

logit — Logistic regression, reporting coefficients

Description	Quick start	Menu	Syntax
Options	Remarks and examples	Stored results	Methods and formulas
References	Also see		

Description

`logit` fits a logit model for a binary response by maximum likelihood; it models the probability of a positive outcome given a set of regressors. *depvar* equal to nonzero and nonmissing (typically *depvar* equal to one) indicates a positive outcome, whereas *depvar* equal to zero indicates a negative outcome.

Also see [\[R\] `logistic`](#); `logistic` displays estimates as odds ratios. Many users prefer the `logistic` command to `logit`. Results are the same regardless of which you use—both are the maximum-likelihood estimator. Several auxiliary commands that can be run after `logit`, `probit`, or `logistic` estimation are described in [\[R\] `logistic postestimation`](#).

Quick start

Logit model of `y` on `x1` and `x2`

```
logit y x1 x2
```

Add indicators for categorical variable `a`

```
logit y x1 x2 i.a
```

With cluster-robust standard errors for clustering by levels of `cvar`

```
logit y x1 x2 i.a, vce(cluster cvar)
```

Save separate coefficient estimates for each level of `cvar` to `myresults.dta`

```
statsby _b, by(cvar) saving(myresults): logit y x1 x2 i.a
```

Adjust for complex survey design using `svyset` data

```
svy: logit y x1 x2 i.a
```

Menu

Statistics > Binary outcomes > Logistic regression

Syntax

```
logit depvvar [indepvars] [if] [in] [weight] [, options]
```

<i>options</i>	Description
----------------	-------------

Model	
<code>noconstant</code>	suppress constant term
<code>offset(<i>varname</i>)</code>	include <i>varname</i> in model with coefficient constrained to 1
<code>asis</code>	retain perfect predictor variables
<code>constraints(<i>constraints</i>)</code>	apply specified linear constraints
SE/Robust	
<code>vce(<i>vcetype</i>)</code>	<i>vcetype</i> may be <code>oim</code> , <code>robust</code> , <code>cluster <i>clustvar</i></code> , <code>bootstrap</code> , or <code>jackknife</code>
Reporting	
<code>level(#)</code>	set confidence level; default is <code>level(95)</code>
<code>or</code>	report odds ratios
<code>nocnsreport</code>	do not display constraints
<code>display_options</code>	control columns and column formats, row spacing, line width, display of omitted variables and base and empty cells, and factor-variable labeling
Maximization	
<code>maximize_options</code>	control the maximization process; seldom used
<code>nocoef</code>	do not display coefficient table; seldom used
<code>collinear</code>	keep collinear variables
<code>coeflegend</code>	display legend instead of statistics
<i>indepvars</i> may contain factor variables; see [U] 11.4.3 Factor variables.	
<i>depvvar</i> and <i>indepvars</i> may contain time-series operators; see [U] 11.4.4 Time-series varlists.	
<code>bayes</code> , <code>bootstrap</code> , <code>by</code> , <code>fmm</code> , <code>fp</code> , <code>jackknife</code> , <code>mfp</code> , <code>mi estimate</code> , <code>nestreg</code> , <code>rolling</code> , <code>statsby</code> , <code>stepwise</code> , and <code>svy</code> are allowed; see [U] 11.1.10 Prefix commands. For more details, see [BAYES] <code>bayes: logit</code> and [FMM] <code>fmm: logit</code> .	
<code>vce(bootstrap)</code> and <code>vce(jackknife)</code> are not allowed with the <code>mi estimate</code> prefix; see [MI] <code>mi estimate</code> .	
Weights are not allowed with the <code>bootstrap</code> prefix; see [R] <code>bootstrap</code> .	
<code>vce()</code> , <code>nocoef</code> , and weights are not allowed with the <code>svy</code> prefix; see [SVY] <code>svy</code> .	
<code>fweights</code> , <code>iweights</code> , and <code>pweights</code> are allowed; see [U] 11.1.6 weight.	
<code>nocoef</code> , <code>collinear</code> , and <code>coeflegend</code> do not appear in the dialog box.	
See [U] 20 Estimation and postestimation commands for more capabilities of estimation commands.	

