

Package ‘ivreg2r’

July 21, 2026

Title Extended Instrumental Variables Estimation with Diagnostics

Version 0.1.0

Description Comprehensive instrumental variables and GMM estimation with automatic diagnostics, inspired by the 'Stata' command 'ivreg2' of Baum, Schaffer, and Stillman (2003) <[doi:10.1177/1536867X0300300101](https://doi.org/10.1177/1536867X0300300101)> and Baum, Schaffer, and Stillman (2007) <[doi:10.1177/1536867X0800700402](https://doi.org/10.1177/1536867X0800700402)>. Supports 2SLS, LIML, Fuller, k-class, two-step efficient GMM, and continuously-updated (CUE) estimators. Provides classical, robust, cluster-robust, HAC, and Driscoll-Kraay standard errors. Reports weak identification, underidentification, overidentification, and endogeneity tests at estimation time. All outputs are verified against 'Stata' within tight numerical tolerances.

License GPL-3

Copyright Adapted from the Stata package 'ivreg2' by Christopher F Baum, Mark E Schaffer, and Steven Stillman, distributed under GPL-3. Component-level provenance for the adapted code and the bundled datasets is in inst/COPYRIGHTS.

URL <https://restatr.com/ivreg2r/>, <https://github.com/restatr/ivreg2r>

BugReports <https://github.com/restatr/ivreg2r/issues>

Encoding UTF-8

Language en-US

Depends R (>= 4.4.0)

Imports Formula, generics, stats, tibble

Suggests dplyr, ivreg, knitr, modelsummary, rmarkdown, sandwich, testthat (>= 3.1.5), tidyr

VignetteBuilder knitr

Config/testthat/edition 3

LazyData true

Config/roxygen2/version 8.0.0

NeedsCompilation no

Author Francis DiTraglia [aut, cre],
 Christopher F. Baum [ctb, cph] (Author of the Stata ivreg2 program from
 which ivreg2r is adapted),
 Mark E. Schaffer [ctb, cph] (Author of the Stata ivreg2 program from
 which ivreg2r is adapted),
 Steven Stillman [ctb, cph] (Author of the Stata ivreg2 program from
 which ivreg2r is adapted)

Maintainer Francis DiTraglia <francis.ditraglia@economics.ox.ac.uk>

Repository CRAN

Date/Publication 2026-07-21 11:00:02 UTC

Contents

abdata	3
augment.ivreg2	4
card	5
cigar	7
coef.ivreg2	8
confint.ivreg2	9
diagnostics	10
first_stage	12
fitted.ivreg2	13
formula.ivreg2	14
glance.ivreg2	15
griliches	17
grunfeld	18
ivreg2	19
ivreg2r-conventions	29
ivreg2r-glossary	31
klein	33
model.matrix.ivreg2	34
mroz	35
nlswork	37
nobs.ivreg2	38
phillips	39
predict.ivreg2	40
print.ivreg2	41
print.summary.ivreg2	42
residuals.ivreg2	43
stockwatson	43
summary.ivreg2	45
terms.ivreg2	46
tidy.ivreg2	46
ts-operators	48
update.ivreg2	49
vcov.ivreg2	50
wagepan	51

abdata

*Arellano–Bond (1991) UK Employment Panel***Description**

Employment, wages, and capital stock for UK companies: an unbalanced annual panel over 1976–1984 (140 companies, 1,031 firm-year observations). This is the companion dataset to Arellano, M. and Bond, S. (1991), "Some Tests of Specification for Panel Data: Monte Carlo Evidence and an Application to Employment Equations," *Review of Economic Studies*, 58(2), 277–297, shipped as distributed (levels plus the log transforms and year dummies used throughout the Stata literature).

Usage

abdata

Format

A data frame with 1,031 observations and 16 variables:

ind Industry code.

year Calendar year (1976–1984). The time variable: use as `tvar`.

emp Employment.

wage Real wage.

cap Gross capital stock.

indoutpt Industry output.

n Log employment, $\log(\text{emp})$.

w Log real wage, $\log(\text{wage})$.

k Log gross capital stock, $\log(\text{cap})$.

ys Log industry output, $\log(\text{indoutpt})$.

yr1980 Year dummy: 1 if `year == 1980`.

yr1981 Year dummy: 1 if `year == 1981`.

yr1982 Year dummy: 1 if `year == 1982`.

yr1983 Year dummy: 1 if `year == 1983`.

yr1984 Year dummy: 1 if `year == 1984`.

id Firm identifier (140 distinct companies). Use as `ivar`.

Details

`n`, `w`, `k`, and `ys` are natural logs of `emp`, `wage`, `cap`, and `indoutpt` respectively. `yr1980`–`yr1984` are year-dummy indicators (1 in the named year, 0 otherwise).

Source

Arellano, M. and Bond, S. (1991). Some tests of specification for panel data: Monte Carlo evidence and an application to employment equations. *Review of Economic Studies*, 58(2), 277–297.

Downloaded from <http://fmwww.bc.edu/ec-p/data/macro/abdata.dta>.

Redistribution basis: `system.file("COPYRIGHTS", package = "ivreg2r")`.

See Also

Other ivreg2r datasets: [card](#), [cigar](#), [griliches](#), [grunfeld](#), [klein](#), [mroz](#), [nlswork](#), [phillips](#), [stockwatson](#), [wagepan](#)

Examples

```
data(abdata)

# ivreg2 help file line 1558: one-way cluster on firm id
fit <- ivreg2(n ~ w + k, data = abdata, clusters = ~id)
summary(fit)
```

augment.ivreg2	<i>Augment data with ivreg2 model predictions and residuals</i>
----------------	---

Description

Adds `.fitted` and `.resid` columns to the model frame (or user-supplied data).

Usage

```
## S3 method for class 'ivreg2'
augment(x, data = NULL, ...)
```

Arguments

<code>x</code>	An object of class "ivreg2".
<code>data</code>	A data frame to augment. If <code>NULL</code> (default), uses the stored model frame (<code>x\$model</code>) and attaches the stored fitted values and residuals. An error is raised if <code>model = FALSE</code> was used at estimation time and data is not supplied. When data is supplied, <code>.fitted</code> is computed fresh via predict() on the supplied rows, which matches columns by name rather than position; the result is therefore correct for data that is reordered, subsetted, or extended with new rows. Rows with <code>NA</code> in a required predictor receive <code>NA</code> in <code>.fitted</code> . <code>.resid</code> is added only when the response column is present in data (broom convention); it is omitted otherwise. Supplying data for a model fit with <code>partial =</code> raises an error, since predict() cannot score new data after partialling. Because <code>.fitted</code> is a prediction, not an estimation-sample indicator, a row of data that was excluded from estimation only because the response (or an instrument) was <code>NA</code> still receives a

real predicted value whenever its regressors are complete; use `residuals()` or the stored model frame (`x$model`) to identify which rows were actually used in estimation.

... Additional arguments (ignored).

Value

A `tibble::tibble()` with all original data columns plus `.fitted`, and `.resid` when the response is available.

See Also

`ivreg2()`

Other broom methods: `glance.ivreg2()`, `tidy.ivreg2()`

Examples

```
data(mroz)
mroz_work <- subset(mroz, inlf == 1)
fit <- ivreg2(lwage ~ exper + expersq | educ | age + kidslt6 + kidsge6,
  data = mroz_work, vcov = "robust")
augment(fit) |> head()

# Residuals vs fitted
aug <- augment(fit)
plot(aug$.fitted, aug$.resid, xlab = "Fitted", ylab = "Residuals")
abline(h = 0, lty = 2)
```

card

Card (1995) College Proximity Dataset

Description

Cross-sectional data from the National Longitudinal Survey of Young Men (1966–1981) used by Card (1995) to estimate the return to schooling using college proximity as an instrument for education.

Usage

card

Format

A data frame with 3,010 observations and 33 variables:

nearc2 Grew up near a 2-year college (binary).

nearc4 Grew up near a 4-year college (binary).

educ Years of education.
age Age in years (1976).
fatheduc Father's years of education.
motheduc Mother's years of education.
weight NLS sampling weight.
momdad14 Lived with both parents at age 14 (binary).
sinmom14 Lived with single mother at age 14 (binary).
step14 Lived with stepparent at age 14 (binary).
reg661 Region dummy: New England (1966).
reg662 Region dummy: Middle Atlantic (1966).
reg663 Region dummy: East North Central (1966).
reg664 Region dummy: West North Central (1966).
reg665 Region dummy: South Atlantic (1966).
reg666 Region dummy: East South Central (1966).
reg667 Region dummy: West South Central (1966).
reg668 Region dummy: Mountain (1966).
reg669 Region dummy: Pacific (1966).
south66 Lived in the South in 1966 (binary).
black Black (binary).
smsa Lives in SMSA (binary, 1976).
south Lives in the South (binary, 1976).
smsa66 Lived in SMSA in 1966 (binary).
wage Hourly wage (cents, 1976).
enroll Enrolled in school in 1976 (binary).
KWW Knowledge of the World of Work test score.
IQ IQ score.
married Marital status (1976): coded 1–6 in the source data (1 = married; the remaining codes distinguish other marital statuses). Not a binary indicator; recode before use as a regressor.
libcrd14 Had a library card at age 14 (binary).
exper Years of labor market experience ($\text{age} - \text{educ} - 6$).
lwage Log hourly wage.
expersq Experience squared (exper^2).

Source

Card, D. (1995). "Using Geographic Variation in College Proximity to Estimate the Return to Schooling." In L.N. Christofides, E.K. Grant, and R. Swidinsky (Eds.), *Aspects of Labour Market Behaviour: Essays in Honour of John Vanderkamp*. University of Toronto Press.

Obtained from Stata's bcuse archive (Boston College), bcuse card.

Redistribution basis: `system.file("COPYRIGHTS", package = "ivreg2r")`.

See Also

Other ivreg2r datasets: [abdata](#), [cigar](#), [griliches](#), [grunfeld](#), [klein](#), [mroz](#), [nlswork](#), [phillips](#), [stockwatson](#), [wagepan](#)

Examples

```
data(card)
# IV regression: instrument education with college proximity.
# Control set from Card (1995, Table 2) / Wooldridge (2020), Example 15.4.
fit <- ivreg2(
  lwage ~ exper + expersq + black + smsa + south + smsa66 +
    reg662 + reg663 + reg664 + reg665 + reg666 + reg667 + reg668 + reg669 |
  educ | nearc4,
  data = card
)
summary(fit)
```

cigar

Cigarette Demand Panel (Baltagi–Levin / Baltagi–Griffin–Xiong)

Description

Annual per-capita cigarette sales for 46 U.S. states over 1963–1992 (a balanced panel, 1,380 state-year observations), along with the price per pack, population, consumer price index, per-capita disposable income, and the minimum price per pack among neighboring states used to instrument for price endogeneity. Baltagi, B. H. and Levin, D. (1992), "Cigarette taxation: raising revenues and reducing consumption," *Structural Change and Economic Dynamics*, 3(2), 321–335. See also Baltagi, B. H., Griffin, J. M. and Xiong, W. (2000), "To pool or not to pool: homogeneous versus heterogeneous estimators applied to cigarette demand," *Review of Economics and Statistics*, 82(1), 117–126.

Usage

```
cigar
```

Format

A data frame with 1,380 observations and 9 variables (46 states times 30 years, 1963–1992):

state State identifier code (46 distinct U.S. states). Use as `ivar`.

year Year, coded 63–92 (i.e., 1963–1992). The time variable: use as `tvar`.

price Price per pack of cigarettes (nominal).

pop Population.

pop16 Population above the age of 16.

cpi Consumer price index (1983 = 100).

ndi Per-capita nominal disposable income.

sales Cigarette sales, in packs per capita.

pimin Minimum price per pack of cigarettes in adjoining states.

Details

price and ndi are nominal; deflate by cpi for real terms.

Source

Baltagi, B. H. and Levin, D. (1992). Cigarette taxation: raising revenues and reducing consumption. *Structural Change and Economic Dynamics*, 3(2), 321–335.

Baltagi, B. H., Griffin, J. M. and Xiong, W. (2000). To pool or not to pool: homogeneous versus heterogeneous estimators applied to cigarette demand. *Review of Economics and Statistics*, 82(1), 117–126.

Distributed via R package plm, data(Cigar).

Redistribution basis: system.file("COPYRIGHTS", package = "ivreg2r").

See Also

Other ivreg2r datasets: [abdata](#), [card](#), [griliches](#), [grunfeld](#), [klein](#), [mroz](#), [nlswork](#), [phillips](#), [stockwatson](#), [wagepan](#)

Examples

```
data(cigar)

# Price-endogeneity IV specification in real (CPI-deflated) terms, with
# the neighboring-state minimum price as the excluded instrument, as in
# the GFIC empirical example.
cigar_real <- transform(
  cigar,
  lsales = log(sales),
  lrprice = log(price / cpi),
  lrndi = log(ndi / cpi),
  lrpimin = log(pimin / cpi)
)
fit <- ivreg2(lsales ~ lrndi | lrprice | lrpimin, data = cigar_real)
summary(fit)
```

coef.ivreg2

Extract coefficients from an ivreg2 object

Description

Extract coefficients from an ivreg2 object

Usage

```
## S3 method for class 'ivreg2'
coef(object, ...)
```

Arguments

object An object of class "ivreg2".
 ... Additional arguments (ignored).

Value

Named numeric vector of coefficient estimates.

See Also

[ivreg2\(\)](#)

Other ivreg2 methods: [confint.ivreg2\(\)](#), [diagnostics\(\)](#), [first_stage\(\)](#), [fitted.ivreg2\(\)](#), [formula.ivreg2\(\)](#), [ivreg2\(\)](#), [model.matrix.ivreg2\(\)](#), [nobs.ivreg2\(\)](#), [predict.ivreg2\(\)](#), [print.ivreg2\(\)](#), [print.summary.ivreg2\(\)](#), [residuals.ivreg2\(\)](#), [summary.ivreg2\(\)](#), [terms.ivreg2\(\)](#), [update.ivreg2\(\)](#), [vcov.ivreg2\(\)](#)

confint.ivreg2	<i>Confidence intervals for ivreg2 coefficients</i>
----------------	---

Description

Computes confidence intervals using the t distribution when small = TRUE was used at estimation time, and the standard normal otherwise.

