

Linear Fixed Effects Models

Basics

In this chapter, we consider fixed effects methods for data in which the dependent variable is measured on an interval scale and is linearly dependent on a set of predictor variables. We have a set of individuals ($i = 1, \dots, n$), each of whom is measured at two or more points in time ($t = 1, \dots, T$). Each time point will be referred to as a “period.”

Here's the notation. We let y_{it} be the dependent variable. We have a set of predictor variables that vary over time, represented by the vector $\mathbf{x}_{it}$, and another set of predictor variables $\mathbf{z}_i$ that do not vary over time. (If you're not comfortable with vectors, you can interpret these as single variables). Our basic model for y is

$$y_{it} = \mu_t + \beta \mathbf{x}_{it} + \gamma \mathbf{z}_i + \alpha_i + \varepsilon_{it} \quad (2.1)$$

where μ_t is an intercept that may be different for each period, and β and γ are vectors of coefficients. Although Equation 2.1 seems to be strictly cross-sectional, there is nothing to prevent the $\mathbf{x}_{it}$ vector from including lagged versions of the x variables, except that one must have at least three periods to estimate a model with a lag of one period.

The two “error” terms, α_i and ε_{it} behave somewhat differently from each other. There is a different ε_{it} for each individual at each point in time, but α_i only varies across individuals, not over time. We regard α_i as representing the combined effect on y of all unobserved variables that are constant over time. On the other hand, ε_{it} represents purely random variation at each point in time.

At this point, I'll make some rather strong assumptions about ε_{it} namely, that each ε_{it} has a mean of zero, has a constant variance (for all i and t), and is statistically independent of everything else (except for y). The assumption of zero mean is not critical as it is only relevant for estimating the intercept. The constant variance assumption can sometimes be relaxed to allow for different variances for different t . Note that the $\mathbf{x}_{it}$ at any one period is independent of $\mathbf{x}_{it}$ at any *other* period, which means that $\mathbf{x}_{it}$ is *strictly exogenous*. This assumption may be relaxed in some situations, but the issues involved are neither trivial nor purely technical. I'll discuss some of those issues in Chapter 6.

As for α_i the traditional approach in fixed effects analysis is to assume that this term represents a set of n fixed parameters that can either be directly estimated or removed in some way from

the estimating equations. As noted in Chapter 1, we'll take a more modern approach in this chapter by assuming that α_j represents a set of random variables. Although we'll assume statistical independence of α_j and $\mathbf{x}_{jt}$, we allow for *any* correlations between α_j and $\mathbf{x}_{jt}$ the vector of time-varying predictors. And if we are not interested in γ , we can also allow for any correlations between α_j and $\mathbf{z}_j$. The inclusion of such correlations distinguishes the fixed effects approach from a random effects approach and allows us to say that the fixed effects method “controls” for time-invariant unobservables. At this point, we don't need to impose any restrictions on the mean and variance of α_j .

The Two-Period Case

Estimation of the model in Equation 2.1 is particularly easy when the variables are observed at only two periods ($T = 2$). The two equations are then

$$\begin{aligned} y_{i1} &= \mu_1 + \beta \mathbf{x}_{i1} + \gamma \mathbf{z}_i + \alpha_i + \varepsilon_{i1} \\ y_{i2} &= \mu_2 + \beta \mathbf{x}_{i2} + \gamma \mathbf{z}_i + \alpha_i + \varepsilon_{i2} \end{aligned} \quad (2.2)$$

If we subtract the first equation from the second, we get the “first difference” equation:

$$y_{i2} - y_{i1} = (\mu_2 - \mu_1) + \beta(\mathbf{x}_{i2} - \mathbf{x}_{i1}) + (\varepsilon_{i2} - \varepsilon_{i1}) \quad (2.3)$$

which can be rewritten as

$$\Delta y_i = \Delta \mu + \beta \Delta \mathbf{x}_i + \Delta \varepsilon_i \quad (2.4)$$

where Δ indicates a difference score. Note that both α_j and $\gamma \mathbf{z}_j$ have been “differenced out” of the equation. Hence, we no longer have to be concerned about α_j and its possible correlation with $\Delta \mathbf{x}_j$. On the other hand, we also lose the possibility of estimating γ . Since $\mathbf{x}_{j1}$ and $\mathbf{x}_{j2}$ are each independent of ε_{j1} and ε_{j2} it follows that $\Delta \mathbf{x}_j$ is independent of $\Delta \varepsilon_j$. This implies that one can get unbiased estimates of β by doing ordinary least squares (OLS) regression on the difference scores.

We now apply this method to some real data. Our sample comes from the National Longitudinal Survey of Youth (NLSY; Center for Human Resource Research, 2002).¹ Drawing a subset of a much larger sample, we have 581 children who were interviewed in 1990, 1992, and 1994. Initially, we will work with just three variables that were measured in each of the three interviews:

ANTI antisocial behavior (scale ranges from 0 to 6)

SELF self-esteem (scale ranges from 6 to 24)

POV coded 1 if family is in poverty, otherwise 0

At this point, we're going to ignore the observations in the middle year, 1992, and use only the data in 1990 and 1994. Our objective is to estimate a linear equation with ANTI as the dependent variable² and SELF and POV as independent variables:

$$ANTI_t = \mu_t + \beta_1 SELF_t + \beta_2 POV_t + \alpha + \varepsilon_t, \quad t = 1, 2 \quad (2.5)$$

By expressing the model in this way, we are assuming a particular direction of causation, specifically, that SELF and POV affect ANTI and not the reverse. We also assume that the effects are contemporaneous (no lagged effects of SELF and POV). Both these assumptions will be relaxed in Chapter 6. Last, we assume that β_1 and β_2 are the same at both periods, an assumption that we will relax very shortly. On the other hand, we let the intercept μ_t be different at each period, allowing for changes in the average level of antisocial behavior that are not a consequence of any changes in SELF or POV.