Options

Model
<code>noconstant</code> , <code>offset(<i>varname</i>)</code> , <code>constraints(<i>constraints</i>)</code> ; see [R] Estimation options.
<code>asis</code> forces retention of perfect predictor variables and their associated perfectly predicted observations and may produce instabilities in maximization; see [R] <code>probit</code> .

SE/Robust

`vce(vcetype)` specifies the type of standard error reported, which includes types that are derived from asymptotic theory (`oim`), that are robust to some kinds of misspecification (`robust`), that allow for intragroup correlation (`cluster clustvar`), and that use bootstrap or jackknife methods (`bootstrap`, `jackknife`); see [R] [vce_option](#).

Reporting

`level(#)`; see [R] [Estimation options](#).

`or` reports the estimated coefficients transformed to odds ratios, that is, e^b rather than b . Standard errors and confidence intervals are similarly transformed. This option affects how results are displayed, not how they are estimated. `or` may be specified at estimation or when replaying previously estimated results.

`nocnsreport`; see [R] [Estimation options](#).

`display_options`: `nocl`, `nopvalues`, `noomitted`, `vsquish`, `noemptycells`, `baselevels`, `allbaselevels`, `novlabel`, `fvwrap(#)`, `fvwrapon(style)`, `cformat(%fmt)`, `pformat(%fmt)`, `sformat(%fmt)`, and `nolstretch`; see [R] [Estimation options](#).

Maximization

`maximize_options`: `difficult`, `technique(algorithm_spec)`, `iterate(#)`, `[no]log`, `trace`, `gradient`, `showstep`, `hessian`, `showtolerance`, `tolerance(#)`, `ltolerance(#)`, `nrtolerance(#)`, `nonrtolerance`, and `from(init_specs)`; see [R] [Maximize](#). These options are seldom used.

The following options are available with `logit` but are not shown in the dialog box:

`nocoef` specifies that the coefficient table not be displayed. This option is sometimes used by program writers but is of no use interactively.

`collinear`, `coeflegend`; see [R] [Estimation options](#).

Remarks and examples

[stata.com](https://www.stata.com)

Remarks are presented under the following headings:

[Basic usage](#)
[Model identification](#)

Basic usage

`logit` fits maximum likelihood models with dichotomous dependent (left-hand-side) variables coded as 0/1 (or, more precisely, coded as 0 and not-0).

For grouped data or data in binomial form, a probit model can be fit using `glm` with the `family(binomial)` and `link(logit)` options.

► Example 1

We have data on the make, weight, and mileage rating of 22 foreign and 52 domestic automobiles. We wish to fit a logit model explaining whether a car is foreign on the basis of its weight and mileage. Here is an overview of our data:

```
. use https://www.stata-press.com/data/r16/auto
(1978 Automobile Data)

. keep make mpg weight foreign

. describe
```

Contains data from <https://www.stata-press.com/data/r16/auto.dta>

obs:	74	1978 Automobile Data
vars:	4	13 Apr 2018 17:45
		(_dta has notes)

variable name	storage type	display format	value label	variable label
make	str18	%-18s		Make and Model
mpg	int	%8.0g		Mileage (mpg)
weight	int	%8.0gc		Weight (lbs.)
foreign	byte	%8.0g	origin	Car type

Sorted by: foreign

Note: Dataset has changed since last saved.

. inspect foreign

foreign: Car type

foreign: Car type		Number of Observations		
		Total	Integers	Nonintegers
#	Negative	-	-	-
#	Zero	52	52	-
#	Positive	22	22	-
#				
#	Total	74	74	-
#	Missing	-		
		74		

foreign is labeled and all values are documented in the label.

The variable `foreign` takes on two unique values, 0 and 1. The value 0 denotes a domestic car, and 1 denotes a foreign car.

The model that we wish to fit is

$$\text{Pr}(\text{foreign} = 1) = F(\beta_0 + \beta_1 \text{weight} + \beta_2 \text{mpg})$$

where $F(z) = e^z / (1 + e^z)$ is the cumulative logistic distribution.

