7 on the x -axis. Scanning a set of such plots can be a useful step in exploratory data analysis. The maximum number of plots is 16; any extra variable in the list will be ignored.

By default the graphs are scatterplots, but if you give the `--with-lines` flag they will be line graphs.

For details on usage of the `--output` option, please see the [gnuplot](#) command.

If a named matrix is specified as the data source the x and y lists should be given as 1-based column numbers; or alternatively, if no such numbers are given, all the columns are plotted against time or an index variable.

If the dataset is time-series, then the second sub-list can be omitted, in which case it will implicitly be taken as "time", so you can plot multiple time series in separated sub-graphs.

Menu path: /View/Multiple graphs

sdiff

Argument: *varlist*

The seasonal difference of each variable in *varlist* is obtained and the result stored in a new variable with the prefix `sd_`. This command is available only for seasonal time series.

Menu path: /Add/Seasonal differences of selected variables

set

Variants: `set variable value`
`set --to-file=filename`
`set --from-file=filename`
`set stopwatch`
`set`

Examples: `set svd on`
`set csv_delim tab`
`set horizon 10`
`set --to-file=mysettings.inp`

The most common use of this command is the first variant shown above, where it is used to set the value of a selected program parameter. This is discussed in detail below. The other uses are: with `--to-file`, to write a script file containing all the current parameter settings; with `--from-file` to read a script file containing parameter settings and apply them to the current session; with `stopwatch` to zero the gretl "stopwatch" which can be used to measure CPU time (see the entry for the `$stopwatch` accessor in the gretl function reference); or, if the word `set` is given alone, to print the current settings.

Values set via this comand remain in force for the duration of the gretl session unless they are changed by a further call to `set`. The parameters that can be set in this way are enumerated below. Note that the settings of `hc_version`, `hac_lag` and `hac_kernel` are used when the `--robust` option is given to an estimation command.

The available settings are grouped under the following categories: program interaction and behavior, numerical methods, random number generation, robust estimation, filtering, time series estimation, and interaction with GNU R.

Program interaction and behavior

These settings are used for controlling various aspects of the way gretl interacts with the user.

- `csv_delim`: either comma (the default), space, tab or semicolon. Sets the column delimiter used when saving data to file in CSV format.
- `csv_write_na`: the string used to represent missing values when writing data to file in CSV format. Maximum 7 characters; the default is NA.
- `csv_read_na`: the string taken to represent missing values (NAs) when reading data in CSV format. Maximum 7 characters. The default depends on whether a data column is found to contain numerical data (mostly) or string values. For numerical data the following are taken as indicating NAs: an empty cell, or any of the strings NA, N.A., na, n.a., N/A, #N/A, NaN, .NaN, ., . ., -999, and -9999. For string-valued data only a blank cell, or a cell containing an empty string, is counted as NA. These defaults can be reimposed by giving `default` as the value for `csv_read_na`. To specify that only empty cells are read as NAs, give a value of `""`. Note that empty cells are always read as NAs regardless of the setting of this variable.
- `csv_digits`: a positive integer specifying the number of significant digits to use when writing data in CSV format. By default up to 15 digits are used depending on the precision of the original data. Note that CSV output employs the C library's `fprintf` function with `"%g"` conversion, which means that trailing zeros are dropped.
- `mwrite_g`: on or off (the default). When writing a matrix to file as text, gretl by default uses scientific notation with 18-digit precision, hence ensuring that the stored values are a faithful representation of the numbers in memory. When writing primary data with no more than 6 digits of precision it may be preferable to use `%g` format for a more compact and human-readable file; you can make this switch via `set mwrite_g on`.
- `echo`: off or on (the default). Suppress or resume the echoing of commands in gretl's output.
- `force_decpoint`: on or off (the default). Force gretl to use the decimal point character, in a locale where another character (most likely the comma) is the standard decimal separator.
- `loop_maxiter`: one non-negative integer value (default 100000). Sets the maximum number of iterations that a `while` loop is allowed before halting (see [loop](#)). Note that this setting only affects the `while` variant; its purpose is to guard against inadvertently infinite loops. Setting this value to 0 has the effect of disabling the limit; use with caution.
- `max_verbose`: on or off (the default). Toggles verbose output for the `BFGSmax` and `NRmax` functions (see the User's Guide for details).
- `messages`: off or on (the default). Suppress or resume the printing of non-error messages associated with various commands, for example when a new variable is generated or when the sample range is changed.
- `warnings`: off or on (the default). Suppress or resume the printing of warning messages issued when arithmetical operations produce non-finite values.
- `debug`: 1, 2 or 0 (the default). This is for use with user-defined functions. Setting `debug` to 1 is equivalent to turning `messages` on within all such functions; setting this variable to 2 has the additional effect of turning on `max_verbose` within all functions.
- `shell_ok`: on or off (the default). Enable launching external programs from gretl via the system shell. This is disabled by default for security reasons, and can only be enabled via the graphical user interface (Tools/Preferences/General). However, once set to on, this setting will remain active for future sessions until explicitly disabled.
- `shell_dir`: *path*. Sets the current working directory for shell commands.
- `use_cwd`: on or off (the default). This setting affects the behavior of the [outfile](#) and [store](#) commands, which write external files. Normally, the file will be written in the user's default data directory; if `use_cwd` is on, on the contrary, the file will be created in the working directory when gretl was started.

- `bfgs_verbskip`: one integer. This setting affects the behavior of the `--verbose` option to those commands that use BFGS as an optimization algorithm and is used to compact output. if `bfgs_verbskip` is set to, say, 3, then the `--verbose` switch will only print iterations 3, 6, 9 and so on.
- `skip_missing`: on (the default) or off. Controls gretl's behavior when constructing a matrix from data series: the default is to skip data rows that contain one or more missing values but if `skip_missing` is set off missing values are converted to NaNs.
- `matrix_mask`: the name of a series, or the keyword `null`. Offers greater control than `skip_missing` when constructing matrices from series: the data rows selected for matrices are those with non-zero (and non-missing) values in the specified series. The selected mask remains in force until it is replaced, or removed via the `null` keyword.
- `huge`: a large positive number (by default, 1.0E100). This setting controls the value returned by the accessor [\\$huge](#).

Numerical methods

These settings are used for controlling the numerical algorithms that gretl uses for estimation.

- `optimizer`: either `auto` (the default), `BFGS` or `newton`. Sets the optimization algorithm used for various ML estimators, in cases where both BFGS and Newton-Raphson are applicable. The default is to use Newton-Raphson where an analytical Hessian is available, otherwise BFGS.
- `bhhh_maxiter`: one integer, the maximum number of iterations for gretl's internal BHHH routine, which is used in the `arma` command for conditional ML estimation. If convergence is not achieved after `bhhh_maxiter`, the program returns an error. The default is set at 500.
- `bhhh_tol`: one floating point value, or the string `default`. This is used in gretl's internal BHHH routine to check if convergence has occurred. The algorithm stops iterating as soon as the increment in the log-likelihood between iterations is smaller than `bhhh_tol`. The default value is 1.0E-06; this value may be re-established by typing `default` in place of a numeric value.
- `bfgs_maxiter`: one integer, the maximum number of iterations for gretl's BFGS routine, which is used for `mle`, `gmm` and several specific estimators. If convergence is not achieved in the specified number of iterations, the program returns an error. The default value depends on the context, but is typically of the order of 500.
- `bfgs_tol`: one floating point value, or the string `default`. This is used in gretl's BFGS routine to check if convergence has occurred. The algorithm stops as soon as the relative improvement in the objective function between iterations is smaller than `bfgs_tol`. The default value is the machine precision to the power 3/4; this value may be re-established by typing `default` in place of a numeric value.
- `bfgs_maxgrad`: one floating point value. This is used in gretl's BFGS routine to check if the norm of the gradient is reasonably close to zero when the `bfgs_tol` criterion is met. A warning is printed if the norm of the gradient exceeds 1; an error is flagged if the norm exceeds `bfgs_maxgrad`. At present the default is the permissive value of 5.0.
- `bfgs_richardson`: on or off (the default). Use Richardson extrapolation when computing numerical derivatives in the context of BFGS maximization.
- `initvals`: either `auto` (the default) or the name of a pre-specified matrix. Allows manual setting of the initial parameter estimates for numerical optimization problems (such as ARMA estimation). For details see the *Gretl User's Guide*.

- `lbfgs`: on or off (the default). Use the limited-memory version of BFGS (L-BFGS-B) instead of the ordinary algorithm. This may be advantageous when the function to be maximized is not globally concave.
- `lbfgs_mem`: an integer value in the range 3 to 20 (with a default value of 8). This determines the number of corrections used in the limited memory matrix when L-BFGS-B is employed.
- `nls_tol`: a floating-point value (the default is the machine precision to the power 3/4). Sets the tolerance used in judging whether or not convergence has occurred in nonlinear least squares estimation using the `nls` command.
- `svd`: on or off (the default). Use SVD rather than Cholesky or QR decomposition in least squares calculations. This option applies to the `mo1s` function as well as various internal calculations, but not to the regular `ols` command.
- `fcg`: on or off (the default). Use the algorithm of Fiorentini, Calzolari and Panattoni rather than native gretl code when computing GARCH estimates.
- `gmm_maxiter`: one integer, the maximum number of iterations for gretl's `gmm` command when in iterated mode (as opposed to one- or two-step). The default value is 250.
- `nadarwat_trim`: one integer, the trim parameter used in the `nadarwat` function.
- `fdjac_quality`: one integer between 0 and 2, the algorithm used by the `fdjac` function.

Random number generation

- `seed`: an unsigned integer. Sets the seed for the pseudo-random number generator. By default this is set from the system time; if you want to generate repeatable sequences of random numbers you must set the seed manually.
- `normal_rand`: `ziggurat` (the default) or `box-muller`. Sets the method for generating random normal samples based on uniform input.

Robust estimation

- `bootrep`: an integer. Sets the number of replications for the `restrict` command with the `--bootstrap` option.
- `garch_vcv`: `unset`, `hessian`, `im` (information matrix), `op` (outer product matrix), `qml` (QML estimator), `bw` (Bollerslev-Wooldridge). Specifies the variant that will be used for estimating the coefficient covariance matrix, for GARCH models. If `unset` is given (the default) then the Hessian is used unless the “robust” option is given for the `garch` command, in which case QML is used.
- `arma_vcv`: `hessian` (the default) or `op` (outer product matrix). Specifies the variant to be used when computing the covariance matrix for ARIMA models.
- `force_hc`: `off` (the default) or `on`. By default, with time-series data and when the `--robust` option is given with `ols`, the HAC estimator is used. If you set `force_hc` to “on”, this forces calculation of the regular Heteroskedasticity Consistent Covariance Matrix (HCCM), which does not take autocorrelation into account. Note that VARs are treated as a special case: when the `--robust` option is given the default method is regular HCCM, but the `--robust-hac` flag can be used to force the use of a HAC estimator.
- `robust_z`: `off` (the default) or `on`. This controls the distribution used when calculating p-values based on robust standard errors in the context of least-squares estimators. By default gretl uses the Student *t* distribution but if `robust_z` is turned on the normal distribution is used.

- **hac_lag**: nw1 (the default), nw2, nw3 or an integer. Sets the maximum lag value or bandwidth, p , used when calculating HAC (Heteroskedasticity and Autocorrelation Consistent) standard errors using the Newey-West approach, for time series data. nw1 and nw2 represent two variant automatic calculations based on the sample size, T : for nw1, $p = 0.75 \times T^{1/3}$, and for nw2, $p = 4 \times (T/100)^{2/9}$. nw3 calls for data-based bandwidth selection. See also **qs_bandwidth** and **hac_prewhiten** below.
- **hac_kernel**: bartlett (the default), parzen, or qs (Quadratic Spectral). Sets the kernel, or pattern of weights, used when calculating HAC standard errors.
- **hac_prewhiten**: on or off (the default). Use Andrews-Monahan prewhitening and re-coloring when computing HAC standard errors. This also implies use of data-based bandwidth selection.
- **hc_version**: 0 (the default), 1, 2, 3 or 3a. Sets the variant used when calculating Heteroskedasticity Consistent standard errors with cross-sectional data. The first four options correspond to the HC0, HC1, HC2 and HC3 discussed by Davidson and MacKinnon in *Econometric Theory and Methods*, chapter 5. HC0 produces what are usually called “White’s standard errors”. Variant 3a is the MacKinnon-White “jackknife” procedure.
- **pcse**: off (the default) or on. By default, when estimating a model using pooled OLS on panel data with the **--robust** option, the Arellano estimator is used for the covariance matrix. If you set **pcse** to “on”, this forces use of the Beck and Katz Panel Corrected Standard Errors (which do not take autocorrelation into account).
- **qs_bandwidth**: Bandwidth for HAC estimation in the case where the Quadratic Spectral kernel is selected. (Unlike the Bartlett and Parzen kernels, the QS bandwidth need not be an integer.)

Time series

- **horizon**: one integer (the default is based on the frequency of the data). Sets the horizon for impulse responses and forecast variance decompositions in the context of vector autoregressions.
- **vecm_norm**: phillips (the default), diag, first or none. Used in the context of VECM estimation via the **vecm** command for identifying the cointegration vectors. See the the *Gretl User’s Guide* for details.

Interaction with R

- **R_lib**: on (the default) or off. When sending instructions to be executed by R, use the R shared library by preference to the R executable, if the library is available.
- **R_functions**: off (the default) or on. Recognize functions defined in R as if they were native functions (the namespace prefix “R.” is required). See the *Gretl User’s Guide* for details on this and the previous item.

setinfo

Argument: *series*

Options: **--description=string** (set description)
--graph-name=string (set graph name)
--discrete (mark series as discrete)
--continuous (mark series as continuous)

Examples: **setinfo x1 --description="Description of x1"**
setinfo y --graph-name="Some string"
setinfo z --discrete

Sets up to three attributes of *series*, given by name or ID number, as follows.

If the `--description` flag is given followed by a string in double quotes, that string is used to set the variable's descriptive label. This label is shown in response to the `labels` command, and is also shown in the main window of the GUI program.

If the `--graph-name` flag is given followed by a quoted string, that string will be used in place of the variable's name in graphs.

If one or other of the `--discrete` or `--continuous` option flags is given, the variable's numerical character is set accordingly. The default is to treat all series as continuous; setting a series as discrete affects the way the variable is handled in frequency plots.

Menu path: /Variable/Edit attributes

Other access: Main window pop-up menu

setmiss

Arguments: *value* [*varlist*]

Examples: `setmiss -1`
`setmiss 100 x2`

Get the program to interpret some specific numerical data value (the first parameter to the command) as a code for “missing”, in the case of imported data. If this value is the only parameter, as in the first example above, the interpretation will be applied to all series in the data set. If *value* is followed by a list of variables, by name or number, the interpretation is confined to the specified variable(s). Thus in the second example the data value 100 is interpreted as a code for “missing”, but only for the variable *x2*.

Menu path: /Data/Set missing value code

setobs

Variants: `setobs periodicity startobs`
`setobs unitvar timevar --panel-vars`

Options: `--cross-section` (interpret as cross section)
`--time-series` (interpret as time series)
`--stacked-cross-section` (interpret as panel data)
`--stacked-time-series` (interpret as panel data)
`--panel-vars` (use index variables, see below)
`--panel-time` (see below)
`--panel-groups` (see below)

Examples: `setobs 4 1990:1 --time-series`
`setobs 12 1978:03`
`setobs 1 1 --cross-section`
`setobs 20 1:1 --stacked-time-series`
`setobs unit year --panel-vars`

This command forces the program to interpret the current data set as having a specified structure.

In the first form of the command the *periodicity*, which must be an integer, represents frequency in the case of time-series data (1 = annual; 4 = quarterly; 12 = monthly; 52 = weekly; 5, 6, or 7 = daily; 24 = hourly). In the case of panel data the periodicity means the number of lines per data block: this corresponds to the number of cross-sectional units in the case of stacked cross-sections, or the number of time periods in the case of stacked time series. In the case of simple cross-sectional data the periodicity should be set to 1.

The starting observation represents the starting date in the case of time series data. Years may be given with two or four digits; subperiods (for example, quarters or months) should be separated from the year with a colon. In the case of panel data the starting observation should be given as 1:1; and in the case of cross-sectional data, as 1. Starting observations for daily or weekly data should be given in the form YYYY-MM-DD (or simply as 1 for undated data).

If no explicit option flag is given to indicate the structure of the data the program will attempt to guess the structure from the information given.

The second form of the command (which requires the `--panel-vars` flag) may be used to impose a panel interpretation when the data set contains variables that uniquely identify the cross-sectional units and the time periods. The data set will be sorted as stacked time series, by ascending values of the units variable, *unitvar*.

Panel-specific options

The `--panel-time` and `--panel-groups` options can only be used with a dataset which has already been defined as a panel.

The purpose of `--panel-time` is to set extra information regarding the time dimension of the panel. This should be given on the pattern of the first form of `setobs` noted above. For example, the following may be used to indicate that the time dimension of a panel is quarterly, starting in the first quarter of 1990.

```
setobs 4 1990:1 --panel-time
```

The purpose of `--panel-groups` is to create a string-valued series holding names for the groups (individuals, cross-sectional units) in the panel. With this option you must supply a name for the series followed by a string variable holding a list of group names. The names should be separated by spaces; if a name includes spaces it should be wrapped in backslash-escaped double-quotes. For example, the following will create a series named `country` in which the names in `cstrs` are each repeated *T* times, *T* being the time-series length of the panel.

```
string cstrs
sprintf cstrs "France Germany Italy \"United Kingdom\""
setobs country cstrs --panel-groups
```

Menu path: /Data/Dataset structure

setopt

Arguments: *command* [*action*] *options*

Examples: `setopt mle --hessian`
`setopt ols persist --quiet`
`setopt ols clear`

This command enables the pre-setting of options for a specified command. Ordinarily this is not required, but it may be useful for the writers of `hansl` functions when they wish to make certain command options conditional on the value of an argument supplied by the caller.

For example, suppose a function offers a boolean “quiet” switch, whose intended effect is to suppress the printing of results from a certain regression executed within the function. In that case one might write:

```
if quiet
    setopt ols --quiet
endif
ols ...
```

The `--quiet` option will then be applied to the next `ols` command if and only if the variable `quiet` has a non-zero value.

By default, options set in this way apply only to the following instance of *command*; they are not persistent. However if you give `persist` as the value for *action* the options will continue to apply to the given command until further notice. The antidote to the `persist` action is `clear`: this erases any stored setting for the specified command.

It should be noted that options set via `setopt` are compounded with any options attached to the target command directly. So for example one might append the `--hessian` option to an `mle` command unconditionally but use `setopt` to add `--quiet` conditionally.

shell

Argument: *shellcommand*

Examples: `! ls -al`

`! notepad`

`! launch notepad`

An exclamation mark, `!`, or the keyword `! launch`, at the beginning of a command line is interpreted as an escape to the user's shell. Thus arbitrary shell commands can be executed from within gretl. When `!` is used, the external command is executed synchronously. That is, gretl waits for it to complete before proceeding. If you want to start another program from within gretl and not wait for its completion (asynchronous operation), use `! launch` instead.

For reasons of security this facility is not enabled by default. To activate it, check the box titled "Allow shell commands" under the File, Preferences menu in the GUI program. This also makes shell commands available in the command-line program (and is the only way to do so).

smp1

Variants: `smp1 startobs endobs`
`smp1 +i -j`
`smp1 dumvar --dummy`
`smp1 condition --restrict`
`smp1 --no-missing [ varlist ]`
`smp1 --no-all-missing [ varlist ]`
`smp1 --contiguous [ varlist ]`
`smp1 n --random`
`smp1 full`

Options: `--dummy` (argument is a dummy variable)
`--restrict` (apply boolean restriction)
`--replace` (replace any existing boolean restriction)
`--no-missing` (restrict to valid observations)
`--no-all-missing` (omit empty observations (see below))
`--contiguous` (see below)
`--random` (form random sub-sample)
`--permanent` (see below)
`--balanced` (panel data: try to retain balanced panel)

Examples: `smp1 3 10`
`smp1 1960:2 1982:4`
`smp1 +1 -1`
`smp1 x > 3000 --restrict`
`smp1 y > 3000 --restrict --replace`
`smp1 100 --random`

Resets the sample range. The new range can be defined in several ways. In the first alternate form (and the first two examples) above, *startobs* and *endobs* must be consistent with the periodicity of the data. Either one may be replaced by a semicolon to leave the value unchanged. In the second form, the integers *i* and *j* (which may be positive or negative, and should be signed) are taken as offsets relative to the existing sample range. In the third form *dummyvar* must be an indicator variable with values 0 or 1 at each observation; the sample will be restricted to observations where the value is 1. The fourth form, using `--restrict`, restricts the sample to observations that satisfy the given Boolean condition (which is specified according to the syntax of the [genr](#) command).

The options `--no-missing` and `--no-all-missing` may be used to exclude from the sample observations for which data are missing. The first variant excludes those rows in the dataset for which at least one variable has a missing value, while the second excludes just those rows on which *all* variables have missing values. In each case the test is confined to the variables in *varlist* if this argument is given, otherwise it is applied to all series—with the qualification that in the case of `--no-all-missing` and no *varlist*, the generic variables *index* and *time* are ignored.

The `--contiguous` form of `smp1` is intended for use with time series data. The effect is to trim any observations at the start and end of the current sample range that contain missing values (either for the variables in *varlist*, or for all data series if no *varlist* is given). Then a check is performed to see if there are any missing values in the remaining range; if so, an error is flagged.

With the `--random` flag, the specified number of cases are selected from the current dataset at random (without replacement). If you wish to be able to replicate this selection you should set the seed for the random number generator first (see the [set](#) command).

The final form, `smp1 full`, restores the full data range.

Note that sample restrictions are, by default, cumulative: the baseline for any `smp1` command is the current sample. If you wish the command to act so as to replace any existing restriction you can add the option flag `--replace` to the end of the command. (But this option is not compatible with the `--contiguous` option.)

The internal variable `obs` may be used with the `--restrict` form of `smp1` to exclude particular observations from the sample. For example

```
smp1 obs!=4 --restrict
```

will drop just the fourth observation. If the data points are identified by labels,

```
smp1 obs!="USA" --restrict
```

will drop the observation with label “USA”.

One point should be noted about the `--dummy`, `--restrict` and `--no-missing` forms of `smp1`: “structural” information in the data file (regarding the time series or panel nature of the data) is likely to be lost when this command is issued. You may reimpose structure with the `setobs` command. A related option, for use with panel data, is the `--balanced` flag: this requests that a balanced panel is reconstituted after sub-sampling, via the insertion of “missing rows” if need be. But note that it is not always possible to comply with this request.

By default, restrictions on the current sample range are undoable: by doing `smp1 full` you can restore the unrestricted dataset. However, the `--permanent` flag can be used to substitute the restricted dataset for the original. This option is only available in conjunction with the `--restrict`, `--dummy`, `--no-missing`, `--no-all-missing` or `--random` forms of `smp1`.

Please see the *Gretl User's Guide* for further details.

Menu path: /Sample

spearman

Arguments: *series1 series2*

Option: `--verbose` (print ranked data)

Prints Spearman's rank correlation coefficient for the series *series1* and *series2*. The variables do not have to be ranked manually in advance; the function takes care of this.

The automatic ranking is from largest to smallest (i.e. the largest data value gets rank 1). If you need to invert this ranking, create a new variable which is the negative of the original. For example:

```
series altx = -x
spearman altx y
```

Menu path: /Model/Robust estimation/Rank correlation

sprintf

Arguments: *stringvar format , args*

This command works exactly like the `printf` command, printing the given arguments under the control of the format string, except that the result is written into the named string, *stringvar*.

square

Argument: *varlist*

Option: `--cross` (generate cross-products as well as squares)

Generates new series which are squares of the series in *varlist* (plus cross-products if the `--cross` option is given). For example, `square x y` will generate `sq_x` = x squared, `sq_y` = y squared and (optionally) `x_y` = x times y . If a particular variable is a dummy variable it is not squared because we will get the same variable.

Menu path: /Add/Squares of selected variables

store

Arguments: *filename* [*varlist*]

Options: `--csv` (use CSV format)
`--omit-obs` (see below, on CSV format)
`--no-header` (see below, on CSV format)
`--gnu-octave` (use GNU Octave format)
`--gnu-R` (use GNU R format)
`--gzipped[=level]` (apply gzip compression)
`--jmulti` (use JMulti ASCII format)
`--dat` (use PcGive ASCII format)
`--decimal-comma` (use comma as decimal character)
`--database` (use gretl database format)
`--overwrite` (see below, on database format)
`--comment=string` (see below)

Save data to *filename*. By default all currently defined series are saved but the optional *varlist* argument can be used to select a subset of series. If the dataset is sub-sampled, only the observations in the current sample range are saved.

The format in which the data are written may be controlled in the first instance by the extension or suffix of *filename*, as follows:

- `.gdt`, or no extension: gretl's native XML data format. (If no extension is provided, "`.gdt`" is added automatically.)
- `.gtdb`: gretl's native binary data format.
- `.csv`: comma-separated values (CSV).
- `.txt` or `.asc`: space-separated values.
- `.R`: GNU R format.
- `.m`: GNU Octave format.

The format-related option flags shown above can be used to force the issue of the save format independently of the filename (or to get gretl to write in the formats of PcGive or JMulTi). However, if *filename* has extension `.gdt` or `.gtdb` this necessarily implies use of native format and the addition of a conflicting option flag will generate an error.

When data are saved in native format (only), the `--gzipped` option may be used for data compression, which can be useful for large datasets. The optional parameter for this flag controls the level of compression (from 0 to 9): higher levels produce a smaller file, but compression takes longer. The default level is 1; a level of 0 means that no compression is applied.

The option flags `--omit-obs` and `--no-header` are applicable only when saving data in CSV format. By default, if the data are time series or panel, or if the dataset includes specific observation markers, the CSV file includes a first column identifying the observations (e.g. by date). If the

`--omit-obs` flag is given this column is omitted. The `--no-header` flag suppresses the usual printing of the names of the variables at the top of the columns.

The option flag `--decimal-comma` is also confined to the case of saving data in CSV format. The effect of this option is to replace the decimal point with the decimal comma; in addition the column separator is forced to be a semicolon.

The option of saving in gretl database format is intended to help with the construction of large sets of series, possibly having mixed frequencies and ranges of observations. At present this option is available only for annual, quarterly or monthly time-series data. If you save to a file that already exists, the default action is to append the newly saved series to the existing content of the database. In this context it is an error if one or more of the variables to be saved has the same name as a variable that is already present in the database. The `--overwrite` flag has the effect that, if there are variable names in common, the newly saved variable replaces the variable of the same name in the original dataset.