Let us begin by estimating Equation 2.5 separately for each period, using ordinary least squares regression. The results are shown in the first two columns of Table 2.1. Not surprisingly, in both years, poverty is associated with higher levels of antisocial behavior, whereas self-esteem is associated with lower levels. The coefficients are quite similar across the two years.

In neither of these two regressions are there any controls for timeinvariant predictors, such as sex and race. Rather than putting such variables in the equation, however, we can control for *all* time-invariant predictors by doing the regression with difference scores. For each child and each variable, we subtract the 1990 value from the 1994 value and then regress the ANTI difference on the SELF difference and the POV difference. Since POV is a dummy variable, it might seem inappropriate to subtract one value from the other. But, in fact, dummy variables can be treated just like any other variables in this regard.

The results are in the last column of Table 2.1. Although the equation was estimated in the form of difference scores, the coefficients can be interpreted as if we had estimated Equation 2.5 directly. That is, they represent the effects of each variable in a given year on the value of the dependent variable in the same year. For self-esteem, the estimated coefficient in the difference score model is in between the coefficients for the two separate years and still highly significant. For poverty, on the other hand, the coefficient is dramatically smaller and no longer statistically significant.

Table 2.1 OLS Regression of Antisocial Behavior on Self-Esteem and Poverty

	1990		1994		Difference Score	
	Coefficient	Standard Error	Coefficient	Standard Error	Coefficient	Standard Error
Intercept	2.375**	0.384	2.888**	0.447	0.209**	0.063
Self-esteem	-0.050**	0.019	-0.064**	0.021	-0.056**	0.015
Poverty	0.595**	0.126	0.547**	0.148	-0.036	0.128
R^2	0.05		0.04		0.02	

It's fairly common for fixed estimates to vary markedly from those produced by other methods. In this case, one possible explanation is that the estimates for the poverty effect in the regressions for the separate years were spurious, reflecting the correlation between poverty and some time-invariant variables that affected antisocial behavior.

One should not be too hasty about that conclusion, however. Whenever conventional regression produces a significant coefficient but fixed effects regression does not, there are two possible explanations: (a) The fixed effects coefficient is substantially smaller in magnitude and/or (b) the fixed effects standard error is substantially larger. As already mentioned, standard errors for fixed effects coefficients are often larger than those for other methods, especially when the predictor variable has little variation over time. In fact, most of the variation in poverty is between girls, with only about 24% of the girls moving into or out of poverty between 1990 and 1994.

Nevertheless, the standard error of the poverty coefficient in the difference score model is about the same as the standard error in 1990 and smaller than the standard error in 1994. The conclusion, then, is that insufficient variation is not a problem here. There seems to be a real and substantial decline in the magnitude of the poverty effect when time-invariant variables are controlled. The general lesson is this: Whenever p values for fixed effects methods are noticeably different from those for other methods, always check both the coefficients and their standard errors.

Note, finally, that the intercept of 0.209 is highly significant. This coefficient represents the estimated change in antisocial behavior from Time 1 to Time 2 for a person who did *not* change in self-esteem or poverty.

Extending the Difference Score Method for the Two-Period Case

The basic fixed effects model of Equation 2.1 can be extended to allow the effects of $\mathbf{x}$ and $\mathbf{z}$ to vary over time. In the two-period case, we can write the equations with distinct coefficients at each period:

$$\begin{aligned} y_{i1} &= \mu_1 + \beta_1 \mathbf{x}_{i1} + \gamma_1 \mathbf{z}_i + \alpha_i + \varepsilon_{i1} \\ y_{i2} &= \mu_2 + \beta_2 \mathbf{x}_{i2} + \gamma_2 \mathbf{z}_i + \alpha_i + \varepsilon_{i2} \end{aligned} \quad (2.6)$$

Taking first differences and rearranging terms produces

$$\begin{aligned} y_{i2} - y_{i1} &= (\mu_2 - \mu_1) + \beta_2(\mathbf{x}_{i2} - \mathbf{x}_{i1}) + (\beta_2 - \beta_1)\mathbf{x}_{i1} \\ &\quad + (\gamma_2 - \gamma_1)\mathbf{z}_i + (\varepsilon_{i2} - \varepsilon_{i1}) \end{aligned} \quad (2.7)$$

which could also be written as

$$\Delta y_i = \Delta\mu + \beta_2 \Delta \mathbf{x}_i + \Delta\beta \mathbf{x}_1 + \Delta\gamma \mathbf{z}_i + \Delta\varepsilon_i$$

There are three things to note about this equation. First, as before, α_i has dropped out, so we don't have to be concerned about its potential confounding effects. Second, $\mathbf{z}$ has *not* dropped out, and its coefficient vector is the difference in the coefficient vectors for the two time points. From this we learn that time-invariant variables whose coefficients vary over time *must* be explicitly included in the regression equation. Fixed effects only controls for time-invariant variables with time-invariant effects. Third, the equation now includes $\mathbf{x}_1$ as a predictor, and its coefficient vector is the difference in the coefficient vectors for the two time periods. Thus, for $\mathbf{z}$ and $\mathbf{x}_1$, tests for whether their coefficients are 0 are equivalent to testing whether $\beta_1 = \beta_2$ or $\gamma_1 = \gamma_2$.