To fit this model, we type

```
. logit foreign weight mpg
Iteration 0:   log likelihood =  -45.03321
Iteration 1:   log likelihood = -29.238536
Iteration 2:   log likelihood = -27.244139
Iteration 3:   log likelihood = -27.175277
Iteration 4:   log likelihood = -27.175156
Iteration 5:   log likelihood = -27.175156

Logistic regression              Number of obs   =          74
                                LR chi2(2)         =          35.72
                                Prob > chi2         =          0.0000
                                Pseudo R2          =          0.3966

Log likelihood = -27.175156
```

foreign	Coef.	Std. Err.	z	P> z	[95% Conf. Interval]	
weight	-.0039067	.0010116	-3.86	0.000	-.0058894	-.001924
mpg	-.1685869	.0919175	-1.83	0.067	-.3487418	.011568
_cons	13.70837	4.518709	3.03	0.002	4.851859	22.56487

We find that heavier cars are less likely to be foreign and that cars yielding better gas mileage are also less likely to be foreign, at least holding the weight of the car constant.

◀

□ Technical note

Stata interprets a value of 0 as a negative outcome (failure) and treats all other values (except missing) as positive outcomes (successes). Thus if your dependent variable takes on the values 0 and 1, then 0 is interpreted as failure and 1 as success. If your dependent variable takes on the values 0, 1, and 2, then 0 is still interpreted as failure, but both 1 and 2 are treated as successes.

If you prefer a more formal mathematical statement, when you type `logit y x`, Stata fits the model

$$\Pr(y_j \neq 0 \mid \mathbf{x}_j) = \frac{\exp(\mathbf{x}_j \boldsymbol{\beta})}{1 + \exp(\mathbf{x}_j \boldsymbol{\beta})}$$

□

Model identification

The `logit` command has one more feature, and it is probably the most useful. `logit` automatically checks the model for identification and, if it is underidentified, drops whatever variables and observations are necessary for estimation to proceed. (`logistic`, `probit`, and `ivprobit` do this as well.)

▷ Example 2

Have you ever fit a logit model where one or more of your independent variables perfectly predicted one or the other outcome?

For instance, consider the following data:

Outcome y	Independent variable x
0	1
0	1
0	0
1	0

Say that we wish to predict the outcome on the basis of the independent variable. The outcome is always zero whenever the independent variable is one. In our data, $\Pr(y = 0 \mid x = 1) = 1$, which means that the logit coefficient on x must be minus infinity with a corresponding infinite standard error. At this point, you may suspect that we have a problem.

Unfortunately, not all such problems are so easily detected, especially if you have a lot of independent variables in your model. If you have ever had such difficulties, you have experienced one of the more unpleasant aspects of computer optimization. The computer has no idea that it is trying to solve for an infinite coefficient as it begins its iterative process. All it knows is that at each step, making the coefficient a little bigger, or a little smaller, works wonders. It continues on its merry way until either 1) the whole thing comes crashing to the ground when a numerical overflow error occurs or 2) it reaches some predetermined cutoff that stops the process. In the meantime, you have been waiting. The estimates that you finally receive, if you receive any at all, may be nothing more than numerical roundoff.

Stata watches for these sorts of problems, alerts us, fixes them, and properly fits the model. Let's return to our automobile data. Among the variables we have in the data is one called `repair`, which takes on three values. A value of 1 indicates that the car has a poor repair record, 2 indicates an average record, and 3 indicates a better-than-average record. Here is a tabulation of our data:

```
. use https://www.stata-press.com/data/r16/repair, clear
(1978 Automobile Data)
. tabulate foreign repair
```

Car type	repair			Total
	1	2	3	
Domestic	10	27	9	46
Foreign	0	3	9	12
Total	10	30	18	58

All the cars with poor repair records (`repair = 1`) are domestic. If we were to attempt to predict `foreign` on the basis of the repair records, the predicted probability for the `repair = 1` category would have to be zero. This in turn means that the logit coefficient must be minus infinity, and that would set most computer programs buzzing.

