

Counterfactual Models of Neighborhood Effects: The Effect of Neighborhood Poverty on Dropping Out and Teenage Pregnancy¹

David J. Harding
Harvard University

This article investigates the causal effects of neighborhood on high school dropping out and teenage pregnancy within a counterfactual framework. It shows that when two groups of children, identical at age 10 on observed factors, experience different neighborhoods during adolescence, those in high-poverty neighborhoods are more likely to drop out of high school and have a teenage pregnancy than those in low-poverty neighborhoods. Causal inferences from such associations have been plagued by the possibility of selection bias. Using a new method for sensitivity analysis, these effects are shown to be robust to selection bias. Unobserved factors would have to be unreasonably strong to account for the associations between neighborhood and the outcomes.

INTRODUCTION

Over 35 years ago, Kenneth Clark described inner-city ghetto communities, isolated by poverty and racial segregation, as beset by a self-

¹ A previous version of this article was presented at the American Sociological Association Annual Meeting, August 16–19, 2002, Chicago. Some of the data used in this analysis are derived from sensitive data files of the PSID, obtained under special contractual arrangements designed to protect the anonymity of respondents. These data are not available from the author. Persons interested in obtaining PSID sensitive data files should contact psidhelp@isr.umich.edu. Christopher Winship, Christopher Jencks, Felix Elwert, Stephen Morgan, Robert Sampson, Stanley Lieberson, Dalton Conley, Katherine Newman, Nicholas Christakis, the *AJS* reviewers, and participants in the Harvard Proseminar on Inequality and Social Policy, the Harvard Applied Statistics Workshop, and the STICERD Work in Progress Seminar, London School of Economics, provided helpful comments on previous versions of this article. Cheri Minton provided excellent programming assistance. I am indebted to Christopher Winship and Felix Elwert for assistance in the development of the sensitivity analysis.

perpetuating “tangle of pathology,” which affected their residents’ physical, emotional, and mental health (Clark 1965). Contemporary social theory holds that concentrated-poverty neighborhoods have serious and lasting consequences for their residents and that, all other things being equal, poor children who grow up in high-poverty neighborhoods will experience significantly worse outcomes than poor children in more affluent communities. Though they come to different conclusions about the creation of concentrated-poverty neighborhoods (see also Jargowsky 1997; Quillian 1999), both Wilson (1987) and Massey and Denton (1993) agree that living in a neighborhood with concentrated poverty has serious consequences above and beyond those of growing up in a poor family because of the absence of role models, social isolation from job networks, weakened social institutions, and other factors. Neither, however, tests the neighborhood-effects hypothesis.

These theoretical advances ignited a sustained empirical search for evidence of neighborhood effects. The results thus far have been inconclusive. Social scientists have been unable to convincingly demonstrate the impact of neighborhood characteristics on individual outcomes using observational data (Jencks and Mayer 1990). Furstenberg and Hughes write, “Despite the intensity of interest in neighborhood influences, the conclusion that Jencks and Mayer reached several years ago remains: quantitative research has not demonstrated a convincing association between neighborhoods and children’s development, much less established the causal pathways between characteristics of neighborhoods and child development” (Furstenberg and Hughes 1997, p. 31). These difficulties are largely but not entirely the result of selection bias: unmeasured factors that affect both neighborhood of residence and individual outcomes that could potentially account for the association between neighborhood poverty level and outcomes.

This article provides a new approach to the study of neighborhood effects that (1) allows for the assessment of the sensitivity of results to selection bias and (2) explicitly defines neighborhood effects within a causal framework. Using the “counterfactual” model of causality, this article shows that neighborhood context has meaningful, statistically significant, and fairly robust effects on two critical adolescent outcomes: high school dropout and teenage pregnancy. I use propensity score matching

The Stata and Mathematica programs used to implement the analysis are available from the author. Preparation of this article was supported in part by a National Science Foundation Graduate Research Fellowship awarded to the author and by a National Science Foundation Integrative Graduate Education and Research Traineeship (IG-ERT) grant (98070661). Direct correspondence to David J. Harding, Department of Sociology, William James Hall, 33 Kirkland Street, Cambridge, Massachusetts 02138. E-mail: dharding@wjh.harvard.edu

estimators rather than traditional regression estimators with statistical controls. The framework allows for semiparametric and nonparametric estimators of causal effects by matching “treated” and “control” individuals on the probability of receiving the treatment: in this case, living in a high-poverty neighborhood. In addition, I use sensitivity analysis to assess the impact of selection bias from unobserved covariates on estimates of neighborhood effects. Intuitively, these methods can be thought of as estimating neighborhood effects by comparing individuals growing up in poor and nonpoor neighborhoods, who are otherwise identical on observable characteristics, and then testing the robustness of the results to the presence of unobserved covariates. This analysis capitalizes on the insight that a substantive understanding of the selection process helps one to address selection bias (Winship and Mare 1992).