The `--comment` option is available when saving data as a database or in CSV format. The required parameter is a double-quoted one-line string, attached to the option flag with an equals sign. The string is inserted as a comment into the database index file or at the top of the CSV output.

The `store` command behaves in a special manner in the context of a “progressive loop”. See the *Gretl User’s Guide* for details.

Menu path: /File/Save data; /File/Export data

summary

Variants: `summary [ varlist ]`
`summary --matrix=matname`

Options: `--simple` (basic statistics only)
`--weight=wvar` (weighting variable)
`--by=byvar` (see below)

In its first form, this command prints summary statistics for the variables in *varlist*, or for all the variables in the data set if *varlist* is omitted. By default, output consists of the mean, standard deviation (sd), coefficient of variation (= sd/mean), median, minimum, maximum, skewness coefficient, and excess kurtosis. If the `--simple` option is given, output is restricted to the mean, minimum, maximum and standard deviation.

If the `--by` option is given (in which case the parameter *byvar* should be the name of a discrete variable), then statistics are printed for sub-samples corresponding to the distinct values taken on by *byvar*. For example, if *byvar* is a (binary) dummy variable, statistics are given for the cases *byvar* = 0 and *byvar* = 1. Note: at present, this option is incompatible with the `--weight` option.

If the alternative form is given, using a named matrix, then summary statistics are printed for each column of the matrix. The `--by` option is not available in this case.

Menu path: /View/Summary statistics

Other access: Main window pop-up menu

system

Variants: `system method=estimator`
`sysname <- system`

Examples: `"Klein Model 1" <- system`
`system method=sur`
`system method=3sls`
 See also `klein.inp`, `kmenta.inp`, `greene14_2.inp`

Starts a system of equations. Either of two forms of the command may be given, depending on whether you wish to save the system for estimation in more than one way or just estimate the system once.

To save the system you should assign it a name, as in the first example (if the name contains spaces it must be surrounded by double quotes). In this case you estimate the system using the `estimate` command. With a saved system of equations, you are able to impose restrictions (including cross-equation restrictions) using the `restrict` command.

Alternatively you can specify an estimator for the system using `method=` followed by a string identifying one of the supported estimators: `ols` (Ordinary Least Squares), `tsls` (Two-Stage Least Squares) `sur` (Seemingly Unrelated Regressions), `3s1s` (Three-Stage Least Squares), `fiml` (Full Information Maximum Likelihood) or `liml` (Limited Information Maximum Likelihood). In this case the system is estimated once its definition is complete.

An equation system is terminated by the line `end system`. Within the system four sorts of statement may be given, as follows.

- `equation`: specify an equation within the system. At least two such statements must be provided.
- `instr`: for a system to be estimated via Three-Stage Least Squares, a list of instruments (by variable name or number). Alternatively, you can put this information into the `equation` line using the same syntax as in the `tsls` command.
- `endog`: for a system of simultaneous equations, a list of endogenous variables. This is primarily intended for use with FIML estimation, but with Three-Stage Least Squares this approach may be used instead of giving an `instr` list; then all the variables not identified as endogenous will be used as instruments.
- `identity`: for use with FIML, an identity linking two or more of the variables in the system. This sort of statement is ignored when an estimator other than FIML is used.

After estimation using the `system` or `estimate` commands the following accessors can be used to retrieve additional information:

- `$uhat`: the matrix of residuals, one column per equation.
- `$yhat`: matrix of fitted values, one column per equation.
- `$coeff`: column vector of coefficients (all the coefficients from the first equation, followed by those from the second equation, and so on).
- `$vcv`: covariance matrix of the coefficients. If there are k elements in the `$coeff` vector, this matrix is k by k .
- `$sigma`: cross-equation residual covariance matrix.
- `$sysGamma`, `$sysA` and `$sysB`: structural-form coefficient matrices (see below).

If you want to retrieve the residuals or fitted values for a specific equation as a data series, select a column from the `$uhat` or `$yhat` matrix and assign it to a series, as in

```
series uh1 = $uhat[,1]
```

The structural-form matrices correspond to the following representation of a simultaneous equations model:

$$\Gamma y_t = A y_{t-1} + B x_t + \epsilon_t$$

If there are n endogenous variables and k exogenous variables, Γ is an $n \times n$ matrix and B is $n \times k$. If the system contains no lags of the endogenous variables then the A matrix is not present. If the maximum lag of an endogenous regressor is p , the A matrix is $n \times np$.

Menu path: /Model/Simultaneous equations

tabprint

Argument: [*-f filename*]

Options: --rtf (Produce RTF instead of $\LaTeX$)
 --csv (Produce CSV instead of $\LaTeX$)
 --complete (Create a complete document)
 --format="f1|f2|f3|f4" (Specify a custom format)

Must follow the estimation of a model. Prints the estimated model in tabular form — by default as $\LaTeX$, but as RTF if the --rtf flag is given or as CSV if the --csv flag is given. If a filename is specified using the -f flag output goes to that file, otherwise it goes to a file with a name of the form model_N followed by the extension tex, rtf or csv, where N is the number of models estimated to date in the current session.

If CSV format is selected, values are comma-separated unless the decimal comma is in force, in which case the separator is the semicolon. Note that CSV output may be less complete than the other formats.

The further options discussed below are available only when printing the model as $\LaTeX$.

If the --complete flag is given the $\LaTeX$ file is a complete document, ready for processing; otherwise it must be included in a document.

If you wish alter the appearance of the tabular output, you can specify a custom row format using the --format flag. The format string must be enclosed in double quotes and must be tied to the flag with an equals sign. The pattern for the format string is as follows. There are four fields, representing the coefficient, standard error, t -ratio and p-value respectively. These fields should be separated by vertical bars; they may contain a printf-type specification for the formatting of the numeric value in question, or may be left blank to suppress the printing of that column (subject to the constraint that you can't leave all the columns blank). Here are a few examples:

```
--format="%.4f|%.4f|%.4f|%.4f"
--format="%.4f|%.4f|%.3f| "
--format="%.5f|%.4f| |%.4f"
--format="%.8g|%.8g| |%.4f"
```

The first of these specifications prints the values in all columns using 4 decimal places. The second suppresses the p-value and prints the t -ratio to 3 places. The third omits the t -ratio. The last one again omits the t , and prints both coefficient and standard error to 8 significant figures.

Once you set a custom format in this way, it is remembered and used for the duration of the gretl session. To revert to the default format you can use the special variant --format=default.

Menu path: Model window, /LaTeX

textplot

Argument: *varlist*

Options: --time-series (plot by observation)
 --one-scale (force a single scale)
 --tall (use 40 rows)

Quick and simple ASCII graphics. Without the `--time-series` flag, *varlist* must contain at least two series, the last of which is taken as the variable for the x axis, and a scatter plot is produced. In this case the `--tall` option may be used to produce a graph in which the y axis is represented by 40 rows of characters (the default is 20 rows).

With the `--time-series`, a plot by observation is produced. In this case the option `--one-scale` may be used to force the use of a single scale; otherwise if *varlist* contains more than one series the data may be scaled. Each line represents an observation, with the data values plotted horizontally.

See also [gnuplot](#).

tobit

Arguments: *depvar indepvars*

Options: `--llimit=lval` (specify left bound)
`--rlimit=rval` (specify right bound)
`--vcv` (print covariance matrix)
`--robust` (robust standard errors)
`--cluster=clustvar` (see [logit](#) for explanation)
`--verbose` (print details of iterations)

Estimates a Tobit model, which may be appropriate when the dependent variable is “censored”. For example, positive and zero values of purchases of durable goods on the part of individual households are observed, and no negative values, yet decisions on such purchases may be thought of as outcomes of an underlying, unobserved disposition to purchase that may be negative in some cases.

By default it is assumed that the dependent variable is censored at zero on the left and is uncensored on the right. However you can use the options `--llimit` and `--rlimit` to specify a different pattern of censoring. Note that if you specify a right bound only, the assumption is then that the dependent variable is uncensored on the left.

The Tobit model is a special case of interval regression, which is supported via the [intreg](#) command.

Menu path: /Model/Limited dependent variable/Tobit

tsls

Arguments: *depvar indepvars ; instruments*

Options: `--no-tests` (don’t do diagnostic tests)
`--vcv` (print covariance matrix)
`--robust` (robust standard errors)
`--cluster=clustvar` (clustered standard errors)
`--liml` (use Limited Information Maximum Likelihood)
`--gmm` (use the Generalized Method of Moments)

Example: `tsls y1 0 y2 y3 x1 x2 ; 0 x1 x2 x3 x4 x5 x6`

Computes Instrumental Variables (IV) estimates, by default using two-stage least squares (TSLS) but see below for further options. The dependent variable is *depvar*, *indepvars* is the list of regressors (which is presumed to include at least one endogenous variable); and *instruments* is the list of instruments (exogenous and/or predetermined variables). If the *instruments* list is not at least as long as *indepvars*, the model is not identified.

In the above example, the *ys* are endogenous and the *xs* are the exogenous variables. Note that exogenous regressors should appear in both lists.

Output for two-stage least squares estimates includes the Hausman test and, if the model is over-identified, the Sargan over-identification test. In the Hausman test, the null hypothesis is that OLS

estimates are consistent, or in other words estimation by means of instrumental variables is not really required. A model of this sort is over-identified if there are more instruments than are strictly required. The Sargan test is based on an auxiliary regression of the residuals from the two-stage least squares model on the full list of instruments. The null hypothesis is that all the instruments are valid, and suspicion is thrown on this hypothesis if the auxiliary regression has a significant degree of explanatory power. For a good explanation of both tests see chapter 8 of [Davidson and MacKinnon \(2004\)](#).

For both TSLS and LIML estimation, an additional test result is shown provided that the model is estimated under the assumption of i.i.d. errors (that is, the `--robust` option is not selected). This is a test for weakness of the instruments. Weak instruments can lead to serious problems in IV regression: biased estimates and/or incorrect size of hypothesis tests based on the covariance matrix, with rejection rates well in excess of the nominal significance level ([Stock et al., 2002](#)). The test statistic is the first-stage F -test if the model contains just one endogenous regressor, otherwise it is the smallest eigenvalue of the matrix counterpart of the first stage F . Critical values based on the Monte Carlo analysis of [Stock and Yogo \(2003\)](#) are shown when available.

The R-squared value printed for models estimated via two-stage least squares is the square of the correlation between the dependent variable and the fitted values.

For details on the effects of the `--robust` and `--cluster` options, please see the help for [ols](#).

As alternatives to TSLS, the model may be estimated via Limited Information Maximum Likelihood (the `--liml` option) or via the Generalized Method of Moments (`--gmm` option). Note that if the model is just identified these methods should produce the same results as TSLS, but if it is over-identified the results will differ in general.

If GMM estimation is selected, the following additional options become available:

- `--two-step`: perform two-step GMM rather than the default of one-step.
- `--iterate`: Iterate GMM to convergence.
- `--weights=Wmat`: specify a square matrix of weights to be used when computing the GMM criterion function. The dimension of this matrix must equal the number of instruments. The default is an appropriately sized identity matrix.

Menu path: /Model/Instrumental variables

var

Arguments: `order ylist [ ; xlist ]`

Options: `--nc` (do not include a constant)
`--trend` (include a linear trend)
`--seasonals` (include seasonal dummy variables)
`--robust` (robust standard errors)
`--robust-hac` (HAC standard errors)
`--quiet` (skip output of individual equations)
`--silent` (don't print anything)
`--impulse-responses` (print impulse responses)
`--variance-decomp` (print variance decompositions)
`--lagselect` (show criteria for lag selection)

Examples: `var 4 x1 x2 x3 ; time mydum`
`var 4 x1 x2 x3 --seasonals`
`var 12 x1 x2 x3 --lagselect`

Sets up and estimates (using OLS) a vector autoregression (VAR). The first argument specifies the lag order — or the maximum lag order in case the `--lagselect` option is given (see below). The order may be given numerically, or as the name of a pre-existing scalar variable. Then follows the setup for the first equation. Do not include lags among the elements of *ylist* — they will be added automatically. The semi-colon separates the stochastic variables, for which *order* lags will be included, from any exogenous variables in *xlist*. Note that a constant is included automatically unless you give the `--nc` flag, a trend can be added with the `--trend` flag, and seasonal dummy variables may be added using the `--seasonals` flag.

While a VAR specification usually includes all lags from 1 to a given maximum, it is possible to select a specific set of lags. To do this, substitute for the regular (scalar) *order* argument either the name of a predefined vector or a comma-separated list of lags, enclosed in braces. We show below two ways of specifying that a VAR should include lags 1, 2 and 4 (but not lag 3):

```
var {1,2,4} ylist
matrix p = {1,2,4}
var p ylist
```

A separate regression is reported for each variable in *ylist*. Output for each equation includes *F*-tests for zero restrictions on all lags of each of the variables, an *F*-test for the significance of the maximum lag, and, if the `--impulse-responses` flag is given, forecast variance decompositions and impulse responses.

Forecast variance decompositions and impulse responses are based on the Cholesky decomposition of the contemporaneous covariance matrix, and in this context the order in which the (stochastic) variables are given matters. The first variable in the list is assumed to be “most exogenous” within-period. The horizon for variance decompositions and impulse responses can be set using the [set](#) command. For retrieval of a specified impulse response function in matrix form, see the [irf](#) function.

If the `--robust` option is given, standard errors are corrected for heteroskedasticity. Alternatively, the `--robust-hac` option can be given to produce standard errors that are robust with respect to both heteroskedasticity and autocorrelation (HAC). In general the latter correction should not be needed if the VAR includes sufficient lags.

If the `--lagselect` option is given, the first parameter to the `var` command is taken as the maximum lag order. Output consists of a table showing the values of the Akaike (AIC), Schwarz (BIC) and Hannan-Quinn (HQC) information criteria computed from VARs of order 1 to the given maximum. This is intended to help with the selection of the optimal lag order. The usual VAR output is not presented. The table of information criteria may be retrieved as a matrix via the `$test` accessor.

Menu path: /Model/Time series/Vector autoregression

varlist

Options: `--scalars` (list scalars)
`--accessors` (list accessor variables)

By default, prints a listing of the (series) variables currently available; `ls` may be used as an alias for this command.

If the `--scalars` option is given, prints a listing of any currently defined scalar variables and their values. Otherwise, if the `--accessors` option is given, prints a list of the internal variables currently available via accessors such as [\\$nobs](#) and [\\$uhat](#).

vartest

Arguments: *series1 series2*

Calculates the F statistic for the null hypothesis that the population variances for the variables *series1* and *series2* are equal, and shows its p-value.

Menu path: /Tools/Test statistic calculator

vecm

Arguments: *order rank ylist* [; *xlist*] [; *rxlist*]

Options: --nc (no constant)
 --rc (restricted constant)
 --uc (unrestricted constant)
 --crt (constant and restricted trend)
 --ct (constant and unrestricted trend)
 --seasonals (include centered seasonal dummies)
 --quiet (skip output of individual equations)
 --silent (don't print anything)
 --impulse-responses (print impulse responses)
 --variance-decomp (print variance decompositions)

Examples: `vecm 4 1 Y1 Y2 Y3`
`vecm 3 2 Y1 Y2 Y3 --rc`
`vecm 3 2 Y1 Y2 Y3 ; X1 --rc`
 See also `denmark.inp`, `hamilton.inp`

A VECM is a form of vector autoregression or VAR (see [var](#)), applicable where the variables in the model are individually integrated of order 1 (that is, are random walks, with or without drift), but exhibit cointegration. This command is closely related to the Johansen test for cointegration (see [coint2](#)).

The *order* parameter to this command represents the lag order of the VAR system. The number of lags in the VECM itself (where the dependent variable is given as a first difference) is one less than *order*.

The *rank* parameter represents the cointegration rank, or in other words the number of cointegrating vectors. This must be greater than zero and less than or equal to (generally, less than) the number of endogenous variables given in *ylist*.

ylist supplies the list of endogenous variables, in levels. The inclusion of deterministic terms in the model is controlled by the option flags. The default if no option is specified is to include an “unrestricted constant”, which allows for the presence of a non-zero intercept in the cointegrating relations as well as a trend in the levels of the endogenous variables. In the literature stemming from the work of Johansen (see for example his 1995 book) this is often referred to as “case 3”. The first four options given above, which are mutually exclusive, produce cases 1, 2, 4 and 5 respectively. The meaning of these cases and the criteria for selecting a case are explained in the *Gretl User's Guide*.

The optional lists *xlist* and *rxlist* allow you to specify sets of exogenous variables which enter the model either unrestrictedly (*xlist*) or restricted to the cointegration space (*rxlist*). These lists are separated from *ylist* and from each other by semicolons.

The --seasonals option, which may be combined with any of the other options, specifies the inclusion of a set of centered seasonal dummy variables. This option is available only for quarterly or monthly data.

The first example above specifies a VECM with lag order 4 and a single cointegrating vector. The endogenous variables are Y1, Y2 and Y3. The second example uses the same variables but specifies a lag order of 3 and two cointegrating vectors; it also specifies a “restricted constant”, which is appropriate if the cointegrating vectors may have a non-zero intercept but the Y variables have no

trend.

Following estimation of a VECM some special accessors are available: `$jalpha`, `$jbeta` and `$jvbeta` retrieve, respectively, the α and β matrices and the estimated variance of β . For retrieval of a specified impulse response function in matrix form, see the `irf` function.

Menu path: /Model/Time series/VECM

vif

Must follow the estimation of a model which includes at least two independent variables. Calculates and displays the Variance Inflation Factors (VIFs) for the regressors. The VIF for regressor j is defined as

$$\frac{1}{1 - R_j^2}$$

where R_j is the coefficient of multiple correlation between regressor j and the other regressors. The factor has a minimum value of 1.0 when the variable in question is orthogonal to the other independent variables. [Neter *et al.* \(1990\)](#) suggest inspecting the largest VIF as a diagnostic for collinearity; a value greater than 10 is sometimes taken as indicating a problematic degree of collinearity.

Menu path: Model window, /Tests/Collinearity

wls

Arguments: *wtvar depvar indepvars*

Options: --vcv (print covariance matrix)
 --robust (robust standard errors)
 --quiet (suppress printing of results)

Computes weighted least squares (WLS) estimates using *wtvar* as the weight, *depvar* as the dependent variable, and *indepvars* as the list of independent variables. Let w denote the positive square root of *wtvar*; then WLS is basically equivalent to an OLS regression of $w * \text{depvar}$ on $w * \text{indepvars}$. The R -squared, however, is calculated in a special manner, namely as

$$R^2 = 1 - \frac{\text{ESS}}{\text{WTSS}}$$

where ESS is the error sum of squares (sum of squared residuals) from the weighted regression and WTSS denotes the “weighted total sum of squares”, which equals the sum of squared residuals from a regression of the weighted dependent variable on the weighted constant alone.

If *wtvar* is a dummy variable, WLS estimation is equivalent to eliminating all observations with value zero for *wtvar*.

Menu path: /Model/Other linear models/Weighted Least Squares

xcrrgm

Arguments: *series1 series2 [order]*

Option: --plot=*mode-or-filename* (see below)

Example: xcrrgm x y 12

Prints and graphs the cross-correlogram for *series1* and *series2*, which may be specified by name or number. The values are the sample correlation coefficients between the current value of *series1* and successive leads and lags of *series2*.

If an *order* value is specified the length of the cross-correlogram is limited to at most that number of leads and lags, otherwise the length is determined automatically, as a function of the frequency of the data and the number of observations.

By default, a plot of the cross-correlogram is produced: a gnuplot graph in interactive mode or an ASCII graphic in batch mode. This can be adjusted via the `--plot` option. The acceptable parameters to this option are `none` (to suppress the plot); `ascii` (to produce a text graphic even when in interactive mode); `display` (to produce a gnuplot graph even when in batch mode); or a file name. The effect of providing a file name is as described for the `--output` option of the [gnuplot](#) command.

Menu path: /View/Cross-correlogram

Other access: Main window pop-up menu (multiple selection)

xtab

Arguments: *ylist* [; *xlist*]

Options: `--row` (display row percentages)
`--column` (display column percentages)
`--zeros` (display zero entries)
`--matrix=matname` (use frequencies from named matrix)

Displays a contingency table or cross-tabulation for each combination of the variables included in *ylist*; if a second list *xlist* is given, each variable in *ylist* is cross-tabulated by row against each variable in *xlist* (by column). Variables in these lists can be referenced by name or by number. Note that all the variables must have been marked as discrete. Alternatively, if the `--matrix` option is given, treat the named matrix as a precomputed set of frequencies and display this as a cross-tabulation.

By default the cell entries are given as frequency counts. The `--row` and `--column` options (which are mutually exclusive), replace the counts with the percentages for each row or column, respectively. By default, cells with a zero count are left blank; the `--zeros` option, which has the effect of showing zero counts explicitly, may be useful for importing the table into another program, such as a spreadsheet.

Pearson's chi-square test for independence is displayed if the expected frequency under independence is at least $1.0e-7$ for all cells. A common rule of thumb for the validity of this statistic is that at least 80 percent of cells should have expected frequencies of 5 or greater; if this criterion is not met a warning is printed.

If the contingency table is 2 by 2, Fisher's Exact Test for independence is computed. Note that this test is based on the assumption that the row and column totals are fixed, which may or may not be appropriate depending on how the data were generated. The left p-value should be used when the alternative to independence is negative association (values tend to cluster in the lower left and upper right cells); the right p-value should be used if the alternative is positive association. The two-tailed p-value for this test is calculated by method (b) in section 2.1 of [Agresti \(1992\)](#): it is the sum of the probabilities of all possible tables having the given row and column totals and having a probability less than or equal to that of the observed table.

1.3 Commands by topic

The following sections show the available commands grouped by topic.

Estimation

ar	Autoregressive estimation	ar1	AR(1) estimation
arbond	Arellano-Bond	arch	ARCH model
arima	ARMA model	biprobit	Bivariate probit
dpanel	Dynamic panel models	duration	Duration models
equation	Define equation within a system	estimate	Estimate system of equations

<code>garch</code>	GARCH model	<code>gmm</code>	GMM estimation
<code>heckit</code>	Heckman selection model	<code>hsk</code>	Heteroskedasticity-corrected estimates
<code>intreg</code>	Interval regression model	<code>kalman</code>	Kalman filter
<code>lad</code>	Least Absolute Deviation estimation	<code>logistic</code>	Logistic regression
<code>logit</code>	Logit regression	<code>mle</code>	Maximum likelihood estimation
<code>mpols</code>	Multiple-precision OLS	<code>negbin</code>	Negative Binomial regression
<code>nls</code>	Nonlinear Least Squares	<code>ols</code>	Ordinary Least Squares
<code>panel</code>	Panel models	<code>poisson</code>	Poisson estimation
<code>probit</code>	Probit model	<code>quantreg</code>	Quantile regression
<code>system</code>	Systems of equations	<code>tobit</code>	Tobit model
<code>tsls</code>	Instrumental variables regression	<code>var</code>	Vector Autoregression
<code>vecm</code>	Vector Error Correction Model	<code>wls</code>	Weighted Least Squares

Tests

<code>add</code>	Add variables to model	<code>adf</code>	Augmented Dickey-Fuller test
<code>chow</code>	Chow test	<code>coeffsum</code>	Sum of coefficients
<code>coint</code>	Engle-Granger cointegration test	<code>coint2</code>	Johansen cointegration test
<code>cusum</code>	CUSUM test	<code>difftest</code>	Nonparametric tests for differences
<code>hausman</code>	Panel diagnostics	<code>kpss</code>	KPSS stationarity test
<code>leverage</code>	Influential observations	<code>levinlin</code>	Levin-Lin-Chu test
<code>meantest</code>	Difference of means	<code>modtest</code>	Model tests
<code>normtest</code>	Normality test	<code>omit</code>	Omit variables
<code>qlrtest</code>	Quandt likelihood ratio test	<code>reset</code>	Ramsey's RESET
<code>restrict</code>	Testing restrictions	<code>runs</code>	Runs test
<code>vartest</code>	Difference of variances	<code>vif</code>	Variance Inflation Factors

Transformations

<code>diff</code>	First differences	<code>discrete</code>	Mark variables as discrete
<code>dummify</code>	Create sets of dummies	<code>lags</code>	Create lags
<code>ldiff</code>	Log-differences	<code>logs</code>	Create logs
<code>orthdev</code>	Orthogonal deviations	<code>sdiff</code>	Seasonal differencing
<code>square</code>	Create squares of variables		

Statistics

<code>anova</code>	ANOVA	<code>corr</code>	Correlation coefficients
<code>corrgm</code>	Correlogram	<code>fractint</code>	Fractional integration
<code>freq</code>	Frequency distribution	<code>hurst</code>	Hurst exponent
<code>mahal</code>	Mahalanobis distances	<code>pca</code>	Principal Components Analysis
<code>pergm</code>	Periodogram	<code>spearman</code>	Spearman's rank correlation
<code>summary</code>	Descriptive statistics	<code>xcorgm</code>	Cross-correlogram
<code>xtab</code>	Cross-tabulate variables		

Dataset

<code>append</code>	Append data	<code>data</code>	Import from database
<code>dataset</code>	Manipulate the dataset	<code>delete</code>	Delete variables
<code>genr</code>	Generate a new variable	<code>info</code>	Information on data set
<code>join</code>	Manage data sources	<code>labels</code>	Labels for variables
<code>markers</code>	Observation markers	<code>nulldata</code>	Creating a blank dataset
<code>open</code>	Open a data file	<code>rename</code>	Rename variables
<code>setinfo</code>	Edit attributes of variable	<code>setmiss</code>	Missing value code
<code>setobs</code>	Set frequency and starting observation	<code>smp1</code>	Set the sample range
<code>store</code>	Save data	<code>varlist</code>	Listing of variables

Graphs

<code>boxplot</code>	Boxplots	<code>gnuplot</code>	Create a gnuplot graph
<code>graphpg</code>	Gretl graph page	<code>plot</code>	
<code>qqplot</code>	Q-Q plot	<code>rmplot</code>	Range-mean plot
<code>scatters</code>	Multiple pairwise graphs	<code>textplot</code>	ASCII plot

Printing

<code>eqnprint</code>	Print model as equation	<code>modprint</code>	Print a user-defined model
<code>outfile</code>	Direct printing to file	<code>print</code>	Print data or strings
<code>printf</code>	Formatted printing	<code>sprintf</code>	Printing to a string
<code>tabprint</code>	Print model in tabular form		

Prediction

`fcast` Generate forecasts

Programming

<code>break</code>	Break from loop	<code>catch</code>	Catch errors
<code>clear</code>		<code>debug</code>	Debugging
<code>elif</code>	Flow control	<code>else</code>	Flow control
<code>end</code>	End block of commands	<code>endif</code>	Flow control
<code>endloop</code>	End a command loop	<code>flush</code>	
<code>foreign</code>	Non-native script	<code>function</code>	Define a function
<code>if</code>	Flow control	<code>include</code>	Include function definitions
<code>loop</code>	Start a command loop	<code>makepkg</code>	Make function package
<code>run</code>	Execute a script	<code>set</code>	Set program parameters
<code>setopt</code>	Set options for next command		

Utilities

<code>eval</code>		<code>help</code>	Help on commands
<code>install</code>		<code>modeltab</code>	The model table
<code>pvalue</code>	Compute p-values	<code>quit</code>	Exit the program
<code>shell</code>	Execute shell commands		

1.4 Short-form command options

As can be seen from section 1.2, the behavior of many gretl commands can be modified via the use of option flags. These take the form of two dashes followed by a string which is somewhat descriptive of the effect of the option.