Let's try Stata on this problem.

```
. logit foreign b3.repair
note: 1.repair != 0 predicts failure perfectly
      1.repair dropped and 10 obs not used

Iteration 0:   log likelihood = -26.992087
Iteration 1:   log likelihood = -22.483187
Iteration 2:   log likelihood = -22.230498
Iteration 3:   log likelihood = -22.229139
Iteration 4:   log likelihood = -22.229138

Logistic regression               Number of obs   =           48
                                LR chi2(1)         =           9.53
                                Prob > chi2         =          0.0020
                                Pseudo R2          =          0.1765

Log likelihood = -22.229138
```

foreign	Coef.	Std. Err.	z	P> z	[95% Conf. Interval]	
repair						
1	0 (empty)					
2	-2.197225	.7698003	-2.85	0.004	-3.706005	-.6884436
_cons	7.94e-17	.4714045	0.00	1.000	-.9239359	.9239359

Remember that all the cars with poor repair records (`repair = 1`) are domestic, so the model cannot be fit, or at least it cannot be fit if we restrict ourselves to finite coefficients. Stata noted that fact “note: 1.repair !=0 predicts failure perfectly”. This is Stata’s mathematically precise way of saying what we said in English. When `repair` is 1, the car is domestic.

Stata then went on to say “1.repair dropped and 10 obs not used”. This is Stata eliminating the problem. First 1.repair had to be removed from the model because it would have an infinite coefficient. Then, the 10 observations that led to the problem had to be eliminated, as well, so as not to bias the remaining coefficients in the model. The 10 observations that are not used are the 10 domestic cars that have poor repair records.

Stata then fit what was left of the model, using the remaining observations. Because no observations remained for cars with poor repair records, Stata reports “(empty)” in the row for `repair = 1`.

◀

□ Technical note

Stata is pretty smart about catching problems like this. It will catch “one-way causation by a dummy variable”, as we demonstrated above.

Stata also watches for “two-way causation”, that is, a variable that perfectly determines the outcome, both successes and failures. Here Stata says, “so-and-so predicts outcome perfectly” and stops. Statistics dictates that no model can be fit.

Stata also checks your data for collinear variables; it will say, “so-and-so omitted because of collinearity”. No observations need to be eliminated in this case, and model fitting will proceed without the offending variable.

It will also catch a subtle problem that can arise with continuous data. For instance, if we were estimating the chances of surviving the first year after an operation, and if we included in our model `age`, and if all the persons over 65 died within the year, Stata would say, “age > 65 predicts failure perfectly”. It would then inform us about the fix-up it takes and fit what can be fit of our model.

logit (and logistic, probit, and ivprobit) will also occasionally display messages such as

Note: 4 failures and 0 successes completely determined.

There are two causes for a message like this. The first—and most unlikely—case occurs when a continuous variable (or a combination of a continuous variable with other continuous or dummy variables) is simply a great predictor of the dependent variable. Consider Stata's `auto.dta` dataset with 6 observations removed.

```
. use https://www.stata-press.com/data/r16/auto
(1978 Automobile Data)

. drop if foreign==0 & gear_ratio > 3.1
(6 observations deleted)

. logit foreign mpg weight gear_ratio, nolog
Logistic regression                                Number of obs   =          68
                                                    LR chi2(3)      =          72.64
                                                    Prob > chi2     =          0.0000
Log likelihood = -6.4874814                        Pseudo R2      =          0.8484
```

foreign	Coef.	Std. Err.	z	P> z	[95% Conf. Interval]	
mpg	-.4944907	.2655508	-1.86	0.063	-1.014961	.0259792
weight	-.0060919	.003101	-1.96	0.049	-.0121698	-.000014
gear_ratio	15.70509	8.166234	1.92	0.054	-.300436	31.71061
_cons	-21.39527	25.41486	-0.84	0.400	-71.20747	28.41694

Note: 4 failures and 0 successes completely determined.

There are no missing standard errors in the output. If you receive the “completely determined” message and have one or more missing standard errors in your output, see the second case discussed below.