Usage

```
## S3 method for class 'ivreg2'
confint(object, parm, level = 0.95, ...)
```

Arguments

object An object of class "ivreg2".
 parm A specification of which parameters to give intervals for, either a numeric vector of positions or a character vector of names. If missing, all parameters are included.
 level The confidence level (default 0.95).
 ... Additional arguments (ignored).

Value

A matrix with columns for the lower and upper confidence limits.

See Also

`ivreg2()`; [ivreg2r-conventions](#) for when the t vs normal distribution is used.

Other `ivreg2` methods: `coef.ivreg2()`, `diagnostics()`, `first_stage()`, `fitted.ivreg2()`, `formula.ivreg2()`, `ivreg2()`, `model.matrix.ivreg2()`, `nobs.ivreg2()`, `predict.ivreg2()`, `print.ivreg2()`, `print.summary.ivreg2()`, `residuals.ivreg2()`, `summary.ivreg2()`, `terms.ivreg2()`, `update.ivreg2()`, `vcov.ivreg2()`

diagnostics

Extract IV diagnostic tests as a tibble

Description

Flattens the IV specification tests stored on a fitted model into a single tibble, one row per test, so they can be filtered, joined, or printed like any other tidyverse output instead of navigated as a nested list (`x$diagnostics$overid$stat` and friends).

Usage

```
diagnostics(x, ...)

## S3 method for class 'ivreg2'
diagnostics(x, ...)
```

Arguments

`x` A fitted model object.

`...` Additional arguments (ignored).

Value

A `tibble::tibble()` with one row per diagnostic test present on the model and the following columns:

`test` Character. A stable machine-readable key, meant for `filter(test == ...)`-style selection (e.g. "overid", "weak_id_robust", "endogeneity").

`test_name` Character. The display name printed by `summary()`.

`statistic` Double. The test statistic, or — for the Stock-Yogo rows — the critical value rather than a statistic.

`df` Integer. Degrees of freedom, NA when not applicable.

`df2` Integer. Denominator degrees of freedom, NA unless the statistic is F-distributed.

`p_value` Double. NA when not defined (e.g. Stock-Yogo critical values, or an exactly identified overidentification row).

`tested_vars` Character. The variable(s) tested; populated only for the endogeneity, orthogonality, and redundancy rows.

note Character. A Stata-parity disclosure, or NA when none applies. A note appears when an explicit option the user set is silently not honored inside an automatically computed diagnostic, matching Stata's `ivreg2`: a non-Bartlett kernel on the identification and redundancy tests (Stata's `ranktest` hard-codes Bartlett); `kiefer` on the identification tests (the Kiefer VCE structure is dropped); `psd` on the identification and redundancy tests (`ranktest` never receives it); `center` on the endogeneity test; and `method = "cue"` (or `"lim1"`, for orthogonality) on the endogeneity and orthogonality C-statistics, which come from a recursive re-estimation that does not use those estimators. When more than one disclosure applies to the same test, the note column joins the sentences with a space; each note is a complete sentence, so splitting on sentence boundaries recovers the individual disclosures programmatically.

Rows appear only for the tests actually computed on the model: an IV fit reports underidentification, weak identification, overidentification, and — for every endogenous regressor, unless narrowed to a subset via `endog =` — the endogeneity test by default, while `orthog =` and `redundant =` add the corresponding orthogonality and redundancy rows only when the user requests them. The Stock-Yogo rows (keys of the form `"sy_iv_size_10"`) report critical values, not test statistics, so their `p_value` is always NA. An OLS fit (single-part formula) has no diagnostics at all and returns a zero-row tibble with the columns above.

The note column is unique to this accessor: `tidy()` and `glance()` keep their numeric-only tidy semantics and carry no disclosure text, so `diagnostics()` is the surface that exposes a note programmatically. The same notes are printed by `summary()` as footnotes below the diagnostics block.

Small-sample correction

the overidentification row (Sargan/Hansen J) and the endogeneity row (C-statistic) are never small-sample corrected, regardless of the main model's `small` argument. The first-stage F-statistics are always small-sample corrected, regardless of `small`; they are not columns of this tibble but are reached via `first_stage()`. All three match Stata's `ivreg2`. See [ivreg2r-conventions](#) for the full statement of which corrections `small` controls.

See Also

[first_stage\(\)](#), [glance.ivreg2\(\)](#)

Other `ivreg2` methods: [coef.ivreg2\(\)](#), [confint.ivreg2\(\)](#), [first_stage\(\)](#), [fitted.ivreg2\(\)](#), [formula.ivreg2\(\)](#), [ivreg2\(\)](#), [model.matrix.ivreg2\(\)](#), [nobs.ivreg2\(\)](#), [predict.ivreg2\(\)](#), [print.ivreg2\(\)](#), [print.summary.ivreg2\(\)](#), [residuals.ivreg2\(\)](#), [summary.ivreg2\(\)](#), [terms.ivreg2\(\)](#), [update.ivreg2\(\)](#), [vcov.ivreg2\(\)](#)

Examples

```
data(card)
fit <- ivreg2(
  lwage ~ exper + expersq + black + south | educ | nearc4 + nearc2,
  data = card, vcov = "robust"
)
diagnostics(fit)
```

first_stage	<i>Extract first-stage regression objects</i>
-------------	---

Description

Returns a named list of `ivreg2_first_stage` objects, one per endogenous variable. Each `ivreg2_first_stage` object supports the following methods:

Usage

```
first_stage(x, ...)

## S3 method for class 'ivreg2'
first_stage(x, ...)
```

Arguments

<code>x</code>	A fitted model object.
<code>...</code>	Additional arguments (ignored).

Details

`coef()` First-stage coefficient vector.
`vcov()` First-stage variance-covariance matrix.
`residuals()` First-stage residuals.
`fitted()` First-stage fitted values.
`nobs()` Number of observations used in the first stage.
`confint()` Confidence intervals for the first-stage coefficients.
`print()` Compact one-line-per-coefficient display.
`summary()` Full first-stage regression table: coefficients, standard errors, t-values, p-values, the first-stage F-statistic, and the partial R-squared (its `print()` method renders the table).
`tidy()` Broom-style tibble with one row per coefficient (term, estimate, std.error, statistic, p.value, and, optionally, confidence limits).
`glance()` Broom-style one-row tibble of first-stage fit statistics (sigma, degrees of freedom, F-statistic, partial R-squared, and related diagnostics).

Value

A named list of `ivreg2_first_stage` objects.

Small-sample correction

the first-stage F-statistic (surfaced via `glance()` and `summary()` on each `ivreg2_first_stage` object) always uses small-sample degrees of freedom, regardless of the main model's `small` argument. This matches Stata's `ivreg2`. See [ivreg2r-conventions](#) for the full statement of which corrections `small` controls.

See Also

`ivreg2()`

Other `ivreg2` methods: `coef.ivreg2()`, `confint.ivreg2()`, `diagnostics()`, `fitted.ivreg2()`, `formula.ivreg2()`, `ivreg2()`, `model.matrix.ivreg2()`, `nobs.ivreg2()`, `predict.ivreg2()`, `print.ivreg2()`, `print.summary.ivreg2()`, `residuals.ivreg2()`, `summary.ivreg2()`, `terms.ivreg2()`, `update.ivreg2()`, `vcov.ivreg2()`

Examples

```
# Mirrors the Stata `ivreg2` help-file example at line 1325, which
# demonstrates the equivalence of the Kleibergen-Paap rk Wald F and the
# first-stage F with a single endogenous regressor (and, in passing,
# the `savefirst` option that `first_stage = TRUE` corresponds to).
data(mroz)
mroz_work <- subset(mroz, inlf == 1)
fit <- ivreg2(lwage ~ exper + expersq | educ | age + kidslt6 + kidsge6,
             data = mroz_work, vcov = "robust", first_stage = TRUE)
fs <- first_stage(fit)
coef(fs$educ)
summary(fs$educ)
# KP rk Wald F = the robust first-stage F (single endogenous regressor)
c(kp_rk_wald_F = fit$diagnostics$weak_id_robust$stat,
  first_stage_F = glance(fs$educ)$f_stat)
```

fitted.ivreg2	<i>Extract fitted values from an ivreg2 object</i>
---------------	--

Description

Respects `na.action`: when the model was fit with `na.exclude`, NAs are reinserted at the omitted row positions so the result aligns with the original data frame (matching base R's `fitted.lm`).

Usage

```
## S3 method for class 'ivreg2'
fitted(object, ...)
```

Arguments

<code>object</code>	An object of class "ivreg2".
<code>...</code>	Additional arguments (ignored).

Value

Numeric vector of fitted values.

See Also

[ivreg2\(\)](#)

Other ivreg2 methods: [coef.ivreg2\(\)](#), [confint.ivreg2\(\)](#), [diagnostics\(\)](#), [first_stage\(\)](#), [formula.ivreg2\(\)](#), [ivreg2\(\)](#), [model.matrix.ivreg2\(\)](#), [nobs.ivreg2\(\)](#), [predict.ivreg2\(\)](#), [print.ivreg2\(\)](#), [print.summary.ivreg2\(\)](#), [residuals.ivreg2\(\)](#), [summary.ivreg2\(\)](#), [terms.ivreg2\(\)](#), [update.ivreg2\(\)](#), [vcov.ivreg2\(\)](#)

formula.ivreg2	<i>Extract formula from an ivreg2 object</i>
----------------	--

Description

Extract formula from an ivreg2 object

Usage

```
## S3 method for class 'ivreg2'
formula(x, ...)
```

Arguments

x	An object of class "ivreg2".
...	Additional arguments passed to Formula::formula.Formula() (e.g., rhs, lhs, collapse).

Value

The original model formula.

See Also

[ivreg2\(\)](#)

Other ivreg2 methods: [coef.ivreg2\(\)](#), [confint.ivreg2\(\)](#), [diagnostics\(\)](#), [first_stage\(\)](#), [fitted.ivreg2\(\)](#), [ivreg2\(\)](#), [model.matrix.ivreg2\(\)](#), [nobs.ivreg2\(\)](#), [predict.ivreg2\(\)](#), [print.ivreg2\(\)](#), [print.summary.ivreg2\(\)](#), [residuals.ivreg2\(\)](#), [summary.ivreg2\(\)](#), [terms.ivreg2\(\)](#), [update.ivreg2\(\)](#), [vcov.ivreg2\(\)](#)

glance.ivreg2	<i>Glance at an ivreg2 object</i>
---------------	-----------------------------------

Description

Returns a single-row tibble of model-level summary statistics and the headline IV diagnostics. The column set is deliberately compact so that table tools such as **modelsummary** render a sensible default.

Usage

```
## S3 method for class 'ivreg2'
glance(x, diagnostics = TRUE, ...)
```

Arguments

x	An object of class "ivreg2".
diagnostics	Logical: include the headline IV diagnostic columns? Default TRUE. Set to FALSE for a goodness-of-fit summary without the test statistics. Follows the same convention as broom's glance.ivreg().
...	Additional arguments (ignored).

Details

glance() returns a fixed set of columns for a given value of diagnostics, using NA for metrics that do not apply to the fitted model. All diagnostic columns are NA for OLS models (single-part formula). overid_stat and overid_p are also NA when the model is exactly identified (the number of excluded instruments equals the number of endogenous regressors), and weak_id_robust_stat is NA under vcov = "iid" (the Cragg-Donald F in weak_id_stat is reported instead of the Kleibergen-Paap F).

Set diagnostics = FALSE for a compact goodness-of-fit summary without the IV test columns.

For the full set of computed tests — including Stock-Yogo critical values, Anderson-Rubin, Stock-Wright, endogeneity, orthogonality, and redundancy — beyond these six headline columns, see [diagnostics\(\)](#).

Value

A single-row [tibble::tibble\(\)](#).

Always present (9 columns): r.squared, adj.r.squared, sigma, statistic (model F or Wald chi-squared), p.value, df (model numerator degrees of freedom), df.residual, nobis, vcov_type.

When diagnostics = TRUE (default, 6 additional columns): the headline IV specification tests — weak_id_stat (Cragg-Donald Wald F), weak_id_robust_stat (Kleibergen-Paap rk Wald F), underid_stat and underid_p (underidentification), overid_stat and overid_p (Sargan/Hansen J overidentification).

The remaining stored quantities and configuration flags are **not** in `glance()` — this keeps the goodness-of-fit block usable in rendered tables. They remain available as named elements on the fitted object: the estimation method, `lambda/kclass_value/fuller_parameter`, `coviv`, `center`, `psd`, `kernel/bw`, `kiefer`, `dkraay`, `sw`, cluster counts, `cue_convergence`, `partial_ct`, `small`, the cross-products `yy/yye`, ranks and condition numbers (`rank`, `rankzz`, `condxx`, `condzz`), the log-likelihood `ll`, and the full diagnostic list `x$diagnostics` (endogeneity, orthogonality, redundancy, Anderson-Rubin, Stock-Wright, and the Cragg-Donald/Kleibergen-Paap eigenvalues).