Some options require a parameter, which must be joined to the option “flag” with an equals sign. Among the options that do *not* require a parameter, certain common ones have a short form—a single dash followed by a single letter—and it is considered idiomatic to use the short forms in hansl scripts. The table below shows the relevant mapping: for any command which supports the long-form option in the first column, the short form in the second column is also supported.

long form	short form
--verbose	-v
--quiet	-q
--robust	-r
--hessian	-h
--window	-w

Chapter 2

Gretl functions

2.1 Introduction

This chapter presents two alphabetical listings: first, the “accessors” which enable the user to retrieve the values of internal variables; and second, the functions proper that are available in the context of the `genr` command.

2.2 Accessors

\$ahat

Output: series

Must follow the estimation of a fixed-effect panel data model. Returns a series containing the estimates of the individual fixed effects (per-unit intercepts).

\$aic

Output: scalar

Returns the Akaike Information Criterion for the last estimated model, if available. See the *Gretl User's Guide* for details of the calculation.

\$bic

Output: scalar

Returns Schwarz's Bayesian Information Criterion for the last estimated model, if available. See the *Gretl User's Guide* for details of the calculation.

\$chisq

Output: scalar

Returns the overall chi-square statistic from the last estimated model, if available.

\$coeff

Output: matrix or scalar

Argument: *s* (name of coefficient, optional)

With no arguments, `$coeff` returns a column vector containing the estimated coefficients for the last model. With the optional string argument it returns a scalar, namely the estimated parameter named *s*. See also [\\$stderr](#), [\\$vcv](#).

Example:

```
open bjg
arima 0 1 1 ; 0 1 1 ; lg
b = $coeff          # gets a vector
```



```
macoef = $coeff(theta_1) # gets a scalar
```

If the “model” in question is actually a system, the result depends on the characteristics of the system: for VARs and VECMs the value returned is a matrix with one column per equation, otherwise it is a column vector containing the coefficients from the first equation followed by those from the second equation, and so on.

\$command

Output: string

Must follow the estimation of a model; returns the command word, for example `ols` or `probit`.

\$companion

Output: matrix

Must follow the estimation of a VAR or a VECM; returns the companion matrix.

\$datatype

Output: scalar

Returns an integer value representing the sort of dataset that is currently loaded: 0 = no data; 1 = cross-sectional (undated) data; 2 = time-series data; 3 = panel data.

\$depvar

Output: string

Must follow the estimation of a single-equation model; returns the name of the dependent variable.

\$df

Output: scalar

Returns the degrees of freedom of the last estimated model. If the last model was in fact a system of equations, the value returned is the degrees of freedom per equation; if this differs across the equations then the value given is the number of observations minus the mean number of coefficients per equation (rounded up to the nearest integer).

\$dwpval

Output: scalar

Returns the p-value for the Durbin-Watson statistic for the model last estimated (if available), computed using the Imhof procedure.

Due to the limited precision of computer arithmetic, the Imhof integral can go negative when the Durbin-Watson statistic is close to its lower bound. In that case the accessor returns NA. Since any other failure mode results in an error being flagged it is probably safe to assume that an NA value means the true p-value is “very small”, although we are unable to quantify it.

\$ec

Output: matrix

Must follow the estimation of a VECM; returns a matrix containing the error correction terms. The number of rows equals the number of observations used and the number of columns equals the cointegration rank of the system.

\$error

Output: scalar

Returns the program's internal error code, which will be non-zero in case an error has occurred but has been trapped using [catch](#). Note that using this accessor causes the internal error code to be reset to zero. If you want to get the error message associated with a given `$error` you need to store the value in a temporary variable, as in

```
err = $error
if (err)
    printf "Got error %d (%s)\n", err, errmsg(err);
endif
```

See also [catch](#), [errmsg](#).

\$ess

Output: scalar

Returns the error sum of squares of the last estimated model, if available.

\$evals

Output: matrix

Must follow the estimation of a VECM; returns a vector containing the eigenvalues that are used in computing the trace test for cointegration.

\$fcast

Output: matrix

Must follow the [fcast](#) forecasting command; returns the forecast values as a matrix. If the model on which the forecast was based is a system of equations the returned matrix will have one column per equation, otherwise it is a column vector.

\$fcerr

Output: matrix

Must follow the [fcast](#) forecasting command; returns the standard errors of the forecasts, if available, as a matrix. If the model on which the forecast was based is a system of equations the returned matrix will have one column per equation, otherwise it is a column vector.

\$fevd

Output: matrix

Must follow estimation of a VAR. Returns a matrix containing the forecast error variance decomposition (FEVD). This matrix has h rows where h is the forecast horizon, which can be chosen using `set horizon` or otherwise is set automatically based on the frequency of the data.

For a VAR with p variables, the matrix has p^2 columns: the first p columns contain the FEVD for the first variable in the VAR; the second p columns the FEVD for the second variable; and so on.

The (decimal) fraction of the forecast error for variable i attributable to innovation in variable j is therefore found in column $(i - 1)p + j$.

\$Fstat

Output: scalar

Returns the overall F-statistic from the last estimated model, if available.

\$gmmcrit

Output: scalar

Must follow a `gmm` block. Returns the value of the GMM objective function at its minimum.

\$h

Output: series

Must follow a `garch` command. Returns the estimated conditional variance series.

\$hausman

Output: row vector

Must follow estimation of a model via either `tsls` or `panel` with the random effects option. Returns a 1×3 vector containing the value of the Hausman test statistic, the corresponding degrees of freedom and the p-value for the test, in that order.

\$hqic

Output: scalar

Returns the Hannan-Quinn Information Criterion for the last estimated model, if available. See the *Gretl User's Guide* for details of the calculation.

\$huge

Output: scalar

Returns a very large positive number. By default this is 1.0E100, but the value can be changed using the `set` command.

\$jalpha

Output: matrix

Must follow the estimation of a VECM, and returns the loadings matrix. It has as many rows as variables in the VECM and as many columns as the cointegration rank.

\$jbeta

Output: matrix

Must follow the estimation of a VECM, and returns the cointegration matrix. It has as many rows as variables in the VECM (plus the number of exogenous variables that are restricted to the cointegration space, if any), and as many columns as the cointegration rank.

\$jvbeta

Output: square matrix

Must follow the estimation of a VECM, and returns the estimated covariance matrix for the elements of the cointegration vectors.

In the case of unrestricted estimation, this matrix has a number of rows equal to the unrestricted elements of the cointegration space after the Phillips normalization. If, however, a restricted system is estimated via the `restrict` command with the `--full` option, a singular matrix with $(n + m)r$ rows will be returned (n being the number of endogenous variables, m the number of exogenous variables that are restricted to the cointegration space, and r the cointegration rank).

Example: the code

```
open denmark.gdt
vecm 2 1 LRM LRY IBO IDE --rc --seasonals -q
s0 = $jvbeta

restrict --full
  b[1,1] = 1
  b[1,2] = -1
  b[1,3] + b[1,4] = 0
end restrict
s1 = $jvbeta

print s0
print s1
```

produces the following output.

```
s0 (4 x 4)

0.019751    0.029816   -0.00044837    -0.12227
0.029816    0.31005    -0.45823    -0.18526
-0.00044837  -0.45823     1.2169    -0.035437
-0.12227    -0.18526    -0.035437     0.76062

s1 (5 x 5)

0.0000    0.0000    0.0000    0.0000    0.0000
0.0000    0.0000    0.0000    0.0000    0.0000
0.0000    0.0000    0.27398   -0.27398   -0.019059
0.0000    0.0000   -0.27398    0.27398    0.019059
0.0000    0.0000   -0.019059   0.019059   0.0014180
```

\$lang

Output: string

Returns a string representing the national language in force currently, if this can be determined. The string is composed of a two-letter ISO 639-1 language code (for example, `en` for English, `jp` for Japanese, `el` for Greek) followed by an underscore plus a two-letter ISO 3166-1 country code. Thus for example Portuguese in Portugal gives `pt_PT` while Portuguese in Brazil gives `pt_BR`.

If the national language cannot be determined, the string “unknown” is returned.

\$llt

Output: series

For selected models estimated via Maximum Likelihood, returns the series of per-observation log-likelihood values. At present this is supported only for binary logit and probit, tobit and heckit.

\$lnl

Output: scalar

Returns the log-likelihood for the last estimated model (where applicable).

\$macheps

Output: scalar

Returns the value of “machine epsilon”, which gives an upper bound on the relative error due to rounding in double-precision floating point arithmetic.

\$mnlprobs

Output: matrix

Following estimation of a multinomial logit model (only), retrieves a matrix holding the estimated probabilities of each possible outcome at each observation in the model’s sample range. Each row represents an observation and each column an outcome.

\$ncoeff

Output: integer

Returns the total number of coefficients estimated in the last model.

\$nobs

Output: integer

Returns the number of observations in the currently selected sample.

\$nvars

Output: integer

Returns the number of variables in the dataset (including the constant).

\$obsdate

Output: series

Applicable when the current dataset is time-series with annual, quarterly, monthly or decennial frequency, or is dated daily or weekly, or when the dataset is a panel with time-series information set appropriately (see the [setobs](#) command). The returned series holds 8-digit numbers on the pattern YYYYMMDD (ISO 8601 “basic” date format), which correspond to the day of the observation, or the first day of the observation period in case of a time-series frequency less than daily.

Such a series can be helpful when using the [join](#) command.

\$obsmajor

Output: series

Applicable when the observations in the current dataset have a major:minor structure, as in quarterly time series (year:quarter), monthly time series (year:month), hourly data (day:hour) and panel data (individual:period). Returns a series holding the major or low-frequency component of each observation (for example, the year).

See also [\\$obsminor](#), [\\$obsmicro](#).

\$obsmicro

Output: series

Applicable when the observations in the current dataset have a major:minor:micro structure, as in dated daily time series (year:month:day). Returns a series holding the micro or highest-frequency component of each observation (for example, the day).

See also [\\$obsmajor](#), [\\$obsminor](#).

\$obsminor

Output: series

Applicable when the observations in the current dataset have a major:minor structure, as in quarterly time series (year:quarter), monthly time series (year:month), hourly data (day:hour) and panel data (individual:period). Returns a series holding the minor or high-frequency component of each observation (for example, the month).

See also [\\$obsmajor](#), [\\$obsmicro](#).

\$pd

Output: integer

Returns the frequency or periodicity of the data (e.g. 4 for quarterly data). In the case of panel data the value returned is the time-series length.

\$pi

Output: scalar

Returns the value of π in double precision.

\$pvalue

Output: scalar or matrix

Returns the p-value of the test statistic that was generated by the last explicit hypothesis-testing command, if any (e.g. `chow`). See the *Gretl User's Guide* for details.

In most cases the return value is a scalar but sometimes it is a matrix (for example, the trace and lambda-max p-values from the Johansen cointegration test); in that case the values in the matrix are laid out in the same pattern as the printed results.

See also [\\$test](#).

\$qlrbreak

Output: scalar

Must follow an invocation of the `qlrtest` command (the QLR test for a structural break at an unknown point). The value returned is the 1-based index of the observation at which the test statistic is maximized.

\$rho

Output: scalar
Argument: *n* (scalar, optional)

Without arguments, returns the first-order autoregressive coefficient for the residuals of the last model. After estimating a model via the `ar` command, the syntax `$rho(n)` returns the corresponding estimate of $\rho(n)$.

\$rsq

Output: scalar

Returns the unadjusted R^2 from the last estimated model, if available.

\$sample

Output: series

Must follow estimation of a single-equation model. Returns a dummy series with value 1 for observations used in estimation, 0 for observations within the currently defined sample range but not used (presumably because of missing values), and NA for observations outside of the current range.

If you wish to compute statistics based on the sample that was used for a given model, you can do, for example:

```
ols y 0 xlist
genr sdum = $sample
smp1 sdum --dummy
```

\$sargan

Output: row vector

Must follow a `ts1s` command. Returns a 1×3 vector, containing the value of the Sargan over-identification test statistic, the corresponding degrees of freedom and p-value, in that order. If the model is exactly identified, the statistic is unavailable, and trying to access it provokes an error.

\$sigma

Output: scalar or matrix

Requires that a model has been estimated. If the last model was a single equation, returns the (scalar) Standard Error of the Regression (or in other words, the standard deviation of the residuals, with an appropriate degrees of freedom correction). If the last model was a system of equations, returns the cross-equation covariance matrix of the residuals.

\$stderr

Output: matrix or scalar
Argument: *s* (name of coefficient, optional)

With no arguments, `$stderr` returns a column vector containing the standard error of the coefficients for the last model. With the optional string argument it returns a scalar, namely the standard error of the parameter named *s*.

If the “model” in question is actually a system, the result depends on the characteristics of the system: for VARs and VECMs the value returned is a matrix with one column per equation, otherwise it is a column vector containing the coefficients from the first equation followed by those from the second equation, and so on.

See also [\\$coeff](#), [\\$vcv](#).

\$stopwatch

Output: scalar

Must be preceded by `set stopwatch`, which activates the measurement of CPU time. The first use of this accessor yields the seconds of CPU time that have elapsed since the `set stopwatch` command. At each access the clock is reset, so subsequent uses of `$stopwatch` yield the seconds of CPU time since the previous access.

\$sysA

Output: matrix

Must follow estimation of a simultaneous equations system. Returns the matrix of coefficients on the lagged endogenous variables, if any, in the structural form of the system. See the [system](#) command.

\$sysB

Output: matrix

Must follow estimation of a simultaneous equations system. Returns the matrix of coefficients on the exogenous variables in the structural form of the system. See the [system](#) command.

\$sysGamma

Output: matrix

Must follow estimation of a simultaneous equations system. Returns the matrix of coefficients on the contemporaneous endogenous variables in the structural form of the system. See the [system](#) command.

\$sysinfo

Output: bundle

Returns a bundle containing information on the capabilities of the gretl build and the system on which gretl is running. The members of the bundle are as follows:

- `mpi`: integer, equals 1 if the system supports MPI (Message Passing Interface), otherwise 0.
- `omp`: integer, equals 1 if gretl is built with support for Open MP, otherwise 0.
- `nproc`: integer, the number of processors available.
- `mpimax`: integer, the maximum number of MPI processes that can be run in parallel. This is zero if MPI is not supported, otherwise it equals the local `nproc` value unless an MPI hosts file has been specified, in which case it is the sum of the number of processors or “slots” across all the machines referenced in that file.

- `wordlen`: integer, either 32 or 64 for 32- and 64-bit systems respectively.
- `os`: string representing the operating system, either `linux`, `osx`, `windows` or `other`.
- `hostname`: the name of the host machine on which the current gretl process is running (with a fallback of `localhost` in case the name cannot be determined).

Note that individual elements in the bundle can be accessed using “dot” notation without any need to copy the whole bundle under a user-specified name. For example,

```
if $sysinfo.os == "linux"
    # do something linux-specific
endif
```

\$T

Output: integer

Returns the number of observations used in estimating the last model.

\$t1

Output: integer

Returns the 1-based index of the first observation in the currently selected sample.

\$t2

Output: integer

Returns the 1-based index of the last observation in the currently selected sample.

\$test

Output: scalar or matrix

Returns the value of the test statistic that was generated by the last explicit hypothesis-testing command, if any (e.g. `chow`). See the *Gretl User's Guide* for details.

In most cases the return value is a scalar but sometimes it is a matrix (for example, the trace and lambda-max statistics from the Johansen cointegration test); in that case the values in the matrix are laid out in the same pattern as the printed results.

See also [\\$pvalue](#).

\$trsq

Output: scalar

Returns TR^2 (sample size times R-squared) from the last model, if available.

\$uhat

Output: series

Returns the residuals from the last model. This may have different meanings for different estimators. For example, after an ARMA estimation `$uhat` will contain the one-step-ahead forecast error; after a probit model, it will contain the generalized residuals.

If the “model” in question is actually a system (a VAR or VECM, or system of simultaneous equations), `$uhat` with no parameters retrieves the matrix of residuals, one column per equation.

\$unit

Output: series

Valid for panel datasets only. Returns a series with value 1 for all observations on the first unit or group, 2 for observations on the second unit, and so on.

\$vcv

Output: matrix or scalar

Arguments: *s1* (name of coefficient, optional)
s2 (name of coefficient, optional)

With no arguments, `$vcv` returns a square matrix containing the estimated covariance matrix for the coefficients of the last model. If the last model was a single equation, then you may supply the names of two parameters in parentheses to retrieve the estimated covariance between the parameters named *s1* and *s2*. See also [\\$coeff](#), [\\$stderr](#).

This accessor is not available for VARs or VECMs; for models of that sort see [\\$sigma](#) and [\\$txtinv](#).

\$vecGamma

Output: matrix

Must follow the estimation of a VECM; returns a matrix in which the Gamma matrices (coefficients on the lagged differences of the cointegrated variables) are stacked side by side. Each row represents an equation; for a VECM of lag order p there are $p - 1$ sub-matrices.

\$version

Output: scalar

Returns an integer value that codes for the program version. The gretl version string takes the form *x.y.z* (for example, 1.9.10). The return value from this accessor is formed as $10000*x + 100*y + z$, so that 1.9.10 translates as 10910.

\$vma

Output: matrix

Must follow the estimation of a VAR or a VECM; returns a matrix containing the VMA representation up to the order specified via the `set horizon` command. See the *Gretl User's Guide* for details.

\$windows

Output: integer

Returns 1 if gretl is running on MS Windows, otherwise 0. By conditioning on the value of this variable you can write shell calls that are portable across different operating systems.

Also see the [shell](#) command.

\$xlist

Output: list

If the last model was a single equation, returns the list of regressors. If the last model was a system of equations, returns the “global” list of exogenous and predetermined variables (in the same order in which they appear in [\\$sysB](#)). If the last model was a VAR, returns the list of exogenous regressors, if any.

\$xtxinv

Output: matrix

Following estimation of a VAR or VECM (only), returns $X'X^{-1}$, where X is the common matrix of regressors used in each of the equations. This accessor is not available for a VECM estimated with a restriction imposed on α , the “loadings” matrix.

\$yhat

Output: series

Returns the fitted values from the last regression.

\$ylist

Output: list

If the last model estimated was a VAR, VECM or simultaneous system, returns the associated list of endogenous variables. If the last model was a single equation, this accessor gives a list with a single element, the dependent variable. In the special case of the biprobit model the list contains two elements.

2.3 Functions proper**abs**

Output: same type as input

Argument: x (scalar, series or matrix)

Returns the absolute value of x .

acos

Output: same type as input

Argument: x (scalar, series or matrix)

Returns the arc cosine of x , that is, the value whose cosine is x . The result is in radians; the input should be in the range -1 to 1 .

acosh

Output: same type as input

Argument: x (scalar, series or matrix)

Returns the inverse hyperbolic cosine of x (positive solution). x should be greater than 1 ; otherwise, NA is returned. See also [cosh](#).

aggregate

Output: matrix

Arguments: x (series or list)

$byvar$ (series or list)

$funcname$ (string)

In the simplest version, both x and $byvar$ are individual series. In that case this function returns a matrix with three columns: the first holds the distinct values of $byvar$, sorted in ascending order; the second holds the count of observations at which $byvar$ takes on each of these values; and the

third holds the values of the statistic specified by *funcname* calculated on series *x*, using only those observations at which *byvar* takes on the value given in the first column.

More generally, if *byvar* is a list with *n* members then the left-hand *n* columns hold the combinations of the distinct values of each of the *n* series and the count column holds the number of observations at which each combination is realized. If *x* is a list with *m* members then the rightmost *m* columns hold the values of the specified statistic for each of the *x* variables, again calculated on the sub-sample indicated in the first column(s).

The following values of *funcname* are supported “natively”: [sum](#), [sumall](#), [mean](#), [sd](#), [var](#), [sst](#), [skewness](#), [kurtosis](#), [min](#), [max](#), [median](#), [nobs](#) and [gini](#). Each of these functions takes a series argument and returns a scalar value, and in that sense can be said to “aggregate” the series in some way. You may give the name of a user-defined function as the aggregator; like the built-ins, such a function must take a single series argument and return a scalar value.

Note that although a count of cases is provided automatically the *nobs* function is not redundant as an aggregator, since it gives the number of valid (non-missing) observations on *x* at each *byvar* combination.

For a simple example, suppose that *region* represents a coding of geographical region using integer values 1 to *n*, and *income* represents household income. Then the following would produce an $n \times 3$ matrix holding the region codes, the count of observations in each region, and mean household income for each of the regions:

```
matrix m = aggregate(income, region, mean)
```

For an example using lists, let *gender* be a male/female dummy variable, let *race* be a categorical variable with three values, and consider the following:

```
list BY = gender race
list X = income age
matrix m = aggregate(X, BY, sd)
```

The *aggregate* call here will produce a 6×5 matrix. The first two columns hold the 6 distinct combinations of gender and race values; the middle column holds the count for each of these combinations; and the rightmost two columns contain the sample standard deviations of *income* and *age*.

Note that if *byvar* is a list, some combinations of the *byvar* values may not be present in the data (giving a count of zero). In that case the value of the statistics for *x* are recorded as NaN (not a number). If you want to ignore such cases you can use the [selifr](#) function to select only those rows that have a non-zero count. The column to test is one place to the right of the number of *byvar* variables, so we can do:

```
matrix m = aggregate(X, BY, sd)
scalar c = nelem(BY)
m = selifr(m, m[,c+1])
```

argname

Output: string

Argument: *s* (string)

For *s* the name of a parameter to a user-defined function, returns the name of the corresponding argument, or an empty string if the argument was anonymous.

asin

Output: same type as input
 Argument: x (scalar, series or matrix)

Returns the arc sine of x , that is, the value whose sine is x . The result is in radians; the input should be in the range -1 to 1 .

asinh

Output: same type as input
 Argument: x (scalar, series or matrix)

Returns the inverse hyperbolic sine of x . See also [sinh](#).

atan

Output: same type as input
 Argument: x (scalar, series or matrix)

Returns the arc tangent of x , that is, the value whose tangent is x . The result is in radians.

atanh

Output: same type as input
 Argument: x (scalar, series or matrix)

Returns the inverse hyperbolic tangent of x . See also [tanh](#).

atof

Output: scalar
 Argument: s (string)

Closely related to the C library function of the same name. Returns the result of converting the string s (or the leading portion thereof, after discarding any initial white space) to a floating-point number. Unlike C's `atof`, however, the decimal character is always assumed (for reasons of portability) to be `."`. Any characters that follow the portion of s that converts to a floating-point number under this assumption are ignored.

If none of s (following any discarded white space) is convertible under the stated assumption, NA is returned.

```
# examples
x = atof("1.234") # gives x = 1.234
x = atof("1,234") # gives x = 1
x = atof("1.2y")  # gives x = 1.2
x = atof("y")     # gives x = NA
x = atof(",234")  # gives x = NA
```

See also [sscanf](#) for more flexible string to numeric conversion.

bessel

Output: same type as input
 Arguments: *type* (character)
 v (scalar)
 x (scalar, series or matrix)

Computes one of the Bessel function variants for order v and argument x . The return value is of the same type as x . The specific function is selected by the first argument, which must be J, Y, I, or K. A good discussion of the Bessel functions can be found on Wikipedia; here we give a brief account.

case J: Bessel function of the first kind. Resembles a damped sine wave. Defined for real v and x , but if x is negative then v must be an integer.

case Y: Bessel function of the second kind. Defined for real v and x but has a singularity at $x = 0$.

case I: Modified Bessel function of the first kind. An exponentially growing function. Acceptable arguments are as for case J.

case K: Modified Bessel function of the second kind. An exponentially decaying function. Diverges at $x = 0$ and is not defined for negative x . Symmetric around $v = 0$.

BFGSmax

Output: scalar
 Arguments: $\&b$ (reference to matrix)
 f (function call)
 g (function call, optional)

Numerical maximization via the method of Broyden, Fletcher, Goldfarb and Shanno. On input the vector b should hold the initial values of a set of parameters, and the argument f should specify a call to a function that calculates the (scalar) criterion to be maximized, given the current parameter values and any other relevant data. If the object is in fact minimization, this function should return the negative of the criterion. On successful completion, BFGSmax returns the maximized value of the criterion, and b holds the parameter values which produce the maximum.

The optional third argument provides a means of supplying analytical derivatives (otherwise the gradient is computed numerically). The gradient function call g must have as its first argument a pre-defined matrix that is of the correct size to contain the gradient, given in pointer form. It also must take the parameter vector as an argument (in pointer form or otherwise). Other arguments are optional.

For more details and examples see the chapter on numerical methods in the *Gretl User's Guide*. See also [NRmax](#), [fdjac](#), [simann](#).

bkfilt

Output: series
 Arguments: y (series)
 $f1$ (integer, optional)
 $f2$ (integer, optional)
 k (integer, optional)

Returns the result from application of the Baxter–King bandpass filter to the series y . The optional parameters $f1$ and $f2$ represent, respectively, the lower and upper bounds of the range of frequencies to extract, while k is the approximation order to be used. If these arguments are not supplied then the following default values are used: $f1 = 8$, $f2 = 32$, $k = 8$.

If $f2$ is greater than or equal to the number of available observations, then the “low-pass” version of the filter will be run and the resulting series should be taken as an estimate of the trend component, rather than the cycle. See also [bwfilt](#), [hpfilt](#).

boxcox

Output: series
 Arguments: *y* (series)
 d (scalar)

Returns the Box-Cox transformation with parameter *d* for the positive series *y*.

$$y_t^{(d)} = \begin{cases} \frac{y_t^d - 1}{d} & \text{if } d \neq 0 \\ \log(y_t) & \text{if } d = 0 \end{cases}$$

bread

Output: bundle
 Arguments: *fname* (string)
 import (boolean, optional)

Reads a bundle from a text file. The string *fname* must contain the name of the file from which the bundle is to be read. If this name has the suffix “.gz” it is assumed that gzip compression has been applied in writing the file.