Note `gear_ratio`'s large coefficient. `logit` thought that the 4 observations with the smallest predicted probabilities were essentially predicted perfectly.

```
. predict p
(option pr assumed; Pr(foreign))

. sort p

. list p in 1/4
```

	P
1.	1.34e-10
2.	6.26e-09
3.	7.84e-09
4.	1.49e-08

If this happens to you, you do not have to do anything. Computationally, the model is sound. The second case discussed below requires careful examination.

The second case occurs when the independent terms are all dummy variables or continuous ones with repeated values (for example, age). Here one or more of the estimated coefficients will have missing standard errors. For example, consider this dataset consisting of 6 observations.


```
. use https://www.stata-press.com/data/r16/logitxmpl, clear
. list, separator(0)
```

	y	x1	x2
1.	0	0	0
2.	0	0	0
3.	0	1	0
4.	1	1	0
5.	0	0	1
6.	1	0	1

```
. logit y x1 x2
```

```
Iteration 0:  log likelihood = -3.819085
Iteration 1:  log likelihood = -2.9527336
Iteration 2:  log likelihood = -2.8110282
Iteration 3:  log likelihood = -2.7811973
Iteration 4:  log likelihood = -2.7746107
Iteration 5:  log likelihood = -2.7730128
(output omitted)
Iteration 15996: log likelihood = -2.7725887 (not concave)
Iteration 15997: log likelihood = -2.7725887 (not concave)
Iteration 15998: log likelihood = -2.7725887 (not concave)
Iteration 15999: log likelihood = -2.7725887 (not concave)
Iteration 16000: log likelihood = -2.7725887 (not concave)
convergence not achieved
```

```
Logistic regression                                Number of obs   =           6
                                                    LR chi2(1)      =           2.09
                                                    Prob > chi2     =          0.1480
Log likelihood = -2.7725887                        Pseudo R2      =          0.2740
```

	y	Coef.	Std. Err.	z	P> z	[95% Conf. Interval]	
	x1	18.3704	2	9.19	0.000	14.45047	22.29033
	x2	18.3704	.	.	.	.	.
	_cons	-18.3704	1.414214	-12.99	0.000	-21.14221	-15.5986

Note: 2 failures and 0 successes completely determined.

```
convergence not achieved
r(430);
```

Three things are happening here. First, `logit` iterates almost forever and then declares nonconvergence. Second, `logit` can fit the outcome ($y = 0$) for the covariate pattern $x1 = 0$ and $x2 = 0$ (that is, the first two observations) perfectly. This observation is the “2 failures and 0 successes completely determined”. Third, if this observation is dropped, then $x1$, $x2$, and the constant are collinear.

This is the cause of the nonconvergence, the message “completely determined”, and the missing standard errors. It happens when you have a covariate pattern (or patterns) with only one outcome and there is collinearity when the observations corresponding to this covariate pattern are dropped.

If this happens to you, confirm the causes. First, identify the covariate pattern with only one outcome. (For your data, replace $x1$ and $x2$ with the independent variables of your model.)

```
. egen pattern = group(x1 x2)
. quietly logit y x1 x2, iterate(100)
. predict p
(option pr assumed; Pr(y))
. summarize p
```

Variable	Obs	Mean	Std. Dev.	Min	Max
p	6	.3333333	.2581989	1.05e-08	.5

If successes were completely determined, that means that there are predicted probabilities that are almost 1. If failures were completely determined, that means that there are predicted probabilities that are almost 0. The latter is the case here, so we locate the corresponding value of `pattern`:

```
. tabulate pattern if p < 1e-7
```

group(x1 x2)	Freq.	Percent	Cum.
1	2	100.00	100.00
Total	2	100.00	

Once we omit this covariate pattern from the estimation sample, `logit` can deal with the collinearity:

```
. logit y x1 x2 if pattern != 1, nolog
note: x2 omitted because of collinearity

Logistic regression               Number of obs   =           4
                                LR chi2(1)         =           0.00
                                Prob > chi2         =          1.0000
Log likelihood = -2.7725887       Pseudo R2      =          0.0000
```

y	Coef.	Std. Err.	z	P> z	[95% Conf. Interval]	
x1	0	2	0.00	1.000	-3.919928	3.919928
x2	0 (omitted)					
_cons	0	1.414214	0.00	1.000	-2.771808	2.771808