Let's try this for the NLSY data. The data set includes several time-invariant variables that we'll examine as possible predictors:

BLACK 1 if child is black, otherwise 0

HISPANIC 1 if child is Hispanic, otherwise 0

CHILDAGE child's age in 1990

MARRIED 1 if mother was currently married in 1990, otherwise 0

GENDER 1 if female, 0 if male

MOMAGE mother's age at birth of child

MOMWORK 1 if mother was employed in 1990, otherwise 0

The first two variables, BLACK and HISPANIC, represent two categories of a three-category variable-the reference category being white, non-Hispanic. These seven variables will now be included in the difference score regression for antisocial behavior, along with the difference scores for self-esteem and poverty. The 1990 measures of self-esteem and poverty are also

included.

Results shown in Table 2.2 are consistent with the findings in Table 2.1. The coefficient for self-esteem (the difference score) is about $-.06$ and is highly significant, while the coefficient for poverty (the difference score) is

Table 2.2 OLS Estimates of Extended Difference Score Model

	<i>Coefficient</i>	<i>Standard Error</i>	<i>p Value</i>
Intercept	−0.550	1.360	.6859
Self-esteem difference	−0.060	0.020	.0024
Poverty difference	0.031	0.156	.8446
Self-esteem 1990	−0.018	0.025	.4826
Poverty 1990	0.121	0.178	.4991
Black	−0.100	0.155	.5185
Hispanic	0.084	0.164	.6109
ChildAge	0.220	0.107	.0409
Married	−0.206	0.154	.1808
Gender	0.101	0.126	.4262
MomAge	−0.040	0.030	.1842
MomWork	−0.153	0.140	.2735

far from significant. Neither the 1990 self-esteem coefficient nor the 1990 poverty coefficient approaches statistical significance, indicating that the effects of self-esteem and poverty were stable over time. Of the seven time-invariant variables, only one-child's age in 1990 is statistically significant (just barely). That doesn't mean that the other six variables are not affecting antisocial behavior. Rather, it means that their effects in 1990 and 1994 are essentially the same.

A First-Difference Method for Three or More Periods per Individual

When each individual is observed at three or more points in time ($T > 2$), it's not so obvious how to extend the methods we just considered. For the NLSY data, we actually have three years of data-1990, 1992, and 1994. One possible approach is to construct and estimate two first-difference equations. Starting with Equation 2.2, we have

$$y_{i2} - y_{i1} = (\mu_2 - \mu_1) + \beta(\mathbf{x}_{i2} - \mathbf{x}_{i1}) + (\varepsilon_{i2} - \varepsilon_{i1})$$

$$y_{i3} - y_{i2} = (\mu_3 - \mu_2) + \beta(\mathbf{x}_{i3} - \mathbf{x}_{i2}) + (\varepsilon_{i3} - \varepsilon_{i2}) \quad (2.8)$$

These equations can be estimated separately by OLS, and each will give unbiased estimates of β . The first two columns of Table 2.3 show the results for the NLSY data. The negative coefficients for self-esteem are highly significant in both difference equations and are roughly of the same magnitude. Poverty is far from significant in both equations. The intercepts represent the change from one period to the next, after adjusting for the two predictor variables. Although there was an increase in antisocial behavior for each 2-year interval, only the change from 1992 to 1994 was statistically significant.

Under the assumption that β doesn't vary over time, the two equations should be estimated simultaneously for optimal efficiency. This can be accomplished by creating a single data set with two records for each person, one with the difference scores for the first equation and the other with the difference scores for the second equation. There should also be a dummy variable distinguishing the first record from the second record. And also, there should be a variable with a common ID number for the two records from each person.

The third column of Table 2.3 shows the results of applying OLS to this combined data set of 1,162 records. Not surprisingly, the coefficients for self-esteem and poverty are in between their respective coefficients in the first two columns. The standard errors are appreciably lower, however,

Table 2.3 First-Difference Regressions of Antisocial Behavior on Self-Esteem and Poverty

	<i>1992–1994 OLS</i>		<i>1990–1992 OLS</i>		<i>Combined OLS</i>		<i>Combined GLS</i>	
	<i>Coefficient</i>	<i>Standard Error</i>	<i>Coefficient</i>	<i>Standard Error</i>	<i>Coefficient</i>	<i>Standard Error</i>	<i>Coefficient</i>	<i>Standard Error</i>
Intercept	0.71**	0.059	0.040	0.053	0.045	0.056	0.05	0.056
Self-esteem	−0.072**	0.016	−0.039**	0.014	−0.055**	0.010	−0.055**	0.010
Poverty	0.216	0.136	0.197	0.133	0.213	0.095	0.139	0.094
Equation dummy					0.122	0.080	0.122	0.094

because more information is being used. The intercept can be interpreted as an estimate of $\mu_2 - \mu_1$ while the coefficient for the equation dummy is an estimate of $(\mu_3 - \mu_2) - (\mu_2 - \mu_1)$. Both are positive, indicating an increase in antisocial behavior from Time 1 to Time 2, and a more rapid increase from Time 2 to Time 3. But neither is statistically significant.