The file in question should be an appropriately defined XML file: it should contain a `gretl-bundle` element, which is used to store zero or more `bundled-item` elements. For example,

```
<gretl-bundle name="temp">
<bundled-item key="s" type="string">moo</bundled-item>
<bundled-item key="x" type="scalar">3</bundled-item>
</gretl-bundle>
```

As you may expect, such files are generated automatically by the companion function [bwrite](#).

If a non-zero value is given for the optional *import* argument, the input file is looked for in the user’s “dot” directory. In this case the *fname* argument should be a plain filename, without any path component.

Should an error occur (such as the file being badly formatted or inaccessible), an error is returned via the [\\$error](#) accessor.

See also [mread](#), [bwrite](#).

bwfilt

Output: series
 Arguments: *y* (series)
 n (integer)
 omega (scalar)

Returns the result from application of a low-pass Butterworth filter with order *n* and frequency cutoff *omega* to the series *y*. The cutoff is expressed in degrees and must be greater than 0 and less than 180. Smaller cutoff values restrict the pass-band to lower frequencies and hence produce a smoother trend. Higher values of *n* produce a sharper cutoff, at the cost of possible numerical instability.

Inspecting the periodogram of the target series is a useful preliminary when you wish to apply this function. See the *Gretl User’s Guide* for details. See also [bkfilt](#), [hpfilt](#).

bwrite

Output: integer
 Arguments: *B* (bundle)
 fname (string)
 export (boolean, optional)

Writes the bundle *B* to an XML file named *fname*. For a summary description of its format, see [bread](#). If file *fname* already exists, it will be overwritten. The return value is 0 on successful completion; if an error occurs, such as the file being unwritable, the return value will be non-zero.

If a non-zero value is given for the *export* argument, the output file will be written into the user's "dot" directory. In this case a plain filename, without any path component, should be given for the second argument.

See also [bread](#), [mwrite](#).

cdemean

Output: matrix
 Argument: *X* (matrix)

Centers the columns of matrix *X* around their means.

cdf

Output: same type as input
 Arguments: *d* (string)
 ... (see below)
 x (scalar, series or matrix)
 Examples: *p1* = cdf(*N*, -2.5)
 p2 = cdf(*X*, 3, 5.67)
 p3 = cdf(*D*, 0.25, -1, 1)

Cumulative distribution function calculator. Returns $P(X \leq x)$, where the distribution of *X* is determined by the string *d*. Between the arguments *d* and *x*, zero or more additional scalar arguments are required to specify the parameters of the distribution, as follows (but note that the normal distribution has its own convenience function, [cnorm](#)).

<i>Distribution</i>	<i>d</i>	<i>Arg 2</i>	<i>Arg 3</i>	<i>Arg 4</i>
Standard normal	z, n or N	-	-	-
Bivariate normal	D	ρ	-	-
Student's <i>t</i> (central)	t	df	-	-
Chi square	c, x or X	df	-	-
Snedecor's <i>F</i>	f or F	df (num.)	df (den.)	-
Gamma	g or G	shape	scale	-
Binomial	b or B	probability	trials	-
Poisson	p or P	mean	-	-
Weibull	w or W	shape	scale	-
Generalized Error	E	shape	-	-
Non-central χ^2	ncX	df	non-centrality	-
Non-central <i>F</i>	ncF	df (num.)	df (den.)	non-centrality
Non-central <i>t</i>	nct	df	non-centrality	-

Note that most cases have aliases to help memorizing the codes. The bivariate normal case is special: the syntax is `x = cdf(D, rho, z1, z2)` where `rho` is the correlation between the variables `z1` and `z2`.

The parametrization `gretl` uses for the Gamma random variate implies that its density function can be written as

$$f(x; k, \theta) = \frac{x^{k-1} e^{-x/\theta}}{\theta^k \Gamma(k)}$$

where $k > 0$ is the shape parameter and $\theta > 0$ is the scale parameter.

See also [pdf](#), [critical](#), [invcdf](#), [pvalue](#).

cdiv

Output: matrix
Arguments: `X` (matrix)
 `Y` (matrix)

Complex division. The two arguments must have the same number of rows, n , and either one or two columns. The first column contains the real part and the second (if present) the imaginary part. The return value is an $n \times 2$ matrix or, if the result has no imaginary part, an n -vector. See also [cmult](#).

ceil

Output: same type as input
Argument: `x` (scalar, series or matrix)

Ceiling function: returns the smallest integer greater than or equal to x . See also [floor](#), [int](#).

cholesky

Output: square matrix
Argument: `A` (positive definite matrix)

Performs a Cholesky decomposition of the matrix A , which is assumed to be symmetric and positive definite. The result is a lower-triangular matrix L which satisfies $A = LL'$. The function will fail if A is not symmetric or not positive definite. See also [psdroot](#).

chowlin

Output: matrix
Arguments: `Y` (matrix)
 `xfac` (integer)
 `X` (matrix, optional)

Expands the input data, Y , to a higher frequency, using the interpolation method of [Chow and Lin \(1971\)](#). It is assumed that the columns of Y represent data series; the returned matrix has as many columns as Y and `xfac` times as many rows.

The second argument represents the expansion factor: it should be 3 for expansion from quarterly to monthly or 4 for expansion from annual to quarterly, these being the only supported factors. The optional third argument may be used to provide a matrix of covariates at the higher (target) frequency.

The regressors used by default are a constant and quadratic trend. If X is provided, its columns are used as additional regressors; it is an error if the number of rows in X does not equal `xfac` times the number of rows in Y .

cmult

Output: matrix
 Arguments: X (matrix)
 Y (matrix)

Complex multiplication. The two arguments must have the same number of rows, n , and either one or two columns. The first column contains the real part and the second (if present) the imaginary part. The return value is an $n \times 2$ matrix, or, if the result has no imaginary part, an n -vector. See also [cdiv](#).

cnorm

Output: same type as input
 Argument: x (scalar, series or matrix)

Returns the cumulative distribution function for a standard normal. See also [dnorm](#), [qnorm](#).

colname

Output: string
 Arguments: M (matrix)
 col (integer)

Retrieves the name for column col of matrix M . If M has no column names attached the value returned is an empty string; if col is out of bounds for the given matrix an error is flagged. See also [colnames](#).

colnames

Output: scalar
 Arguments: M (matrix)
 S (array of strings or list)

Attaches names to the columns of the $T \times k$ matrix M . If S is a named list, the names are taken from the names of the listed series; the list must have k members. If S is an array of strings, it should contain k elements. For backward compatibility, a single string may also be given as the second argument; in that case it should contain k space-separated substrings.

The return value is 0 on successful completion, non-zero on error. See also [rownames](#).

Example:

```
matrix M = {1, 2; 2, 1; 4, 1}
strings S = array(2)
S[1] = "Col1"
S[2] = "Col2"
colnames(M, S)
print M
```

cols

Output: integer
 Argument: X (matrix)

Returns the number of columns of X . See also [mshape](#), [rows](#), [unvech](#), [vec](#), [vech](#).

corr

Output: scalar
 Arguments: *y1* (series or vector)
 y2 (series or vector)

Computes the correlation coefficient between *y1* and *y2*. The arguments should be either two series, or two vectors of the same length. See also [cov](#), [mcov](#), [mcorr](#).

corrgm

Output: matrix
 Arguments: *x* (series, matrix or list)
 p (integer)
 y (series or vector, optional)

If only the first two arguments are given, computes the correlogram for *x* for lags 1 to *p*. Let *k* represent the number of elements in *x* (1 if *x* is a series, the number of columns if *x* is a matrix, or the number of list-members if *x* is a list). The return value is a matrix with *p* rows and *2k* columns, the first *k* columns holding the respective autocorrelations and the remainder the respective partial autocorrelations.

If a third argument is given, this function computes the cross-correlogram for each of the *k* elements in *x* and *y*, from lead *p* to lag *p*. The returned matrix has *2p + 1* rows and *k* columns. If *x* is series or list and *y* is a vector, the vector must have just as many rows as there are observations in the current sample range.

cos

Output: same type as input
 Argument: *x* (scalar, series or matrix)

Returns the cosine of *x*.

cosh

Output: same type as input
 Argument: *x* (scalar, series or matrix)

Returns the hyperbolic cosine of *x*.

$$\cosh x = \frac{e^x + e^{-x}}{2}$$

See also [acosh](#), [sinh](#), [tanh](#).

cov

Output: scalar
 Arguments: *y1* (series or vector)
 y2 (series or vector)

Returns the covariance between *y1* and *y2*. The arguments should be either two series, or two vectors of the same length. See also [corr](#), [mcov](#), [mcorr](#).

critical

Output: same type as input
 Arguments: *c* (character)
 ... (see below)
 p (scalar, series or matrix)
 Examples: `c1 = critical(t, 20, 0.025)`
 `c2 = critical(F, 4, 48, 0.05)`

Critical value calculator. Returns x such that $P(X > x) = p$, where the distribution X is determined by the character c . Between the arguments c and p , zero or more additional scalar arguments are required to specify the parameters of the distribution, as follows.

<i>Distribution</i>	<i>c</i>	<i>Arg 2</i>	<i>Arg 3</i>
Standard normal	z, n or N	-	-
Student's t (central)	t	degrees of freedom	-
Chi square	c, x or X	degrees of freedom	-
Snedecor's F	f or F	df (num.)	df (den.)
Binomial	b or B	p	n
Poisson	p or P	λ	-

See also [cdf](#), [invcdf](#), [pvalue](#).

cum

Output: same type as input
 Argument: x (series or matrix)

Cumulates x . When x is a series, produces a series $y_t = \sum_{s=m}^t x_s$; the starting point of the summation, m , is the first non-missing observation of the currently selected sample. If any missing values are encountered in x , subsequent values of y will be set to missing. When x is a matrix, its elements are cumulated by columns.

See also [diff](#).

curl

Output: scalar
 Argument: $\&b$ (reference to bundle)

Provides a somewhat flexible means of obtaining a text buffer containing data from an internet server, using libcurl. On input the bundle b must contain a string named URL which gives the full address of the resource on the target host. Other optional elements are as follows.

- “header”: a string specifying an HTTP header to be sent to the host.
- “postdata”: a string holding data to be sent to the host.

The `header` and `postdata` fields are intended for use with an HTTP POST request; if `postdata` is present the POST method is implicit, otherwise the GET method is implicit. (But note that for straightforward GET requests [readfile](#) offers a simpler interface.)

One other optional bundle element is recognized: if a scalar named `include` is present and has a non-zero value, this is taken as a request to include the header received from the host with the output body.

On completion of the request, the text received from the server is added to the bundle under the key “output”.

If an error occurs in formulating the request (for example there’s no URL on input) the function fails, otherwise it returns 0 if the request succeeds or non-zero if it fails, in which case the error message from the curl library is added to the bundle under the key “errmsg”. Note, however, that “success” in this sense does not necessarily mean you got the data you wanted; all it means is that some response was received from the server. You must check the content of the output buffer (which may in fact be a message such as “Page not found”).

Here is an example of use: downloading some data from the US Bureau of Labor Statistics site, which requires sending a JSON query. Note the use of `sprintf` to embed double-quotes in the POST data.

```
bundle req
req.URL = "http://api.bls.gov/publicAPI/v1/timeseries/data/"
req.include = 1
req.header = "Content-Type: application/json"
string s
sprintf s, "{\"seriesid\": [\"LEU0254555900\"]}"
req.postdata = s
err = curl(&req)
if err == 0
    s = req.output
    string line
    loop while getline(s, line) --quiet
        printf "%s\\n", line
    endloop
endif
```

See also the `jsonget` function for means of processing JSON data received.

deseas

Output: series
 Arguments: *x* (series)
 c (character, optional)

Depends on having TRAMO/SEATS or X-12-ARIMA installed. Returns a deseasonalized (seasonally adjusted) version of the input series *x*, which must be a quarterly or monthly time series. To use X-12-ARIMA give *X* as the second argument; to use TRAMO give *T*. If the second argument is omitted then X-12-ARIMA is used.

Note that if the input series has no detectable seasonal component this function will fail. Also note that both TRAMO/SEATS and X-12-ARIMA offer numerous options; `deseas` calls them with all options at their default settings. For both programs, the seasonal factors are calculated on the basis of an automatically selected ARIMA model. One difference between the programs which can sometimes make a substantial difference to the results is that by default TRAMO performs a prior adjustment for outliers while X-12-ARIMA does not.

det

Output: scalar
 Argument: *A* (square matrix)

Returns the determinant of *A*, computed via the LU factorization. See also `ldet`, `rcond`.

diag

Output: matrix
 Argument: X (matrix)

Returns the principal diagonal of X in a column vector. Note: if X is an $m \times n$ matrix, the number of elements of the output vector is $\min(m, n)$. See also [tr](#).

diagcat

Output: matrix
 Arguments: A (matrix)
 B (matrix)

Returns the direct sum of A and B , that is a matrix holding A in its north-west corner and B in its south-east corner. If both A and B are square, the resulting matrix is block-diagonal.

diff

Output: same type as input
 Argument: y (series, matrix or list)

Computes first differences. If y is a series, or a list of series, starting values are set to NA. If y is a matrix, differencing is done by columns and starting values are set to 0.

When a list is returned, the individual variables are automatically named according to the template `d_ varname` where *varname* is the name of the original series. The name is truncated if necessary, and may be adjusted in case of non-uniqueness in the set of names thus constructed.

See also [cum](#), [ldiff](#), [sdiff](#).

digamma

Output: same type as input
 Argument: x (scalar, series or matrix)

Returns the digamma (or Psi) function of x , that is $\frac{d\Gamma(x)}{dx}$.

dnorm

Output: same type as input
 Argument: x (scalar, series or matrix)

Returns the density of the standard normal distribution at x . To get the density for a non-standard normal distribution at x , pass the z -score of x to the `dnorm` function and multiply the result by the Jacobian of the z transformation, namely $1/\sigma$, as illustrated below:

```
mu = 100
sigma = 5
x = 109
fx = (1/sigma) * dnorm((x-mu)/sigma)
```

See also [cnorm](#), [qnorm](#).

dsort

Output: same type as input
 Argument: x (series or vector)

Sorts x in descending order, skipping observations with missing values when x is a series. See also [sort](#), [values](#).

dummify

Output: list
 Arguments: x (series)
 $omitval$ (scalar, optional)

The argument x should be a discrete series. This function creates a set of dummy variables coding for the distinct values in the series. By default the smallest value is taken as the omitted category and is not explicitly represented.

The optional second argument represents the value of x which should be treated as the omitted category. The effect when a single argument is given is equivalent to `dummify(x, min(x))`. To produce a full set of dummies, with no omitted category, use `dummify(x, NA)`.

The generated variables are automatically named according to the template `Dvarname_i` where *varname* is the name of the original series and i is a 1-based index. The original portion of the name is truncated if necessary, and may be adjusted in case of non-uniqueness in the set of names thus constructed.

easterday

Output: same type as input
 Argument: x (scalar, series or matrix)

Given the year in argument x , returns the date of Easter in the Gregorian calendar as *month + day/100*. Note that April the 10th, is, under this convention, 4.1; hence, 4.2 is April the 20th, not April the 2nd (which would be 4.02).

```
scalar e = easterday(2014)
scalar m = floor(e)
scalar d = 100*(e-m)
```

eigengen

Output: matrix
 Arguments: A (square matrix)
 $\&U$ (reference to matrix, or null)

Computes the eigenvalues, and optionally the right eigenvectors, of the $n \times n$ matrix A . If all the eigenvalues are real an $n \times 1$ matrix is returned; otherwise the result is an $n \times 2$ matrix, the first column holding the real components and the second column the imaginary components. The eigenvalues are not guaranteed to be sorted in any particular order.

The second argument must be either the name of an existing matrix preceded by `&` (to indicate the “address” of the matrix in question), in which case an auxiliary result is written to that matrix, or the keyword `null`, in which case the auxiliary result is not produced.

If a non-null second argument is given, the specified matrix will be over-written with the auxiliary result. (It is not required that the existing matrix be of the right dimensions to receive the result.) It will be organized as follows:

- If the i -th eigenvalue is real, the i -th column of U will contain the corresponding eigenvector;
- If the i -th eigenvalue is complex, the i -th column of U will contain the real part of the corresponding eigenvector and the next column the imaginary part. The eigenvector for the conjugate eigenvalue is the conjugate of the eigenvector.

In other words, the eigenvectors are stored in the same order as the eigenvalues, but the real eigenvectors occupy one column, whereas complex eigenvectors take two (the real part comes first); the total number of columns is still n , because the conjugate eigenvector is skipped.

See also [eigensym](#), [eigsolve](#), [qrdecomp](#), [svd](#).

eigensym

Output: matrix
 Arguments: A (symmetric matrix)
 $\&U$ (reference to matrix, or null)

Works just as [eigen](#), but the argument A must be symmetric (in which case the calculations can be reduced). Unlike [eigen](#), eigenvalues are returned in ascending order.

Note: if you're interested in the eigen-decomposition of a matrix of the form $X'X$, where X is a large matrix, it is preferable to compute it via the prime operator $X'X$ rather than using the more general syntax $X'*X$. The former expression uses a specialized algorithm which has the double advantage of being more efficient computationally and of ensuring that the result will be free by construction of machine precision artifacts that may render it numerically non-symmetric.

eigsolve

Output: matrix
 Arguments: A (symmetric matrix)
 B (symmetric matrix)
 $\&U$ (reference to matrix, or null)

Solves the generalized eigenvalue problem $|A - \lambda B| = 0$, where both A and B are symmetric and B is positive definite. The eigenvalues are returned directly, arranged in ascending order. If the optional third argument is given it should be the name of an existing matrix preceded by $\&$; in that case the generalized eigenvectors are written to the named matrix.

epochday

Output: scalar or series
 Arguments: $year$ (scalar or series)
 $month$ (scalar or series)
 day (scalar or series)

Returns the number of the day in the current epoch specified by year, month and day. The epoch day equals 1 for the first of January in the year 1 AD; it stood at 733786 on 2010-01-01. If any of the arguments are given as series the value returned is a series, otherwise it is a scalar.

For the inverse function, see [isodate](#).

errmsg

Output: string
 Argument: $errno$ (integer)

Retrieves the gretl error message associated with $errno$. See also [\\$error](#).

exp

Output: same type as input
 Argument: x (scalar, series or matrix)

Returns e^x . Note that in case of matrices the function acts element by element. For the matrix exponential function, see [mexp](#).

fcstats

Output: matrix
 Arguments: y (series or vector)
 f (series or vector)

Produces a column vector holding several statistics which may be used for evaluating the series f as a forecast of the series y over the current sample range. Two vectors of the same length may be given in place of two series arguments.

The layout of the returned vector is as follows:

- 1 Mean Error (ME)
- 2 Mean Squared Error (MSE)
- 3 Mean Absolute Error (MAE)
- 4 Mean Percentage Error (MPE)
- 5 Mean Absolute Percentage Error (MAPE)
- 6 Theil's U
- 7 Bias proportion, UM
- 8 Regression proportion, UR
- 9 Disturbance proportion, UD

For details on the calculation of these statistics, and the interpretation of the U values, please see the *Gretl User's Guide*.

fdjac

Output: matrix
 Arguments: b (column vector)
 $fcall$ (function call)

Calculates a numerical approximation to the Jacobian associated with the n -vector b and the transformation function specified by the argument $fcall$. The function call should take b as its first argument (either straight or in pointer form), followed by any additional arguments that may be needed, and it should return an $m \times 1$ matrix. On successful completion `fdjac` returns an $m \times n$ matrix holding the Jacobian. Example:

```
matrix J = fdjac(theta, myfunc(&theta, X))
```

The function can use three different methods: simple forward-difference, bilateral difference or 4-nodes Richardson extrapolation. Respectively:

$$J_0 = \frac{f(x+h) - f(x)}{h}$$

$$J_1 = \frac{f(x+h) - f(x-h)}{2h}$$

$$J_2 = \frac{8(f(x+h) - f(x-h)) - (f(x+2h) - f(x-2h))}{12h}$$

The three alternatives above provide, generally, a trade-off between accuracy and speed. You can choose among methods by using the `set` command and specify the value 0, 1 or 2 for the `fdjac_quality` variable.

For more details and examples see the chapter on numerical methods in the *Gretl User's Guide*.

See also [BFGSmax](#), [set](#).

fft

Output: matrix

Argument: X (matrix)

Discrete real Fourier transform. If the input matrix X has n columns, the output has $2n$ columns, where the real parts are stored in the odd columns and the complex parts in the even ones.

Should it be necessary to compute the Fourier transform on several vectors with the same number of elements, it is numerically more efficient to group them into a matrix rather than invoking `fft` for each vector separately. See also [ffti](#).

ffti

Output: matrix

Argument: X (matrix)

Inverse discrete real Fourier transform. It is assumed that X contains n complex column vectors, with the real part in the odd columns and the imaginary part in the even ones, so the total number of columns should be $2n$. A matrix with n columns is returned.

Should it be necessary to compute the inverse Fourier transform on several vectors with the same number of elements, it is numerically more efficient to group them into a matrix rather than invoking `ffti` for each vector separately. See also [fft](#).

filter

Output: series

Arguments: x (series or matrix)

a (scalar or vector, optional)

b (scalar or vector, optional)

$y0$ (scalar, optional)

Computes an ARMA-like filtering of the argument x . The transformation can be written as

$$y_t = \sum_{i=0}^q a_i x_{t-i} + \sum_{i=1}^p b_i y_{t-i}$$

If argument x is a series, the result will be itself a series. Otherwise, if x is a matrix with T rows and k columns, the result will be a matrix of the same size, in which the filtering is performed column by column.

The two arguments a and b are optional. They may be scalars, vectors or the keyword `null`.

If a is a scalar, this is used as a_0 and implies $q = 0$; if it is a vector of $q + 1$ elements, they contain the coefficients from a_0 to a_q . If a is `null` or omitted, this is equivalent to setting $a_0 = 1$ and $q = 0$.

If b is a scalar, this is used as b_1 and implies $p = 1$; if it is a vector of p elements, they contain the coefficients from b_1 to b_p . If b is null or omitted, this is equivalent to setting $B(L) = 1$.

The optional scalar argument $y0$ is taken to represent all values of y prior to the beginning of sample (used only when $p > 0$). If omitted, it is understood to be 0. Pre-sample values of x are always assumed zero.

See also [bkfilt](#), [bwfilt](#), [fracdiff](#), [hpfilt](#), [movavg](#), [varsimul](#).

Example:

```

nulldata 5
y = filter(index, 0.5, -0.9, 1)
print index y --byobs
x = seq(1,5)' ~ (1 | zeros(4,1))
w = filter(x, 0.5, -0.9, 1)
print x w

```

produces

index		y
1	1	-0.40000
2	2	1.36000
3	3	0.27600
4	4	1.75160
5	5	0.92356

x (5 x 2)

1	1
2	0
3	0
4	0
5	0

w (5 x 2)

-0.40000	-0.40000
1.3600	0.36000
0.27600	-0.32400
1.7516	0.29160
0.92356	-0.26244

firstobs

Output: integer

Argument: y (series)

Returns the 1-based index of the first non-missing observation for the series y . Note that if some form of subsampling is in effect, the value returned may be smaller than the dollar variable [\\$t1](#). See also [lastobs](#).

fixname

Output: string

Argument: *rawname* (string)

Intended for use in connection with the [join](#) command. Returns the result of converting *rawname* to a valid gretl identifier, which must start with a letter, contain nothing but (ASCII) letters, digits and the underscore character, and must not exceed 31 characters. The rules used in conversion are:

1. Skip any leading non-letters.
2. Until the 31-character limit is reached or the input is exhausted: transcribe “legal” characters; skip “illegal” characters apart from spaces; and replace one or more consecutive spaces with an underscore, unless the previous character transcribed is an underscore in which case space is skipped.

floor

Output: same type as input
 Argument: *y* (scalar, series or matrix)

Returns the greatest integer less than or equal to *x*. Note: [int](#) and `floor` differ in their effect for negative arguments: `int(-3.5)` gives -3 , while `floor(-3.5)` gives -4 .

fracdiff

Output: series
 Arguments: *y* (series)
 d (scalar)

$$\Delta^d y_t = y_t - \sum_{i=1}^{\infty} \psi_i y_{t-i}$$

where

$$\psi_i = \frac{\Gamma(i-d)}{\Gamma(-d)\Gamma(i+1)}$$

Note that in theory fractional differentiation is an infinitely long filter. In practice, presample values of y_t are assumed to be zero.

gammafun

Output: same type as input
 Argument: *x* (scalar, series or matrix)

Returns the gamma function of *x*.

genseries

Output: scalar
 Arguments: *varname* (string)
 rhs (series)

Provides the script writer with a convenient means of generating series whose names are not known in advance, and/or creating a series and appending it to a list in a single operation.

The first argument gives the name of the series to create (or modify); this can be a string literal, a string variable, or an expression that evaluates to a string. The second argument, *rhs* (“right-hand side”), defines the source series: this can be the name of an existing series or an expression that evaluates to a series, as would appear to the right of the equals sign when defining a series in the usual way.

The return value from this function is the ID number of the series in the dataset, a value suitable for inclusion in a list (or -1 on failure).

For example, suppose you want to add n random normal series to the dataset and put them all into a named list. The following will do the job:

```
list Normals = null
loop i=1..n --quiet
    Normals += genseries(sprintf("norm%d", i), normal())
endloop
```

On completion `Normals` will contain the series `norm1`, `norm2` and so on.

getenv

Output: string
Argument: *s* (string)

If an environment variable by the name of *s* is defined, returns the string value of that variable, otherwise returns an empty string. See also [ngetenv](#).

getline

Output: scalar
Arguments: *source* (string)
target (string)

This function is used to read successive lines from *source*, which should be a named string variable. On each call a line from the source is written to *target* (which must also be a named string variable), with the newline character stripped off. The value returned is 1 if there was anything to be read (including blank lines), 0 if the source has been exhausted.

Here is an example in which the content of a text file is broken into lines:

```
string s = readfile("data.txt")
string line
scalar i = 1
loop while getline(s, line)
    printf "line %d = '%s'\n", i++, line
endloop
```

In this example we can be sure that the source is exhausted when the loop terminates. If the source might not be exhausted you should follow your regular call(s) to `getline` with a “clean up” call, in which *target* is replaced by `null` (or omitted altogether) as in

```
getline(s, line) # get a single line
getline(s, null) # clean up
```

Note that although the reading position advances at each call to `getline`, *source* is not modified by this function, only *target*.

ghk

Output: matrix
 Arguments: C (matrix)
 A (matrix)
 B (matrix)
 U (matrix)
 &mathit{dP} (reference to matrix, or null)

Computes the GHK (Geweke, Hajivassiliou, Keane) approximation to the multivariate normal distribution function; see for example [Geweke \(1991\)](#). The value returned is an $n \times 1$ vector of probabilities.