See Also

[ivreg2\(\)](#)

Other broom methods: [augment.ivreg2\(\)](#), [tidy.ivreg2\(\)](#)

Examples

```
data(mroz)
mroz_work <- subset(mroz, inlf == 1)
fit <- ivreg2(lwage ~ exper + expersq | educ | age + kidslt6 + kidsge6,
             data = mroz_work, vcov = "robust")

# Full output with diagnostics
glance(fit)

# Compact output without diagnostics
glance(fit, diagnostics = FALSE)

# Extract specific diagnostics
glance(fit)[, c("overid_stat", "overid_p")]
glance(fit)[, c("weak_id_stat", "weak_id_robust_stat")]

# Diagnostics dropped from glance() remain on the fitted object
fit$diagnostics$endogeneity

# Compare Sargan (IID) vs Hansen J (robust)
fit_iid <- ivreg2(lwage ~ exper + expersq | educ |
                 age + kidslt6 + kidsge6, data = mroz_work)
data.frame(
  vcov = c("iid", "robust"),
  overid = c(glance(fit_iid)$overid_stat, glance(fit)$overid_stat),
  overid_p = c(glance(fit_iid)$overid_p, glance(fit)$overid_p)
)

# A compact modelsummary table built from the curated glance() columns
if (requireNamespace("modelsummary", quietly = TRUE)) {
  modelsummary::modelsummary(
    list("2SLS" = fit),
    statistic = "std.error",
    gof_map = c("nobs", "r.squared", "weak_id_stat", "overid_stat")
  )
}
```

griliches*Griliches (1976) Young Men's Wages Dataset*

Description

Wages and characteristics of young men from the National Longitudinal Survey (NLS), originally analyzed in Griliches (1976). This dataset is a standard IV example used in Hayashi (2000, Ch. 3) and Baum, Schaffer & Stillman (2007, pp. 492–494) to illustrate weak identification tests, Anderson–Rubin inference, Stock–Wright S statistics, and instrument redundancy.

Usage

griliches

Format

A data frame with 758 observations and 20 variables:

rns Residency in the South (1 = yes).

rns80 Residency in the South in 1980.

mrt Marital status (1 = married).

mrt80 Marital status in 1980.

smsa Resides in metropolitan area (1 = urban).

smsa80 Resides in metropolitan area in 1980.

med Mother's education (years).

iq IQ score.

kww Score on Knowledge of the World of Work test.

year Survey year (66–73 or 80).

age Age at survey.

age80 Age in 1980.

s Completed years of schooling.

s80 Completed years of schooling in 1980.

expr Work experience (years).

expr80 Work experience in 1980 (years).

tenure Job tenure (years).

tenure80 Job tenure in 1980 (years).

lw Log wage.

lw80 Log wage in 1980.

Details

Two cross-sections are pooled (survey years 66–73 and 80). The Baum, Schaffer & Stillman (2007) examples use year dummies via `factor(year)`.

Source

Griliches, Z. (1976). Wages of very young men. *Journal of Political Economy*, 84(4), S69–S85.

Downloaded from <http://fmwww.bc.edu/ec-p/data/hayashi/griliches76.dta>.

Redistribution basis: `system.file("COPYRIGHTS", package = "ivreg2r")`.

References

Hayashi, F. (2000). *Econometrics*. Princeton University Press. (Chapter 3, p. 255.)

Baum, C.F., Schaffer, M.E. and Stillman, S. (2007). Enhanced routines for instrumental variables/generalized method of moments estimation and testing. *The Stata Journal*, 7(4), 465–506. (pp. 492–494.)

See Also

Other ivreg2r datasets: [abdata](#), [card](#), [cigar](#), [grunfeld](#), [klein](#), [mroz](#), [nlswork](#), [phillips](#), [stockwatson](#), [wagepan](#)

Examples

```
data(griliches)

# Baum, Schaffer & Stillman (2007), p. 493: IV with robust SEs, weak-ID
# diagnostics, redundancy test
fit <- ivreg2(lw ~ s + expr + tenure + rns + smsa + factor(year) |
             iq | age + mrt,
             data = griliches, vcov = "robust",
             redundant = "mrt")

summary(fit)
```

grunfeld

Grunfeld (1958) Corporate Investment Panel

Description

Corporate investment, market value, and capital stock for 10 large U.S. firms observed annually over 1935–1954: a balanced panel of 10 firms times 20 years (200 observations). Grunfeld, Y. (1958). *The Determinants of Corporate Investment*. PhD dissertation, University of Chicago. See Kleiber, C. and Zeileis, A. (2010) for the definitive account of the dataset’s provenance and its many circulating variants.

Usage

`grunfeld`

Format

A data frame with 200 observations and 6 variables (10 firms times 20 years, 1935–1954):

company Firm identifier (1–10). Use as ivar.

year Calendar year (1935–1954). The time variable: use as tvar.

invest Gross investment, current year.

mvalue Market value of the firm, prior year.

kstock Capital stock, prior year.

time Time index (1–20), a linear trend.

Source

Grunfeld, Y. (1958). *The Determinants of Corporate Investment*. PhD dissertation, University of Chicago.

Kleiber, C. and Zeileis, A. (2010). The Grunfeld data at 50. *German Economic Review*, 11(4), 404–417.

Distributed via Stata's webuse grunfeld.

Redistribution basis: `system.file("COPYRIGHTS", package = "ivreg2r")`.

See Also

Other ivreg2r datasets: [abdata](#), [card](#), [cigar](#), [griliches](#), [klein](#), [mroz](#), [nlswork](#), [phillips](#), [stockwatson](#), [wagepan](#)

Examples

```
data(grunfeld)

# ivreg2 help file line 1595: Driscoll-Kraay standard errors (bw = 2),
# small-sample correction
fit <- ivreg2(
  invest ~ mvalue + kstock, data = grunfeld,
  dkraay = 2, small = TRUE, tvar = "year", ivar = "company"
)
summary(fit)
```

Description

Estimate models by OLS, two-stage least squares (2SLS), LIML, Fuller, or k-class with automatic diagnostic tests. Uses a three-part formula for IV: `y ~ exog | endo | instruments`.

Usage

```

ivreg2(
  formula,
  data,
  weights,
  subset,
  na.action = stats::na.omit,
  vcov = "iid",
  clusters = NULL,
  endog = NULL,
  orthog = NULL,
  redundant = NULL,
  method = "2sls",
  kclass = NULL,
  fuller = 0,
  coviv = FALSE,
  small = FALSE,
  dofminus = 0L,
  sdofminus = 0L,
  weight_type = "aweight",
  kernel = NULL,
  bw = NULL,
  tvar = NULL,
  ivar = NULL,
  kiefer = FALSE,
  dkraay = NULL,
  wmatrix = NULL,
  smatrix = NULL,
  b0 = NULL,
  noid = FALSE,
  partial = NULL,
  nopartialsmall = FALSE,
  center = FALSE,
  psd = NULL,
  sw = FALSE,
  reduced_form = "none",
  first_stage = FALSE,
  model = TRUE,
  x = FALSE,
  y = TRUE
)

```

Arguments

formula	A formula: $y \sim \text{exog}$ (OLS), $y \sim \text{exog} \mid \text{endo} \mid \text{instruments}$ (IV), or $y \sim \text{exog} \mid 0 \mid \text{instruments}$ (no endogenous regressors, with the excluded instruments contributing surplus moment conditions — Stata's (=z1 z2) form). The latter is estimated by OLS; the surplus conditions drive the Sargan/Hansen J overiden-
---------	--

	tification test and orthog C-tests, and method = "gmm2s" with vcov = "robust" gives Cragg's (1983) heteroskedastic OLS (HOLS) estimator (with the default iid VCE the two-step weighting matrix is proportional to $(Z'Z)^{-1}$ and the estimates equal OLS exactly). The response must be numeric; a factor or character response is rejected with an error.
data	A data frame containing the variables in the formula.
weights	Optional analytic weights expression (evaluated in data), equivalent to Stata's [aw=varname]. Must be strictly positive. For aweights and pweights, the weights are normalized internally to sum to N, following Stata's convention. This makes sigma (RMSE) scale-invariant for those types: multiplying all weights by a constant changes no coefficients, standard errors, or test statistics. Frequency and importance weights are <i>not</i> normalized — they redefine N as sum(weights) (see weight_type), so their scale matters by construction. See ivreg2r-conventions for how weights interact with sigma and for the comparison to <code>lm()</code> (coefficients match <code>lm()</code> for OLS; SEs and the VCV match only when small = TRUE).
subset	Optional subset expression (evaluated in data).
na.action	Function for handling NAs (default na.omit). Use na.exclude to drop incomplete cases during estimation but reinsert NAs at the omitted positions in <code>fitted()</code> , <code>residuals()</code> , and <code>predict()</code> output, so results align with the original data frame. Stata drops missing observations silently (equivalent to na.omit) and has no analogue of na.exclude .
vcov	Character: covariance type. One of "iid" (classical), "robust" (White heteroskedasticity-robust), "HAC" (heteroskedasticity and autocorrelation consistent), or "AC" (autocorrelation consistent). Matching is case-insensitive ("hac", "Hac", and "HAC" are all accepted), but the value is always normalized to and printed as the canonical spelling shown above. Use small = TRUE to apply finite-sample corrections. To match Stata's ivreg2, robust: use vcov = "robust". To match Stata's ivreg2, robust small: use vcov = "robust", small = TRUE.
clusters	One-sided formula specifying one or two cluster variables (e.g. ~ firmid for one-way, ~ firmid + year for two-way). Two-way clustering uses the Cameron-Gelbach-Miller (2006) formula. The effective cluster count is min(M1, M2) per Stata convention. The small argument controls whether the finite-sample correction $(N-1)/(N-K) * M/(M-1)$ is applied (matching Stata's cluster() small combination). Unlike Stata's ivreg2, which silently drops observations with missing cluster values from the estimation sample, ivreg2r raises an error; the user should drop or handle those rows explicitly (e.g., by filtering them out) so the estimation sample is never changed silently.
endog	Character vector of endogenous regressor names to test for exogeneity (endogeneity test / C-statistic). If NULL (default), tests all endogenous regressors. Names must match variables in the endogenous part of the formula. Ignored for OLS models. Note: Unlike Stata's ivreg2, which only computes the endogeneity test when the endog() option is explicitly specified, ivreg2r computes it automatically for all IV models.
orthog	Character vector of instrument names to test for orthogonality (instrument-subset C-statistic). Names must be included or excluded instruments (not endogenous

	regressors or the intercept). If NULL (default), no orthogonality test is computed. Ignored for OLS models. Equivalent to Stata's <code>orthog()</code> option.
redundant	Character vector of excluded instrument names to test for redundancy (zero first-stage explanatory power). The test is a KP rk LM test of $H_0: \text{rank}=0$ on the first-stage coefficient matrix for the tested instruments, conditional on maintained instruments. If NULL (default), no redundancy test is computed. On a model with no endogenous regressors (a one-part OLS formula, or the IV form $y \sim \text{exog} \mid 0 \mid \text{instruments}$), a warning reports that the test is skipped; Stata's <code>redundant()</code> is silently ignored in those cases. <code>noid = TRUE</code> and <code>b0</code> also suppress the redundancy test; supplying <code>redundant</code> alongside either is dropped with a warning rather than silently ignored. Equivalent to Stata's <code>redundant()</code> option.
method	Character: estimation method. One of "2sls" (default), "liml", "kclass", "gmm2s" (two-step efficient GMM), or "cue" (continuously updated GMM estimator). For OLS models (1-part formula), this is ignored. When <code>fuller > 0</code> is specified, method is automatically promoted to "liml". "gmm2s" uses the inverse of the moment covariance matrix as the optimal weighting matrix, yielding efficient estimates under the specified error structure. Incompatible with <code>fuller</code> and <code>kclass</code>. "cue" continuously updates both moment conditions and moment covariance at each candidate beta during optimization. Under iid errors (no robust/cluster/kernel VCE), CUE with no other options is equivalent to LIML with <code>coviv = TRUE</code>. Under non-iid errors, CUE generalizes LIML to produce efficient estimates. Incompatible with <code>fuller</code>, <code>kclass</code>, <code>wmatrix</code>, and <code>smatrix</code>. CUE optimizer note: This package uses R's <code>optim()</code> (BFGS + Nelder-Mead, scaled by the starting values) while Stata uses Mata's <code>optimize()</code> (modified Newton-Raphson). The CUE objective is non-convex and can have multiple local minima, and it can possess economically degenerate minima (e.g., negative R-squared) – a documented property of the CUE ratio objective (Hausman, Lewis, Menzel & Newey, 2011), not an optimizer failure. On every benchmarked model the two implementations agree on the minimum, including one such pathological case, where both converge to the same degenerate basin; because those basins are extremely flat, coefficient agreement across implementations there is limited to a few digits. Inspect the J-statistic and R-squared to judge whether a CUE solution is economically sensible.
kclass	Numeric scalar: user-supplied k value for k-class estimation. When supplied, method is automatically set to "kclass". Must be non-negative. Cannot be combined with <code>method = "liml"</code> or <code>fuller</code> .
fuller	Numeric scalar: Fuller (1977) modification parameter. Non-negative. 0 (the default) applies no Fuller adjustment; a positive value activates the Fuller-modified LIML estimator, which automatically sets <code>method</code> to "liml" and <code>k = lambda - fuller / (N - L)</code> . <code>fuller = 1</code> gives the bias-corrected LIML estimator; <code>fuller = 4</code> targets MSE. Cannot be combined with <code>kclass</code> .
coviv	Logical: if TRUE, use the 2SLS bread $(X_{\text{hat}}'X_{\text{hat}})^{-1}$ instead of the k-class bread for VCV computation in LIML/k-class estimation. This gives the "COVIV" (covariance at the IV estimates) VCV that is robust to misspecification of the LIML model. Applies only to <code>method = "liml"</code> and <code>method = "kclass"</code> ; for every other method (OLS, 2SLS, GMM2S, CUE) it is reset to FALSE with a warning. Default FALSE.