Although the combined OLS estimates should be unbiased, they ignore the fact that $\varepsilon_2 - \varepsilon_1$ is likely to be negatively correlated with $\varepsilon_3 - \varepsilon_2$ because they share a common component, ε_2 , with opposite signs. This implies that the coefficient estimates may not be fully efficient and the standard error estimates may be biased. We can solve this problem by estimating the correlation between the error terms and then using generalized least squares (GLS) to take account of that correlation.

Most comprehensive statistical packages have routines to do GLS. Such programs typically require the specification of an ID variable so that observations from the same individual can be identified. I used the **xtreg** command in Stata with the **pa** option, which estimates the linear model using GLS.³ The GLS results, shown in the last column of Table 2.3, are very similar to the OLS coefficient estimates and standard errors in the preceding column.

The first-difference method can be easily extended to more than three periods per individual. For T periods per individual, $T - 1$ records are created, each with difference scores between adjacent time points for all variables. Additionally, there should be a variable containing a common ID number for all observations from the same individual, and a variable or a set of dummy variables to distinguish the different records. The regression is then estimated on the entire set of records, using GLS to adjust for correlations among the error terms. Unless T is large, for example, greater than 10, it's probably best to allow the error correlation matrix to be unstructured. That is, the matrix would allow for a different correlation between each pair of error terms. With larger T , it may be preferable to impose a simplified structure to reduce the number of distinct correlations that need to be estimated. For more details, see Greene (2000).

Dummy Variable Method for Two or More Periods per Individual

Although the multiple-difference-score method is a reasonable way to estimate a fixed effects model for the multiperiod case, the name “fixed effects” is usually reserved for a different method, one that can be implemented either by dummy variables or by constructing mean deviations. The results produced by the fixed effects method are not identical to those produced by the difference-score method, although they will usually be very similar. In the two-period case, the two methods give identical results.

The dummy variable method requires a data set with a rather different structure: one record for each period for each individual. For the NLSY data, for example, the required data set has three records for each of the 581 children, for a total of 1,743 records. The time-varying variables have the same variable names on each record but different values. For any time-invariant variables, their values are simply replicated across the multiple records for each

individual. There should also be an ID variable with a common value for all the records for each individual. Last, there should be a variable distinguishing the different periods for each individual. For the NLSY data, for example, the variable TIME has values of 1, 2, and 3, corresponding to 1990, 1992, and 1994. Table 2.4 shows the first 15 records of this data set, corresponding to the first five persons.

Table 2.4 Data Set With Three Observations per Person (First Five Persons)

<i>ID</i>	<i>TIME</i>	<i>ANTI</i>	<i>SELF</i>	<i>POV</i>	<i>GENDER</i>
1	1	1	21	1	1
1	2	1	24	1	1
1	3	1	23	1	1
2	1	0	20	0	1
2	2	0	24	0	1
2	3	0	24	0	1
3	1	5	21	0	0
3	2	5	24	0	0
3	3	5	24	0	0
4	1	2	23	0	0
4	2	3	21	0	0
4	3	1	21	0	0
5	1	1	22	0	1
5	2	0	23	0	1
5	3	0	24	0	1

To implement the method, it's necessary to construct a set of dummy variables to distinguish the individuals in the data set. In our example, that means 580 dummy variables to represent the 581 children. Many statistical packages can do this automatically by specifying the ID variable as a categorical variable. If the TIME variable is also specified as categorical, two dummy variables will be created to distinguish the three different years. One can then use OLS to estimate the coefficients. The coefficients for the dummy variables created from the ID variable are actually the estimates of the α_j in Equation 2.1, under the constraint that one of them is equal to 0.

I did this in Stata using the **reg** command,⁴ with results shown in the left-hand panel of Table 2.5. Only the coefficients for the first nine dummy variables are shown.

Table 2.5 Regression of Antisocial Behavior on Self-Esteem and Poverty, Dummy Variable Method

	<i>Fixed Effects</i>			<i>Conventional OLS</i>		
	<i>Coefficient</i>	<i>Standard Error</i>	<i>p</i>	<i>Coefficient</i>	<i>Standard Error</i>	<i>p</i>
SELF	−0.055	0.010	.00	−0.067	0.011	.00
POV	0.112	0.093	.23	0.518	0.079	.00
TIME_2	0.044	0.059	.45	0.051	0.090	.58
TIME_3	0.211	0.059	.00	0.223	0.091	.01
ID_2	−0.887	0.819	.28			
ID_3	4.131	0.811	.00			
ID_4	1.057	0.819	.20			
ID_5	−0.536	0.819	.51			
ID_6	0.040	0.820	.96			
ID_7	2.170	0.821	.01			
ID_8	0.910	0.820	.27			
ID_9	−0.276	0.819	.74			

Comparing the results in Table 2.5 with those in the last column of Table 2.3 (obtained with the first-difference method), we see that the coefficient and standard error for self-esteem are virtually identical. The coefficient for poverty is a little smaller using the dummy variable method, but it is far from significant with either method. The coefficients for TIME_2 and TIME_3 represent contrasts with the omitted category (TIME_1). We see that antisocial behavior increases, on average, over time, and that TIME_3 is significantly higher than TIME_1.