The argument C ($m \times m$) should give the Cholesky factor (lower triangular) of the covariance matrix of m normal variates. The arguments A and B should both be $n \times m$, giving respectively the lower and upper bounds applying to the variates at each of n observations. Where variates are unbounded, this should be indicated using the built-in constant [Shuge](#) or its negative.

The matrix U should be $m \times r$, with r the number of pseudo-random draws from the uniform distribution; suitable functions for creating U are [muniform](#) and [halton](#).

We illustrate below with a relatively simple case where the multivariate probabilities can be calculated analytically. The series P and Q should be numerically very similar to one another, P being the “true” probability and Q its GHK approximation:

```

nulldata 20
series inf1 = -2*uniform()
series sup1 = 2*uniform()
series inf2 = -2*uniform()
series sup2 = 2*uniform()

scalar rho = 0.25
matrix V = {1, rho; rho, 1}

series P = cdf(D, rho, inf1, inf2) - cdf(D, rho, sup1, inf2) \
- cdf(D, rho, inf1, sup2) + cdf(D, rho, sup1, sup2)

C = cholesky(V)
U = halton(2, 100)

series Q = ghk(C, {inf1, inf2}, {sup1, sup2}, U)

```

The optional dP argument can be used to retrieve the $n \times k$ matrix of derivatives of the probabilities, where k equals $2m + m(m + 1)/2$. The first m columns hold the derivatives with respect to the lower bounds, the next m those with respect to the upper bounds, and the remainder the derivatives with respect to the unique elements of the C matrix in “vech” order.

gini

Output: scalar
 Argument: y (series)

Returns Gini’s inequality index for the series y .

ginv

Output: matrix
 Argument: A (matrix)

Returns A^+ , the Moore–Penrose or generalized inverse of A , computed via the singular value decomposition.

This matrix has the properties

$$\begin{aligned} AA^+A &= A \\ A^+AA^+ &= A^+ \end{aligned}$$

Moreover, the products A^+A and AA^+ are symmetric by construction.

See also [inv](#), [svd](#).

halton

Output: matrix
 Arguments: m (integer)
 r (integer)
 $offset$ (integer, optional)

Returns an $m \times r$ matrix containing m Halton sequences of length r ; m is limited to a maximum of 40. The sequences are constructed using the first m primes. By default the first 10 elements of each sequence are discarded, but this figure can be adjusted via the optional *offset* argument, which should be a non-negative integer. See [Halton and Smith \(1964\)](#).

hdprod

Output: matrix
 Arguments: X (matrix)
 Y (matrix)

Horizontal direct product. The two arguments must have the same number of rows, r . The return value is a matrix with r rows, in which the i -th row is the Kronecker product of the corresponding rows of X and Y .

In other words, if X is an $r \times k$ matrix, Y is an $r \times m$ matrix and Z is the result matrix of the horizontal direct product of X times Y , then Z will have r rows and $k \cdot m$ columns; moreover,

$$Z_{in} = X_{ij}Y_{il}$$

where $n = (j - 1)m + l$.

As far as we know there isn't an established name for this operation in matrix algebra; it is called "horizontal direct product" in the GAUSS programming language.

Example: the code

```
A = {1,2,3; 4,5,6}
B = {0,1; -1,1}
C = hdprod(A, B)
```

produces the following matrix:

```
0   1   0   2   0   3
-4  4  -5   5  -6   6
```

hpfilt

Output: series
 Arguments: y (series)
 λ (scalar, optional)

Returns the cycle component from application of the Hodrick–Prescott filter to series y . If the smoothing parameter, λ , is not supplied then a data-based default is used, namely 100 times the square of the periodicity (100 for annual data, 1600 for quarterly data, and so on). See also [bkfilt](#), [bwfilt](#).

I

Output: square matrix
 Argument: n (integer)

Returns an identity matrix with n rows and columns.

imaxc

Output: row vector
 Argument: X (matrix)

Returns the row indices of the maxima of the columns of X .

See also [imaxr](#), [iminc](#), [maxc](#).

imaxr

Output: column vector
 Argument: X (matrix)

Returns the column indices of the maxima of the rows of X .

See also [imaxc](#), [iminr](#), [maxr](#).

imhof

Output: scalar
 Arguments: M (matrix)
 x (scalar)

Computes $\text{Prob}(u' Au < x)$ for a quadratic form in standard normal variates, u , using the procedure developed by [Imhof \(1961\)](#).

If the first argument, M , is a square matrix it is taken to specify A , otherwise if it's a column vector it is taken to be the precomputed eigenvalues of A , otherwise an error is flagged.

See also [pvalue](#).

iminc

Output: row vector
 Argument: X (matrix)

Returns the row indices of the minima of the columns of X .

See also [iminr](#), [imaxc](#), [minc](#).

iminr

Output: column vector

Argument: X (matrix)

Returns the column indices of the minima of the rows of X .

See also [iminc](#), [imaxr](#), [minr](#).

inbundle

Output: integer

Arguments: b (bundle)

key (string)

Checks whether bundle b contains a data-item with name key . The value returned is an integer code for the type of the item: 0 for no match, 1 for scalar, 2 for series, 3 for matrix, 4 for string, 5 for bundle and 6 for array. The function [typestr](#) may be used to get the string corresponding to this code.

infnorm

Output: scalar

Argument: X (matrix)

Returns the ∞ -norm of the $r \times c$ matrix X , namely,

$$\|X\|_{\infty} = \max_i \sum_{j=1}^c |X_{ij}|$$

See also [onenorm](#).

inlist

Output: integer

Arguments: L (list)

y (series)

Returns the (1-based) position of y in list L , or 0 if y is not present in L . The second argument may be given as the name of a series or alternatively as an integer ID number.

int

Output: same type as input

Argument: x (scalar, series or matrix)

Returns the integer part of x , truncating the fractional part. Note: `int` and [floor](#) differ in their effect for negative arguments: `int(-3.5)` gives -3 , while `floor(-3.5)` gives -4 . See also [ceil](#).

inv

Output: matrix

Argument: A (square matrix)

Returns the inverse of A . If A is singular or not square, an error message is produced and nothing is returned. Note that `gretl` checks automatically the structure of A and uses the most efficient numerical procedure to perform the inversion.

The matrix types `gretl` checks for are: identity; diagonal; symmetric and positive definite; symmetric but not positive definite; and triangular.

Note: it makes sense to use this function only if you plan to use the inverse of A more than once. If you just need to compute an expression of the form $A^{-1}B$, you'll be much better off using the "division" operators `\` and `/`. See the *Gretl User's Guide* for details.

See also [ginv](#), [invpd](#).

invcdf

Output: same type as input
 Arguments: d (string)
 ... (see below)
 p (scalar, series or matrix)

Inverse cumulative distribution function calculator. Returns x such that $P(X \leq x) = p$, where the distribution of X is determined by the string d . Between the arguments d and p , zero or more additional scalar arguments are required to specify the parameters of the distribution, as follows.

<i>Distribution</i>	<i>d</i>	<i>Arg 2</i>	<i>Arg 3</i>	<i>Arg 4</i>
Standard normal	z, n or N	–	–	–
Gamma	g or G	shape	scale	–
Student's t (central)	t	degrees of freedom	–	–
Chi square	c, x or X	degrees of freedom	–	–
Snedecor's F	f or F	df (num.)	df (den.)	–
Binomial	b or B	p	n	–
Poisson	p or P	λ	–	–
Standardized GED	E	shape	–	–
Non-central χ^2	ncX	df	non-centrality	–
Non-central F	ncF	df (num.)	df (den.)	non-centrality
Non-central t	nct	df	non-centrality	–

See also [cdf](#), [critical](#), [pvalue](#).

invmls

Output: same type as input
 Argument: x (scalar, series or matrix)

Returns the inverse Mills ratio at x , that is the ratio between the standard normal density and the complement to the standard normal distribution function, both evaluated at x .

This function uses a dedicated algorithm which yields greater accuracy compared to calculation using [dnorm](#) and [cnorm](#), but the difference between the two methods is appreciable only for very large negative values of x .

See also [cdf](#), [cnorm](#), [dnorm](#).

invpd

Output: square matrix
 Argument: A (positive definite matrix)

Returns the inverse of the symmetric, positive definite matrix A . This function is slightly faster than [inv](#) for large matrices, since no check for symmetry is performed; for that reason it should be used with care.

Note: if you're interested in the inversion of a matrix of the form $X'X$, where X is a large matrix, it is preferable to compute it via the prime operator $X'X$ rather than using the more general syntax $X'*X$. The former expression uses a specialized algorithm which has the double advantage of being more efficient computationally and of ensuring that the result will be free by construction of machine precision artifacts that may render it numerically non-symmetric.

irf

Output: matrix
Arguments: *target* (integer)
 shock (integer)
 alpha (scalar between 0 and 1, optional)

This function is available only when the last model estimated was a VAR or VECM. It returns a matrix containing the estimated response of the *target* variable to an impulse of one standard deviation in the *shock* variable. These variables are identified by their position in the VAR specification: for example, if *target* and *shock* are given as 1 and 3 respectively, the returned matrix gives the response of the first variable in the VAR for a shock to the third variable.

If the optional *alpha* argument is given, the returned matrix has three columns: the point estimate of the responses, followed by the lower and upper limits of a $1 - \alpha$ confidence interval obtained via bootstrapping. (So *alpha* = 0.1 corresponds to 90 percent confidence.) If *alpha* is omitted or set to zero, only the point estimate is provided.

The number of periods (rows) over which the response is traced is determined automatically based on the frequency of the data, but this can be overridden via the [set](#) command, as in `set horizon 10`.

irr

Output: scalar
Argument: *x* (series or vector)

Returns the Internal Rate of Return for *x*, considered as a sequence of payments (negative) and receipts (positive). See also [npv](#).

isconst

Output: integer
Arguments: *y* (series or vector)
 panel-code (integer, optional)

Without the optional second argument, returns 1 if *y* has a constant value over the current sample range (or over its entire length if *y* is a vector), otherwise 0.

The second argument is accepted only if the current dataset is a panel and *y* is a series. In that case a *panel-code* value of 0 calls for a check for time-invariance, while a value of 1 means check for cross-sectional invariance (that is, in each time period the value of *y* is the same for all groups).

If *y* is a series, missing values are ignored in checking for constancy.

isdiscrete

Output: integer
Argument: *name* (string)

If *name* is the identifier for a currently defined series, returns 1 if the series is marked as discrete-valued, otherwise 0. If *name* does not identify a series, returns NA.

isdummy

Output: integer
Argument: *x* (series or vector)

If all the values contained in *x* are 0 or 1 (or missing), returns the number of ones, otherwise 0.

isnan

Output: same type as input
Argument: *x* (scalar or matrix)

Given a scalar argument, returns 1 if *x* is “Not a Number” (NaN), otherwise 0. Given a matrix argument, returns a matrix of the same dimensions with 1s in positions where the corresponding element of the input is NaN and 0s elsewhere.

isnull

Output: integer
Argument: *name* (string)

Returns 0 if *name* is the identifier for a currently defined object, be it a scalar, a series, a matrix, list, string or bundle; otherwise returns 1.

isoconv

Output: scalar
Arguments: *date* (series)
 &*year* (reference to series)
 &*month* (reference to series)
 &*day* (reference to series, optional)

Given a series *date* holding dates in ISO 8601 “basic” format (YYYYMMDD), this function writes the year, month and (optionally) day components into the series named by the second and subsequent arguments. An example call, assuming the series *dates* contains suitable 8-digit values:

```
series y, m, d
isoconv(dates, &y, &m, &d)
```

The return value from this function is 0 on successful completion, non-zero on error.

isodate

Output: see below
Arguments: *ed* (scalar or series)
 as-string (boolean, optional)

The argument *ed* is interpreted as an epoch day (which equals 1 for the first of January in the year 1 AD). The default return value — of the same type as *ed* — is an 8-digit number, or a series of such numbers, on the pattern YYYYMMDD (ISO 8601 “basic” format), giving the calendar date corresponding to the epoch day.

If *ed* is a scalar (only) and the optional second argument *as-string* is non-zero, the return value is not numeric but rather a string on the pattern YYYY-MM-DD (ISO 8601 “extended” format).

For the inverse function, see [epochday](#).

iwishart

Output: matrix
Arguments: *S* (symmetric matrix)
 v (integer)

Given *S* (a positive definite $p \times p$ scale matrix), returns a drawing from the Inverse Wishart distribution with *v* degrees of freedom. The returned matrix is also $p \times p$. The algorithm of [Odell and Feiveson \(1966\)](#) is used.

jsonget

Output: string
Arguments: *buf* (string)
 path (string)

The argument *buf* should be a JSON buffer, as may be retrieved from a suitable website via the [curl](#) function, and the *path* argument should be a JsonPath specification.

This function returns a string representing the data found in the buffer at the specified path. Data types of double (floating-point), int (integer) and string are supported. In the case of doubles or ints, their string representation is returned (using the “C” locale for doubles). If the object to which *path* refers is an array, the members are printed one per line in the returned string.

For an account of JsonPath syntax, see for example <http://goessner.net/articles/JsonPath/>.

kdensity

Output: matrix
Arguments: *x* (series or vector)
 scale (scalar, optional)
 control (boolean, optional)

Computes a kernel density estimate for the series or vector *x*. The returned matrix has two columns, the first holding a set of evenly spaced abscissae and the second the estimated density at each of these points.

The optional *scale* parameter can be used to adjust the degree of smoothing relative to the default of 1.0 (higher values produce a smoother result). The *control* parameter acts as a boolean: 0 (the default) means that the Gaussian kernel is used; a non-zero value switches to the Epanechnikov kernel.

A plot of the results may be obtained using the [gnuplot](#) command, as in

```
matrix d = kdensity(x)
gnuplot 2 1 --matrix=d --with-lines --suppress-fitted
```

kfilter

Output: scalar
 Arguments: $\&E$ (reference to matrix, or null)
 $\&V$ (reference to matrix, or null)
 $\&S$ (reference to matrix, or null)
 $\&P$ (reference to matrix, or null)
 $\&G$ (reference to matrix, or null)

Requires that a Kalman filter be set up. Performs a forward, filtering pass and returns 0 on successful completion or 1 if numerical problems are encountered.

The optional matrix arguments can be used to retrieve the following information: E gets the matrix of one-step ahead prediction errors and V gets the variance matrix for these errors; S gets the matrix of estimated values of the state vector and P the variance matrix of these estimates; G gets the Kalman gain. All of these matrices have T rows, corresponding to T observations. For the column dimensions and further details see the *Gretl User's Guide*.

See also [kalman](#), [ksmooth](#), [ksimul](#).

kpsscrit

Output: matrix
 Arguments: T (scalar)
 $trend$ (boolean)

Returns a row vector containing critical values at the 10, 5 and 1 percent levels for the KPSS test for stationarity of a time series. T should give the number of observations and $trend$ should be 1 if the test includes a trend, 0 otherwise.

The critical values given are based on response surfaces estimated in the manner set out by [Sephton \(1995\)](#). See also the [kpss](#) command.

ksimul

Output: matrix
 Arguments: v (matrix)
 w (matrix)
 $\&S$ (reference to matrix, or null)

Requires that a Kalman filter be set up. Performs a simulation and returns a matrix holding simulated values of the observable variables.

The argument v supplies artificial disturbances for the state transition equation and w supplies disturbances for the observation equation, if applicable. The optional argument S may be used to retrieve the simulated state vector. For details see the *Gretl User's Guide*.

See also [kalman](#), [kfilter](#), [ksmooth](#).

ksmooth

Output: matrix
 Argument: $\&P$ (reference to matrix, or null)

Requires that a Kalman filter be set up. Performs a backward, smoothing pass and returns a matrix holding smoothed estimates of the state vector. The optional argument P may be used to retrieve the MSE of the smoothed state. For details see the *Gretl User's Guide*.

See also [kalman](#), [kfilter](#), [ksimul](#).

kurtosis

Output: scalar
Argument: *x* (series)

Returns the excess kurtosis of the series *x*, skipping any missing observations.

lags

Output: list
Arguments: *p* (integer)
y (series or list)
bylag (boolean, optional)

Generates lags 1 to *p* of the series *y*, or if *y* is a list, of all series in the list. If *p* = 0, the maximum lag defaults to the periodicity of the data; otherwise *p* must be positive.

The generated variables are automatically named according to the template *varname _ i* where *varname* is the name of the original series and *i* is the specific lag. The original portion of the name is truncated if necessary, and may be adjusted in case of non-uniqueness in the set of names thus constructed.

When *y* is a list and the lag order is greater than 1, the default ordering of the terms in the returned list is by variable: all lags of the first series in the input list followed by all lags of the second series, and so on. The optional third argument can be used to change this: if *bylag* is non-zero then the terms are ordered by lag: lag 1 of all the input series, then lag 2 of all the series, and so on.

lastobs

Output: integer
Argument: *y* (series)

Returns the 1-based index of the last non-missing observation for the series *y*. Note that if some form of subsampling is in effect, the value returned may be larger than the dollar variable \$t2. See also [firstobs](#).

ldet

Output: scalar
Argument: *A* (square matrix)

Returns the natural log of the determinant of *A*, computed via the LU factorization. See also [det](#), [rcond](#).

ldiff

Output: same type as input
Argument: *y* (series or list)

Computes log differences; starting values are set to NA.

When a list is returned, the individual variables are automatically named according to the template *ld_varname* where *varname* is the name of the original series. The name is truncated if necessary, and may be adjusted in case of non-uniqueness in the set of names thus constructed.

See also [diff](#), [sdiff](#).

lincomb

Output: series
Arguments: L (list)
 b (vector)

Computes a new series as a linear combination of the series in the list L . The coefficients are given by the vector b , which must have length equal to the number of series in L .

See also [wmean](#).

linearize

Output: series
Argument: x (series)

Depends on having TRAMO installed. Returns a “linearized” version of the input series; that is, a series in which any missing values are replaced by interpolated values and outliers are adjusted. TRAMO’s fully automatic mechanism is used; consult the TRAMO documentation for details.

Note that if the input series has no missing values and no values that TRAMO regards as outliers, this function will return a copy of the original series.

ljungbox

Output: scalar
Arguments: y (series)
 p (integer)

Computes the Ljung–Box Q' statistic for the series y using lag order p , over the currently defined sample range. The lag order must be greater than or equal to 1 and less than the number of available observations.

This statistic may be referred to the chi-square distribution with p degrees of freedom as a test of the null hypothesis that the series y is not serially correlated. See also [pvalue](#).

lngamma

Output: same type as input
Argument: x (scalar, series or matrix)

Returns the log of the gamma function of x .

log

Output: same type as input
Argument: x (scalar, series, matrix or list)

Returns the natural logarithm of x ; produces NA for non-positive values. Note: `ln` is an acceptable alias for `log`.

When a list is returned, the individual variables are automatically named according to the template `l_varname` where *varname* is the name of the original series. The name is truncated if necessary, and may be adjusted in case of non-uniqueness in the set of names thus constructed.

log10

Output: same type as input
Argument: x (scalar, series or matrix)

Returns the base-10 logarithm of x ; produces NA for non-positive values.

log2

Output: same type as input
 Argument: x (scalar, series or matrix)

Returns the base-2 logarithm of x ; produces NA for non-positive values.

loess

Output: series
 Arguments: y (series)
 x (series)
 d (integer, optional)
 q (scalar, optional)
 $robust$ (boolean, optional)

Performs locally-weighted polynomial regression and returns a series holding predicted values of y for each non-missing value of x . The method is as described by [Cleveland \(1979\)](#).

The optional arguments d and q specify the order of the polynomial in x and the proportion of the data points to be used in local estimation, respectively. The default values are $d = 1$ and $q = 0.5$. The other acceptable values for d are 0 and 2. Setting $d = 0$ reduces the local regression to a form of moving average. The value of q must be greater than 0 and cannot exceed 1; larger values produce a smoother outcome.

If a non-zero value is given for the *robust* argument the local regressions are iterated twice, with the weights being modified based on the residuals from the previous iteration so as to give less influence to outliers.

See also [nadarwat](#), and in addition see the *Gretl User's Guide* for details on nonparametric methods.

logistic

Output: same type as input
 Argument: x (scalar, series or matrix)

Returns the logistic function of the argument x , that is, $\Lambda(x) = e^x / (1 + e^x)$. If x is a matrix, the function is applied element by element.

lower

Output: square matrix
 Argument: A (matrix)

Returns an $n \times n$ lower triangular matrix B for which $B_{ij} = A_{ij}$ if $i \geq j$, and 0 otherwise.

See also [upper](#).

lrvar

Output: scalar
 Arguments: y (series or vector)
 k (integer)

Returns the long-run variance of y , calculated using a Bartlett kernel with window size k . The default window size, namely the integer part of the cube root of the sample size, can be selected by giving a negative value for k .

In formulae:

$$\hat{\omega}^2(k) = \frac{1}{T} \sum_{t=k}^{T-k} \left[\sum_{i=-k}^k w_i (y_t - \bar{X})(y_{t-i} - \bar{Y}) \right]$$

with

$$w_i = 1 - \frac{|i|}{k+1}$$

max

Output: scalar or series

Argument: y (series or list)

If the argument y is a series, returns the (scalar) maximum of the non-missing observations in the series. If the argument is a list, returns a series each of whose elements is the maximum of the values of the listed variables at the given observation.

See also [min](#), [xmax](#), [xmin](#).

maxc

Output: row vector

Argument: X (matrix)

Returns a row vector containing the maxima of the columns of X .

See also [imaxc](#), [maxr](#), [minc](#).

maxr

Output: column vector

Argument: X (matrix)

Returns a column vector containing the maxima of the rows of X .

See also [imaxc](#), [maxc](#), [minr](#).

mcorr

Output: matrix

Argument: X (matrix)

Computes a correlation matrix treating each column of X as a variable. See also [corr](#), [cov](#), [mcov](#).

mcov

Output: matrix

Argument: X (matrix)

Computes a covariance matrix treating each column of X as a variable. See also [corr](#), [cov](#), [mcorr](#).

mcovg

Output: matrix
 Arguments: X (matrix)
 u (vector, optional)
 w (vector, optional)
 p (integer)

Returns the matrix covariogram for a $T \times k$ matrix X (typically containing regressors), an (optional) T -vector u (typically containing residuals), an (optional) $(p+1)$ -vector of weights w , and a lag order p , which must be greater than or equal to 0.

The returned matrix is given by

$$\sum_{j=-p}^p \sum_j w_{|j|} (X_t' u_t u_{t-j} X_{t-j})$$

If u is given as null the u terms are omitted, and if w is given as null all the weights are taken to be 1.0.

mean

Output: scalar or series
 Argument: x (series or list)

If x is a series, returns the (scalar) sample mean, skipping any missing observations.

If x is a list, returns a series y such that y_t is the mean of the values of the variables in the list at observation t , or NA if there are any missing values at t .

meanc

Output: row vector
 Argument: X (matrix)

Returns the means of the columns of X . See also [meanr](#), [sumc](#), [sdc](#).

meanr

Output: column vector
 Argument: X (matrix)

Returns the means of the rows of X . See also [meanc](#), [sumr](#).

median

Output: scalar
 Argument: y (series)

The median of the non-missing observations in series y . See also [quantile](#).

mexp

Output: square matrix
 Argument: A (square matrix)

Computes the matrix exponential,

$$e^A = \sum_{k=0}^{\infty} \frac{A^k}{k!} = \frac{I}{0!} + \frac{A}{1!} + \frac{A^2}{2!} + \frac{A^3}{3!} + \dots$$

(This series is sure to converge.) The algorithm used is 11.3.1 from [Golub and Van Loan \(1996\)](#).

min

Output: scalar or series

Argument: y (series or list)

If the argument y is a series, returns the (scalar) minimum of the non-missing observations in the series. If the argument is a list, returns a series each of whose elements is the minimum of the values of the listed variables at the given observation.

See also [max](#), [xmax](#), [xmin](#).

minc

Output: row vector

Argument: X (matrix)

Returns the minima of the columns of X .

See also [iminc](#), [maxc](#), [minr](#).

minr

Output: column vector

Argument: X (matrix)

Returns the minima of the rows of X .

See also [iminr](#), [maxr](#), [minc](#).

missing

Output: same type as input

Argument: x (scalar, series or list)

Returns a binary variable holding 1 if x is NA. If x is a series, the comparison is done element by element; if x is a list of series, the output is a series with 1 at observations for which at least one series in the list has a missing value, and 0 otherwise.

See also [misszero](#), [ok](#), [zeromiss](#).

misszero

Output: same type as input

Argument: x (scalar or series)

Converts NAs to zeros. If x is a series, the conversion is done element by element. See also [missing](#), [ok](#), [zeromiss](#).

mlag

Output: matrix
 Arguments: X (matrix)
 p (scalar or vector)
 m (scalar, optional)

Shifts up or down the rows of X . If p is a positive scalar, the returned matrix Y has typical element $Y_{i,j} = X_{i-p,j}$ for $i \geq p$ and zero otherwise. In other words, the columns of X are shifted down by p rows and the first p rows are filled with the value m . If p is a negative number, X is shifted up and the last rows are filled with the value m . If m is omitted, it is understood to be zero.

If p is a vector, the above operation is carried out for each element in p , joining the resulting matrices horizontally.

mnormal

Output: matrix
 Arguments: r (integer)
 c (integer)

Returns a matrix with r rows and c columns, filled with standard normal pseudo-random variates. See also [normal](#), [muniform](#).

mols

Output: matrix
 Arguments: Y (matrix)
 X (matrix)
 $\&U$ (reference to matrix, or null)
 $\&V$ (reference to matrix, or null)

Returns a $k \times n$ matrix of parameter estimates obtained by OLS regression of the $T \times n$ matrix Y on the $T \times k$ matrix X .

If the third argument is not null, the $T \times n$ matrix U will contain the residuals. If the final argument is given and is not null then the $k \times k$ matrix V will contain (a) the covariance matrix of the parameter estimates, if Y has just one column, or (b) $X'X^{-1}$ if Y has multiple columns.

By default, estimates are obtained via Cholesky decomposition, with a fallback to QR decomposition if the columns of X are highly collinear. The use of SVD can be forced via the command `set svd on`.

See also [mpols](#), [mrls](#).

monthlen

Output: integer
 Arguments: *month* (integer)
 year (integer)
 weeklen (integer)

Returns the number of (relevant) days in the specified month in the specified year; *weeklen*, which must equal 5, 6 or 7, gives the number of days in the week that should be counted (a value of 6 omits Sundays, and a value of 5 omits both Saturdays and Sundays).

movavg

Output: series
 Arguments: x (series)
 p (scalar)
 $control$ (integer, optional)

Depending on the value of the parameter p , returns either a simple or an exponentially weighted moving average of the input series x .