I focus on teenage pregnancy and high school dropping out, two critical outcomes for any young person’s future adult life because of rising returns to schooling and the serious consequences of teenage pregnancy for a teen mother’s children and for her own future outcomes (Freeman 1997; Ashenfelter and Rouse 2000; Maynard 1997). Adolescence is the period in the life course in which neighborhood effects would become visible. It is the developmental stage in which a young person’s social world begins to incorporate peers and the larger community (Darling and Steinberg 1997). The theoretical mechanisms by which neighborhoods might influence adolescent outcomes have been extensively discussed and summarized elsewhere (e.g., Wilson 1987, 1996; Jencks and Mayer 1990; Small and Newman 2001). I focus here on the complexities of evaluating whether neighborhood context affects youth outcomes.

I begin by reviewing past research on neighborhood effects on youth outcomes from both experiments and observational studies, focusing on the problem of selection bias, and then consider recent research on the predictors of living in a high-poverty neighborhood—knowledge that will be important for constructing matched groups of treatment and control individuals. I then discuss the counterfactual causal framework and its application to neighborhood effects research, including the use of propensity score matching estimators and sensitivity analysis. Estimates of the effects of neighborhood poverty on high school dropping out and teenage pregnancy follow.

ESTIMATES OF NEIGHBORHOOD EFFECTS

Two types of studies have attempted to estimate the impact of neighborhoods on children’s outcomes: experiments and observational studies. Experiments assign subjects to poor and nonpoor neighborhoods at random

with respect to family and individual characteristics and therefore avoid the problems of selection bias that plague observational studies of neighborhood effects.

Mobility Experiments

Studies of two housing mobility programs, Gatreaux and Moving to Opportunity (MTO), provide the most convincing evidence to date of the existence of neighborhood effects. Gatreaux might be more accurately described as a “quasi-experiment” because random assignment to treatment and control was not explicitly built into the study design. Results suggest that those moving to middle-class suburbs subsequently have higher rates of labor force participation, earn higher wages and benefits, and are more likely to graduate high school, attend college, and attend a four-year college than those moving to city neighborhoods (Rosenbaum et al. 1993; Rosenbaum and Popkin 1991; J. Rosenbaum 1995; Rubinowitz and Rosenbaum 2000). However, as Katz, Kling, and Liebman (2001) point out, because Gatreaux participants were contacted years after their initial moves and because only a small sample was studied, selective retention could have created an unrepresentative sample, potentially biasing results upward.

Encouraged by the success of the Gatreaux program and desiring a better test of its findings, the Department of Housing and Urban Development started the MTO demonstration project in 1994. Preliminary results from Baltimore, Boston, and Los Angeles suggest that in some cities treatment groups experienced higher labor force participation, better health care and child care, better mental and physical health, fewer behavior problems among children, less juvenile crime, higher test scores, and better schools, but they also were more likely to be held back in school and develop school discipline problems, and adults had fewer contacts with neighborhood friends (Ladd and Ludwig 1997; Hanratty, McLanahan, and Pettit 1998; Katz et al. 2001; Ludwig, Duncan, and Hirschfield 2001; Ludwig, Ladd, and Duncan 2001).

Observational Studies

While social experiments are undoubtedly a valuable tool in measuring neighborhood effects, their high cost and the difficulty of implementation ensure they are rare. In addition, since experiments can, ethically, only move people from less desirable to more desirable neighborhoods, they are restricted to low-income samples. Experiments such as MTO and Gatreaux are also based on unrepresentative samples of those living in high-poverty neighborhoods. Only those who self-select into the program

are randomized into treatment and control groups. Hence, studies using observational data are also needed for research on neighborhood effects.