We omit the collinear variable. Then we must decide whether to include or omit the observations with `pattern = 1`. We could include them,

```
. logit y x1, nolog
Logistic regression               Number of obs   =           6
                                LR chi2(1)         =           0.37
                                Prob > chi2         =          0.5447
Log likelihood = -3.6356349       Pseudo R2      =          0.0480
```

y	Coef.	Std. Err.	z	P> z	[95% Conf. Interval]	
x1	1.098612	1.825742	0.60	0.547	-2.479776	4.677001
_cons	-1.098612	1.154701	-0.95	0.341	-3.361784	1.164559

or exclude them,

```
. logit y x1 if pattern != 1, nolog
```

Logistic regression	Number of obs	=	4
	LR chi2(1)	=	0.00
	Prob > chi2	=	1.0000
Log likelihood = -2.7725887	Pseudo R2	=	0.0000

y	Coef.	Std. Err.	z	P> z	[95% Conf. Interval]	
x1	0	2	0.00	1.000	-3.919928	3.919928
_cons	0	1.414214	0.00	1.000	-2.771808	2.771808

If the covariate pattern that predicts outcome perfectly is meaningful, you may want to exclude these observations from the model. Here you would report that covariate pattern such and such predicted outcome perfectly and that the best model for the rest of the data is But, more likely, the perfect prediction was simply the result of having too many predictors in the model. Then you would omit the extraneous variables from further consideration and report the best model for all the data.



Stored results

logit stores the following in `e()`:

Scalars

<code>e(N)</code>	number of observations
<code>e(N_cds)</code>	number of completely determined successes
<code>e(N_cdf)</code>	number of completely determined failures
<code>e(k)</code>	number of parameters
<code>e(k_eq)</code>	number of equations in <code>e(b)</code>
<code>e(k_eq_model)</code>	number of equations in overall model test
<code>e(k_dv)</code>	number of dependent variables
<code>e(df_m)</code>	model degrees of freedom
<code>e(r2_p)</code>	pseudo- R^2
<code>e(ll)</code>	log likelihood
<code>e(ll_0)</code>	log likelihood, constant-only model
<code>e(N_clust)</code>	number of clusters
<code>e(chi2)</code>	χ^2
<code>e(p)</code>	p -value for model test
<code>e(rank)</code>	rank of <code>e(V)</code>
<code>e(ic)</code>	number of iterations
<code>e(rc)</code>	return code
<code>e(converged)</code>	1 if converged, 0 otherwise

Macros

<code>e(cmd)</code>	logit
<code>e(cmdline)</code>	command as typed
<code>e(depvar)</code>	name of dependent variable
<code>e(wtype)</code>	weight type
<code>e(wexp)</code>	weight expression
<code>e(title)</code>	title in estimation output
<code>e(clustvar)</code>	name of cluster variable
<code>e(offset)</code>	linear offset variable
<code>e(chi2type)</code>	Wald or LR; type of model χ^2 test
<code>e(vce)</code>	<i>vcetype</i> specified in <code>vce()</code>
<code>e(vcetype)</code>	title used to label Std. Err.
<code>e(opt)</code>	type of optimization
<code>e(which)</code>	max or min; whether optimizer is to perform maximization or minimization
<code>e(ml_method)</code>	type of ml method
<code>e(user)</code>	name of likelihood-evaluator program

<code>e(technique)</code>	maximization technique
<code>e(properties)</code>	b V
<code>e(estat_cmd)</code>	program used to implement estat
<code>e(predict)</code>	program used to implement predict
<code>e(marginsok)</code>	predictions allowed by margins
<code>e(marginsnotok)</code>	predictions disallowed by margins
<code>e(asbalanced)</code>	factor variables fvset as asbalanced
<code>e(asobserved)</code>	factor variables fvset as asobserved