If $p > 1$, a simple p -term moving average is computed, that is, $\frac{1}{p} \sum_{i=0}^{p-1} x_{t-i}$. If a non-zero value is supplied for the optional $control$ parameter the MA is centered, otherwise it is “trailing”.

If $0 < p < 1$, an exponential moving average is computed: $y_t = px_t + (1 - p)y_{t-1}$. By default the output series y is initialized using the first valid value of x , but the $control$ parameter may be used to specify the number of initial observations that should be averaged to produce y_0 . A zero value for $control$ indicates that all the observations should be used.

mpols

Output: matrix
 Arguments: Y (matrix)
 X (matrix)
 $\&U$ (reference to matrix, or null)

Works exactly as [mols](#), except that the calculations are done in multiple precision using the GMP library.

By default GMP uses 256 bits for each floating point number, but you can adjust this using the environment variable GRETL_MP_BITS, e.g. GRETL_MP_BITS=1024.

mrangden

Output: matrix
 Arguments: d (string)
 $p1$ (scalar)
 $p2$ (scalar, conditional)
 $p3$ (scalar, conditional)
 $rows$ (integer)
 $cols$ (integer)
 Examples: `matrix mx = mrangden(u, 0, 100, 50, 1)`
`matrix mt14 = mrangden(t, 14, 20, 20)`

Works like [rangden](#) except that the return value is a matrix rather than a series. The initial arguments to this function (the number of which depends on the selected distribution) are as described for [rangden](#), but they must be followed by two integers to specify the number of rows and columns of the desired random matrix.

The first example above calls for a uniform random column vector of length 50, while the second example specifies a 20×20 random matrix with drawings from the t distribution with 14 degrees of freedom.

See also [mnormal](#), [muniform](#).

mread

Output: matrix
 Arguments: *fname* (string)
 import (boolean, optional)

Reads a matrix from a file named *fname*. If the filename has the suffix “.gz” it is assumed that gzip compression has been applied in writing the data; if it has the suffix “.bin” the file is assumed to be in binary format (see [mwrite](#) for details). Otherwise the file is assumed to be plain text, conforming to the following specification:

- It may start with any number of comments, defined as lines that start with the hash mark, #; such lines are ignored.
- The first non-comment line must contain two integers, separated by a space or a tab, indicating the number of rows and columns, respectively.
- The columns must be separated by spaces or tab characters.
- The decimal separator must be the dot character, “.”.

If a non-zero value is given for the optional *import* argument, the input file is looked for in the user’s “dot” directory. This is intended for use with the matrix-exporting functions offered in the context of the [foreign](#) command. In this case the *fname* argument should be a plain filename, without any path component.

Should an error occur (such as the file being badly formatted or inaccessible), an empty matrix is returned.

See also [bread](#), [mwrite](#).

mreverse

Output: matrix
 Argument: *X* (matrix)

Returns a matrix containing the rows of *X* in reverse order. If you wish to obtain a matrix in which the columns of *X* appear in reverse order you can do:

```
matrix Y = mreverse(X')'
```

mrls

Output: matrix
 Arguments: *Y* (matrix)
 X (matrix)
 R (matrix)
 q (column vector)
 &*U* (reference to matrix, or null)
 &*V* (reference to matrix, or null)

Restricted least squares: returns a $k \times n$ matrix of parameter estimates obtained by least-squares regression of the $T \times n$ matrix *Y* on the $T \times k$ matrix *X* subject to the linear restriction $RB = q$, where *B* denotes the stacked coefficient vector. *R* must have kn columns; each row of this matrix represents a linear restriction. The number of rows in *q* must match the number of rows in *R*.

If the fifth argument is not null, the $T \times n$ matrix U will contain the residuals. If the final argument is given and is not null then the $k \times k$ matrix V will hold the restricted counterpart to the matrix $X'X^{-1}$. The variance matrix of the estimates for equation i can be constructed by multiplying the appropriate sub-matrix of V by an estimate of the error variance for that equation.

mshape

Output: matrix
 Arguments: X (matrix)
 r (integer)
 c (integer)

Rearranges the elements of X into a matrix with r rows and c columns. Elements are read from X and written to the target in column-major order. If X contains fewer than $k = rc$ elements, the elements are repeated cyclically; otherwise, if X has more elements, only the first k are used.

See also [cols](#), [rows](#), [unvech](#), [vec](#), [vech](#).

msortby

Output: matrix
 Arguments: X (matrix)
 j (integer)

Returns a matrix in which the rows of X are reordered by increasing value of the elements in column j . This is a stable sort: rows that share the same value in column j will not be interchanged.

muniform

Output: matrix
 Arguments: r (integer)
 c (integer)

Returns a matrix with r rows and c columns, filled with uniform (0,1) pseudo-random variates. Note: the preferred method for generating a scalar uniform r.v. is to use the [randgen1](#) function.

See also [mnormal](#), [uniform](#).

mwrite

Output: integer
 Arguments: X (matrix)
 $fname$ (string)
 $export$ (boolean, optional)

Writes the matrix X to a file named $fname$. By default this file will be plain text; the first line will hold two integers, separated by a tab character, representing the number of rows and columns; on the following lines the matrix elements appear, in scientific notation, separated by tabs (one line per row). See below for alternative formats.

If a file $fname$ already exists, it will be overwritten. The return value is 0 on successful completion; if an error occurs, such as the file being unwritable, the return value will be non-zero.

If a non-zero value is given for the $export$ argument, the output file will be written into the user's "dot" directory, where it is accessible by default via the matrix-loading functions offered in the context of the [foreign](#) command. In this case a plain filename, without any path component, should be given for the second argument.

Matrices stored via the `mwrite` function in its default form can be easily read by other programs; see the *Gretl User's Guide* for details.

Two mutually exclusive inflections of this function are available, as follows:

- If *fname* has the suffix “.gz” the file is written with gzip compression. The format is basically as described above except that the matrix elements are written in a single column, in column-major order.
- If *fname* has the suffix “.bin” the file is written in binary format. In this case the first 19 bytes contain the characters `gretl_binary_matrix`, the next 8 bytes contain two 32-bit integers giving the number of rows and columns, and the remainder of the file contains the matrix elements as little-endian “doubles”, in column-major order. If `gretl` is run on a big-endian system, the binary values are converted to little endian on writing, and converted to big endian on reading.

Note that if the matrix file is to be read by a third-party program it is not advisable to use the gzip or binary options. But if the file is intended for reading by `gretl` the alternative formats save space, and the binary format allows for much faster reading of large matrices.

See also [mread](#).

mxtab

Output: matrix
 Arguments: *x* (series or vector)
 y (series or vector)

Returns a matrix holding the cross tabulation of the values contained in *x* (by row) and *y* (by column). The two arguments should be of the same type (both series or both column vectors), and because of the typical usage of this function, are assumed to contain integer values only.

See also [values](#).

nadarwat

Output: series
 Arguments: *y* (series)
 x (series)
 h (scalar)

Returns the Nadaraya-Watson nonparametric estimator of the conditional mean of *y* given *x*. It returns a series holding the nonparametric estimate of $E(y_i|x_i)$ for each nonmissing element of the series *x*.

$$m(x_i) = \frac{\sum_{j=1}^n y_j \cdot K_h(x_i - x_j)}{\sum_{j=1}^n K_h(x_i - x_j)}$$

where the kernel function $K_h(\cdot)$ is given by

$$K_h(x) = \exp\left(-\frac{x^2}{2h}\right)$$

for $|x| < \tau$ and zero otherwise.

The argument *h*, known as the *bandwidth*, is a parameter (a positive real number) given by the user. This is usually a small number: larger values of *h* make $m(x)$ smoother; a popular choice is $n^{-0.2}$. More details are given in the *Gretl User's Guide*.

The scalar τ is used to prevent numerical problems when the kernel function is evaluated too far away from zero and is called the trim parameter.

The trim parameter can be adjusted via the `nadarwat_trim` setting, as a multiple of h . The default value is 4.

The user may provide a negative value for the bandwidth: this is interpreted as conventional syntax to obtain the leave-one-out estimator, that is a variant of the estimator that does not use the i -th observation for evaluating $m(x_i)$. This makes the Nadaraya-Watson estimator more robust numerically and its usage is normally advised when the estimator is computed for inference purposes. Of course, the bandwidth actually used is the absolute value of h .

In formulae, the leave-one-out estimator is

$$m(x_i) = \frac{\sum_{j \neq i} y_j \cdot K_h(x_i - x_j)}{\sum_{j \neq i} K_h(x_i - x_j)}$$

nelem

Output: integer

Argument: L (list, matrix, bundle or array)

Returns the number of elements in the argument, which may be a list, a matrix, a bundle, or an array (but not a series).

ngetenv

Output: scalar

Argument: s (string)

If an environment variable by the name of s is defined and has a numerical value, returns that value; otherwise returns NA. See also [getenv](#).

nlines

Output: scalar

Argument: buf (string)

Returns a count of the complete lines (that is, lines that end with the newline character) in buf .

nobs

Output: integer

Argument: y (series)

Returns the number of non-missing observations for the variable y in the currently selected sample.

normal

Output: series

Arguments: μ (scalar)

σ (scalar)

Generates a series of Gaussian pseudo-random variates with mean μ and standard deviation σ . If no arguments are supplied, standard normal variates $N(0,1)$ are produced. The values are produced using the Ziggurat method ([Marsaglia and Tsang, 2000](#)).

See also [randgen](#), [mnormal](#), [muniform](#).

npv

Output: scalar
 Arguments: x (series or vector)
 r (scalar)

Returns the Net Present Value of x , considered as a sequence of payments (negative) and receipts (positive), evaluated at annual discount rate r , which must be expressed as a decimal fraction, not a percentage (0.05 rather than 5%). The first value is taken as dated “now” and is not discounted. To emulate an NPV function in which the first value is discounted, prepend zero to the input sequence. Supported data frequencies are annual, quarterly, monthly, and undated (undated data are treated as if annual).

See also [irr](#).

NRmax

Output: scalar
 Arguments: $\&b$ (reference to matrix)
 f (function call)
 g (function call, optional)
 h (function call, optional)

Numerical maximization via the Newton–Raphson method. On input the vector b should hold the initial values of a set of parameters, and the argument f should specify a call to a function that calculates the (scalar) criterion to be maximized, given the current parameter values and any other relevant data. If the object is in fact minimization, this function should return the negative of the criterion. On successful completion, NRmax returns the maximized value of the criterion, and b holds the parameter values which produce the maximum.

The optional third and fourth arguments provide means of supplying analytical derivatives and an analytical (negative) Hessian, respectively. The functions referenced by g and h must take as their first argument a pre-defined matrix that is of the correct size to contain the gradient or Hessian, respectively, given in pointer form. They also must take the parameter vector as an argument (in pointer form or otherwise). Other arguments are optional. If either or both of the optional arguments are omitted, a numerical approximation is used.

For more details and examples see the chapter on numerical methods in the *Gretl User's Guide*. See also [BFGSmax](#), [fdjac](#).

nullspace

Output: matrix
 Argument: A (matrix)

Computes the right nullspace of A , via the singular value decomposition: the result is a matrix B such that $AB = [0]$, except when A has full column rank, in which case an empty matrix is returned. Otherwise, if A is $m \times n$, B will be an $n \times (n - r)$ matrix, where r is the rank of A .

See also [rank](#), [svd](#).

obs

Output: series

Returns a series of consecutive integers, setting 1 at the start of the dataset. Note that the result is invariant to subsampling. This function is especially useful with time-series datasets. Note: you can write t instead of obs with the same effect.

See also [obsnum](#).

obslabel

Output: string
Argument: t (integer)

Returns the observation label for observation t , where t is a 1-based index. The inverse function is provided by [obsnum](#).

obsnum

Output: integer
Argument: s (string)

Returns an integer corresponding to the observation specified by the string s . Note that the result is invariant to subsampling. This function is especially useful with time-series datasets. For example, the following code

```
open denmark
k = obsnum(1980:1)
```

yields $k = 25$, indicating that the first quarter of 1980 is the 25th observation in the denmark dataset.

See also [obs](#), [obslabel](#).

ok

Output: see below
Argument: x (scalar, series, matrix or list)

If x is a scalar, returns 1 if x is not NA, otherwise 0. If x is a series, returns a series with value 1 at observations with non-missing values and zeros elsewhere. If x is a list, the output is a series with 0 at observations for which at least one series in the list has a missing value, and 1 otherwise.

If x is a matrix the behavior is a little different, since matrices cannot contain NAs: the function returns a matrix of the same dimensions as x , with 1s in positions corresponding to finite elements of x and 0s in positions where the elements are non-finite (either infinities or not-a-number, as per the IEEE 754 standard).

See also [missing](#), [misszero](#), [zeromiss](#). But note that these functions are not applicable to matrices.

onenorm

Output: scalar
Argument: X (matrix)

Returns the 1-norm of the $r \times c$ matrix X :

$$\|X\|_1 = \max_j \sum_{i=1}^r |X_{ij}|$$

See also [infnorm](#), [rcond](#).

ones

Output: matrix
 Arguments: r (integer)
 c (integer)

Outputs a matrix with r rows and c columns, filled with ones.

See also [seq](#), [zeros](#).

orthdev

Output: series
 Argument: y (series)

Only applicable if the currently open dataset has a panel structure. Computes the forward orthogonal deviations for variable y , that is

$$\tilde{y}_{i,t} = \sqrt{\frac{T_i - t}{T_i - t + 1}} \left(y_{i,t} - \frac{1}{T_i - t} \sum_{s=t+1}^{T_i} y_{i,s} \right)$$

This transformation is sometimes used instead of differencing to remove individual effects from panel data. For compatibility with first differences, the deviations are stored one step ahead of their true temporal location (that is, the value at observation t is the deviation that, strictly speaking, belongs at $t - 1$). That way one loses the first observation in each time series, not the last.

See also [diff](#).

pdf

Output: same type as input
 Arguments: d (string)
 ... (see below)
 x (scalar, series or matrix)
 Examples: $f1 = \text{pdf}(N, -2.5)$
 $f2 = \text{pdf}(X, 3, y)$
 $f3 = \text{pdf}(W, \text{shape}, \text{scale}, y)$

Probability density function calculator. Returns the density at x of the distribution identified by the code d . See [cdf](#) for details of the required (scalar) arguments. The distributions supported by the `pdf` function are the normal, Student's t , chi-square, F , Gamma, Weibull, Generalized Error, Binomial and Poisson. Note that for the Binomial and the Poisson what's calculated is in fact the probability mass at the specified point. For Student's t , chi-square, F the noncentral variants are supported too.

For the normal distribution, see also [dnorm](#).

pergm

Output: matrix
 Arguments: x (series or vector)
 bandwidth (scalar, optional)

If only the first argument is given, computes the sample periodogram for the given series or vector. If the second argument is given, computes an estimate of the spectrum of x using a Bartlett lag window of the given bandwidth, up to a maximum of half the number of observations ($T/2$).

Returns a matrix with two columns and $T/2$ rows: the first column holds the frequency, ω , from $2\pi/T$ to π , and the second the corresponding spectral density.

pexpand

Output: series
Argument: v (vector)

Only applicable if the currently open dataset has a panel structure. Performs the inverse operation of [pshrink](#). That is, given a vector of length equal to the number of individuals in the current panel sample, it returns a series in which each value is repeated T times, for T the time-series length of the panel. The resulting series is therefore non-time varying.

pmax

Output: series
Arguments: y (series)
 $mask$ (series, optional)

Only applicable if the currently open dataset has a panel structure. Returns a series holding the maxima of variable y for each cross-sectional unit (repeated for each time period).

If the optional second argument is provided then observations for which the value of $mask$ is zero are ignored.

See also [pmin](#), [pmean](#), [pnobs](#), [psd](#), [pxsum](#), [pshrink](#), [psum](#).

pmean

Output: series
Arguments: y (series)
 $mask$ (series, optional)

Only applicable if the currently open dataset has a panel structure. Computes the time-mean of variable y for each cross-sectional unit; that is,

$$\bar{y}_i = \frac{1}{T_i} \sum_{t=1}^{T_i} y_{i,t}$$

where T_i is the number of valid observations for unit i .

If the optional second argument is provided then observations for which the value of $mask$ is zero are ignored.

See also [pmax](#), [pmin](#), [pnobs](#), [psd](#), [pxsum](#), [pshrink](#), [psum](#).

pmin

Output: series
Arguments: y (series)
 $mask$ (series, optional)

Only applicable if the currently open dataset has a panel structure. Returns a series holding the minima of variable y for each cross-sectional unit (repeated for each time period).

If the optional second argument is provided then observations for which the value of $mask$ is zero are ignored.

See also [pmax](#), [pmean](#), [pnobs](#), [psd](#), [pshrink](#), [psum](#).

pnobs

Output: series
 Arguments: y (series)
 $mask$ (series, optional)

Only applicable if the currently open dataset has a panel structure. Returns a series holding the number of valid observations of variable y for each cross-sectional unit (repeated for each time period).

If the optional second argument is provided then observations for which the value of $mask$ is zero are ignored.

See also [pmax](#), [pmin](#), [pmean](#), [psd](#), [pshrink](#), [psum](#).

polroots

Output: matrix
 Argument: a (vector)

Finds the roots of a polynomial. If the polynomial is of degree p , the vector a should contain $p + 1$ coefficients in ascending order, i.e. starting with the constant and ending with the coefficient on x^p .

If all the roots are real they are returned in a column vector of length p , otherwise a $p \times 2$ matrix is returned, the real parts in the first column and the imaginary parts in the second.

polyfit

Output: series
 Arguments: y (series)
 q (integer)

Fits a polynomial trend of order q to the input series y using the method of orthogonal polynomials. The series returned holds the fitted values.

princomp

Output: matrix
 Arguments: X (matrix)
 p (integer)
 $covmat$ (boolean, optional)

Let the matrix X be $T \times k$, containing T observations on k variables. The argument p must be a positive integer less than or equal to k . This function returns a $T \times p$ matrix, P , holding the first p principal components of X .

The optional third argument acts as a boolean switch: if it is non-zero the principal components are computed on the basis of the covariance matrix of the columns of X (the default is to use the correlation matrix).

The elements of P are computed as

$$P_{tj} = \sum_{i=1}^k Z_{ti} v_i^{(j)}$$

where Z_{ti} is the standardized value of variable i at observation t , $Z_{ti} = (X_{ti} - \bar{X}_i) / \hat{\sigma}_i$, and $v^{(j)}$ is the j th eigenvector of the correlation (or covariance) matrix of the X_i s, with the eigenvectors ordered by decreasing value of the corresponding eigenvalues.

See also [eigensym](#).

prodc

Output: row vector
Argument: X (matrix)

Returns the product of the elements of X , by column. See also [prodr](#), [meanc](#), [sdc](#), [sumc](#).

prodr

Output: column vector
Argument: X (matrix)

Returns the product of the elements of X , by row. See also [prodc](#), [meanr](#), [sumr](#).

psd

Output: series
Arguments: y (series)
 $mask$ (series, optional)

Only applicable if the currently open dataset has a panel structure. Computes the per-unit sample standard deviation for variable y , that is

$$\sigma_i = \sqrt{\frac{1}{T_i - 1} \sum_{t=1}^{T_i} (y_{i,t} - \bar{y}_i)^2}$$

The above formula holds for $T_i \geq 2$, where T_i is the number of valid observations for unit i ; if $T_i = 0$, NA is returned; if $T_i = 1$, 0 is returned.

If the optional second argument is provided then observations for which the value of $mask$ is zero are ignored.

Note: this function makes it possible to check whether a given variable (say, X) is time-invariant via the condition $\max(\text{psd}(X)) == 0$.

See also [pmax](#), [pmin](#), [pmean](#), [pnobs](#), [pshrink](#), [psum](#).

psdroot

Output: square matrix
Argument: A (symmetric matrix)

Performs a generalized variant of the Cholesky decomposition of the matrix A , which must be positive semidefinite (but which may be singular). If the input matrix is not square an error is flagged, but symmetry is assumed and not tested; only the lower triangle of A is read. The result is a lower-triangular matrix L which satisfies $A = LL'$. Indeterminate elements in the solution are set to zero.

For the case where A is positive definite, see [cholesky](#).

pshrink

Output: matrix
Argument: y (series)

Only applicable if the currently open dataset has a panel structure. Returns a column vector holding the first valid observation for the series y for each cross-sectional unit in the panel, over the current sample range. If a unit has no valid observations for the input series it is skipped.

This function provides a means of compacting the series returned by functions such as [pmax](#) and [pmean](#), in which a value pertaining to each cross-sectional unit is repeated for each time period.

See [pexpand](#) for the inverse operation.

psum

Output: series
 Arguments: y (series)
 $mask$ (series, optional)

This function is applicable only if the current dataset has a panel structure. It computes the sum over time of variable y for each cross-sectional unit; that is,

$$S_i = \sum_{t=1}^{T_i} y_{i,t}$$

where T_i is the number of valid observations for unit i .

If the optional second argument is provided then observations for which the value of $mask$ is zero are ignored.

See also [pmax](#), [pmean](#), [pmin](#), [pnoobs](#), [psd](#), [pxsum](#), [pshrink](#).

pvalue

Output: same type as input
 Arguments: c (character)
 ... (see below)
 x (scalar, series or matrix)
 Examples: $p1 = pvalue(z, 2.2)$
 $p2 = pvalue(X, 3, 5.67)$
 $p2 = pvalue(F, 3, 30, 5.67)$

P -value calculator. Returns $P(X > x)$, where the distribution of X is determined by the character c . Between the arguments c and x , zero or more additional arguments are required to specify the parameters of the distribution; see [cdf](#) for details. The distributions supported by the `pval` function are the standard normal, t , Chi square, F , gamma, binomial, Poisson, Weibull and Generalized Error.

See also [critical](#), [invcdf](#), [urcpval](#), [imhof](#).

pxsum

Output: series
 Arguments: y (series)
 $mask$ (series, optional)

Only applicable if the currently open dataset has a panel structure. Computes the cross-sectional sum for variable y in each period; that is,

$$\tilde{y}_t = \sum_{i=1}^N y_{i,t}$$

where N is the number of cross-sectional units.

If the optional second argument is provided then observations for which the value of *mask* is zero are ignored.

Note that this function works in a different dimension from the [pmean](#) function.

qform

Output: matrix
 Arguments: *x* (matrix)
 A (symmetric matrix)

Computes the quadratic form $Y = xAx'$. Using this function instead of ordinary matrix multiplication guarantees more speed and better accuracy, when *A* is a generic symmetric matrix. However, in the special case $A = I$, the simple expression $x'x$ performs much better than `qform(x', I(rows(x)))`.

If *x* and *A* are not conformable, or *A* is not symmetric, an error is returned.

qlrpval

Output: scalar
 Arguments: *X2* (scalar)
 df (integer)
 p1 (scalar)
 p2 (scalar)

P-values for the test statistic from the QLR sup-Wald test for a structural break at an unknown point (see [qlrtest](#)), as per [Hansen \(1997\)](#).

The first argument, *X2*, denotes the (chi-square form of) the maximum Wald test statistic and *df* denotes its degrees of freedom. The third and fourth arguments represent, as decimal fractions of the overall estimation range, the starting and ending points of the central range of observations over which the successive Wald tests are calculated. For example if the standard approach of 15 percent trimming is adopted, you would set *p1* to 0.15 and *p2* to 0.85.

See also [pvalue](#), [urcpval](#).

qnorm

Output: same type as input
 Argument: *x* (scalar, series or matrix)

Returns quantiles for the standard normal distribution. If *x* is not between 0 and 1, NA is returned. See also [cnorm](#), [dnorm](#).

qrdecomp

Output: matrix
 Arguments: *X* (matrix)
 &*R* (reference to matrix, or null)

Computes the QR decomposition of an $m \times n$ matrix *X*, that is $X = QR$ where *Q* is an $m \times n$ orthogonal matrix and *R* is an $n \times n$ upper triangular matrix. The matrix *Q* is returned directly, while *R* can be retrieved via the optional second argument.

See also [eigengen](#), [eigensym](#), [svd](#).

quadtable

Output: matrix
 Arguments: n (integer)
 $type$ (integer, optional)
 a (scalar, optional)
 b (scalar, optional)

Returns an $n \times 2$ matrix for use with Gaussian quadrature (numerical integration). The first column holds the nodes or abscissae, the second the weights.

The first argument specifies the number of points (rows) to compute. The second argument codes for the type of quadrature: use 1 for Gauss-Hermite (the default); 2 for Gauss-Legendre; or 3 for Gauss-Laguerre. The significance of the optional parameters a and b depends on the selected $type$, as explained below.

Gaussian quadrature is a method of approximating numerically the definite integral of some function of interest. Let the function be represented as the product $f(x)W(x)$. The types of quadrature differ in the specification of the component $W(x)$: in the Hermite case we have $W(x) = \exp(-x^2)$; in the Laguerre case, $W(x) = \exp(-x)$; and in the Legendre case simply $W(x) = 1$.

For each specification of $W(x)$, one can compute a set of nodes, x_i , and weights, w_i , such that $\sum_{i=1}^n f(x_i)w_i$ approximates the desired integral. The method of [Golub and Welsch \(1969\)](#) is used.

When the Gauss-Legendre type is selected, the optional arguments a and b can be used to control the lower and upper limits of integration, the default values being -1 and 1 . (In Hermite quadrature the limits are fixed at $-\infty$ and $+\infty$, while in the Laguerre case they are fixed at 0 and ∞ .)

In the Hermite case a and b play a different role: they can be used to replace the default form of $W(x)$ with the (closely related) normal distribution with mean a and standard deviation b . Supplying values of 0 and 1 for these parameters, for example, has the effect of making $W(x)$ into the standard normal pdf, which is equivalent to multiplying the default x_i values by $\sqrt{2}$ and dividing the default w_i by $\sqrt{\pi}$.

quantile

Output: scalar or matrix
 Arguments: y (series or matrix)
 p (scalar between 0 and 1)

If y is a series, returns the p -quantile for the series. For example, when $p = 0.5$, the median is returned.

If y is a matrix, returns a row vector containing the p -quantiles for the columns of y ; that is, each column is treated as a series.

In addition, for matrix y an alternate form of the second argument is supported: p may be given as a vector. In that case the return value is an $m \times n$ matrix, where m is the number of elements in p and n is the number of columns in y .