For comparison, the right-hand panel of Table 2.5 gives the OLS estimates of the coefficients without the inclusion of the 580 dummy variables. As we saw in the two-period case, the big difference in the results for the two methods is that the coefficient for POV is much larger for conventional OLS, and highly significant. Thus, the apparent effect of poverty on self-esteem disappears when we adjust for all between-person differences and focus only on within-person changes. It's also of some interest to compare the standard errors. The standard error for the POV coefficient is larger for the fixed effects estimate, a fairly typical result that stems from not using the between-person variation. On the other hand, for the coefficients for SELF and the two TIME dummies, the fixed effects standard errors are actually smaller than those for conventional OLS. Why the difference? It's all a matter of the relative magnitudes of within- and between-person variation. For POV, 70% of the variation is between persons, while for SELF

the figure is only 53%.⁵ For the TIME dummies, all the variation is within person and none is between. The best situation for a fixed effects analysis is when all the variation on a time-varying predictor is within persons, but there's still a lot of between-person variation on the response variable.

The problem with the dummy variable method is that the computational requirement of estimating coefficients for all the dummy variables can be quite burdensome, especially in large samples where it may be beyond the capacity of the software or the machine memory. Fortunately, there is an alternative algorithm-the mean deviation method-that produces exactly the same results. The one drawback is that it doesn't give estimates for the coefficients of the dummy variables representing different persons, but those are rarely of interest anyway.

The mean deviation algorithm works like this. For each person and for each time-varying variable (both response and predictor variables), we compute the means over time for that person:

$$\bar{y}_i = \frac{1}{n_i} \sum_t y_{it}$$

$$\bar{x}_i = \frac{1}{n_i} \sum_t x_{it}$$

where n_i is the number of measurements for person i . Then we subtract the person-specific means from the observed values of each variable:

$$y_{it}^* = y_{it} - \bar{y}_i$$

$$x_{it}^* = x_{it} - \bar{x}_i$$

Finally, we regress y^* on x^* , plus variables to represent the effect of time. This is sometimes called a “conditional” method because it conditions out the coefficients for the fixed effects dummy variables.

If you construct the deviation scores yourself and then use an ordinary regression program to estimate the coefficients, you will get the correct OLS estimates for all the coefficients. However, the standard errors and p values will not be correct. That's because the calculation of the degrees of freedom is based on the number of variables in the specified model, when it should actually include the number of dummy variables implicitly used to represent different persons in the sample (580 for the NLSY data). Formulas are available to correct the standard errors and p values (Judge, Hill, Griffiths, & Lee, 1985), but it's much easier to let the software do it for you. For example, the **xtreg** command⁶ in Stata does the correct calculations for a fixed effects model; SAS does it with the ABSORB statement in PROC GLM.

Using the **xtreg** command, I specified a fixed effects model (FE option) with ID as the variable that identifies records from the same person. Results were identical to the first five lines of Table 2.5. **xtreg** also reports several additional statistics that are specific to fixed effects models:

1.

An F test of the null hypothesis that all the coefficients for the fixed effects dummy variables are zero. In this case, the p value is less than .0001, so the null hypothesis can be confidently rejected. This is equivalent to saying that there is evidence for person-level unobserved heterogeneity. That is, there are stable differences in antisocial behavior between persons that are not fully accounted for by the measured predictor variables.

2.

An estimate of the proportion of variance in the dependent variable that is attributable to the fixed effects (the α_i s), labeled “rho (fraction of variance due to u_i).” In this case, the estimate is 0.64.

3.

An estimate of the correlation between the fixed effects at and $\alpha_i \hat{\beta} \mathbf{x}_{it}$, the estimated linear combination of the time-varying predictors. In a random effects model, this correlation is assumed to be 0. For these data, the correlation is .068.

4.

Three R^2 s: within, between, and overall. The within R^2 is just the usual R^2 calculated for the regression using the mean deviation variables. Here it's .033. The between R^2 is the squared correlation between the person-specific mean of y and the *predicted* person-specific mean of y . In this case, it's .041. Finally, the overall R^2 (.036 in this case) is the squared correlation between y itself and the predicted value of y . All three of these R^2 s are calculated using predicted values based on the estimated regression coefficients but *not* using the coefficients for the fixed effects dummy variables. If you include those (using the dummy variable method), the R^2 goes up to .73 for these data.

As noted before, one characteristic of this method is the inability to estimate coefficients for time-invariant predictors. This is evident from the fact that subtracting the person-specific mean of a time-invariant predictor from the individual values (which are the same at all periods) yields a value of 0 for all persons. Keep in mind, however, that we are still controlling all time-invariant predictors even though they drop out of the equation. In the next section, moreover, we'll see how to test whether the effects of those variables are themselves time invariant.