small	Logical: if TRUE, use small-sample corrections (t/F instead of z/chi-squared, N-K denominator for sigma).
dofminus	Non-negative integer: large-sample degrees-of-freedom adjustment. Subtracted from N in large-sample variance formulas (e.g., $\sigma^2 = \text{rss}/(N - \text{dofminus})$). Useful when fixed effects have been partialled out. Equivalent to Stata's <code>dofminus()</code> option.
sdofminus	Non-negative integer: small-sample degrees-of-freedom adjustment. Subtracted from the residual degrees of freedom alongside K (e.g., $\text{df.residual} = N - K - \text{dofminus} - \text{sdofminus}$). Useful when partialling out regressors. Equivalent to Stata's <code>sdofminus()</code> option.
weight_type	Character: type of weights. One of "aweight" (analytic weights, default), "fweight" (frequency weights), "pweight" (probability/sampling weights), or "iweight" (importance weights). aweight: Normalized to sum to N. Standard for WLS. fweight: Integer-valued; N is redefined as <code>sum(weights)</code>. HC meat uses $w * e^2$ (linear, not quadratic). pweight: Normalized to sum to N. Forces robust VCE, with a warning (overrides <code>vcov = "iid"</code> to "robust" and <code>vcov = "AC"</code> — explicit or kiefer-implied — to "HAC"). iweight: Not normalized; N is redefined as <code>sum(weights)</code> (float). All intermediate calculations use float N; posted <code>nobs</code> and <code>df.residual</code> use <code>floor(N)</code>. Restricted to IID VCE only — incompatible with robust, cluster, HAC, and GMM estimation.
kernel	Character: kernel function for HAC/AC standard errors. One of "bartlett", "parzen", "truncated", "tukey-hanning", "tukey-hamming", "qs" (quadratic spectral), "daniell", or "tent". When specified without an explicit <code>vcov</code> change, <code>vcov</code> is automatically set: "iid" becomes "AC", "robust" becomes "HAC". Requires <code>bw</code> and <code>tvar</code>. Note: Kleibergen-Paap identification tests (underidentification and weak identification) always use the Bartlett kernel internally, regardless of the user-specified kernel. This matches a known behavior in Stata's <code>ivreg2</code>, where <code>ranktest</code> hard-codes the Bartlett kernel.
bw	Numeric or "auto": bandwidth for kernel estimation. Must be positive numeric, or "auto" for automatic selection via Newey-West (1994). Auto selection is available for Bartlett, Parzen, and Quadratic Spectral kernels only, and is not supported for panel data (<code>ivar</code>). Required when <code>kernel</code> is specified.
tvar	Character: name of the time variable in data. Required for HAC/AC estimation and for time-series operators <code>l()</code> / <code>d()</code> in the formula (see ts-operators). Lags are resolved by time <i>value</i> (gap-aware), mirroring Stata's <code>tsset</code> .
ivar	Character: name of the panel identifier variable in data. Optional; only needed for panel data (HAC/AC estimation, or so that time-series operators lag within panel units).
kiefer	Logical: if TRUE, use the Kiefer (1980) VCE — autocorrelation-consistent with <code>kernel = Truncated</code> and <code>bandwidth = T</code> (the full time span). Requires panel data (<code>tvar + ivar</code>). Incompatible with robust VCE, clustering, explicit kernel or bandwidth. Equivalent to Stata's <code>ivreg2 ... , kiefer</code> .

<code>dkraay</code>	Positive numeric scalar: bandwidth for Driscoll-Kraay (1998) VCE. When specified, clusters on the time variable and applies kernel smoothing across time lags, producing standard errors robust to cross-sectional dependence. Requires panel data (<code>tvar + ivar</code>). Incompatible with explicit <code>bw</code> . If <code>kernel</code> is not specified, defaults to Bartlett. Equivalent to Stata's <code>ivreg2 ... , dkraay(3)</code> .
<code>wmatrix</code>	Numeric matrix: user-supplied $L \times L$ weighting matrix for GMM estimation, where L is the number of instruments (including exogenous regressors). When supplied without <code>method = "gmm2s"</code> , produces inefficient GMM with full sandwich VCV (method becomes "gmmw"). When combined with <code>method = "gmm2s"</code> , used as the first-step weighting matrix (second step uses the optimal weighting matrix). Must be symmetric. Incompatible with <code>method = "liml"</code> , <code>kclass</code> , and <code>fuller</code> . When <code>method</code> is not "gmm2s", ignored without robust VCE, clustering, or HAC (with a warning). Equivalent to Stata's <code>wmatrix()</code> option. When used (not ignored), it is echoed back as <code>fit\$W</code> , matching Stata's <code>e(W)</code> . A matrix with <code>dimnames</code> is matched to the instrument columns by name (reordered as needed, mirroring Stata's <code>matsort</code> ; names that do not cover the instrument list are an error); an unnamed matrix is used positionally, in instrument column order.
<code>smatrix</code>	Numeric matrix: user-supplied $L \times L$ moment covariance matrix for GMM estimation. When supplied, the GMM estimation uses this matrix instead of computing Omega from residuals. Implies efficient GMM (method is promoted to "gmm2s" if "2sls"). Must be symmetric. Incompatible with <code>method = "liml"</code> , <code>kclass</code> , and <code>fuller</code> . Equivalent to Stata's <code>smatrix()</code> option. Echoed back as <code>fit\$S</code> (never recomputed), matching Stata's <code>e(S)</code> . A matrix with <code>dimnames</code> is matched to the instrument columns by name (reordered as needed, mirroring Stata's <code>matsort</code>); an unnamed matrix is used positionally, in instrument column order.
<code>b0</code>	Numeric vector: evaluate the CUE objective at this fixed parameter vector without optimization. When supplied, <code>method</code> is promoted to "cue" and the <code>J(b0)</code> statistic is computed and stored. Identification diagnostics (<code>underid</code> , <code>weakid</code> , <code>first-stage</code> , <code>AR</code> , <code>SW</code> , <code>endogeneity</code> , <code>orthog</code> , and the redundancy test) are suppressed (matching Stata's <code>b0()</code> option); <code>reduced_form</code> is not suppressed. An explicit <code>endog</code> , <code>orthog</code> , <code>redundant</code> , or <code>first_stage = TRUE</code> request supplied alongside <code>b0</code> is dropped with a warning, rather than silently ignored. Length must equal the number of regressors K . Named vectors are reordered to match model matrix columns; unnamed vectors are used as-is in model matrix column order. Incompatible with <code>method = "gmm2s"</code> , "liml", <code>kclass</code> , <code>fuller</code> , and <code>wmatrix</code> .
<code>noid</code>	Logical: if <code>TRUE</code> , suppress computation of underidentification, weak identification, and redundancy test statistics. Overidentification, Anderson-Rubin, Stock-Wright, <code>endogeneity</code> , and <code>first-stage</code> diagnostics are still computed. Default <code>FALSE</code> . Useful for speeding up simulation loops where rank-based identification tests are expensive. An explicit <code>redundant</code> request supplied alongside <code>noid = TRUE</code> is dropped with a warning, rather than silently ignored. Equivalent to Stata's <code>noid</code> option.
<code>partial</code>	Character vector: exogenous regressors to partial out via Frisch-Waugh-Lovell projection before estimation. Coefficients on partialled variables are not recoverable and are not reported. Special values:

- `"_cons"` or `"(Intercept)"`: partial only the constant (demean).
- `"_all"`: partial all included exogenous regressors.
- NULL (default): no partialling.

By default, `sdofminus` is automatically incremented by the number of partialled variables (including the constant). Use `nopartialsml` = TRUE to suppress this adjustment.

Note: FWL invariance holds for OLS, 2SLS, LIML, and two-step GMM, but **not** for CUE. The CUE objective recomputes the moment covariance at each iteration, so partialled residuals produce a different optimization surface. A warning is issued when `partial` is used with `method = "cue"`.

After partialling, `predict()` is restricted to residuals only (no `newdata`). Summary output notes that total SS, model F, and R-squared are partial-model values.

Equivalent to Stata's `partial()` option.

`nopartialsml` Logical: if TRUE, suppress the automatic `sdofminus` adjustment from partialling. Default FALSE. Equivalent to Stata's `nopartialsml` option.

`center` Logical: if TRUE, subtract the mean of moment conditions (scores) before computing the S matrix (meat of the sandwich VCE). Default FALSE. Centering only affects non-homoskedastic VCE types (HC, cluster, HAC); a warning is issued if used with IID or AC VCE.

Note: Centering is applied to the main model's VCE and to diagnostic tests that use the main model's S matrix (overidentification, orthogonality). However, the endogeneity test (`endog`) computes its own S matrix from the restricted model *without* centering, even when `center = TRUE`. This matches Stata's `ivreg2`, where `center` is not forwarded to the recursive call for the endogeneity test.

The Kleibergen-Paap identification statistics (underidentification, weak identification) and the instrument redundancy test are likewise always computed *without* centering, regardless of `center`. This matches Stata, whose `ranktest` calls (`underid`, `weak-id`, and `redundancy`) never receive the `center` option.

When a user-supplied `smatrix` is given, `center` has no effect on the estimator, the GMM weighting matrix, the stored `$S`, or any statistic computed from the supplied matrix (the overidentification and orthogonality tests reuse it): the supplied matrix is used as-is (mirroring the `psd` argument's identical interaction with `smatrix`, documented below). Centering still applies only to diagnostics that rebuild their own moment covariance, such as the Stock-Wright S statistic.

`psd` Character or NULL: PSD correction for the moment covariance matrix S. NULL (default) applies no correction. `"psd0"` zeroes negative eigenvalues (Politis 2007). `"psda"` replaces negative eigenvalues with their absolute values (Stock & Watson 2008). The correction is applied to S *before* the variance-covariance matrix is assembled from it, matching Stata's `m_omega`; the corrected S is stored as `$S`. The conventional (`vcov = "iid"`) VCV and the Kleibergen-Paap identification statistics are never corrected, also matching Stata (its `ranktest` does not receive the `psd` option). A warning is emitted whenever negative eigenvalues are detected and corrected (Stata corrects silently). Ignored when a user-supplied `smatrix` is given: the user matrix is used as-is for estimation, the VCV, and `$S`, matching Stata (whose `m_omega` is never invoked for a supplied S). Equivalent to Stata's `psd0` and `psda` options.

Under an exactly-singular psd0-corrected VCV, the model F statistic uses a conditioning-guarded (swept) inverse and can differ from Stata, which inverts the near-singular block directly; neither is a well-defined Wald statistic in that degenerate case.

sw	Logical: if TRUE, compute the Stock-Watson (2008, Econometrica) panel-robust VCE. Requires panel data (<code>ivar</code>). Incompatible with clustering, HAC kernels, <code>kiefer</code>, <code>dkraay</code>, <code>fweights</code>, and <code>iweights</code>. Forces <code>vcov = "robust"</code> automatically, with a warning when <code>vcov = "iid"</code> is thereby overridden. Equivalent to Stata's <code>sw</code> option (labeled "BETA VERSION" in Stata). Precondition: the Stock-Watson (2008) correction is derived for a within-transformed (fixed-effects) panel regression with fixed T and large N. <code>ivreg2()</code> does not perform the within transformation for you: within-transform the data yourself (de-mean every variable by panel unit) before fitting, and pass <code>dofminus</code> equal to the number of panel units, which charges the absorbed unit means to the degrees of freedom. See the worked example in the "Fixed-effects panels: the Stock-Watson correction" section of <code>vignette("time-series-gmm")</code>.
reduced_form	Character: what reduced-form output to store. "none" (default) stores nothing. "rf" stores the $y \sim Z$ regression (equivalent to Stata's <code>saverf</code>; Stata's <code>rf</code> option <i>displays</i> the reduced form without storing it). "system" stores the full system of $y +$ all endogenous variables regressed on Z, with cross-equation VCV (equivalent to Stata's <code>savesfirst</code>, an option present in <code>ivreg2.ado</code> but absent from its help file; Stata's <code>sfirst</code> displays the system without storing it). On a model with no endogenous regressors (a one-part OLS formula, or the IV form $y \sim \text{exog} \mid \emptyset \mid \text{instruments}$), no reduced form is computed — matching Stata, which skips reduced-form estimation whenever the endogenous list is empty — and a warning reports that the request was ignored. Note: the reduced-form variance-covariance matrix always applies OLS-style small-sample degrees-of-freedom scaling, regardless of the main model's <code>small</code> argument. This matches Stata's <code>ivreg2</code> (<code>ivreg2.ado:3091-3097</code>). See ivreg2r-conventions for the full statement of which corrections small controls.
first_stage	Logical: if TRUE, store extractable first-stage regression objects on the fitted model. Access them via <code>first_stage()</code>; see <code>first_stage()</code> for the full list of supported S3 methods (<code>coef()</code>, <code>vcov()</code>, <code>summary()</code>, <code>tidy()</code>, <code>glance()</code>, and others). Equivalent to Stata's <code>savefirst</code> option. Default FALSE. On a model with no endogenous regressors (a one-part OLS formula, or the IV form $y \sim \text{exog} \mid \emptyset \mid \text{instruments}$), nothing is stored and a warning reports that the request was ignored.
model	Logical: if TRUE (default), store the model frame in the return object.
x	Logical: if TRUE, store model matrices (X, Z) in the return object.
y	Logical: if TRUE (default), store the response vector in the return object.

Value

An object of class "ivreg2". Beyond the usual components (coefficients, residuals, `vcov`, diagnostics, ...), the object stores the estimated moment-condition covariance matrix as `$S` and the GMM weighting matrix as `$W`, corresponding to Stata's `e(S)` and `e(W)`:

- S is present for every fit. Normalization matches Stata's m_omega : $iid\ sigma^2\ (Z'Z)/N$ with $\sigma^2 = RSS/(N - dofminus)$; robust and HAC, meat divided by $N - dofminus$; cluster, meat divided by N ; psd-corrected when psd is set; built from centered moments when center = TRUE. A user-supplied smatrix is echoed, never recomputed and never psd-corrected.
- W is NULL for method = "liml" and "kclass" (as in Stata, no weighting matrix is defined for k-class estimators). For two-step GMM, CUE, and b0 fits it is the inverse of the defining S ; for one-step OLS/2SLS it is $(1/\sigma^2)\ (Z'Z/N)^{-1}$; a user-supplied wmatrix is echoed (it is the first-step weighting matrix under method = "gmm2s").