For a series of length n , the p -quantile, q , is defined as:

$$q = y_{[k]} + [(n+1) \cdot p - k](y_{[k+1]} - y_{[k]})$$

where k is the integer part of $(n+1) \cdot p$ and $y_{[i]}$ is the i -th element of the series when sorted from smallest to largest.

randgen

Output: series
 Arguments: *d* (string)
 p1 (scalar or series)
 p2 (scalar or series, conditional)
 p3 (scalar, conditional)
 Examples: `series x = randgen(u, 0, 100)`
 `series t14 = randgen(t, 14)`
 `series y = randgen(B, 0.6, 30)`
 `series g = randgen(G, 1, 1)`
 `series P = randgen(P, mu)`

All-purpose random number generator. The argument *d* is a string (in most cases just a single character) which specifies the distribution from which the pseudo-random numbers should be drawn. The arguments *p1* to *p3* specify the parameters of the selected distribution; the number of such parameters depends on the distribution. For distributions other than the beta-binomial, the parameters *p1* and (if applicable) *p2* may be given as either scalars or series: if they are given as scalars the output series is identically distributed, while if a series is given for *p1* or *p2* the distribution is conditional on the parameter value at each observation. In the case of the beta-binomial all the parameters must be scalars.

Specifics are given below: the string code for each distribution is shown in parentheses, followed by the interpretation of the argument *p1* and, where applicable, *p2* and *p3*.

Distribution	<i>d</i>	<i>p1</i>	<i>p2</i>	<i>p3</i>
Uniform (continuous)	u or U	minimum	maximum	-
Uniform (discrete)	i	minimum	maximum	-
Normal	z, n or N	mean	standard deviation	-
Student's <i>t</i>	t	degrees of freedom	-	-
Chi square	c, x or X	degrees of freedom	-	-
Snedecor's <i>F</i>	f or F	df (num.)	df (den.)	-
Gamma	g or G	shape	scale	-
Binomial	b or B	<i>p</i>	<i>n</i>	-
Poisson	p or P	mean	-	-
Weibull	w or W	shape	scale	-
Generalized Error	e or E	shape	-	-
Beta	beta	shape1	shape2	-
Beta-Binomial	bb	<i>n</i>	shape1	shape2

See also [normal](#), [uniform](#), [mrandgen](#), [randgen1](#).

randgen1

Output: scalar
 Arguments: *d* (character)
 p1 (scalar)
 p2 (scalar, conditional)
 Examples: `scalar x = randgen1(z, 0, 1)`
 `scalar g = randgen1(g, 3, 2.5)`

Works like [randgen](#) except that the return value is a scalar rather than a series.

The first example above calls for a value from the standard normal distribution, while the second specifies a drawing from the Gamma distribution with shape 3 and scale 2.5.

See also [mrandgen](#).

randint

Output: integer
 Arguments: *min* (integer)
 max (integer)

Returns a pseudo-random integer in the closed interval [*min*, *max*]. See also [randgen](#).

rank

Output: integer
 Argument: *X* (matrix)

Returns the rank of *X*, numerically computed via the singular value decomposition. See also [svd](#).

ranking

Output: same type as input
 Argument: *y* (series or vector)

Returns a series or vector with the ranks of *y*. The rank for observation *i* is the number of elements that are less than y_i plus one half the number of elements that are equal to y_i . (Intuitively, you may think of chess points, where victory gives you one point and a draw gives you half a point.) One is added so the lowest rank is 1 instead of 0.

Formally,

$$\text{rank}(y_i) = 1 + \sum_{j \neq i} \left[I(y_j < y_i) + 0.5 \cdot I(y_j = y_i) \right]$$

where *I* denotes the indicator function.

See also [sort](#), [sortby](#).

rcond

Output: scalar
 Argument: *A* (square matrix)

Returns the reciprocal condition number for *A* with respect to the 1-norm. In many circumstances, this is a better measure of the sensitivity of *A* to numerical operations such as inversion than the determinant.

Given that *A* is non-singular, we may define

$$\kappa(A) = \|A\|_1 \cdot \|A^{-1}\|_1$$

This function returns $\kappa(A)^{-1}$.

See also [det](#), [ldet](#), [onenorm](#).

readfile

Output: string
 Arguments: *fname* (string)
 codeset (string, optional)

If a file by the name of *fname* exists and is readable, returns a string containing the content of this file, otherwise flags an error.

If *fname* starts with the identifier of a supported internet protocol (<http://>, <ftp://> or <https://>), libcurl is invoked to download the resource. See also [curl](#) for more elaborate downloading operations.

If the text to be read is not encoded in UTF-8, gretl will try recoding it from the current locale codeset if that is not UTF-8, or from ISO-8859-15 otherwise. If this simple default does not meet your needs you can use the optional second argument to specify a codeset. For example, if you want to read text in Microsoft codepage 1251 and that is not your locale codeset, you should give a second argument of "cp1251".

Also see the [sscanf](#) and [getline](#) functions.

regsub

Output: string
 Arguments: *s* (string)
 match (string)
 repl (string)

Returns a copy of *s* in which all occurrences of the pattern *match* are replaced using *repl*. The arguments *match* and *repl* are interpreted as Perl-style regular expressions.

See also [strsub](#) for simple substitution of literal strings.

remove

Output: integer
 Argument: *fname* (string)

If a file by the name of *fname* exists and is writable by the user, removes (deletes) the named file. Returns 0 on successful completion, non-zero if there is no such file or the file cannot be removed.

replace

Output: same type as input
 Arguments: *x* (series or matrix)
 find (scalar or vector)
 subst (scalar or vector)

Replaces each element of *x* equal to the *i*-th element of *find* with the corresponding element of *subst*.

If *find* is a scalar, *subst* must also be a scalar. If *find* and *subst* are both vectors, they must have the same number of elements. But if *find* is a vector and *subst* a scalar, then all matches will be replaced by *subst*.

Example:

```
a = {1,2,3;3,4,5}
find = {1,3,4}
```

```
subst = {-1,-8, 0}
b = replace(a, find, subst)
print a b
```

produces

a (2 x 3)

```
1  2  3
3  4  5
```

b (2 x 3)

```
-1  2 -8
-8  0  5
```

resample

Output: same type as input
 Arguments: x (series or matrix)
 blocksize (integer, optional)

The initial description of this function pertains to cross-sectional or time-series data; see below for the case of panel data.

Resamples from x with replacement. In the case of a series argument, each value of the returned series, y_t , is drawn from among all the values of x_t with equal probability. When a matrix argument is given, each row of the returned matrix is drawn from the rows of x with equal probability.

The optional argument *blocksize* represents the block size for resampling by moving blocks. If this argument is given it should be a positive integer greater than or equal to 2. The effect is that the output is composed by random selection with replacement from among all the possible contiguous sequences of length *blocksize* in the input. (In the case of matrix input, this means contiguous rows.) If the length of the data is not an integer multiple of the block size, the last selected block is truncated to fit.

If the argument x is a series and the dataset takes the form of a panel, resampling by moving blocks is not supported. The basic form of resampling is supported, but has this specific interpretation: the data are resampled “by individual”. Suppose you have a panel in which 100 individuals are observed over 5 periods. Then the returned series will again be composed of 100 blocks of 5 observations: each block will be drawn with equal probability from the 100 individual time series, with the time-series order preserved.

round

Output: same type as input
 Argument: x (scalar, series or matrix)

Rounds to the nearest integer. Note that when x lies halfway between two integers, rounding is done “away from zero”, so for example 2.5 rounds to 3, but `round(-3.5)` gives -4. This is a common convention in spreadsheet programs, but other software may yield different results. See also [ceil](#), [floor](#), [int](#).

rownames

Output: integer
 Arguments: M (matrix)
 S (array of strings or list)

Attaches names to the rows of the $m \times n$ matrix M . If S is a named list, the names are taken from the names of the listed series; the list must have m members. If S is an array of strings, it should contain m elements. For backward compatibility, a single string may also be given as the second argument; in that case it should contain m space-separated substrings.

The return value is 0 on successful completion, non-zero on error. See also [colnames](#).

Example:

```
matrix M = {1, 2; 2, 1; 4, 1}
strings S = array(3)
S[1] = "Row1"
S[2] = "Row2"
S[3] = "Row3"
rownames(M, S)
print M
```

rows

Output: integer
 Argument: X (matrix)

Returns the number of rows of the matrix X . See also [cols](#), [mshape](#), [unvech](#), [vec](#), [vech](#).

sd

Output: scalar or series
 Argument: x (series or list)

If x is a series, returns the (scalar) sample standard deviation, skipping any missing observations.

If x is a list, returns a series y such that y_t is the sample standard deviation of the values of the variables in the list at observation t , or NA if there are any missing values at t .

See also [var](#).

sdc

Output: row vector
 Arguments: X (matrix)
 df (scalar, optional)

Returns the standard deviations of the columns of X . If df is positive it is used as the divisor for the column variances, otherwise the divisor is the number of rows in X (that is, no degrees of freedom correction is applied). See also [meanc](#), [sumc](#).

sdiff

Output: same type as input
 Argument: y (series or list)

Computes seasonal differences: $y_t - y_{t-k}$, where k is the periodicity of the current dataset (see [\\$pd](#)). Starting values are set to NA.

When a list is returned, the individual variables are automatically named according to the template `sd_varname` where *varname* is the name of the original series. The name is truncated if necessary, and may be adjusted in case of non-uniqueness in the set of names thus constructed.

See also [diff](#), [ldiff](#).

seasonals

Output: list
 Arguments: *baseline* (integer, optional)
 center (boolean, optional)

Applicable only if the dataset has a time-series structure with periodicity greater than 1. Returns a list of dummy variables coding for the period or season, named S1, S2 and so on.

The optional *baseline* argument can be used to exclude one period from the set of dummies. For example, if you give a baseline value of 1 with quarterly data the returned list will hold dummies for quarters 2, 3 and 4 only. If this argument is omitted or set to zero a full set of dummies is generated; if non-zero, it must be an integer from 1 to the periodicity of the data.

The *center* argument, if non-zero, calls for the dummies to be centered; that is, to have their population mean subtracted. For example, with quarterly data centered seasonals will have values -0.25 and 0.75 rather than 0 and 1.

selifc

Output: matrix
 Arguments: *A* (matrix)
 b (row vector)

Selects from *A* only the columns for which the corresponding element of *b* is non-zero. *b* must be a row vector with the same number of columns as *A*.

See also [selifr](#).

selifr

Output: matrix
 Arguments: *A* (matrix)
 b (column vector)

Selects from *A* only the rows for which the corresponding element of *b* is non-zero. *b* must be a column vector with the same number of rows as *A*.

See also [selifc](#), [trimr](#).

seq

Output: row vector
 Arguments: *a* (integer)
 b (integer)
 k (integer, optional)

Given only two arguments, returns a row vector filled with consecutive integers, with *a* as first element and *b* last. If *a* is greater than *b* the sequence will be decreasing.

If the third argument is given, returns a row vector containing a sequence of integers starting with *a* and incremented (or decremented, if *a* is greater than *b*) by *k* at each step. The final value is the

largest member of the sequence that is less than or equal to b (or *mutatis mutandis* for a greater than b). The argument k must be a positive integer.

See also [ones](#), [zeros](#).

setnote

Output: integer
 Arguments: b (bundle)
 key (string)
 $note$ (string)

Sets a descriptive note for the object identified by key in the bundle b . This note will be shown when the `print` command is used on the bundle. This function returns 0 on success or non-zero on failure (for example, if there is no object in b under the given key).

simann

Output: scalar
 Arguments: $\&b$ (reference to matrix)
 f (function call)
 $maxit$ (integer, optional)

Implements simulated annealing, which may be helpful in improving the initialization for a numerical optimization problem.

On input the first argument holds the initial value of a parameter vector and the second argument specifies a function call which returns the (scalar) value of the maximand. The optional third argument specifies the maximum number of iterations (which defaults to 1024). On successful completion, `simann` returns the final value of the maximand and b holds the associated parameter vector.

For more details and an example see the chapter on numerical methods in the *Gretl User's Guide*. See also [BFGSmax](#), [NRmax](#).

sin

Output: same type as input
 Argument: x (scalar, series or matrix)

Returns the sine of x . See also [cos](#), [tan](#), [atan](#).

sinh

Output: same type as input
 Argument: x (scalar, series or matrix)

Returns the hyperbolic sine of x .

$$\sinh x = \frac{e^x - e^{-x}}{2}$$

See also [asinh](#), [cosh](#), [tanh](#).

skewness

Output: scalar
 Argument: x (series)

Returns the skewness value for the series x , skipping any missing observations.

sort

Output: same type as input

Argument: x (series or vector)

Sorts x in ascending order, skipping observations with missing values when x is a series. See also [dsort](#), [values](#). For matrices specifically, see [msortby](#).

sortby

Output: series

Arguments: $y1$ (series)

$y2$ (series)

Returns a series containing the elements of $y2$ sorted by increasing value of the first argument, $y1$. See also [sort](#), [ranking](#).

sprintf

Output: string

Arguments: *format* (string)

. . . (see below)

The returned string is constructed by printing the values of the trailing arguments, indicated by the dots above, under the control of *format*. See the help for the [printf](#) command for details regarding the format and other arguments.

square

Output: list

Arguments: L (list)

cross-products (boolean, optional)

Returns a list that references the squares of the variables in the list L , named on the pattern *sq_varname*. If the optional second argument is present and has a non-zero value, the returned list also includes the cross-products of the elements of L ; these are named on the pattern *var1_var2*. In these patterns the input variable names are truncated if need be, and the output names may be adjusted in case of duplication of names in the returned list.

sqrt

Output: same type as input

Argument: x (scalar, series or matrix)

Returns the positive square root of x ; produces NA for negative values.

Note that if the argument is a matrix the operation is performed element by element and, since matrices cannot contain NA, negative values generate an error. For the “matrix square root” see [cholesky](#).

sscanf

Output: integer
 Arguments: *src* (string)
 format (string)
 ... (see below)

Reads values from *src* under the control of *format* and assigns these values to one or more trailing arguments, indicated by the dots above. Returns the number of values assigned. This is a simplified version of the `sscanf` function in the C programming language.

src may be either a literal string, enclosed in double quotes, or the name of a predefined string variable. *format* is defined similarly to the format string in `printf` (more on this below). *args* should be a comma-separated list containing the names of pre-defined variables: these are the targets of conversion from *src*. (For those used to C: one can prefix the names of numerical variables with `&` but this is not required.)

Literal text in *format* is matched against *src*. Conversion specifiers start with `%`, and recognized conversions include `%f`, `%g` or `%lf` for floating-point numbers; `%d` for integers; `%s` for strings; and `%m` for matrices. You may insert a positive integer after the percent sign: this sets the maximum number of characters to read for the given conversion (or the maximum number of rows in the case of matrix conversion). Alternatively, you can insert a literal `*` after the percent to suppress the conversion (thereby skipping any characters that would otherwise have been converted for the given type). For example, `%3d` converts the next 3 characters in *source* to an integer, if possible; `.*g` skips as many characters in *source* as could be converted to a single floating-point number.

Matrix conversion works thus: the scanner reads a line of input and counts the (space- or tab-separated) number of numeric fields. This defines the number of columns in the matrix. By default, reading then proceeds for as many lines (rows) as contain the same number of numeric columns, but the maximum number of rows to read can be limited as described above.

In addition to `%s` conversion for strings, a simplified version of the C format `%N[chars]` is available. In this format *N* is the maximum number of characters to read and *chars* is a set of acceptable characters, enclosed in square brackets: reading stops if *N* is reached or if a character not in *chars* is encountered. The function of *chars* can be reversed by giving a circumflex, `^`, as the first character; in that case reading stops if a character in the given set is found. (Unlike C, the hyphen does not play a special role in the *chars* set.)

If the source string does not (fully) match the format, the number of conversions may fall short of the number of arguments given. This is not in itself an error so far as `gretl` is concerned. However, you may wish to check the number of conversions performed; this is given by the return value.

Some examples follow:

```
scalar x
scalar y
sscanf("123456", "%3d%3d", x, y)

sprintf S, "1 2 3 4\n5 6 7 8"
S
matrix m
sscanf(S, "%m", m)
print m
```

sst

Output: scalar
 Argument: *y* (series)

Returns the sum of squared deviations from the mean for the non-missing observations in series *y*. See also [var](#).

stringify

Output: integer
Arguments: *y* (series)
S (array of strings)

Provides a means of defining string values for the series *y*. Two conditions must be satisfied for this to work: the target series must have nothing but integer values, none of them less than 1, and the array *S* must have at least *n* elements where *n* is the largest value in *y*. In addition each element of *S* must be valid UTF-8. See also [strvals](#).

The value returned is zero on success or a positive error code on error.

strlen

Output: integer
Argument: *s* (string)

Returns the number of characters in the string *s*. Note that this does not necessarily equal the number of bytes if some characters are outside of the printable-ASCII range.

strncmp

Output: integer
Arguments: *s1* (string)
s2 (string)
n (integer, optional)

Compares the two string arguments and returns an integer less than, equal to, or greater than zero if *s1* is found, respectively, to be less than, to match, or be greater than *s2*, up to the first *n* characters. If *n* is omitted the comparison proceeds as far as possible.

Note that if you just want to compare two strings for equality, that can be done without using a function, as in `if (s1 == s2) ...`

strsplit

Output: string or array of strings
Arguments: *s* (string)
i (integer, optional)

With no second argument, returns the array of strings that results from the splitting of *s* on white space.

If the second argument is provided, returns space-separated element *i* from the string *s*. The index *i* is 1-based, and it is an error if *i* is less than 1. In case *s* contains no spaces and *i* equals 1, a copy of the entire input string is returned; otherwise, in case *i* exceeds the number of space-separated elements an empty string is returned.

strstr

Output: string
Arguments: *s1* (string)
s2 (string)

Searches *s1* for an occurrence of the string *s2*. If a match is found, returns a copy of the portion of *s1* that starts with *s2*, otherwise returns an empty string.

rstrip

Output: string
Argument: *s* (string)

Returns a copy of the argument *s* from which leading and trailing white space have been removed.

strsub

Output: string
Arguments: *s* (string)
 find (string)
 subst (string)

Returns a copy of *s* in which all occurrences of *find* are replaced by *subst*. See also [regsub](#) for more complex string replacement via regular expressions.

strvals

Output: array of strings
Argument: *y* (series)

If the series *y* is string-valued, returns an array containing all its distinct values, ordered by the associated numerical values starting at 1. If *y* is not string-valued an empty strings array is returned. See also [stringify](#).

substr

Output: string
Arguments: *s* (string)
 start (integer)
 end (integer)

Returns a substring of *s*, from the character with (1-based) index *start* to that with index *end*, inclusive.

sum

Output: scalar or series
Argument: *x* (series, matrix or list)

If *x* is a series, returns the (scalar) sum of the non-missing observations in *x*. See also [sumall](#).

If *x* is a matrix, returns the sum of the elements of the matrix.

If *x* is a list, returns a series *y* such that y_t is the sum of the values of the variables in the list at observation *t*, or NA if there are any missing values at *t*.

sumall

Output: scalar
Argument: *x* (series)

Returns the sum of the observations of *x* over the current sample range, or NA if there are any missing values.

sumc

Output: row vector
 Argument: X (matrix)

Returns the sums of the columns of X . See also [meanc](#), [sumr](#).

sumr

Output: column vector
 Argument: X (matrix)

Returns the sums of the rows of X . See also [meanr](#), [sumc](#).

svd

Output: row vector
 Arguments: X (matrix)
 & U (reference to matrix, or null)
 & V (reference to matrix, or null)

Performs the singular values decomposition of the $r \times c$ matrix X :

$$X = U \begin{bmatrix} \sigma_1 & & & \\ & \sigma_2 & & \\ & & \ddots & \\ & & & \sigma_n \end{bmatrix} V$$

where $n = \min(r, c)$. U is $r \times n$ and V is $n \times c$, with $U'U = I$ and $VV' = I$.

The singular values are returned in a row vector. The left and/or right singular vectors U and V may be obtained by supplying non-null values for arguments 2 and 3, respectively. For any matrix A , the code

```
s = svd(A, &U, &V)
B = (U .* s) * V
```

should yield B identical to A (apart from machine precision).

See also [eigengen](#), [eigensym](#), [qrdecomp](#).

tan

Output: same type as input
 Argument: x (scalar, series or matrix)

Returns the tangent of x . See also [atan](#), [cos](#), [sin](#).

tanh

Output: same type as input
 Argument: x (scalar, series or matrix)

Returns the hyperbolic tangent of x .

$$\tanh x = \frac{e^{2x} - 1}{e^{2x} + 1}$$

See also [atanh](#), [cosh](#), [sinh](#).

toepsolv

Output: column vector
Arguments: c (vector)
 r (vector)
 b (vector)

Solves a Toeplitz system of linear equations, that is $Tx = b$ where T is a square matrix whose element $T_{i,j}$ equals c_{i-j} for $i \geq j$ and r_{j-i} for $i \leq j$. Note that the first elements of c and r must be equal, otherwise an error is returned. Upon successful completion, the function returns the vector x .

The algorithm used here takes advantage of the special structure of the matrix T , which makes it much more efficient than other unspecialized algorithms, especially for large problems. Warning: in certain cases, the function may spuriously issue a singularity error when in fact the matrix T is nonsingular; this problem, however, cannot arise when T is positive definite.

tolower

Output: string
Argument: s (string)

Returns a copy of s in which any upper-case characters are converted to lower case.

toupper

Output: string
Argument: s (string)

Returns a copy of s in which any lower-case characters are converted to upper case.

tr

Output: scalar
Argument: A (square matrix)

Returns the trace of the square matrix A , that is, the sum of its diagonal elements. See also [diag](#).

transp

Output: matrix
Argument: X (matrix)

Returns the transpose of X . Note: this is rarely used; in order to get the transpose of a matrix, in most cases you can just use the prime operator: X' .

trimr

Output: matrix
Arguments: X (matrix)
 $ttop$ (integer)
 $tbot$ (integer)

Returns a matrix that is a copy of X with $ttop$ rows trimmed at the top and $tbot$ rows trimmed at the bottom. The latter two arguments must be non-negative, and must sum to less than the total rows of X .

See also [selifr](#).

typestr

Output: string

Argument: *typecode* (integer)

Returns the name of the gretl data-type corresponding to *typecode*. This may be used in conjunction with the function [inbundle](#). The value returned is one of “scalar”, “series”, “matrix”, “string”, “bundle”, “array” or “null”.

uniform

Output: series

Arguments: a (scalar)

b (scalar)

Generates a series of uniform pseudo-random variates in the interval (a, b) , or, if no arguments are supplied, in the interval $(0,1)$. The algorithm used by default is the SIMD-oriented Fast Mersenne Twister developed by [Saito and Matsumoto \(2008\)](#).

See also [randgen](#), [normal](#), [mnormal](#), [muniform](#).

uniq

Output: column vector

Argument: x (series or vector)

Returns a vector containing the distinct elements of x , not sorted but in their order of appearance. See [values](#) for a variant that sorts the elements.

unvech

Output: square matrix

Argument: v (vector)

Returns an $n \times n$ symmetric matrix obtained by rearranging the elements of v . The number of elements in v must be a triangular integer — i.e., a number k such that an integer n exists with the property $k = n(n + 1)/2$. This is the inverse of the function [vech](#).

See also [mshape](#), [vech](#).

upper

Output: square matrix

Argument: A (square matrix)

Returns an $n \times n$ upper triangular matrix B for which $B_{ij} = A_{ij}$ if $i \leq j$ and 0 otherwise.

See also [lower](#).

urcpval

Output: scalar
 Arguments: *tau* (scalar)
 n (integer)
 niv (integer)
 itv (integer)

P-values for the test statistic from the Dickey-Fuller unit-root test and the Engle-Granger cointegration test, as per [MacKinnon \(1996\)](#).

The arguments are as follows: *tau* denotes the test statistic; *n* is the number of observations (or 0 for an asymptotic result); *niv* is the number of potentially cointegrated variables when testing for cointegration (or 1 for a univariate unit-root test); and *itv* is a code for the model specification: 1 for no constant, 2 for constant included, 3 for constant and linear trend, 4 for constant and quadratic trend.

Note that if the test regression is “augmented” with lags of the dependent variable, then you should give an *n* value of 0 to get an asymptotic result.

See also [pvalue](#), [qlrpval](#).

values

Output: column vector
 Argument: *x* (series or vector)

Returns a vector containing the distinct elements of *x* sorted in ascending order. If you wish to truncate the values to integers before applying this function, use the expression `values(int(x))`.

See also [uniq](#), [dsort](#), [sort](#).

var

Output: scalar or series
 Argument: *x* (series or list)

If *x* is a series, returns the (scalar) sample variance, skipping any missing observations.

If *x* is a list, returns a series y such that y_t is the sample variance of the values of the variables in the list at observation *t*, or NA if there are any missing values at *t*.

In each case the sum of squared deviations from the mean is divided by $(n - 1)$ for $n > 1$. Otherwise the variance is given as zero if $n = 1$, or as NA if $n = 0$.

See also [sd](#).

varname

Output: string
 Argument: *v* (integer or list)

If given an integer argument, returns the name of the variable with ID number *v*, or generates an error if there is no such variable.

If given a list argument, returns a string containing the names of the variables in the list, separated by commas. If the supplied list is empty, so is the returned string.

varnum

Output: integer
 Argument: *varname* (string)

Returns the ID number of the variable called *varname*, or NA if there is no such variable.

varsimul

Output: matrix
 Arguments: *A* (matrix)
 U (matrix)
 y0 (matrix)

Simulates a p -order n -variable VAR, that is $y_t = \sum_{i=1}^p A_i y_{t-i} + u_t$. The coefficient matrix A is composed by stacking the A_i matrices horizontally; it is $n \times np$, with one row per equation. This corresponds to the first n rows of the matrix `$compan` provided by gretl's `var` and `vecm` commands.

The u_t vectors are contained (as rows) in U ($T \times n$). Initial values are in $y0$ ($p \times n$).

If the VAR contains deterministic terms and/or exogenous regressors, these can be handled by folding them into the U matrix: each row of U then becomes $u_t = B'x_t + e_t$.

The output matrix has $T + p$ rows and n columns; it holds the initial p values of the endogenous variables plus T simulated values.

See also `$compan`, `var`, `vecm`.

vec

Output: column vector
 Argument: X (matrix)

Stacks the columns of X as a column vector. See also `mshape`, `unvech`, `vech`.

vech

Output: column vector
 Argument: A (square matrix)

Returns in a column vector the elements of A on and above the diagonal. Typically, this function is used on symmetric matrices; in this case, it can be undone by the function `unvech`. See also `vec`.

weekday

Output: same type as input
 Arguments: *year* (scalar or series)
 month (scalar or series)
 day (scalar or series)

Returns the day of the week (Sunday = 0, Monday = 1, etc.) for the date(s) specified by the three arguments, or NA if the date is invalid. Note that all three arguments must be of the same type, either scalars (integers) or series.

wmean

Output: series
 Arguments: Y (list)
 W (list)

Returns a series y such that y_t is the weighted mean of the values of the variables in list Y at observation t , the respective weights given by the values of the variables in list W at t . The weights can therefore be time-varying. The lists Y and W must be of the same length and the weights must be non-negative.