Interactions With Time in the Fixed Effects Method

In the two-period case, we saw how to extend the method to allow the coefficients for predictor variables to change over time. For a time-varying predictor, we simply added the variable measured at time 1 to the model. For a time-invariant predictor, we added the variable itself to the model. For the dummy variable method (or the equivalent mean deviation method), these extensions are accomplished by including interactions with time.

For the three-period NLSY data, Table 2.6 shows the results of including interactions between TIME (treated as categorical) and both time-varying and time-invariant predictors. Since TIME has three categories, there are two interactions with each predictor. Note that the time-invariant predictors do not have main effects included in the model. If we had tried to include them, the software would have dropped them from the model because they have no variation within persons.

For each of the interactions, the t statistic tests whether a coefficient at Time 2 or Time 3 is different from the coefficient at Time 1. Of the 18 interactions, only one (TIME_3*CHILDAGE) is statistically significant ($p = .024$). For that interaction, the coefficient of 0.227 indicates that the

Table 2.6 Interactions With Time

	<i>Coefficient</i>	<i>Standard Error</i>	<i>t</i>	<i>p</i>
TIME_2	0.291	1.245	0.23	.82
TIME_3	−0.444	1.258	−0.35	.72
SELF	−0.034	0.016	−2.08	.04
POV	0.097	0.130	0.75	.46
TIME_2 * SELF	−0.026	0.020	−1.28	.20
TIME_3 * SELF	−0.023	0.021	−1.09	.28
TIME_2 * POV	−0.112	0.152	−0.74	.46
TIME_3 * POV	0.099	0.155	0.64	.52
TIME_2 * BLACK	0.250	0.144	1.74	.08
TIME_3 * BLACK	−0.110	0.144	−0.77	.44
TIME_2 * HISPANIC	0.190	0.154	1.23	.22
TIME_3 * HISPANIC	0.075	0.153	0.49	.62
TIME_2 * CHILDAGE	0.076	0.100	0.76	.45
TIME_3 * CHILDAGE	0.227	0.100	2.26	.02
TIME_2 * MARRIED	−0.095	0.143	−0.67	.51
TIME_3 * MARRIED	−0.176	0.143	−1.23	.22
TIME_2 * GENDER	0.041	0.118	0.35	.73
TIME_3 * GENDER	0.107	0.118	0.91	.37
TIME_2 * MOMAGE	−0.027	0.028	−0.96	.34
TIME_3 * MOMAGE	−0.042	0.028	−1.52	.13
TIME_2 * MOMWORK	0.0137	0.131	1.05	.29
TIME_3 * SMOMWORK	−0.144	0.130	−1.11	.27

coefficient for CHILDAGE is 0.227 higher at Time 3 than at Time 1. Of course, with 18 tests it's a pretty good bet that at least one of them is statistically significant at the .05 level even if there's nothing really going on. A simultaneous test that all 18 interaction coefficients are equal to 0 yields a *p* value of .15.

Comparison With Random Effects Models

A popular alternative to the linear fixed effects model is the random effects or mixed model. This

model is based on the same equation that we used for the fixed effects model:

$$y_{it} = \mu_i + \beta x_{it} + \gamma z_i + \alpha_i + \varepsilon_{it} \quad (2.9)$$

The crucial difference is that now, instead of treating α_j as a set of fixed numbers (which is equivalent to treating α_j as random but with all possible correlations with x_{jt}), we assume that α_j is a set of random variables with a specified probability distribution. For example, it is typical to assume that each α_j is normally distributed with a mean of 0, constant variance, and is independent of all the other variables on the right-hand side of the equation.

There's a lot of software available to estimate the random effects model. SAS can do it with the MIXED procedure. The **xtreg** command in Stata does GLS estimation of the random effects model by default. Table 2.7 presents the estimates produced by **xtreg**, both with and without time-invariant predictors.

The fact that the random effects method can include time-invariant predictors is the most apparent difference between the fixed and the random effects models. In this case, however, we see that the inclusion of those variables doesn't make much difference in the coefficients for the time-varying predictors, self-esteem, and poverty.

Like the conventional OLS estimates in Table 2.5 and unlike the fixed effects estimates, both variables have coefficients that are highly significant. The similarity of random effects estimates and OLS estimates is not surprising. If the random effects assumption that α is uncorrelated with all other variables is correct, both methods produce consistent (and therefore approximately unbiased estimates) of the coefficients in Equation 2.9. But when that assumption is incorrect, both methods will yield biased estimates.