The Panel Study of Income Dynamics (PSID) is the most commonly used longitudinal data set for investigating neighborhood effects. Table 1 shows unadjusted rates of high school dropping out and teenage pregnancy by race and neighborhood poverty among a sample of PSID respondents. For both race groups and both outcomes, high rates of neighborhood poverty *appear* to result in more negative outcomes. Base rates of high school dropping out and teenage pregnancy approximately double moving from low- to moderate-poverty neighborhoods and increase again by one-quarter to one-half moving from moderate- to high-poverty neighborhoods.

The crucial question is whether these differences by neighborhood can be causally attributed to neighborhood context or whether they are simply due to differences between individuals living in different neighborhoods. Because differences between individuals and families in high-poverty neighborhoods and those in low-poverty neighborhoods will bias estimates of neighborhood effects, it is necessary to control for such differences. This problem is variously referred to in different disciplines as selection bias, omitted variable bias, hidden bias, confounding, or unobserved heterogeneity.

Estimates of neighborhood effects on high school graduation, years of schooling, and teenage nonmarital childbearing from standard regression models are extremely sensitive to the individual and family characteristics for which one controls, with strong effects when no individual and family characteristics are controlled and smaller and often nonsignificant effects when an extensive set of individual and family attributes are controlled (Ginther, Haveman, and Wolfe 2000). It is no surprise then that studies with different model specifications come to contradictory conclusions about the importance of neighborhoods for child development (e.g., compare Brooks-Gunn et al. [1993] with Evans, Oates, and Schwab [1992]).

Clearly a more theoretically informed model of the relationship between family and individual characteristics and neighborhood effects is needed. In particular, controlling for attributes that may themselves be affected by neighborhood characteristics may bias estimates of neighborhood effects. For instance, theory suggests that neighborhood poverty and segregation negatively influence family income (Sampson and Morenoff 1997), so controlling for family income measured during or after neighborhood context is measured likely downwardly biases estimates of neighborhood effects.

In addition to basic regression models with statistical controls for covariates, two other techniques have been used to remove selection bias from observational studies of neighborhood effects. Aaronson (1997, 1998)

TABLE 1
HIGH SCHOOL DROPPING OUT AND TEENAGE PREGNANCY RATES BY RACE AND
NEIGHBORHOOD POVERTY RATE AMONG PSID RESPONDENTS BORN 1958–77

NEIGHBORHOOD POVERTY RATE	HIGH SCHOOL DROPPING OUT (%)		TEENAGE PREGNANCY (%)	
	Nonblack	Black	Nonblack	Black
Low:				
0%–9.99%	12 (1,574)	13 (206)	5 (911)	13 (135)
Moderate:				
10%–19.99% ...	23 (667)	27 (587)	13 (434)	27 (386)
High:				
20%–29.99% ...	34 (241)	32 (704)	18 (139)	30 (452)
30%–39.99% ...	36 (67)	33 (485)	20 (49)	28 (312)
40%+	36 (22)	40 (421)	21 (14)	32 (293)
Total	18 (2,571)	31 (2,403)	9 (1,547)	28 (1,578)

NOTE.—Unweighted numbers of cases appear in parentheses.

uses sibling fixed-effects models, finding that a 10% increase in proportion of young adults in the neighborhood who are high school dropouts increases the likelihood of dropout by 3.6%. As Aaronson points out, this method has two drawbacks. First, fixed effects estimators often have large standard errors. Second, and more important, the fixed effect estimator does not control for unobserved family characteristics that vary over time and therefore differ between siblings.

Duncan, Connell, and Klebanov (1997) use instrumental variables (IVs) to estimate neighborhood effects on schooling (see also Evans et al. 1992), with mother's neighborhood characteristics after the child leaves home serving as an instrument for child's neighborhood characteristics. This estimator relies on the assumption that mother's neighborhood is correlated with child's neighborhood but not with unobserved dimensions of mother's parenting that may be correlated with the child's neighborhood.

However, IV estimators have three drawbacks. First, like fixed-effect estimators, they tend to have large standard errors. Second, they are often based on assumptions that are either implausible or untestable. For example, Duncan et al. (1997) note that their analysis assumes that parents are very responsive to changing housing needs—parents will change neighborhoods after their children leave home if their neighborhood preferences have changed. Third, IV estimators only capture the effect of the

treatment on the subset of the sample that is on the margin, those whose treatment assignment is changed by the instrument. This will bias estimates of neighborhood effects if this subset is different from the entire sample (Imbens and Angrist 1994). However, if one is interested in policy effects, those on the margin are in fact often the group of substantive interest.