See also [wsd](#), [wvar](#).

wsd

Output: series
Arguments: Y (list)
 W (list)

Returns a series y such that y_t is the weighted sample standard deviation of the values of the variables in list Y at observation t , the respective weights given by the values of the variables in list W at t . The weights can therefore be time-varying. The lists Y and W must be of the same length and the weights must be non-negative.

See also [wmean](#), [wvar](#).

wvar

Output: series
Arguments: X (list)
 W (list)

Returns a series y such that y_t is the weighted sample variance of the values of the variables in list X at observation t , the respective weights given by the values of the variables in list W at t . The weights can therefore be time-varying. The lists Y and W must be of the same length and the weights must be non-negative.

The weighted sample variance is computed as

$$s_w^2 = \frac{n'}{n' - 1} \frac{\sum_{i=1}^n w_i (x_i - \bar{x}_w)^2}{\sum_{i=1}^n w_i}$$

where n' is the number of non-zero weights and $\bar{x}_w$ is the weighted mean.

See also [wmean](#), [wsd](#).

xmax

Output: scalar
Arguments: x (scalar)
 y (scalar)

Returns the greater of x and y , or NA if either value is missing.

See also [xmin](#), [max](#), [min](#).

xmin

Output: scalar
Arguments: x (scalar)
 y (scalar)

Returns the lesser of x and y , or NA if either value is missing.

See also [xmax](#), [max](#), [min](#).

zeromiss

Output: same type as input

Argument: x (scalar or series)

Converts zeros to NAs. If x is a series, the conversion is done element by element. See also [missing](#), [misszero](#), [ok](#).

zeros

Output: matrix

Arguments: r (integer)

c (integer)

Outputs a zero matrix with r rows and c columns. See also [ones](#), [seq](#).

Chapter 3

Operators

3.1 Precedence

Table 3.1 lists the operators available in `gretl` in order of decreasing precedence: the operators on the first row have the highest precedence, those on the second row have the second highest precedence and so on. Operators on any given row have equal precedence. Where successive operators have the same precedence the order of evaluation is in general left to right. The exception is exponentiation; the expression `a^b^c` is equivalent to `a^(b^c)`, not `(a^b)^c`.

Table 3.1: Operator precedence

()	[]	{}		
!	++	--	^	'
*	/	%	\	**
+	-	~		
>	<	>=	<=	..
==	!=			
&&				
?:				

In addition to the basic forms shown in the Table, several operators also have a “dot form” (as in “`.*`” which is read as “dot plus”). These are element-wise versions of the basic operators, for use with matrices exclusively; they have the same precedence as their basic counterparts. The available dot operators are as follows.

`.^` `.*` `./` `.+` `.-` `.>` `.<` `.>=` `.<=` `.=`

Each basic operator is shown once again in the following list along with a brief account of its meaning. Apart from the first three sets of grouping symbols, all operators are binary except where otherwise noted.

()	Function call
[]	Subscripting
{}	Matrix definition
!	Unary logical NOT
++	Increment (unary)
--	Decrement (unary)
^	Exponentiation
'	Matrix transpose (unary) or transpose-multiply (binary)
*	Multiplication
/	Division, matrix “right division”

%	Modulus
\	Matrix “left division”
**	Kronecker product
+	Addition
-	Subtraction
~	Matrix horizontal concatenation
	Matrix vertical concatenation
>	Boolean greater than
<	Boolean less than
>=	Greater than or equal
<=	Less than or equal
..	Range from-to (in constructing lists)
==	Boolean equality test
!=	Boolean inequality test
&&	Logical AND
	Logical OR
?:	Conditional expression

Details on the use of the matrix-related operators (including the dot operators) can be found in the chapter on matrices in the *Gretl User’s Guide*.

3.2 Assignment

The operators mentioned above are all intended for use on the right-hand side of an expression which assigns a value to a variable (or which just computes and displays a value — see the `eval` command). In addition we have the assignment operator itself, “=”. In effect this has the lowest precedence of all: the entire right-hand side is evaluated before assignment takes place.

Besides plain “=” several “inflected” versions of assignment are available. These may be used only when the left-hand side variable is already defined. The inflected assignment yields a value that is a function of the prior value on the left and the computed value on the right. Such operators are formed by prepending a regular operator symbol to the equals sign. For example,

```
y += x
```

The new value assigned to `y` by the statement above is the prior value of `y` plus `x`. The other available inflected operators, which work in an exactly analogous fashion, are as follows.

```
-=   *=   /=   %=   ^=   ~=   |=
```

In addition, a special form of inflected assignment is provided for matrices. Say matrix `M` is 2×2 . If you execute `M = 5` this has the effect of replacing `M` with a 1×1 matrix with single element 5. But if you do `M .= 5` this assigns the value 5 to all elements of `M` without changing its dimensions.

3.3 Increment and decrement

The unary operators `++` and `--` follow their operand,¹ which must be a variable of scalar type. Their simplest use is in stand-alone expressions, such as

¹The C programming language also supports prefix versions of `++` and `--`, which increment or decrement their operand before yielding its value. Only the postfix form is supported by `gretl`.

```
j++ # shorthand for j = j + 1  
k-- # shorthand for k = k - 1
```

However, they can also be embedded in more complex expressions, in which case they first yield the original value of the variable in question, then have the side-effect of incrementing or decrementing the variable's value. For example:

```
scalar i = 3  
k = i++  
matrix M = zeros(10, 1)  
M[i++] = 1
```

After the second line, `k` has the value 3 and `i` has value 4. The last line assigns the value 1 to element 4 of matrix `M` and sets `i` = 5.

Chapter 4

Comments in scripts

When a script does anything non-obvious, it's a good idea to add comments explaining what's going on. This is particularly useful if you plan to share the script with others, but it's also useful as a reminder to yourself — when you revisit a script some months later and wonder what it was supposed to be doing.

The comment mechanism can also be helpful when you're developing a script. There may come a point where you want to execute a script, but bypass execution of some portion of it. Obviously you could delete the portion you wish to bypass, but rather than lose that section you can “comment it out” so that it is ignored by gretl.

Two sorts of comments are supported by gretl. The simpler one is this:

- If a hash mark, #, is encountered in a gretl script, everything from that point to the end of the current line is treated as a comment, and ignored.

If you wish to “comment out” several lines using this mode, you'll have to place a hash mark at the start of each line.

The second sort of comment is patterned after the C programming language:

- If the sequence /* is encountered in a script, all the following input is treated as a comment until the sequence */ is found.

Comments of this sort can extend over several lines. Using this mode it is easy to add lengthy explanatory text, or to get gretl to ignore substantial blocks of commands. As in C, comments of this type cannot be nested.

How do these two comment modes interact? You can think of gretl as starting at the top of a script and trying to decide at each point whether it should or should not be in “ignore mode”. In doing so it follows these rules:

- If we're not in ignore mode, then # puts us into ignore mode till the end of the current line.
- If we're not in ignore mode, then /* puts us into ignore mode until */ is found.

This means that each sort of comment can be masked by the other.

- If /* follows # on a given line which does not already start in ignore mode, then there's nothing special about /*, it's just part of a #-style comment.
- If # occurs when we're already in ignore mode, it is just part of a comment.

A few examples follow.

```
/* multi-line comment
# hello
# hello */
```

In the above example the hash marks are not special; in particular the hash mark on the third line does not prevent the multi-line comment from terminating at */.

```
# single-line comment /* hello
```

Assuming we were not in ignore mode before the line shown above, it is just a single-line comment: the /* is masked, and does not open a multi-line comment.

You can append a comment to a command:

```
ols 1 0 2 3 # estimate the baseline model
```

Example of “commenting out”:

```
/*  
# let's skip this for now  
ols 1 0 2 3 4  
omit 3 4  
*/
```

Chapter 5

Options, arguments and path-searching

5.1 Invoking gretl

`gretl` (under MS Windows, `gretl.exe`)¹.

— Opens the program and waits for user input.

`gretl datafile`

— Starts the program with the specified datafile in its workspace. The data file may be in any of several formats (see the *Gretl User's Guide*); the program will try to detect the format of the file and treat it appropriately. See also Section 5.4 below for path-searching behavior.

`gretl --help` (or `gretl -h`)

— Print a brief summary of usage and exit.

`gretl --version` (or `gretl -v`)

— Print version identification for the program and exit.

`gretl --english` (or `gretl -e`)

— Force use of English instead of translation.

`gretl --run scriptfile` (or `gretl -r scriptfile`)

— Start the program and open a window displaying the specified script file, ready to run. See Section 5.4 below for path-searching behavior.

`gretl --db database` (or `gretl -d database`)

— Start the program and open a window displaying the specified database. If the database files (the `.bin` file and its accompanying `.idx` file) are not in the default system database directory, you must specify the full path. See also the *Gretl User's Guide* for details on databases.

`gretl --dump` (or `gretl -c`)

— Dump the program's configuration information to a plain text file (the name of the file is printed on standard output). May be useful for trouble-shooting.

`gretl --debug` (or `gretl -g`)

— (MS Windows only) Open a console window to display any messages sent to the “standard output” or “standard error” streams. Such messages are not usually visible on Windows; this may be useful for trouble-shooting.

5.2 Preferences dialog

Various things in `gretl` are configurable under the “Tools, Preferences” menu. Separate menu items are devoted to the choice of the monospaced font to be used in `gretl` screen output, and, on some platforms, the font used for menus and other messages. The other options are organized under five tabs, as follows.

¹On Linux, a “wrapper” script named `gretl` is installed. This script checks whether the `DISPLAY` environment variable is set; if so, it launches the GUI program, `gretl_x11`, and if not it launches the command-line program, `gretlcli`.

General: Here you can configure the base directory for gretl's shared files. In addition there are several check boxes. If your native language setting is not English and the local decimal point character is not the period ("."), unchecking "Use locale setting for decimal point" will make gretl use the period regardless. Checking "Allow shell commands" makes it possible to invoke shell commands in scripts and in the gretl console (this facility is disabled by default for security reasons).

Databases tab: You can select the directory in which to start looking for native gretl databases; the directory in which to start looking for RATS databases; the host name of the gretl database server to access; and the IP number and port number of the HTTP proxy server to use when contacting the database server (if you're behind a firewall).

Programs tab: You can specify the names or paths to various third-party programs that may be called by gretl under certain conditions. Note that the item "Command to compile T_EX files" can be set to either latex or pdflatex; if latex is selected, T_EX output will be previewed in DVI format; if pdflatex is selected, the preview will be in PDF format.

HCCME tab: Set preferences regarding robust covariance matrix estimation. See the *Gretl User's Guide* for details.

Manuals tab: Select your preferred language for the full gretl documentation in PDF format (currently only English and Italian are supported). When using the English documentation you can also choose between US letter paper and A4 paper.

Settings chosen via the Preferences dialog are stored from one gretl session to the next. Under MS Windows they are stored in the Windows registry; on other platforms they are stored in a plain text file named `.gretl2rc` in the user's home directory.

5.3 Invoking gretlcli

`gretlcli`

— Opens the program and waits for user input.

`gretlcli datafile`

— Starts the program with the specified datafile in its workspace. The data file may be in any format supported by gretl (see the *Gretl User's Guide* for details). The program will try to detect the format of the file and treat it appropriately. See also Section 5.4 for path-searching behavior.

`gretlcli --help` (or `gretlcli -h`)

— Prints a brief summary of usage.

`gretlcli --version` (or `gretlcli -v`)

— Prints version identification for the program.

`gretlcli --english` (or `gretlcli -e`)

— Force use of English instead of translation.

`gretlcli --run scriptfile` (or `gretlcli -r scriptfile`)

— Execute the commands in *scriptfile* then hand over input to the command line. See Section 5.4 for path-searching behavior.

`gretlcli --batch scriptfile` (or `gretlcli -b scriptfile`)

— Execute the commands in *scriptfile* then exit. When using this option you will probably want to redirect output to a file. See Section 5.4 for path-searching behavior.

When using the `--run` and `--batch` options, the script file in question must call for a data file to be opened. This can be done using the `open` command within the script.

5.4 Path searching

When the name of a data file or script file is supplied to `gretl` or `gretlcli` on the command line, the file is looked for as follows:

1. “As is”. That is, in the current working directory or, if a full path is specified, at the specified location.
2. In the user’s `gretl` directory (see Table 5.1 for the default values; note that `PERSONAL` is a placeholder that is expanded by Windows in a user- and language-specific way, typically involving “My Documents” on English-language systems).
3. In any immediate sub-directory of the user’s `gretl` directory.
4. In the case of a data file, search continues with the main `gretl` data directory. In the case of a script file, the search proceeds to the system script directory. See Table 5.1 for the default settings. (`PREFIX` denotes the base directory chosen at the time `gretl` is installed.)
5. In the case of data files the search then proceeds to all immediate sub-directories of the main data directory.

Table 5.1: Default path settings

	<i>Linux</i>	<i>MS Windows</i>
User directory	<code>\$HOME/gretl</code>	<code>PERSONAL\gretl</code>
System data directory	<code>PREFIX/share/gretl/data</code>	<code>PREFIX\gretl\data</code>
System script directory	<code>PREFIX/share/gretl/scripts</code>	<code>PREFIX\gretl\scripts</code>

Thus it is not necessary to specify the full path for a data or script file unless you wish to override the automatic searching mechanism. (This also applies within `gretlcli`, when you supply a filename as an argument to the `open` or `run` commands.)

When a command script contains an instruction to open a data file, the search order for the data file is as stated above, except that the directory containing the script is also searched, immediately after trying to find the data file “as is”.

MS Windows

Under MS Windows configuration information for `gretl` and `gretlcli` is stored in the Windows registry. A suitable set of registry entries is created when `gretl` is first installed, and the settings can be changed under `gretl`’s “Tools, Preferences” menu. In case anyone needs to make manual adjustments to this information, the entries can be found (using the standard Windows program `regedit.exe`) under `Software\gretl` in `HKEY_LOCAL_MACHINE` (the main `gretl` directory and the paths to various auxiliary programs) and `HKEY_CURRENT_USER` (all other configurable variables).

Chapter 6

Reserved Words

Reserved words, which cannot be used as the names of variables, fall into the following categories:

- Names of constants and data types, plus a few specials: `const`, `NA`, `null`, `obs`, `scalar`, `series`, `matrix`, `string`, `list`, `bundle`, `array`, `kalman`, `void`, `continue`, `next`, `to`.
- Names of gretl commands (see section [1.2](#)).

User-defined functions cannot have names which collide with built-in functions, the names of which are shown in Table [6.1](#).

Table 6.1: Function names

BFGSmax	I	NRmax	abs	acos	acosh	aggregate	argname
array	asin	asinh	atan	atanh	atof	bessel	bkfilt
bootci	bootpval	boxcox	bread	bwfilt	bwrite	cdemean	cdf
cdiv	ceil	cholesky	chowlin	cmult	cnorm	colname	colnames
cols	corr	corrgm	cos	cosh	cov	critical	cum
curl	deseas	det	diag	diagcat	diff	digamma	dnorm
dotwrite	dsort	dummify	easterday	eigengen	eigensym	eigsolve	epochday
errmsg	exp	fcstats	fdjac	fft	ffti	filter	firstobs
fixname	floor	fracdiff	fraclag	freq	gammafun	genseries	getenv
getline	ghk	gini	ginv	grab	halton	hdprod	hpfilt
imaxc	imaxr	imhof	iminc	iminr	inbundle	infnorm	inlist
int	inv	invcdf	invmls	invpd	irf	irr	isconst
isdiscrete	isdummy	isnan	isnull	isoconv	isodate	isstring	iwishart
jsonget	kdensity	kfilter	kpsscrt	ksimul	ksmooth	kurtosis	lags
lastobs	ldet	ldiff	lincomb	linearize	ljungbox	ln	lngamma
loess	log	log10	log2	logistic	lower	lrvar	max
maxc	maxr	mcorr	mcov	mcovg	mean	meanc	meanr
median	mexp	min	minc	minr	missing	misszero	mlag
mnormal	mols	monthlen	movavg	mpiallred	mpibcast	mpirecv	mpireduce
mpiscatter	mpisend	mpols	mrangden	mread	mreverse	mrls	mshape
msortby	muniform	mwrite	mxtab	nadarwat	nelem	ngetenv	nlines
nobs	normal	npv	nullspace	obslabel	obsnum	ok	onenorm
ones	orthdev	pdf	pergm	pexpand	pmax	pmean	pmin
pnobs	polroots	polyfit	princomp	printf	prodc	prodr	psd
psdroot	pshrink	psum	putarray	pvalue	pxsum	qform	qlrpval
qnorm	qrdecomp	quadtable	quantile	randgen	randgen1	randint	rank
ranking	rcond	readfile	regsub	remove	replace	resample	round
rownames	rows	sd	sdc	sdiff	seasonals	selifc	selifr
seq	setnote	simann	sin	sinh	skewness	sort	sortby
sprintf	sqrt	square	sscanf	sst	stringify	strlen	strncmp
strsplit	strstr	strstrip	strsub	strvals	substr	sum	sumall
sumc	sumr	svd	tan	tanh	toepsolv	tolower	toupper
tr	transp	trimr	typeof	typestr	uniform	uniq	unvech
upper	urcpval	values	var	varname	varnum	varsimul	vec
vech	weekday	wmean	wsd	wvar	xmax	xmin	zeromiss
zeros							

Bibliography

- Agresti, A. (1992) 'A survey of exact inference for contingency tables', *Statistical Science* 7: 131–153.
- Akaike, H. (1974) 'A new look at the statistical model identification', *IEEE Transactions on Automatic Control* AC-19: 716–723.
- Arellano, M. and S. Bond (1991) 'Some tests of specification for panel data: Monte carlo evidence and an application to employment equations', *The Review of Economic Studies* 58: 277–297.
- Belsley, D., E. Kuh and R. Welsch (1980) *Regression Diagnostics*, New York: Wiley.
- Breusch, T. S. and A. R. Pagan (1979) 'A simple test for heteroscedasticity and random coefficient variation', *Econometrica* 47: 1287–1294.
- Choi, I. (2001) 'Unit root tests for panel data', *Journal of International Money and Finance* 20(2): 249–272.
- Chow, G. C. and A.-I. Lin (1971) 'Best linear unbiased interpolation, distribution, and extrapolation of time series by related series', *The Review of Economics and Statistics* 53(4): 372–375.
- Cleveland, W. S. (1979) 'Robust locally weighted regression and smoothing scatterplots', *Journal of the American Statistical Association* 74(368): 829–836.
- Davidson, R. and J. G. MacKinnon (1993) *Estimation and Inference in Econometrics*, New York: Oxford University Press.
- (2004) *Econometric Theory and Methods*, New York: Oxford University Press.
- Doornik, J. A. (1998) 'Approximations to the asymptotic distribution of cointegration tests', *Journal of Economic Surveys* 12: 573–593. Reprinted with corrections in McAleer and Oxley (1999).
- Edgerton, D. and C. Wells (1994) 'Critical values for the cusumsq statistic in medium and large sized samples', *Oxford Bulletin of Economics and Statistics* 56: 355–365.
- Elliott, G., T. J. Rothenberg and J. H. Stock (1996) 'Efficient tests for an autoregressive unit root', *Econometrica* 64: 813–836.
- Engle, R. F. and C. W. J. Granger (1987) 'Co-integration and error correction: Representation, estimation, and testing', *Econometrica* 55: 251–276.
- Fiorentini, G., G. Calzolari and L. Panattoni (1996) 'Analytic derivatives and the computation of GARCH estimates', *Journal of Applied Econometrics* 11: 399–417.
- Geweke, J. (1991) 'Efficient simulation from the multivariate normal and student-t distributions subject to linear constraints'. In *Computer Science and Statistics: Proceedings of the Twenty-third symposium on the Interface*, pp. 571–578. Alexandria, VA: American Statistical Association.
- Geweke, J. and S. Porter-Hudak (1983) 'The estimation and application of long memory time series models', *Journal of Time Series Analysis* 4: 221–238.
- Godfrey, L. G. (1994) 'Testing for serial correlation by variable addition in dynamic models estimated by instrumental variables', *The Review of Economics and Statistics* 76(3): 550–559.
- Golub, G. H. and C. F. Van Loan (1996) *Matrix Computations*, Baltimore and London: The John Hopkins University Press, third edn.

- Golub, G. H. and J. H. Welsch (1969) 'Calculation of Gauss quadrature rules', *Mathematics of Computation* 23: 221–230.
- Greene, W. H. (2000) *Econometric Analysis*, Upper Saddle River, NJ: Prentice-Hall, fourth edn.
- Halton, J. H. and G. B. Smith (1964) 'Algorithm 247: Radical-inverse quasi-random point sequence', *Communications of the ACM* 7: 701–702.
- Hamilton, J. D. (1994) *Time Series Analysis*, Princeton, NJ: Princeton University Press.
- Hansen, B. E. (1997) 'Approximate asymptotic p values for structural-change tests', *Journal of Business & Economic Statistics* 15: 60–67.
- Heckman, J. (1979) 'Sample selection bias as a specification error', *Econometrica* 47: 153–161.
- Im, K. S., M. H. Pesaran and Y. Shin (2003) 'Testing for unit roots in heterogeneous panels', *Journal of Econometrics* 115: 53–74.
- Imhof, J. P. (1961) 'Computing the distribution of quadratic forms in normal variables', *Biometrika* 48: 419–426.
- Kiviet, J. F. (1986) 'On the rigour of some misspecification tests for modelling dynamic relationships', *Review of Economic Studies* 53: 241–261.
- Koenker, R. (1981) 'A note on studentizing a test for heteroscedasticity', *Journal of Econometrics* 17: 107–112.
- _____ (1994) 'Confidence intervals for regression quantiles'. In P. Mandl and M. Huskova (eds.), *Asymptotic Statistics*, pp. 349–359. New York: Springer-Verlag.
- Koenker, R. and G. Bassett (1978) 'Regression quantiles', *Econometrica* 46: 33–50.
- Koenker, R. and J. Machado (1999) 'Goodness of fit and related inference processes for quantile regression', *Journal of the American Statistical Association* 94: 1296–1310.
- Koenker, R. and Q. Zhao (1994) 'L-estimation for linear heteroscedastic models', *Journal of Non-parametric Statistics* 3: 223–235.
- Kwiatkowski, D., P. C. B. Phillips, P. Schmidt and Y. Shin (1992) 'Testing the null of stationarity against the alternative of a unit root: How sure are we that economic time series have a unit root?', *Journal of Econometrics* 54: 159–178.
- Levin, A., C.-F. Lin and J. Chu (2002) 'Unit root tests in panel data: asymptotic and finite-sample properties', *Journal of Econometrics* 108: 1–24.
- Locke, C. (1976) 'A test for the composite hypothesis that a population has a gamma distribution', *Communications in Statistics — Theory and Methods* A5: 351–364.
- MacKinnon, J. G. (1996) 'Numerical distribution functions for unit root and cointegration tests', *Journal of Applied Econometrics* 11: 601–618.
- Maddala, G. S. (1992) *Introduction to Econometrics*, Englewood Cliffs, NJ: Prentice-Hall.
- Marsaglia, G. and W. W. Tsang (2000) 'The ziggurat method for generating random variables', *Journal of Statistical Software* 5: 3–30.
- McAleer, M. and L. Oxley (1999) *Practical Issues in Cointegration Analysis*, Oxford: Blackwell.
- Nerlove, M. (1971) 'Further evidence on the estimation of dynamic economic relations from a time series of cross sections', *Econometrica* 39: 359–382.
- Neter, J., W. Wasserman and M. H. Kutner (1990) *Applied Linear Statistical Models*, Boston: Irwin, third edn.

- Ng, S. and P. Perron (2001) 'Lag length selection and the construction of unit root tests with good size and power', *Econometrica* 69(6): 1519–1554.
- Odell, P. L. and A. H. Feiveson (1966) 'A numerical procedure to generate a sample covariance matrix', *Journal of the American Statistical Association* 61: 199–203.
- Perron, P. and Z. Qu (2007) 'A simple modification to improve the finite sample properties of Ng and Perron's unit root tests', *Economics Letters* 94(1): 12–19.
- Pesaran, M. H. and L. W. Taylor (1999) 'Diagnostics for IV regressions', *Oxford Bulletin of Economics and Statistics* 61(2): 255–281.
- Phillips, P. C. B. and K. Shimotsu (2004) 'Local Whittle estimation in nonstationary and unit root cases', *The Annals of Statistics* 32(2): 659–692. URL <http://arxiv.org/pdf/math/0406462>.
- Ramanathan, R. (2002) *Introductory Econometrics with Applications*, Fort Worth: Harcourt, fifth edn.
- Robinson, P. (1995) 'Gaussian semiparametric estimation of long range dependence', *Annals of Statistics* 22: 1630–1661.
- Saito, M. and M. Matsumoto (2008) 'SIMD-oriented Fast Mersenne Twister: a 128-bit pseudorandom number generator'. In A. Keller, S. Heinrich and H. Niederreiter (eds.), *Monte Carlo and Quasi-Monte Carlo Methods 2006*, pp. 607–622. Berlin: Springer.
- Satterthwaite, F. E. (1946) 'An approximate distribution of estimates of variance components', *Biometrics Bulletin* 2(6): 110–114.
- Schwert, G. W. (1989) 'Tests for unit roots: A Monte Carlo investigation', *Journal of Business and Economic Statistics* 7(2): 5–17.
- Sephton, P. S. (1995) 'Response surface estimates of the KPSS stationarity test', *Economics Letters* 47: 255–261.
- Shapiro, S. and L. Chen (2001) 'Composite tests for the gamma distribution', *Journal of Quality Technology* 33: 47–59.
- Stock, J. H., J. H. Wright and M. Yogo (2002) 'A survey of weak instruments and weak identification in generalized method of moments', *Journal of Business & Economic Statistics* 20(4): 518–529.
- Stock, J. H. and M. Yogo (2003) 'Testing for weak instruments in linear IV regression'. NBER Technical Working Paper 284. URL <http://www.nber.org/~myogo/papers/published/RothenbergCh5.pdf>.
- Swamy, P. A. V. B. and S. S. Arora (1972) 'The exact finite sample properties of the estimators of coefficients in the error components regression models', *Econometrica* 40: 261–275.
- Welch, B. L. (1951) 'On the comparison of several mean values: An alternative approach', *Biometrika* 38: 330–336.
- Windmeijer, F. (2005) 'A finite sample correction for the variance of linear efficient two-step GMM estimators', *Journal of Econometrics* 126: 25–51.