Why is it that POV is highly significant in the random effects model but far from significant in the fixed effects model? As explained earlier, whenever a coefficient is significant in a random effects model but not in a

Table 2.7 GLS Estimates of Random Effects Models

	<i>Coefficient</i>	<i>Standard Error</i>	<i>p</i>	<i>Coefficient</i>	<i>Standard Error</i>	<i>p</i>
SELF	−0.062	0.009	.00	−0.060	0.009	.00
POV	0.247	0.080	.00	0.296	0.077	.00
TIME_2	0.047	0.059	.42	0.047	0.059	.42
TIME_3	0.216	0.059	.00	0.216	0.059	.00
BLACK	0.227	0.126	.07			
HISPANIC	−0.218	0.138	.11			
CHILDAGE	0.088	0.091	.33			
MARRIED	−0.049	0.126	.70			
GENDER	−0.483	0.106	.00			
MOMAGE	−0.022	0.025	.39			
MOMWORK	0.261	0.115	.02			

fixed effects model, the first thing to do is compare the standard error estimates. Because fixed effects standard errors are often substantially higher than random effects standard errors, that alone could explain the higher p values. In this case, however, although the fixed effects standard error for POV is a little higher (0.09 vs. 0.08), that's not enough to account for the difference. Even if we substitute 0.08 for 0.09, the POV coefficient would still not be significant in the fixed effects model. Clearly, the main difference is in the magnitude of the two coefficients, 0.11 for fixed effects, and 0.25 (or 0.30) for random effects (depending on whether the time-invariant predictors are controlled or not). The most plausible explanation for this difference is that there are unobserved variables that “explain away” the observed association between poverty and antisocial behavior. When these variables are controlled, via fixed effects, the relationship disappears.

The key point here is that, contrary to popular belief, estimating a random effects model does not really “control” for unobserved heterogeneity. That's because the conventional random effects model assumes no correlation between the unobserved variables and the observed variables. The fixed effects model, on the other hand, allows for any correlations between time-invariant predictors and the time-varying predictors. It does so, however, at the cost of some efficiency in the event that those correlations are really zero.

It has been shown that the random effects model is actually a special case of the fixed effects model (Mundlak, 1978). That is, if you start with the conventional random effects model of Equation 2.9 and then allow for all possible correlations between the α_j term and the $\mathbf{x}_{jt}$

variables, you end up with something equivalent to the fixed effects model. In general, whenever there is a choice between two nested models, one being a restricted version of the other, there is a trade-off between bias and efficiency. The simpler model (the random effects model) will lead to more efficient estimates, but those estimates may be biased if the restrictions of the model are wrong. The more complex model (the fixed effects model) is less prone to bias but at the expense of greater sampling variability.

Given these trade-offs, it would be useful to have a statistical test that compared the random effects and fixed effects models. Such a test would help determine whether the biases inherent in the random effects method are small enough to ignore, or whether we need to move to the less restrictive fixed effects model. The best-known test is the Hausman (1978) test of the null hypothesis that the random effects coefficients are identical to the fixed effects coefficients.⁷ This test is available in several software packages. For the example at hand, the fairest test would be to compare the fixed effects coefficients in Table 2.5 with the random effects coefficients in the left-hand panel of Table 2.7, which control for several time-invariant variables. When I used this test with the **xtreg** command in Stata, I got a *p* value of .04, suggesting some evidence against the random effects model and in favor of the fixed effects model. In the next section, I consider an alternative test that may have somewhat better properties than the Hausman test currently implemented in Stata.

A Hybrid Method

We shall now consider an approach that combines some of the virtues of the fixed effects and random effects models. As we saw earlier, one way to estimate the fixed effects model is to express all variables as deviations from their person-specific means, then run an OLS regression on those mean deviation variables. In the hybrid method, the time-varying **x** variables are again transformed into deviations from their person-specific means, but the response variable *y* is not. Furthermore, unlike our previous fixed effects methods, we now include the time-invariant **z** variables in the regression model. In addition, we also include variables that are the person-specific means for each of the time-varying variables (which are also time invariant). Finally, instead of doing OLS regression, we estimate a random effects model to ensure that the standard errors reflect the dependence among the multiple observations for each person.⁸

Table 2.8 shows the results for the NLSY data. DSELF and DPOV are the deviation variables. MSELF and MPOV are the person-specific means. The first thing to notice in this table is that coefficients and standard errors for DSELF, DPOV, and the two time dummies are *identical* to

those we saw for the fixed effects method in Table 2.5. Thus, we have yet another method for producing fixed effects estimates.⁹ Actually, for DSELF and DPOV, it doesn't matter what time-invariant variables are in the equation. Even if we delete MSELF, MPOV, and the other time-invariant variables, the coefficients and standard errors for the time-varying variables will remain the same. Of course, what we get from this hybrid approach are estimates for the effects of the time-invariant variables, something that we couldn't obtain from the conventional fixed effects method.

Table 2.8 Hybrid Estimates

	<i>Coefficient</i>	<i>Standard Error</i>	<i>p</i>
DSELF	−0.055	0.010	.00
DPOV	0.112	0.093	.22
MSELF	−0.090	0.022	.00
MPOV	0.616	0.157	.00
BLACK	0.111	0.132	.40
HISPANIC	−0.280	0.139	.04
CHILDAGE	0.086	0.091	.35
MARRIED	−0.128	0.128	.32
GENDER	−0.508	0.107	.00
MOMAGE	−0.011	0.025	.65
MOMWORK	0.164	0.119	.17
TIME_2	0.044	0.059	.45
TIME_3	0.211	0.059	.00

In the multilevel model literature (Bryk & Raudenbusch, 1992; Goldstein, 1987; Kreft & De Leeuw, 1995), the practice of subtracting person-specific means from each time-varying variable is called *group mean centering*. Although it is well-known that group mean centering can produce substantially different results, this literature has not made the connection to fixed effects models, nor has it been recognized that group mean centering controls for all time-invariant predictors.