None of these methods (controls in regression models, sibling fixed effects, and IV) is entirely satisfactory for dealing with selection bias in neighborhood effects. This study uses another method, propensity score matching with sensitivity analysis, to address the selection bias problem. It capitalizes on knowledge of how subjects are assigned to treatment (in this case, how people come to live in a high-poverty neighborhood) and assesses sensitivity of results to selection bias. Propensity score matching offers advantages over more traditional regression techniques: matching estimators are nonparametric, generally more efficient, largely avoid problems of multicollinearity, and ensure that treatment and control cases are reasonable comparisons (these are discussed further below). However, none of these advantages directly deal with selection bias. Indeed, matching estimators make the same assumption as regression estimators, that treatment assignment is “strongly ignorable” net of the variables controlled for or matched. However, propensity score matching sets up a relatively simple method of assessing the sensitivity of results to the presence of an unobserved covariate, allowing the analyst to examine how results change given various assumptions about the magnitude of selection bias.

The next section reviews recent research on the determinants of living in a high-poverty neighborhood, which will inform the estimation of neighborhood effects below.

DETERMINANTS OF LIVING IN A HIGH-POVERTY NEIGHBORHOOD

There has been considerable research on residential mobility, neighborhood choice, and the determinants of living in a poor neighborhood. This knowledge will be critical in constructing matched groups of individuals in low- and high-poverty neighborhoods in the analysis below, so I briefly summarize the literature here.

First, life cycle plays an important role in mobility and housing decisions (Rossi 1980). Mobility has been linked to marital status (Speare and Goldscheider 1987; South and Crowder 1998a); age of parents and age of children (South and Deane 1993; South and Crowder 1997b, 1998a); home ownership (South and Deane 1993; South and Crowder 1998a); and ed-

ucation, income, and employment (South and Crowder 1997*a*, 1997*b*, 1998*a*).

Second, households choose to reside in neighborhoods, communities, and jurisdictions based on an assessment of the taxes and services of the area and their own needs (Tiebout 1956). Though the evidence on the Tiebout hypothesis is mixed (Dowding, John, and Biggs 1994; Rhode and Strumpf 2000), mobility decisions do seem to reflect some key area factors such as school quality (Figlio and Lucas 2000), crime rates and victimization (Cullen and Levitt 1996; Dugan 1999), and changes in neighborhood composition (Lee, Oropesa, and Kanan 1994; Crowder 2000).

Third, race limits residential choice among minorities, especially blacks. Audit studies document discrimination in all sectors of the real estate market (Yinger 1995). Net of background factors, blacks are less likely than whites to move and less likely to improve their housing and neighborhood environments when they do move (South and Deane 1993; Cook and Bruin 1993; St. John, Edwards, and Wenk 1995). Whites are especially likely to leave neighborhoods with growing minority populations and neighborhoods with multiple minorities (Denton and Massey 1991; Crowder 2000), and housing values are considerably lower in black neighborhoods than in white neighborhoods (Harris 1999). However, there is also some evidence that the role of race is declining (Tolnay, Crowder, and Adelman 2000; Farley and Frey 1994; Frey and Farley 1996).

Finally, a small but growing body of literature has focused on the mobility behavior of poor persons and those living in high-poverty neighborhoods, highlighting the importance of race, education, and changes in marital status and labor market status. Blacks and Hispanics of all income levels are considerably more likely than whites to live in high-poverty neighborhoods (Jargowsky 1997). Approximately one-fourth of persistently poor adults enter or leave a census tract with a poverty rate over 20% each year, and whites and families without children are more likely than blacks and families with children to leave poor neighborhoods (Gramlich, Laren, and Sealand 1992; South and Crowder 1997*a*). Among both blacks and whites, those who are recently married, have more education, draw higher income, are renters, and do not receive public assistance are more likely to move from poor to nonpoor neighborhoods, while those who are recently divorced, are renters, have less education, receive lower income, and have recently become unemployed are more likely to move from nonpoor to poor neighborhoods (South and Crowder 1997*a*).

Metropolitan-level factors also play an important role in determining whether individuals live in high-poverty neighborhoods. An urban area's mean income, income inequality, and degree of spatial segregation by race and by income are associated with the extent of neighborhood poverty

(Jargowsky 1997) and therefore also with the probability that an area's residents will live in a high-poverty neighborhood.² Movement between poor and nonpoor neighborhoods by both blacks and whites is related to a metropolitan area's racial and economic segregation, vacant housing, and new housing (South and Crowder 1997*a*, 1998*b*).