Rows and columns are named by the R instrument columns — (Intercept) rather than Stata's $_cons$, and R column order (exogenous regressors before excluded instruments) rather than Stata's. Passing $fit\$S$ back via smatrix reproduces the corresponding efficient-GMM fit, mirroring the $e(S)$ reuse workflow in Stata's ivreg2 help file.

Overidentification and endogeneity tests

When the model carries more instruments than regressors, ivreg2() automatically reports an overidentification test: the Sargan (1958) statistic under $vcov = "iid"$, and the Hansen J statistic under robust, cluster, or HAC errors. The test evaluates the $L - K$ overidentifying restrictions jointly, where L is the number of instruments and K the number of regressors. A rejection indicates that at least one instrument violates its orthogonality condition, without identifying which one. A non-rejection is necessary but not sufficient for instrument validity. The test has no power against a bias shared by every instrument, so orthogonality conditions that fail in the same direction can still appear jointly satisfied. The test is also uninformative about the K exactly-identifying moment conditions, whose validity must be defended on substantive grounds. An exactly identified model has no overidentifying restrictions, and therefore no overidentification test.

The endogeneity test, also called the C-statistic or difference test, asks whether regressors treated as endogenous could instead be treated as exogenous. It is the overidentification statistic of the model that adds the suspect regressors to the instrument set, treating them as exogenous, minus that of the model that treats them as endogenous. Under the null of exogeneity, this difference follows a chi-square distribution with degrees of freedom equal to the number of regressors tested. Both overidentification statistics are formed from the same estimated moment covariance, which guarantees that the difference is non-negative (Hayashi 2000). ivreg2() computes this test automatically for every IV model, whereas Stata computes it only when the endog() option is given explicitly. The endog argument selects which endogenous regressors are tested. The statistic takes the difference-in-Sargan form under iid errors and the difference-in-Hansen-J form under robust, cluster, or HAC errors.

LIML, Fuller, and k-class estimators

LIML, Fuller's modification, and the general k-class estimators are alternatives to 2SLS motivated by weak instruments. LIML is approximately median-unbiased, and its Stock-Yogo weak-identification thresholds are more forgiving than those for 2SLS. LIML has no finite-sample moments of any order, however. Its sampling distribution can therefore be dispersed, and individual estimates can lie far from the truth. Fuller's (1977) modification sets $k = \lambda - \text{fuller} / (N - L)$, which yields an estimator with finite moments. The choice $\text{fuller} = 1$ is approximately best-unbiased, and $\text{fuller} = 4$ minimizes mean squared error, both within the standard weak-instrument framework. Weak-identification pretesting for LIML and Fuller uses estimator-specific critical values rather than the 2SLS values. LIML coincides with 2SLS when the model is exactly identified,

because the LIML eigenvalue equals one in that case. Fuller does not, since it subtracts $\text{fuller} / (N - L)$ from that eigenvalue.

See Also

[ivreg2r-conventions](#) for the statistical conventions (sigma normalization, small, weights, degrees of freedom) and how they differ from `lm()`; [ivreg2r-glossary](#) for the diagnostic-output acronyms; [summary.ivreg2\(\)](#) and [print.ivreg2\(\)](#) for console output; [tidy.ivreg2\(\)](#), [glance.ivreg2\(\)](#), and [augment.ivreg2\(\)](#) for broom-style output; [first_stage\(\)](#) for first-stage regressions. The vignettes `vignette("introduction", "ivreg2r")`, `vignette("advanced-iv", "ivreg2r")`, and `vignette("time-series-gmm", "ivreg2r")` work through applied examples.

Other ivreg2 methods: [coef.ivreg2\(\)](#), [confint.ivreg2\(\)](#), [diagnostics\(\)](#), [first_stage\(\)](#), [fitted.ivreg2\(\)](#), [formula.ivreg2\(\)](#), [model.matrix.ivreg2\(\)](#), [nobs.ivreg2\(\)](#), [predict.ivreg2\(\)](#), [print.ivreg2\(\)](#), [print.summary.ivreg2\(\)](#), [residuals.ivreg2\(\)](#), [summary.ivreg2\(\)](#), [terms.ivreg2\(\)](#), [update.ivreg2\(\)](#), [vcov.ivreg2\(\)](#)

Examples

```
data(mroz)
mroz_work <- subset(mroz, inlf == 1)

# --- Just-identified IV: return to schooling ---
# Wooldridge (2020), Example 15.1: father's education instruments
# education in a simple wage equation (replicates the published
# estimates: educ = 0.059, SE = 0.035). With one instrument for one
# endogenous variable, no overid test is possible; instrument validity
# must be defended on substantive grounds.
fit <- ivreg2(lwage ~ 1 | educ | fatheduc,
             data = mroz_work)
summary(fit)

# --- Overidentified IV: testing instrument validity ---
# The Stata ivreg2 help file's illustrative baseline (line 1274):
# age, kidslt6, and kidsge6 give two overidentifying restrictions,
# so the Sargan test can detect misspecification. Note its weak
# first stage (Cragg-Donald F ~ 4.3) -- see the LIML example below.
fit_overid <- ivreg2(lwage ~ exper + expersq | educ |
                   age + kidslt6 + kidsge6, data = mroz_work)
summary(fit_overid)

# --- Robust standard errors ---
# With heteroskedasticity, Kleibergen-Paap diagnostics replace
# Anderson/Cragg-Donald, and Hansen J replaces Sargan.
fit_robust <- ivreg2(lwage ~ exper + expersq | educ |
                   age + kidslt6 + kidsge6, data = mroz_work,
                   vcov = "robust")
summary(fit_robust)

# --- LIML ---
# Same baseline as above (a documented variation on the help-file
```

```

# spec). Its weak first stage (Cragg-Donald F ~ 4.3) is the setting
# where 2SLS bias toward OLS grows with the overidentification degree
# relative to instrument strength, and LIML's approximate
# median-unbiasedness -- established within the iid Stock-Yogo
# framework -- is attractive. LIML equals 2SLS when just-identified.
fit_liml <- ivreg2(lwage ~ exper + expersq | educ |
                  age + kidslt6 + kidsge6, data = mroz_work,
                  method = "liml")
summary(fit_liml)

# --- Endogeneity test ---
# Is education actually endogenous? The C-statistic tests
# H0: OLS is consistent (educ is exogenous).
# ivreg2 help file line 1283: the endog(educ) example on this equation.
fit_endog <- ivreg2(lwage ~ exper + expersq | educ |
                   age + kidslt6 + kidsge6, data = mroz_work,
                   endog = "educ")
if (requireNamespace("dplyr", quietly = TRUE)) {
  diagnostics(fit_endog) |> dplyr::filter(test == "endogeneity")
}

# --- Clustering ---
# Griliches (1976) wage equation, cluster on year.
# Adapted from Stata `ivreg2` help-file example (line 1250): the help
# command also includes year dummies (xi i.year), omitted here, and no
# small option. Either way the fit warns -- 7 year clusters < the
# instrument count makes the moment covariance singular, which is the
# help file's own lead-in to partialling; see
# vignette("time-series-gmm").
data(griliches)
fit_cl <- ivreg2(lw ~ s + expr + tenure + rns + smsa |
                iq | med + kww + age,
                data = griliches, clusters = ~year, small = TRUE)
summary(fit_cl)

# --- Two-step efficient GMM ---
# ivreg2 help file line 1154: efficient GMM with robust weighting.
# The optimal weighting matrix uses the first-step robust moment
# covariance; Hansen J replaces Sargan.
fit_gmm <- ivreg2(lw ~ s + expr + tenure + rns + smsa + factor(year) |
                 iq | med + kww + age + mrt,
                 data = griliches, method = "gmm2s", vcov = "robust")
summary(fit_gmm)

```

Description

ivreg2r is a translation of Stata's ivreg2, and it follows Stata's statistical conventions rather than those of R's `lm()/vcov()` ecosystem. This topic collects those conventions in one place. The API itself is idiomatic R (formula interface, tidyverse-friendly output); only the *statistical* choices below lean on Stata. See `ivreg2()` for the estimation function and [ivreg2r-glossary](#) for the diagnostic acronyms printed by `summary()`.

Default VCE and small-sample corrections

`vcov` defaults to "iid" (classical) and `small` defaults to FALSE, matching Stata's ivreg2 defaults. With `small = FALSE`, inference uses z and chi-squared statistics and the large-sample variance normalization. Set `small = TRUE` for t and F statistics and the finite-sample $N - K$ correction, matching Stata's `small` option. `small` is the single, uniform control for finite-sample corrections across every VCE type ("iid", "robust", "HAC", "AC"); it is not a VCE selector. There are no "HC0"/"HC1" values — use `vcov = "robust"` and toggle `small`.

Three statistics sit outside that uniform control, all matching Stata: the first-stage F statistics (`first_stage()`) always use small-sample degrees of freedom regardless of `small`; the overidentification (Sargan/Hansen J) and endogeneity (C-statistic) tests are never small-sample corrected, even when `small = TRUE`; and the stored reduced-form variance-covariance matrix (`reduced_form`) always applies OLS-style small-sample scaling regardless of the main model's `small`.

Sigma (RMSE) normalization

Under `small = FALSE`, $\sigma^2 = RSS/(N - dof_{minus})$. Under `small = TRUE`, $\sigma^2 = RSS/(N - K - dof_{minus} - sdof_{minus})$. This differs from `lm()`, which always divides the residual sum of squares by its residual degrees of freedom. Because ivreg2r reports the large-sample sigma by default, its RMSE and any quantity built from sigma differ from `lm()` unless `small = TRUE`.

Standard errors relative to lm()

Coefficients match `lm()` for OLS. Standard errors and the full variance-covariance matrix match `lm()` only when `small = TRUE`; the default `small = FALSE` reports Stata-convention standard errors without the $N - K$ finite-sample correction. Under analytic ("aweight") or probability ("pweight") weights, sigma additionally differs from `lm(..., weights = w)` by a factor of $\sqrt{N / \sum(w)}$, because `lm()` uses raw (unnormalized) weights while ivreg2r normalizes them (see below).

Weight normalization

Analytic and probability weights are normalized internally to sum to N , following Stata. This makes sigma scale-invariant for those types: multiplying every weight by a constant leaves all coefficients, standard errors, and test statistics unchanged. Frequency ("fweight") and importance ("iweight") weights are *not* normalized; they redefine N as $\sum(\text{weights})$, so their scale is meaningful by construction. Probability weights force a robust VCE. See the `weight_type` argument of `ivreg2()`.

Degrees-of-freedom adjustments

`dofminus` is a large-sample adjustment subtracted from N in large-sample variance formulas (for example when fixed effects have been partialled out upstream). `sdofminus` is a small-sample adjustment subtracted from the residual degrees of freedom alongside K (for example the count

of regressors removed by partial). Both mirror Stata's `dofminus()` and `sdofofminus()` options. Stata's own source labels `dofminus` a large-sample adjustment (e.g. number of fixed effects) and `sdofofminus` a small-sample one (e.g. number of partialled-out regressors).

Confidence intervals and test statistics

`confint()` and the model Wald test use the standard normal and chi-squared distributions when `small = FALSE`, and the t and F distributions when `small = TRUE`. The reported model test is therefore a chi-squared Wald statistic by default and an F statistic under `small = TRUE`.

See Also

`ivreg2()` for the estimation function and the full argument list; [ivreg2r-glossary](#) for the diagnostic-output acronyms.

Other ivreg2r reference: [ivreg2r-glossary](#)

ivreg2r-glossary

Glossary of diagnostic-output acronyms

Description

`summary()` and `print()` report a battery of identification, weak-instrument, and overidentification diagnostics, many named by acronym. This topic glosses every acronym that reaches the console. The underlying statistics are stored on the fitted object under `fit$diagnostics`; the headline ones are also in `glance()`.

Identification and weak instruments

Underidentification test Tests whether the excluded instruments have rank high enough to identify the endogenous regressors. Reported as the Anderson (1951) canonical-correlations LM statistic under `vcov = "iid"`, and the Kleibergen-Paap rk LM statistic under robust, cluster, or HAC errors.

Weak identification test Gauges instrument strength against the Stock-Yogo (2005) critical values. Reported as the Cragg-Donald (1993) Wald F under `vcov = "iid"`, and the Kleibergen-Paap rk Wald F under robust, cluster, or HAC errors.

Cragg-Donald (CD) The iid weak-identification Wald F (and the underlying eigenvalue statistic); see "Weak identification test".

Kleibergen-Paap (KP) rk Rank-based LM and Wald statistics that replace the Anderson and Cragg-Donald statistics when errors are not iid.

Anderson-Rubin (AR) A weak-instrument-robust test of the joint significance of the endogenous regressors: valid regardless of instrument strength.

Stock-Wright (S) A weak-instrument-robust statistic for the same hypothesis, the score (LM) counterpart of the Anderson-Rubin Wald test. It is distinct from the Hansen J overidentification statistic: S constrains the endogenous coefficients to zero and has degrees of freedom equal to the number of excluded instruments.

First-stage diagnostics

Shea partial R-squared A partial R-squared for the first stage that accounts for correlation among instruments; the relevant measure with more than one endogenous regressor.

SW F (Sanderson-Windmeijer) A conditional first-stage F statistic for the identification of an individual endogenous regressor, given the others.

AP (Angrist-Pischke) A conditional first-stage F (or chi-squared, depending on the VCE) for an individual endogenous regressor.

Overidentification and endogeneity

Sargan / Hansen J The overidentification test: Sargan (1958) under `vcov = "iid"`, and the Hansen J statistic under robust, cluster, or HAC errors. Tests the joint validity of the overidentifying restrictions.