Matrices

<code>e(b)</code>	coefficient vector
<code>e(Cns)</code>	constraints matrix
<code>e(ilog)</code>	iteration log (up to 20 iterations)
<code>e(gradient)</code>	gradient vector
<code>e(mns)</code>	vector of means of the independent variables
<code>e(rules)</code>	information about perfect predictors
<code>e(V)</code>	variance–covariance matrix of the estimators
<code>e(V_modelbased)</code>	model-based variance

Functions

<code>e(sample)</code>	marks estimation sample
------------------------	-------------------------

In addition to the above, the following is stored in `r()`:

Matrices

<code>r(table)</code>	matrix containing the coefficients with their standard errors, test statistics, <i>p</i> -values, and confidence intervals
-----------------------	--

Note that results stored in `r()` are updated when the command is replayed and will be replaced when any `r`-class command is run after the estimation command.

Methods and formulas

Cramer (2003, chap. 9) surveys the prehistory and history of the logit model. The word “logit” was coined by Berkson (1944) and is analogous to the word “probit”. For an introduction to probit and logit, see, for example, Aldrich and Nelson (1984), Cameron and Trivedi (2010), Jones (2007), Long (1997), Long and Freese (2014), Pampel (2000), or Powers and Xie (2008).

The likelihood function for logit is

$$\ln L = \sum_{j \in S} w_j \ln F(\mathbf{x}_j \mathbf{b}) + \sum_{j \notin S} w_j \ln \{1 - F(\mathbf{x}_j \mathbf{b})\}$$

where S is the set of all observations j , such that $y_j \neq 0$, $F(z) = e^z / (1 + e^z)$, and w_j denotes the optional weights. $\ln L$ is maximized as described in [R] **Maximize**.

This command supports the Huber/White/sandwich estimator of the variance and its clustered version using `vce(robust)` and `vce(cluster clustvar)`, respectively. See [P] **_robust**, particularly *Maximum likelihood estimators* and *Methods and formulas*. The scores are calculated as $\mathbf{u}_j = \{1 - F(\mathbf{x}_j \mathbf{b})\} \mathbf{x}_j$ for the positive outcomes and $-F(\mathbf{x}_j \mathbf{b}) \mathbf{x}_j$ for the negative outcomes.

`logit` also supports estimation with survey data. For details on VCEs with survey data, see [SVY] **Variance estimation**.

Joseph Berkson (1899–1982) was born in New York City and studied at the College of the City of New York, Columbia, and Johns Hopkins, earning both an MD and a doctorate in statistics. He then worked at Johns Hopkins before moving to the Mayo Clinic in 1931 as a biostatistician. Among many other contributions, his most influential one drew upon a long-sustained interest in the logistic function, especially his 1944 paper on bioassay, in which he introduced the term “logit”. Berkson was a frequent participant in controversy—sometimes humorous, sometimes bitter—on subjects such as the evidence for links between smoking and various diseases and the relative merits of probit and logit methods and of different calculation methods.