The estimates for the mean variables, MSELF and MPOV, are not particularly enlightening in themselves. But it's important to include these variables for two reasons. First, they help us get better estimates of the effects of the other time-invariant variables. Excluding MSELF and MPOV would mean that we are not fully controlling for these variables. Second, comparing their coefficients with those of the mean deviation variables, DSELF and DPOV, give us some

insight into what's going on here. If the assumptions of the random effects model are correct (i.e., the α_j term is uncorrelated with the $\mathbf{x}$ variables), then the deviation coefficient should be the same as the mean coefficient for each variable (apart from sampling variability). This isn't far from true for DSELF and MSELF. But the coefficient for MPOV is much larger than the coefficient for DPOV. When we estimated the conventional random effects model, what we got for SELF and POV were weighted averages of these “within” and “between” coefficients. This further suggests that we can test the random effects model against the fixed effects model by testing equality for these two pairs of coefficients (an alternative to the Hausman test discussed earlier). This is easy to do in Stata (using a Wald test) and yields a p value of .007, fairly clear-cut evidence against the random effects model.

Another attraction of the hybrid approach is that it can be extended in interesting ways that are not easily handled with the conventional methods for fixed effects estimation. The random effects models that we have been considering so far are all random *intercept* models. It's also possible, to estimate random *slope* models. For example, instead of assuming that the coefficient for DSELF is the same for everyone, we could assume that it is a random variable, and then estimate its mean and standard deviation. Such models are easily handled with the MIXED procedure in SAS, or with the **xtmixed** command in Stata. The latter produced an estimated mean of -0.055 for the coefficient of DSELF. The estimated standard deviation for the coefficient was 0.070. That's more than twice its standard error of 0.024, which strongly suggests that the effect of DSELF varies across persons.

Using the hybrid method, it's also possible to estimate models with more complex error structures (e.g., three-level structures or autoregressive structures) than the rather simple structure implied by the conventional fixed effects model. For further information on such models, see Singer and Willett (2003).

Summary

There are several equivalent ways to estimate linear fixed effects models for quantitative response variables:

- 1.

When each individual is observed at exactly two periods, compute the difference scores for all time-varying predictors. Then, do OLS regression of the difference score for the response variable on the difference scores for the predictors.

- 2.

For any number of periods, structure the data so that there is one record for each period for

each individual. Then do OLS regression with the inclusion of dummy variables for all individuals (excluding one).

3.

For the data structure in Method 2, transform all the time-varying variables into deviations from individual-specific means. Then do OLS regression on the deviation scores (with corrections for standard errors test statistics and *p* values). This is conveniently done in Stata with the **xtreg** command.

4.

For the data structure in Method 2, transform only the *predictor* variables into deviations from individual-specific means. Then estimate a random effects model, with predictor variables that include both the means and deviations from the means.

Of these methods, the fourth is the most flexible. It offers the following capabilities not shared with one or more of the other methods:

- The inclusion of predictor variables that do not vary within individuals
- A test of the fixed effects versus random effects assumption
- Random coefficients for those predictors that vary within individuals
- Less restrictive error structures

Regardless of which computational method is used, the fixed effects method effectively controls for all time-invariant predictors, both measured and unmeasured. This is its principal attraction as compared with random effects methods. A key assumption of the fixed effects method, however, is that the time-invariant predictors must have the same effects at all occasions. Variables whose effects are not constant across occasions must be explicitly included in the model. And, of course, the fixed effects approach does nothing to control for unmeasured predictors that vary over time.

Notes

1. I am indebted to Peter Tice for preparing this data set and making it available to me.
2. Because ANTI only takes on integer values and is positively skewed, an argument could be made for estimating ordered logit models rather than linear models. Indeed, we shall estimate such models for these data in Chapter 3. However, the qualitative results of the logit models are nearly the same as those for the linear models in this chapter.
3. A random effects model (the default in **xtreg**) is not appropriate in this case because the model allows only positive correlations between the error term. With multiple first-difference

equations, the correlations are often negative.

4. The **reg** command (and all other Stata commands discussed in this chapter) was used with the **xi** prefix in order to treat TIME and ID as categorical variables.

5. These numbers can be obtained by running an analysis of variance with each variable as the dependent variable and the ID variable as a categorical predictor.

6. Another Stata command that implements the mean-deviation method is **areg** with the **absorb(id)** option.

7. The Hausman test is computed as follows. Let **b** be the vector of fixed effects coefficients (excluding the constant) and let **β** be the corresponding vector of random effects coefficients. Let $\Sigma = \text{var}(\mathbf{b}) - \text{var}(\beta)$, where $\text{var}(\mathbf{b})$ is the estimated covariance matrix for **b** and similarly for **β**. The statistic is then $m = (\mathbf{b} - \beta)' \Sigma^{-1} (\mathbf{b} - \beta)$, which has a chi-square distribution under the null hypothesis.

8. The hybrid method described here is similar, but not identical, to methods proposed by Mundlak (1978) and Hausman (1978).

9. The estimates are identical only if the data set is balanced—that is, there is an equal number of observed periods for every individual. Otherwise the hybrid estimates will be slightly different from the conventional fixed effects estimates.

<http://dx.doi.org/10.4135/9781412993869.d4>