C-statistic (C-stat) A difference-of-J statistic used for the endogeneity test (`endog`), the orthogonality test (`orthog`), and the redundancy test (`redundant`).

VCE and estimator labels

Kiefer The Kiefer (1980) VCE: autocorrelation-consistent with a Truncated kernel at bandwidth equal to the full time span.

Driscoll-Kraay (DK) The Driscoll-Kraay (1998) panel VCE, robust to cross-sectional dependence.

dofminus / sdofminus Large- and small-sample degrees-of-freedom adjustments; see [ivreg2r-conventions](#).

psd0 / psda Positive-semidefinite corrections for the moment covariance matrix, zeroing (`psd0`) or taking the absolute value of (`psda`) any negative eigenvalues.

COVIV The "covariance at the IV estimates" variance matrix for LIML and k-class estimators, robust to misspecification of the LIML model.

HOLS Cragg's (1983) heteroskedastic OLS estimator, obtained from a no-endogenous-regressor model under `method = "gmm2s"` with a robust VCE.

k-class The family of IV estimators indexed by a scalar k (2SLS at $k = 1$, OLS at $k = 0$, LIML at k equal to the LIML eigenvalue).

See Also

[summary.ivreg2\(\)](#) for the console output; [ivreg2r-conventions](#) for the statistical conventions; [ivreg2\(\)](#) for the arguments that request each test.

Other `ivreg2r` reference: [ivreg2r-conventions](#)

klein

*Klein (1950) Model I Macroeconomic Data***Description**

Annual U.S. national-accounts aggregates for 1920–1941, from Klein’s Model I of the U.S. economy: consumption, profits, wages, investment, capital stock, government spending, and taxes. Klein, L.R. (1950). *Economic Fluctuations in the United States, 1921–1941*. Cowles Commission Monograph No. 11. Wiley.

Usage

klein

Format

A data frame with 22 observations and 12 variables:

yr Calendar year (1920–1941). The time variable: use as `tvar` with `l()`/`d()`.

consump Consumption.

profits Private profits.

wagepriv Private wage bill.

invest Investment.

capital1 Lagged value of capital stock (a primitive; no contemporaneous capital column exists to derive it from).

totinc Total income/demand.

wagegovt Government wage bill.

govt Government spending.

taxnetx Indirect business taxes plus net exports.

wagetot Total U.S. wage bill.

year Calendar year minus 1931 (a linear trend, -11 to 10), used as an instrument – not the time variable (see `yr`).

Details

Two columns both encode "year" – do not confuse them. `yr` is the calendar year (1920–1941) and is the time variable to pass as `tvar` when using the time-series operators `l()` / `d()`. `year` is calendar year minus 1931 (a linear trend ranging from -11 to 10) and is used as an *instrument* (a trend term) in the help-file examples below – it is not a second copy of the time index.

The upstream Stata dataset (webuse `klein`) also ships two precomputed one-period lags, `profits1` (= `L.profits`) and `totinc1` (= `L.totinc`). Those columns are deliberately dropped here: this package’s `l()` time-series operator computes the same lags directly from `profits` and `totinc` given `tvar = "yr"`, so shipping precomputed lag columns would be redundant and could drift out of sync with `l()`. `capital1` (the lagged capital stock) is kept because it is a primitive in this dataset – there is no contemporaneous capital column to lag it from – and it appears directly in the instrument lists below.

Source

Klein, L.R. (1950). *Economic Fluctuations in the United States, 1921–1941*. Cowles Commission Monograph No. 11. New York: Wiley.

Distributed via Stata's webuse klein.

Redistribution basis: `system.file("COPYRIGHTS", package = "ivreg2r")`.

See Also

Other ivreg2r datasets: [abdata](#), [card](#), [cigar](#), [griliches](#), [grunfeld](#), [mroz](#), [nlswork](#), [phillips](#), [stockwatson](#), [wagepan](#)

Examples

```
data(klein)

# ivreg2 help file line 1462: LIML, consumption on lagged profits, with
# profits and wagetot treated as endogenous
fit <- ivreg2(
  consump ~ l(profits, 1) | profits + wagetot |
  govt + taxnetx + year + wagegovt + capital1 + l(totinc, 1),
  data = klein, tvar = "yr", method = "liml"
)
summary(fit)
```

model.matrix.ivreg2	<i>Extract design matrices from an ivreg2 model</i>
---------------------	---

Description

Returns the regressor matrix (X), instrument matrix (Z), or projected regressors ($X_{\text{hat}} = P_Z X$), as used in estimation. For models estimated with `partial`, these are the post-partialling matrices that `coef()`, `residuals()`, and `vcov()` correspond to. The returned matrices do not depend on whether the model was fitted with `x = TRUE`: when the matrices were not stored, they are reconstructed from the model frame and the same partialling projection is re-applied.

Usage

```
## S3 method for class 'ivreg2'
model.matrix(
  object,
  component = c("regressors", "projected", "instruments"),
  ...
)
```

Arguments

<code>object</code>	An object of class "ivreg2".
<code>component</code>	Character: which matrix to return. "regressors" (default) returns the regressor matrix X; "instruments" returns the full instrument matrix Z (NULL for OLS); "projected" returns the projected regressors X_{hat} (equals X for OLS).
<code>...</code>	Additional arguments (ignored).

Value

A numeric matrix, or NULL if `component = "instruments"` for an OLS model.

See Also

`ivreg2()`

Other `ivreg2` methods: `coef.ivreg2()`, `confint.ivreg2()`, `diagnostics()`, `first_stage()`, `fitted.ivreg2()`, `formula.ivreg2()`, `ivreg2()`, `nobs.ivreg2()`, `predict.ivreg2()`, `print.ivreg2()`, `print.summary.ivreg2()`, `residuals.ivreg2()`, `summary.ivreg2()`, `terms.ivreg2()`, `update.ivreg2()`, `vcov.ivreg2()`

mroz

Mroz (1987) Female Labor Supply Dataset

Description

Cross-sectional data on 753 married white women from the Panel Study of Income Dynamics (PSID), 1975. Of these, 428 were working (`inlf == 1`) with observed wages. Used by Mroz (1987) to study married women's labor force participation and hours of work. A classic IV application instruments education with parental education (`motheduc`, `fatheduc`).

Usage

```
mroz
```

Format

A data frame with 753 observations and 22 variables:

inlf In the labor force in 1975 (binary).

hours Hours worked in 1975.

kidslt6 Number of children younger than 6.

kidsge6 Number of children aged 6–18.

age Age in years.

educ Years of education.

wage Estimated hourly wage (1975 dollars).

repwage Reported hourly wage at interview (1976).
hushrs Husband's hours worked in 1975.
husage Husband's age.
huseduc Husband's years of education.
huswage Husband's hourly wage (1975 dollars).
faminc Family income (1975 dollars).
mtr Federal marginal tax rate facing the woman.
motheduc Mother's years of education.
fatheduc Father's years of education.
unem Unemployment rate in county of residence.
city Lives in SMSA (binary).
exper Years of labor market experience.
nwifeinc Non-wife household income (faminc - wage * hours, in thousands of 1975 dollars).
lwage Log estimated hourly wage.
expersq Experience squared (exper^2).

Source

Mroz, T.A. (1987). "The Sensitivity of an Empirical Model of Married Women's Hours of Work to Economic and Statistical Assumptions." *Econometrica*, 55(4), 765–799.

Obtained from Stata's bcuse archive (Boston College), bcuse mroz.

Redistribution basis: `system.file("COPYRIGHTS", package = "ivreg2r")`.

See Also

Other ivreg2r datasets: [abdata](#), [card](#), [cigar](#), [griliches](#), [grunfeld](#), [klein](#), [nlswork](#), [phillips](#), [stockwatson](#), [wagepan](#)

Examples

```
data(mroz)
# Restrict to working women (observed wages)
mroz_work <- subset(mroz, inlf == 1)
# IV regression: instrument education with parental education (Example 15.5,
# "Return to Education for Working Women", in Wooldridge (2020),
# Introductory Econometrics).
# These instruments (both parents' education) differ from the specification
# in ?ivreg2, which instruments educ with fatheduc alone (just-identified)
# and with age + kidslt6 + kidsge6 (overidentified).
fit <- ivreg2(lwage ~ exper + expersq | educ | motheduc + fatheduc,
              data = mroz_work)
summary(fit)
```

nlswork

*NLS Young Women Panel (Stata [XT] Manual Extract)***Description**

Panel data from the U.S. Bureau of Labor Statistics' National Longitudinal Survey of Young Women, 14–24 years of age in 1968, interviewed in survey years 1968–1988. This is the extract distributed with Stata's [XT] manual (webuse nlswork): 28,534 person-year observations.

Usage

nlswork

Format

A data frame with 28,534 observations and 21 variables:

idcode Person identifier. Use as `ivar`.

year Interview year (two-digit, e.g. 70 = 1970). The time variable: use as `tvar`.

birth_yr Birth year (two-digit).

age Age in current year.

race Race: 1 = White, 2 = Black, 3 = Other.

msp 1 if married, spouse present.

nev_mar 1 if never married.

grade Current grade completed.

collgrad 1 if college graduate.

not_smsa 1 if not in an SMSA (metropolitan area).

c_city 1 if central city.

south 1 if in the South.

ind_code Industry of employment (code).

occ_code Occupation (code).

union 1 if union member.

wks_ue Weeks unemployed, last year.

ttl_exp Total work experience (years).

tenure Job tenure (years).

hours Usual hours worked.

wks_work Weeks worked, last year.

ln_wage Log wage (deflated by the GNP deflator).

Details

`race` is coded 1 = White, 2 = Black, 3 = Other (faithful to the upstream numeric coding; not converted to a factor). Many columns contain missing values, consistent with the source survey.

Source

U.S. Bureau of Labor Statistics. National Longitudinal Survey of Young Women, 14–24 years old in 1968. Center for Human Resource Research.

Distributed via Stata's webuse nlswork.

Redistribution basis: `system.file("COPYRIGHTS", package = "ivreg2r")`.

See Also

Other ivreg2r datasets: [abdata](#), [card](#), [cigar](#), [griliches](#), [grunfeld](#), [klein](#), [mroz](#), [phillips](#), [stockwatson](#), [wagepan](#)

Examples

```
data(nlswork)

# ivreg2 help file line 1618: one-way cluster on person id
fit <- ivreg2(ln_wage ~ grade + age + ttl_exp + tenure, data = nlswork,
             clusters = ~idcode)
summary(fit)

# ivreg2 help file line 1628: two-way cluster on person id and year
fit2 <- ivreg2(ln_wage ~ grade + age + ttl_exp + tenure, data = nlswork,
              clusters = ~idcode + year)
summary(fit2)
```

nobs.ivreg2

Extract number of observations from an ivreg2 object

Description

Extract number of observations from an ivreg2 object

Usage

```
## S3 method for class 'ivreg2'
nobs(object, ...)
```

Arguments

`object` An object of class "ivreg2".

`...` Additional arguments (ignored).

Value

Integer: number of observations.

See Also`ivreg2()`

Other `ivreg2` methods: `coef.ivreg2()`, `confint.ivreg2()`, `diagnostics()`, `first_stage()`, `fitted.ivreg2()`, `formula.ivreg2()`, `ivreg2()`, `model.matrix.ivreg2()`, `predict.ivreg2()`, `print.ivreg2()`, `print.summary.ivreg2()`, `residuals.ivreg2()`, `summary.ivreg2()`, `terms.ivreg2()`, `update.ivreg2()`, `vcov.ivreg2()`

phillips

*Phillips Curve Macroeconomic Dataset***Description**

U.S. macroeconomic time series (1948–1996) with inflation and unemployment rates plus lags and first differences. This is the same dataset used in Stata's `ivreg2` help file for HAC and AC examples, originally from Wooldridge (2020).

Usage`phillips`**Format**

A data frame with 49 observations and 11 variables:

year Calendar year (1948–1996).

inf Inflation rate (percent).

unem Unemployment rate (percent).

cinf Change in inflation ($\text{inf} - \text{inf}_1$).

cunem Change in unemployment ($\text{unem} - \text{unem}_1$).

unem_1 Lagged unemployment (one year).

inf_1 Lagged inflation (one year).

unem_2 Lagged unemployment (two years).

inf_2 Lagged inflation (two years).

cinf_1 Lagged change in inflation.

cunem_1 Lagged change in unemployment.

Source

Wooldridge, J.M. (2020). *Introductory Econometrics: A Modern Approach*, 7th ed. Cengage Learning. The underlying series are from the *Economic Report of the President* (a U.S. government work).

Downloaded from <http://fmwww.bc.edu/ec-p/data/wooldridge/phillips.dta>.

Redistribution basis: `system.file("COPYRIGHTS", package = "ivreg2r")`.

See Also

Other ivreg2r datasets: [abdata](#), [card](#), [cigar](#), [griliches](#), [grunfeld](#), [klein](#), [mroz](#), [nlswork](#), [stockwatson](#), [wagepan](#)

Examples

```
data(phillips)
# OLS Phillips curve with AC standard errors: exact replication of the
# Stata ivreg2 help-file example at line 1501 (`ivreg2 cinf unem, bw(3)`;
# Bartlett is Stata's default kernel). Without vcov = "robust", a kernel
# gives AC (autocorrelation-consistent) inference, not HAC.
fit_ac <- ivreg2(cinf ~ unem, data = phillips,
                kernel = "bartlett", bw = 3, tvar = "year")
summary(fit_ac)

# Adding vcov = "robust" gives HAC (Newey-West), replicating the
# help-file example at line 1511.
fit_hac <- ivreg2(cinf ~ unem, data = phillips, vcov = "robust",
                kernel = "bartlett", bw = 3, tvar = "year",
                small = TRUE)
summary(fit_hac)
```

predict.ivreg2	<i>Predict from an ivreg2 model</i>
----------------	-------------------------------------

Description

Predict from an ivreg2 model

Usage

```
## S3 method for class 'ivreg2'
predict(object, newdata, se.fit = FALSE, na.action = stats::na.pass, ...)
```

Arguments

object	An object of class "ivreg2".
newdata	An optional data frame for prediction. If omitted, fitted values from the original data are returned. If the model uses time-series operators (see ts-operators) among the regressors, newdata must contain the time variable (and panel variable, if any) plus the history rows needed to compute the lags; rows whose lags are missing within newdata get NA predictions (matching Stata). Operators confined to the instrument part impose no such requirement — prediction scores only the regressors.
se.fit	Logical: if TRUE, return prediction standard errors alongside fitted values. Standard errors are computed as $\sqrt{\text{diag}(X' V X)}$ where $V = \text{vcov}(\text{object})$, so they reflect the VCE used at estimation time (IID, robust, cluster, HAC, etc.). Not available after partialling (matching Stata's predict, stdp).

na.action	Function for handling NAs in newdata.
...	Additional arguments (ignored).