METHODS

This study uses the counterfactual causal framework, employing propensity score matching and sensitivity analysis.

The Counterfactual Causal Framework

I focus here on providing an intuitive rather than formal explanation of the counterfactual causal framework (Rubin 1974, 1977, 1991; Rosenbaum and Rubin 1983*a*, 1983*b*, 1984, 1985; Rosenbaum 1984*a*, 1984*b*, 1995; Holland 1986; formal reviews can be found in Winship and Morgan [1999] or Winship and Sobel [in press]). This framework borrows both the logic and language of experiments. A causal effect is defined as the difference in outcome between the world in which the subject receives the treatment and the "counterfactual" world in which the same individual does not. In the case of effects of neighborhood poverty, the treatment is exposure to a high-poverty neighborhood. Clearly, a single subject cannot simultaneously experience and not experience the treatment, so it is necessary to use additional data to fill this information gap. This problem is referred to as the "fundamental problem of causal inference" (Holland 1986).

The solution used here is to match each treated subject with one or more control subjects such that the treated subjects are, on average, identical to the control subjects on observable characteristics prior to treatment.³ Control subjects serve as the counterfactual. For example, in a randomized experiment, randomly assigning individuals to treatment and control groups greatly increases the probability that the two groups will be well matched in any particular sample and ensures that the difference between treatment and control groups will be an unbiased estimator of the population average treatment effect. Treatment and control groups are well matched when subject characteristics that affect the outcome are "balanced" in the treatment and control groups. These characteristics are

² This is, of course, true by definition.

³ One might also match multiple treatment subjects to single controls when treatment subjects are more numerous than control subjects.

often referred to as “covariates” or “controls.”⁴ For example, in the literature on neighborhood effects, family income is one covariate that might bias estimates of neighborhood effects since it affects the probability of living in a high-poverty neighborhood and is thought to affect most outcomes experienced by young people. An important difference between a true randomized experiment and an observational study in which matched treatment and control groups are constructed is that with randomization we can expect balance on *both* observed and unobserved characteristics. When we construct matched groups, we can only expect balance on observed and matched covariates.

Under all but trivial conditions, estimates of effects in this framework are the “effect of treatment on the treated” rather than the effect of treatment for the entire population.⁵ In other words, in estimating the effect of growing up in a high-poverty neighborhood, we are estimating the effect on those individuals in the data who actually grew up in a high-poverty neighborhood, not the hypothetical effect of a high-poverty neighborhood on anyone who could conceivably live in one. In addition, the estimates from the counterfactual model are estimates of the average treatment effect, rather than the effect on each individual. As with regression analysis, one should not assume that all individuals experiencing the treatment are affected equally.

Employing the counterfactual causal framework requires careful consideration of the definition of the treatment and the appropriate covariates on which to match. Figure 1 illustrates the identification strategy of matching within the counterfactual framework.⁶ I define the treatment as neighborhood poverty during adolescence (between the ages of 11 and 20) and match on covariates measured at age 10 or before. One should only match on covariates measured before the treatment is experienced. The exception is a covariate that the treatment could not conceivably affect, such as a covariate that remains constant over time, like birth year, race, or gender. *If the treatment and control groups are truly identical prior to treatment, then any differences between the two groups after treatment must be an effect of the treatment.* In other words, if the treatment and control groups differ on a covariate *during the treatment period*, then, under the assumption of no unobserved covariates, that (statistically significant) dif-

⁴ Technically, a covariate is balanced when treatment and control groups have identical distributions on the control variables.

⁵ This is also the case for regression estimates and, with some strong assumptions, for IV estimates as well; see Angrist, Imbens, and Rubin (1996).

⁶ Note that carefully defining the treatment and selecting and appropriately measuring the covariates should help to produce better analyses no matter which estimation methods are used.