References

- Aldrich, J. H., and F. D. Nelson. 1984. *Linear Probability, Logit, and Probit Models*. Newbury Park, CA: SAGE.
- Archer, K. J., and S. A. Lemeshow. 2006. Goodness-of-fit test for a logistic regression model fitted using survey sample data. *Stata Journal* 6: 97–105.
- Berkson, J. 1944. Application of the logistic function to bio-assay. *Journal of the American Statistical Association* 39: 357–365.
- Buis, M. L. 2010a. Direct and indirect effects in a logit model. *Stata Journal* 10: 11–29.
- . 2010b. Stata tip 87: Interpretation of interactions in nonlinear models. *Stata Journal* 10: 305–308.
- Cameron, A. C., and P. K. Trivedi. 2010. *Microeconometrics Using Stata*. Rev. ed. College Station, TX: Stata Press.
- Corral, P., and M. Terbish. 2015. Generalized maximum entropy estimation of discrete choice models. *Stata Journal* 15: 512–522.
- Cramer, J. S. 2003. *Logit Models from Economics and Other Fields*. Cambridge: Cambridge University Press.
- Drukker, D. M. 2016. Probability differences and odds ratios measure conditional-on-covariate effects and population-parameter effects. *The Stata Blog: Not Elsewhere Classified*. <http://blog.stata.com/2016/07/26/probability-differences-and-odds-ratios-measure-conditional-on-covariate-effects-and-population-parameter-effects/>.
- Hilbe, J. M. 2009. *Logistic Regression Models*. Boca Raton, FL: Chapman & Hill/CRC.
- Hosmer, D. W., Jr., S. A. Lemeshow, and R. X. Sturdivant. 2013. *Applied Logistic Regression*. 3rd ed. Hoboken, NJ: Wiley.
- Jones, A. M. 2007. *Applied Econometrics for Health Economists: A Practical Guide*. 2nd ed. Abingdon, UK: Radcliffe.
- Judge, G. G., W. E. Griffiths, R. C. Hill, H. Lütkepohl, and T.-C. Lee. 1985. *The Theory and Practice of Econometrics*. 2nd ed. New York: Wiley.
- Long, J. S. 1997. *Regression Models for Categorical and Limited Dependent Variables*. Thousand Oaks, CA: SAGE.
- Long, J. S., and J. Freese. 2014. *Regression Models for Categorical Dependent Variables Using Stata*. 3rd ed. College Station, TX: Stata Press.
- Miranda, A., and S. Rabe-Hesketh. 2006. Maximum likelihood estimation of endogenous switching and sample selection models for binary, ordinal, and count variables. *Stata Journal* 6: 285–308.
- Mitchell, M. N., and X. Chen. 2005. Visualizing main effects and interactions for binary logit models. *Stata Journal* 5: 64–82.
- O’Fallon, W. M. 1998. Berkson, Joseph. In Vol. 1 of *Encyclopedia of Biostatistics*, ed. P. Armitage and T. Colton, 290–295. Chichester, UK: Wiley.
- Orsini, N., R. Bellocco, and P. C. Sjölander. 2013. Doubly robust estimation in generalized linear models. *Stata Journal* 13: 185–205.
- Pampel, F. C. 2000. *Logistic Regression: A Primer*. Thousand Oaks, CA: SAGE.
- Pedace, R. 2013. *Econometrics for Dummies*. Hoboken, NJ: Wiley.
- Pollock, P. H., III, and B. C. Edwards. 2019. *A Stata Companion to Political Analysis*. 4th ed. Thousand Oaks, CA: CQ Press.
- Powers, D. A., and Y. Xie. 2008. *Statistical Methods for Categorical Data Analysis*. 2nd ed. Bingley, UK: Emerald.

Pregibon, D. 1981. Logistic regression diagnostics. *Annals of Statistics* 9: 705–724.

Schonlau, M. 2005. [Boosted regression \(boosting\): An introductory tutorial and a Stata plugin](#). *Stata Journal* 5: 330–354.

Xu, J., and J. S. Long. 2005. [Confidence intervals for predicted outcomes in regression models for categorical outcomes](#). *Stata Journal* 5: 537–559.

Also see

[R] [logit postestimation](#) — Postestimation tools for logit

[R] [brier](#) — Brier score decomposition

[R] [cloglog](#) — Complementary log-log regression

[R] [exlogistic](#) — Exact logistic regression

[R] [logistic](#) — Logistic regression, reporting odds ratios

[R] [npregress kernel](#) — Nonparametric kernel regression

[R] [npregress series](#) — Nonparametric series regression

[R] [probit](#) — Probit regression

[R] [roc](#) — Receiver operating characteristic (ROC) analysis

[BAYES] [bayes: logit](#) — Bayesian logistic regression, reporting coefficients

[FMM] [fmm: logit](#) — Finite mixtures of logistic regression models

[LASSO] [Lasso intro](#) — Introduction to lasso

[ME] [melogit](#) — Multilevel mixed-effects logistic regression

[MI] [Estimation](#) — Estimation commands for use with mi estimate

[SVY] [svy estimation](#) — Estimation commands for survey data

[XT] [xtlogit](#) — Fixed-effects, random-effects, and population-averaged logit models

[U] [20 Estimation and postestimation commands](#)