Value

When `se.fit = FALSE` (default), a numeric vector of predicted values. When `se.fit = TRUE`, a list with components `fit` (predicted values) and `se.fit` (standard errors of prediction).

See Also

[ivreg2\(\)](#)

Other `ivreg2` methods: [coef.ivreg2\(\)](#), [confint.ivreg2\(\)](#), [diagnostics\(\)](#), [first_stage\(\)](#), [fitted.ivreg2\(\)](#), [formula.ivreg2\(\)](#), [ivreg2\(\)](#), [model.matrix.ivreg2\(\)](#), [nobs.ivreg2\(\)](#), [print.ivreg2\(\)](#), [print.summary.ivreg2\(\)](#), [residuals.ivreg2\(\)](#), [summary.ivreg2\(\)](#), [terms.ivreg2\(\)](#), [update.ivreg2\(\)](#), [vcov.ivreg2\(\)](#)

Examples

```
fit <- ivreg2(lwage ~ exper | educ | nearc4, data = card)

# Fitted values on the estimation sample
head(predict(fit))

# Predictions with standard errors for new data
nd <- data.frame(exper = c(5, 10), educ = c(12, 16), nearc4 = c(0, 1))
predict(fit, newdata = nd, se.fit = TRUE)
```

print.ivreg2	<i>Print an ivreg2 object</i>
--------------	-------------------------------

Description

Print an `ivreg2` object

Usage

```
## S3 method for class 'ivreg2'
print(x, digits = max(3L, getOption("digits") - 3L), ...)
```

Arguments

x	An object of class "ivreg2".
digits	Minimum number of significant digits to print.
...	Additional arguments (ignored).

Value

x, invisibly.

See Also

[ivreg2\(\)](#)

Other ivreg2 methods: [coef.ivreg2\(\)](#), [confint.ivreg2\(\)](#), [diagnostics\(\)](#), [first_stage\(\)](#), [fitted.ivreg2\(\)](#), [formula.ivreg2\(\)](#), [ivreg2\(\)](#), [model.matrix.ivreg2\(\)](#), [nobs.ivreg2\(\)](#), [predict.ivreg2\(\)](#), [print.summary.ivreg2\(\)](#), [residuals.ivreg2\(\)](#), [summary.ivreg2\(\)](#), [terms.ivreg2\(\)](#), [update.ivreg2\(\)](#), [vcov.ivreg2\(\)](#)

`print.summary.ivreg2` *Print a summary.ivreg2 object*

Description

Formats and displays the full estimation output, including coefficient table, fit statistics, and IV diagnostic tests.

Usage

```
## S3 method for class 'summary.ivreg2'
print(
  x,
  digits = max(3L, getOption("digits") - 3L),
  signif.stars = getOption("show.signif.stars", TRUE),
  ...
)
```

Arguments

x	An object of class "summary.ivreg2".
digits	Minimum number of significant digits.
signif.stars	Logical: print significance stars? Default TRUE.
...	Additional arguments passed to printCoefmat() .

Value

x, invisibly.

See Also

[ivreg2\(\)](#)

Other ivreg2 methods: [coef.ivreg2\(\)](#), [confint.ivreg2\(\)](#), [diagnostics\(\)](#), [first_stage\(\)](#), [fitted.ivreg2\(\)](#), [formula.ivreg2\(\)](#), [ivreg2\(\)](#), [model.matrix.ivreg2\(\)](#), [nobs.ivreg2\(\)](#), [predict.ivreg2\(\)](#), [print.summary.ivreg2\(\)](#), [residuals.ivreg2\(\)](#), [summary.ivreg2\(\)](#), [terms.ivreg2\(\)](#), [update.ivreg2\(\)](#), [vcov.ivreg2\(\)](#)

residuals.ivreg2	<i>Extract residuals from an ivreg2 object</i>
------------------	--

Description

Respects na.action: when the model was fit with na.exclude, NAs are reinserted at the omitted row positions so the result aligns with the original data frame (matching base R's residuals.lm).

Usage

```
## S3 method for class 'ivreg2'
residuals(object, ...)
```

Arguments

object	An object of class "ivreg2".
...	Additional arguments (ignored).

Value

Numeric vector of residuals.

See Also

[ivreg2\(\)](#)

Other ivreg2 methods: [coef.ivreg2\(\)](#), [confint.ivreg2\(\)](#), [diagnostics\(\)](#), [first_stage\(\)](#), [fitted.ivreg2\(\)](#), [formula.ivreg2\(\)](#), [ivreg2\(\)](#), [model.matrix.ivreg2\(\)](#), [nobs.ivreg2\(\)](#), [predict.ivreg2\(\)](#), [print.ivreg2\(\)](#), [print.summary.ivreg2\(\)](#), [summary.ivreg2\(\)](#), [terms.ivreg2\(\)](#), [update.ivreg2\(\)](#), [vcov.ivreg2\(\)](#)

stockwatson	<i>Stock and Watson U.S. Macroeconomic Dataset</i>
-------------	--

Description

Quarterly U.S. macroeconomic time series (1959Q1–2000Q4) used in Baum, Schaffer & Stillman (2007) to illustrate HAC standard errors, two-step GMM, and CUE estimation. Original source: Stock and Watson (2003).

Usage

```
stockwatson
```

Format

A data frame with 168 observations and 17 variables:

year Calendar year.

quarter Quarter (1–4).

date Stata quarterly date (0 = 1960Q1).

UR Unemployment rate (percent).

CPI Consumer Price Index.

FFIR Federal funds interest rate.

TBILL Treasury bill rate.

TBON Treasury bond rate.

ER Trade-weighted exchange rate.

GDP Nominal GDP (billions of dollars).

inf Annualized CPI inflation: $100 * \log(\text{CPI} / \text{L4.CPI})$. NA for the first 4 observations.

ggdp Annualized GDP growth: $100 * \log(\text{GDP} / \text{L4.GDP})$. NA for the first 4 observations.

dinf First difference of inflation ($\text{inf} - \text{L.inf}$). NA for the first 5 observations.

ggdp_2 Second lag of GDP growth (L2.ggdp).

TBILL_1 First lag of Treasury bill rate (L.TBILL).

ER_1 First lag of exchange rate (L.ER).

TBON_1 First lag of Treasury bond rate (L.TBON).

Details

The derived variables `inf`, `ggdp`, `dinf`, and the lagged instrument columns are pre-computed following Baum, Schaffer & Stillman (2007, p. 474) so that users can replicate their examples directly without time-series operators.

Source

Stock, J.H. and Watson, M.W. (2003). *Introduction to Econometrics*. Addison-Wesley.

Downloaded from <http://fmwww.bc.edu/ec-p/data/stockwatson/macrodta.dta>.

Redistribution basis: `system.file("COPYRIGHTS", package = "ivreg2r")`.

References

Baum, C.F., Schaffer, M.E. and Stillman, S. (2007). Enhanced routines for instrumental variables/generalized method of moments estimation and testing. *The Stata Journal*, 7(4), 465–506.

See Also

Other `ivreg2r` datasets: [abdata](#), [card](#), [cigar](#), [griliches](#), [grunfeld](#), [klein](#), [mroz](#), [nlswork](#), [phillips](#), [wagepan](#)

Examples

```
data(stockwatson)

# Baum, Schaffer & Stillman (2007), p. 475: 2SLS with IID errors
fit_iv <- ivreg2(dinf ~ 1 | UR | ggdp_2 + TBILL_1 + ER_1 + TBON_1,
               data = stockwatson)
summary(fit_iv)

# Baum, Schaffer & Stillman (2007), p. 476: two-step GMM with HAC (Bartlett, bw=5)
fit_gmm <- ivreg2(dinf ~ 1 | UR | ggdp_2 + TBILL_1 + ER_1 + TBON_1,
               data = stockwatson, method = "gmm2s",
               vcov = "robust", kernel = "bartlett", bw = 5,
               tvar = "date")
summary(fit_gmm)
```

summary.ivreg2	<i>Summary for ivreg2 objects</i>
----------------	-----------------------------------

Description

Builds a coefficient table (estimates, standard errors, test statistics, p-values) and collects model diagnostics for display by [print.summary.ivreg2\(\)](#).

Usage

```
## S3 method for class 'ivreg2'
summary(object, ...)
```

Arguments

object	An object of class "ivreg2".
...	Additional arguments (ignored).

Value

An object of class "summary.ivreg2".

See Also

[ivreg2\(\)](#); [ivreg2r-glossary](#) for the diagnostic acronyms printed here; [ivreg2r-conventions](#) for the statistical conventions.

Other ivreg2 methods: [coef.ivreg2\(\)](#), [confint.ivreg2\(\)](#), [diagnostics\(\)](#), [first_stage\(\)](#), [fitted.ivreg2\(\)](#), [formula.ivreg2\(\)](#), [ivreg2\(\)](#), [model.matrix.ivreg2\(\)](#), [nobs.ivreg2\(\)](#), [predict.ivreg2\(\)](#), [print.ivreg2\(\)](#), [print.summary.ivreg2\(\)](#), [residuals.ivreg2\(\)](#), [terms.ivreg2\(\)](#), [update.ivreg2\(\)](#), [vcov.ivreg2\(\)](#)

terms.ivreg2	<i>Extract terms object from an ivreg2 model</i>
--------------	--

Description

Extract terms object from an ivreg2 model

Usage

```
## S3 method for class 'ivreg2'
terms(x, component = c("regressors", "instruments", "full"), ...)
```

Arguments

x	An object of class "ivreg2".
component	Character: which terms object to return. "regressors" (default) returns terms for all regressors (exogenous + endogenous); "instruments" returns terms for excluded instruments (NULL for OLS); "full" returns terms for the complete formula.
...	Additional arguments (ignored).

Value

A [terms](#) object, or NULL if component = "instruments" for an OLS model.

See Also

[ivreg2\(\)](#)

Other ivreg2 methods: [coef.ivreg2\(\)](#), [confint.ivreg2\(\)](#), [diagnostics\(\)](#), [first_stage\(\)](#), [fitted.ivreg2\(\)](#), [formula.ivreg2\(\)](#), [ivreg2\(\)](#), [model.matrix.ivreg2\(\)](#), [nobs.ivreg2\(\)](#), [predict.ivreg2\(\)](#), [print.ivreg2\(\)](#), [print.summary.ivreg2\(\)](#), [residuals.ivreg2\(\)](#), [summary.ivreg2\(\)](#), [update.ivreg2\(\)](#), [vcov.ivreg2\(\)](#)

tidy.ivreg2	<i>Tidy an ivreg2 object</i>
-------------	------------------------------

Description

Constructs a tibble summarizing coefficient estimates, standard errors, test statistics, and p-values.

Usage

```
## S3 method for class 'ivreg2'
tidy(x, conf.int = TRUE, conf.level = 0.95, exponentiate = FALSE, ...)
```

Arguments

<code>x</code>	An object of class "ivreg2".
<code>conf.int</code>	Logical: include confidence intervals? Default TRUE.
<code>conf.level</code>	Confidence level for intervals. Default 0.95.
<code>exponentiate</code>	Logical: exponentiate the coefficient estimates and confidence interval bounds? Default FALSE. Useful for log-linear models where <code>exp(estimate)</code> gives a multiplicative effect. Standard errors remain on the original (log) scale, following broom convention.
<code>...</code>	Additional arguments (ignored).

Value

A `tibble::tibble()` with columns `term`, `estimate`, `std.error`, `statistic`, `p.value`, and optionally `conf.low`, `conf.high`.

See Also

[ivreg2\(\)](#)

Other broom methods: [augment.ivreg2\(\)](#), [glance.ivreg2\(\)](#)

Examples

```
data(mroz)
mroz_work <- subset(mroz, inlf == 1)
fit <- ivreg2(lwage ~ exper + expersq | educ | age + kidslt6 + kidsge6,
             data = mroz_work, vcov = "robust")

tidy(fit)
tidy(fit, conf.int = FALSE)
tidy(fit, exponentiate = TRUE)

# Compare 2SLS and LIML side-by-side on the help-file baseline spec
# (weak first stage; see the LIML example in ?ivreg2 for the framing)
fit_liml <- ivreg2(lwage ~ exper + expersq | educ |
                  age + kidslt6 + kidsge6,
                  data = mroz_work, method = "liml")
comparison <- rbind(
  cbind(method = "2SLS", tidy(fit)),
  cbind(method = "LIML", tidy(fit_liml))
)
comparison[comparison$term == "educ", c("method", "estimate", "std.error")]
```

ts-operators	<i>Time-series operators for ivreg2 formulas</i>
--------------	--

Description

`l()` (lag) and `d()` (difference) may be used inside an `ivreg2()` formula to create lags and differences of a variable on the fly, resolved against the time variable `tvar` (and panel variable `ivar` for panel data). They mirror Stata's `tsset`-aware `L.` and `D.` operators:

Usage

`l(x, k = 1)`

`d(x, k = 1)`

Arguments

- `x` A numeric variable in the model data.
- `k` Lag order(s) for `l()`; difference order(s) for `d()`. Non-negative integer (vector allowed, expanded to one term per element).

Details

R syntax	Meaning	Stata equivalent
<code>l(x), l(x, 1)</code>	lag 1	<code>L . x</code>
<code>l(x, 1:3)</code>	lags 1, 2, 3 (three terms)	<code>L (1/3) . x</code>
<code>l(x, 0)</code>	<code>x</code> itself	<code>L 0 . x</code>
<code>d(x), d(x, 1)</code>	first difference $x - l(x, 1)$	<code>D . x</code>
<code>d(x, 2)</code>	second-order difference $x - 2\ l(x, 1) + l(x, 2)$	<code>D2 . x</code>

Value

When used inside an `ivreg2()` formula: a numeric vector aligned to the model data. Calling these functions anywhere else is an error.

Lag semantics

Lags are computed by time **value**, not row position: `l(x, k)` at time `t` is the value of `x` at time `t - k` within the same panel unit. If no observation exists at `t - k` (start of the series, or a gap in the time variable), the lag is NA and the row is dropped from the estimation sample — exactly Stata's behavior. Lags never cross panel-unit boundaries. The data does not need to be sorted.

Differences from fixest

Following Stata (not fixest), $d(x, k)$ is the k -th **order** (iterated) difference $(1-L)^k x$, matching Stata's $Dk.x$; fixest's $d(x, k)$ is the distance- k difference $x - 1(x, k)$. The two agree for $k = 1$.

Usage constraints

Operators require `tvar` (and `ivar` for panel data) to be passed to `ivreg2()`; `l()` and `d()` are not standalone lag utilities and signal an error when called outside an `ivreg2()` formula (use, e.g., `dplyr::lag()` with a consecutive-periods check, or `fixest/collapse/plm`, for general-purpose panel lagging). The time variable must be integer-valued (greater than -2^{53} , at most 2^{53}), like Stata's `tsset` — lag arithmetic on fractional or astronomically large time values is not exact in floating point. The lag/difference order k must be a constant non-negative integer (or integer vector for ranges); leads (negative k), nested operators (`l(d(x))`), and ranges on the left-hand side (the model has a single dependent variable) are not supported. When operators appear among the *regressors*, `predict(fit, newdata)` recomputes them within `newdata`, so `newdata` must contain the history rows needed for the lags; rows whose lags are missing get NA predictions (again matching Stata). Operators confined to the instrument part impose no requirement on `newdata`.

See Also

[ivreg2\(\)](#)

Examples

```
data(phillips)
# Stata: ivreg2 cinf (unem = l(1/3).unem), bw(3)
fit <- ivreg2(cinf ~ 1 | unem | l(unem, 1:3), data = phillips,
             tvar = "year", vcov = "AC", kernel = "bartlett", bw = 3)
coef(fit)
```

update.ivreg2

Update and re-fit an ivreg2 model

Description

Updates the formula and/or other arguments of an `ivreg2` call and (optionally) re-fits the model.

Usage

```
## S3 method for class 'ivreg2'
update(object, formula., ..., evaluate = TRUE)
```

Arguments

object	An object of class "ivreg2".
formula.	A formula to update the model formula (see update.formula). Multi-part formula updates are supported.
...	Additional arguments to update in the call (e.g., vcov = "robust", data = new_data).
evaluate	Logical: if TRUE (default), evaluate the updated call; if FALSE, return the unevaluated call.

Value

If evaluate = TRUE, a new ivreg2 object. If evaluate = FALSE, the unevaluated call.

See Also

[ivreg2\(\)](#)

Other ivreg2 methods: [coef.ivreg2\(\)](#), [confint.ivreg2\(\)](#), [diagnostics\(\)](#), [first_stage\(\)](#), [fitted.ivreg2\(\)](#), [formula.ivreg2\(\)](#), [ivreg2\(\)](#), [model.matrix.ivreg2\(\)](#), [nobs.ivreg2\(\)](#), [predict.ivreg2\(\)](#), [print.ivreg2\(\)](#), [print.summary.ivreg2\(\)](#), [residuals.ivreg2\(\)](#), [summary.ivreg2\(\)](#), [terms.ivreg2\(\)](#), [vcov.ivreg2\(\)](#)

vcov.ivreg2

Extract variance-covariance matrix from an ivreg2 object

Description

Extract variance-covariance matrix from an ivreg2 object

Usage

```
## S3 method for class 'ivreg2'
vcov(object, ...)
```

Arguments

object	An object of class "ivreg2".
...	Additional arguments (ignored).

Value

The variance-covariance matrix of the coefficient estimates.

See Also[ivreg2\(\)](#)

Other `ivreg2` methods: [coef.ivreg2\(\)](#), [confint.ivreg2\(\)](#), [diagnostics\(\)](#), [first_stage\(\)](#), [fitted.ivreg2\(\)](#), [formula.ivreg2\(\)](#), [ivreg2\(\)](#), [model.matrix.ivreg2\(\)](#), [nobs.ivreg2\(\)](#), [predict.ivreg2\(\)](#), [print.ivreg2\(\)](#), [print.summary.ivreg2\(\)](#), [residuals.ivreg2\(\)](#), [summary.ivreg2\(\)](#), [terms.ivreg2\(\)](#), [update.ivreg2\(\)](#)

wagepan

*Wooldridge Wage Panel Dataset***Description**

Balanced panel data from the National Longitudinal Survey of Youth (NLSY), 1980–1987. Contains 4,360 observations on 545 young men observed over 8 years (the panel is balanced: every man contributes exactly one observation per year). Useful for demonstrating panel methods, including two-way clustering by individual and year.

Usage

wagepan

Format

A data frame with 4,360 observations and 11 variables:

nr Person identifier.

year Calendar year (1980–1987).

lwage Log hourly wage.

educ Years of education.

black Black (binary).

hisp Hispanic (binary).

exper Years of labor market experience.

expersq Experience squared (exper^2).

married Married (binary).

union Union member (binary).

hours Annual hours worked.

Details

The bundled data are in raw levels, not within-transformed. Fixed-effects use (e.g. the Stock-Watson panel-robust VCE, `sw = TRUE`) requires demeaning each variable by panel unit before fitting; see the worked example in the "Fixed-effects panels: the Stock-Watson correction" section of `vignette("time-series-gmm")`.

Source

Vella, F. and Verbeek, M. (1998). "Whose Wages Do Unions Raise? A Dynamic Model of Unionism and Wage Rate Determination for Young Men." *Journal of Applied Econometrics*, 13(2), 163–183.

Wooldridge, J.M. (2010). *Econometric Analysis of Cross Section and Panel Data*, 2nd ed. MIT Press.

Obtained from Stata's bcuse archive (Boston College), bcuse wagepan.

Redistribution basis: `system.file("COPYRIGHTS", package = "ivreg2r")`.

See Also

Other ivreg2r datasets: [abdata](#), [card](#), [cigar](#), [griliches](#), [grunfeld](#), [klein](#), [mroz](#), [nlswork](#), [phillips](#), [stockwatson](#)

Examples

```
data(wagepan)
# OLS with two-way clustering by person and year
# The canonical two-way clustering example from the ivreg2 help file lives
# in ?nlswork; this block illustrates the same VCE on a different panel.
fit <- ivreg2(lwage ~ educ + black + hisp + exper + expersq + married + union,
             data = wagepan, clusters = ~ nr + year, small = TRUE)
summary(fit)
```

Index

- * **broom methods**
 - augment.ivreg2, 4
 - glance.ivreg2, 15
 - tidy.ivreg2, 46
 - * **datasets**
 - abdata, 3
 - card, 5
 - cigar, 7
 - griliches, 17
 - grunfeld, 18
 - klein, 33
 - mroz, 35
 - nlswork, 37
 - phillips, 39
 - stockwatson, 43
 - wagepan, 51
 - * **ivreg2 methods**
 - coef.ivreg2, 8
 - confint.ivreg2, 9
 - diagnostics, 10
 - first_stage, 12
 - fitted.ivreg2, 13
 - formula.ivreg2, 14
 - ivreg2, 19
 - model.matrix.ivreg2, 34
 - nobs.ivreg2, 38
 - predict.ivreg2, 40
 - print.ivreg2, 41
 - print.summary.ivreg2, 42
 - residuals.ivreg2, 43
 - summary.ivreg2, 45
 - terms.ivreg2, 46
 - update.ivreg2, 49
 - vcov.ivreg2, 50
 - * **ivreg2r datasets**
 - abdata, 3
 - card, 5
 - cigar, 7
 - griliches, 17
 - grunfeld, 18
 - klein, 33
 - mroz, 35
 - nlswork, 37
 - phillips, 39
 - stockwatson, 43
 - wagepan, 51
 - * **ivreg2r reference**
 - ivreg2r-conventions, 29
 - ivreg2r-glossary, 31
- abdata, 3, 7, 8, 18, 19, 34, 36, 38, 40, 44, 52
- augment.ivreg2, 4
- augment.ivreg2(), 16, 28, 47
- card, 4, 5, 8, 18, 19, 34, 36, 38, 40, 44, 52
- cigar, 4, 7, 7, 18, 19, 34, 36, 38, 40, 44, 52
- coef(), 12
- coef.ivreg2, 8
- coef.ivreg2(), 10, 11, 13, 14, 28, 35, 39, 41–43, 45, 46, 50, 51
- confint(), 12, 31
- confint.ivreg2, 9
- confint.ivreg2(), 9, 11, 13, 14, 28, 35, 39, 41–43, 45, 46, 50, 51
- d(ts-operators), 48
- diagnostics, 10, 15
- diagnostics(), 9, 10, 13, 14, 28, 35, 39, 41–43, 45, 46, 50, 51
- first_stage, 12
- first_stage(), 9–11, 14, 26, 28, 30, 35, 39, 41–43, 45, 46, 50, 51
- fitted(), 12, 21
- fitted.ivreg2, 13
- fitted.ivreg2(), 9–11, 13, 14, 28, 35, 39, 41–43, 45, 46, 50, 51
- formula.ivreg2, 14
- formula.ivreg2(), 9–11, 13, 14, 28, 35, 39, 41–43, 45, 46, 50, 51

- Formula::formula.Formula(), [14](#)
- glance(), [12, 31](#)
- glance.ivreg2, [15](#)
- glance.ivreg2(), [5, 11, 28, 47](#)
- griliches, [4, 7, 8, 17, 19, 34, 36, 38, 40, 44, 52](#)
- grunfeld, [4, 7, 8, 18, 18, 34, 36, 38, 40, 44, 52](#)
- ivreg2, [19](#)
- ivreg2(), [5, 9–11, 13, 14, 16, 30–32, 35, 39, 41–43, 45–51](#)
- ivreg2r-conventions, [10–12, 21, 26, 28, 29, 32, 45](#)
- ivreg2r-glossary, [28, 30, 31, 31, 45](#)
- klein, [4, 7, 8, 18, 19, 33, 36, 38, 40, 44, 52](#)
- l (ts-operators), [48](#)
- lm(), [21, 28, 30](#)
- model.matrix.ivreg2, [34](#)
- model.matrix.ivreg2(), [9–11, 13, 14, 28, 39, 41–43, 45, 46, 50, 51](#)
- mroz, [4, 7, 8, 18, 19, 34, 35, 38, 40, 44, 52](#)
- na.exclude, [21](#)
- na.omit, [21](#)
- nlswork, [4, 7, 8, 18, 19, 34, 36, 37, 40, 44, 52](#)
- nobs(), [12](#)
- nobs.ivreg2, [38](#)
- nobs.ivreg2(), [9–11, 13, 14, 28, 35, 41–43, 45, 46, 50, 51](#)
- phillips, [4, 7, 8, 18, 19, 34, 36, 38, 39, 44, 52](#)
- predict(), [4, 21](#)
- predict.ivreg2, [40](#)
- predict.ivreg2(), [9–11, 13, 14, 28, 35, 39, 42, 43, 45, 46, 50, 51](#)
- print(), [12, 31](#)
- print.ivreg2, [41](#)
- print.ivreg2(), [9–11, 13, 14, 28, 35, 39, 41–43, 45, 46, 50, 51](#)
- print.summary.ivreg2, [42](#)
- print.summary.ivreg2(), [9–11, 13, 14, 28, 35, 39, 41–43, 45, 46, 50, 51](#)
- printCoefmat(), [42](#)
- residuals(), [5, 12, 21](#)
- residuals.ivreg2, [43](#)
- residuals.ivreg2(), [9–11, 13, 14, 28, 35, 39, 41, 42, 45, 46, 50, 51](#)
- stockwatson, [4, 7, 8, 18, 19, 34, 36, 38, 40, 43, 52](#)
- summary(), [12, 30, 31](#)
- summary.ivreg2, [45](#)
- summary.ivreg2(), [9–11, 13, 14, 28, 32, 35, 39, 41–43, 46, 50, 51](#)
- terms, [46](#)
- terms.ivreg2, [46](#)
- terms.ivreg2(), [9–11, 13, 14, 28, 35, 39, 41–43, 45, 50, 51](#)
- tibble::tibble(), [5, 10, 15, 47](#)
- tidy(), [12](#)
- tidy.ivreg2, [46](#)
- tidy.ivreg2(), [5, 16, 28](#)
- ts-operators, [48](#)
- update.formula, [50](#)
- update.ivreg2, [49](#)
- update.ivreg2(), [9–11, 13, 14, 28, 35, 39, 41–43, 45, 46, 51](#)
- vcov(), [12, 30](#)
- vcov.ivreg2, [50](#)
- vcov.ivreg2(), [9–11, 13, 14, 28, 35, 39, 41–43, 45, 46, 50](#)
- wagepan, [4, 7, 8, 18, 19, 34, 36, 38, 40, 44, 51](#)