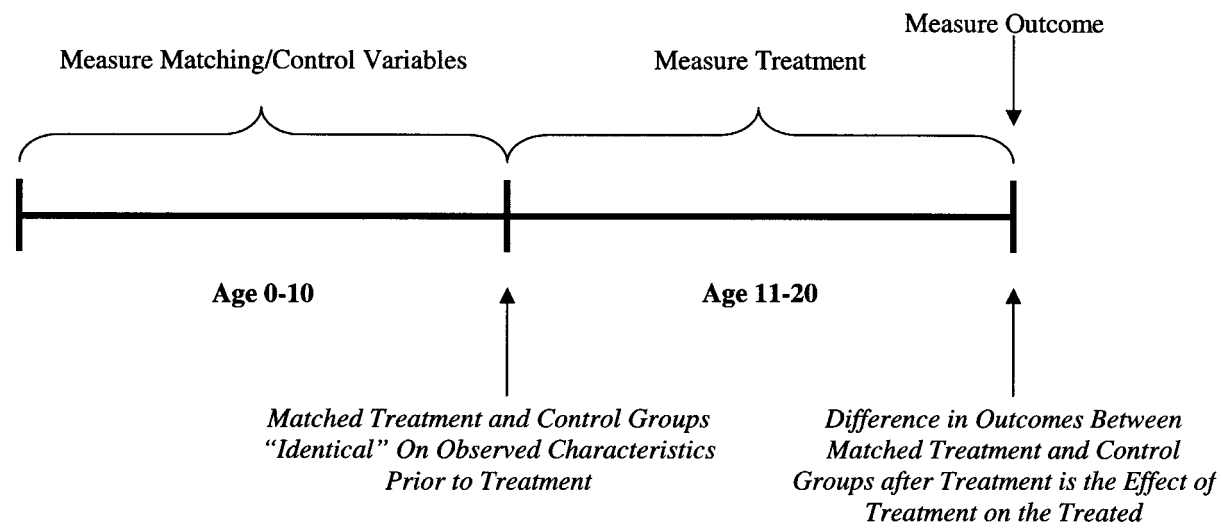


FIG. 1.—Identifying the causal effect of neighborhood poverty during adolescence in the counterfactual framework

ference must be a result of the treatment since the only difference between the treatment and control groups prior to treatment is the treatment assignment itself. In the sensitivity analysis described below, I relax the assumption that there are no unobserved covariates.

Matching on a covariate that is also affected by the treatment biases the treatment effect estimate (Rosenbaum 1984a). In the estimation of neighborhood effects, controlling for covariates such as family income during the treatment period will bias downward the neighborhood effect if neighborhood conditions affect family income and family income affects the outcome.⁷ In other words, controlling for (or matching on) family income during the treatment period removes the portion of the effect of neighborhood poverty on the outcome that operates through family income.⁸ Note that the counterfactual definition of the causal effect of neighborhood poverty during adolescence leads to the policy-relevant estimate and the estimate that is most comparable to estimates of effects from MTO and Gatreux. Policies to move adolescents from high-poverty to low-poverty neighborhoods would certainly include moving an adolescent's entire family, and therefore any neighborhood effects that operate through the family should be considered in assessing the benefits of such policies.⁹

One potential disadvantage of the counterfactual framework is that it is not always possible to find suitable matches for all treatment cases. When this occurs, one can only estimate the treatment effect for the subset of the treated cases that are matched. There is a very real trade-off between estimating the effect for the entire treatment group but having little knowledge about the robustness of the results and estimating the effect for only some of the treatment group but understanding the robustness of the results.

Propensity Score Matching

In an ideal situation, one would match control cases to treatment cases by matching exactly on all observed covariates, such that each treatment-control pair has the same values on all observed covariates. This is known

⁷ One should, however, match on family income prior to treatment. The direction of the bias associated with controlling for an intervening mechanism depends on the relationships between the mechanism and the treatment and the outcome.

⁸ If one were only interested in the portion of the effect that operates through mechanisms other than family income, then matching on family income during adolescence might be appropriate. However, it is not clear how such an estimate should be interpreted theoretically.

⁹ To my knowledge, South and Crowder (1999) is the only other neighborhood effects article using observational data that measures controls prior to treatment.

as exact matching. In practice, the limited size of social science data sets and the need to match on many continuous covariates makes it impossible to find control group cases to match exactly to every treatment case. An alternative is to create a summary measure of the probability of receiving treatment. This is known as the estimated propensity score. Below, I use a combination of these two methods, matching exactly on gender and race and then matching on estimated propensity scores within gender-race groups. Previous empirical work in sociology using propensity score matching can be found in Smith (1997) and Morgan (2001).

The true propensity score is defined as the (unknown) probability that a particular subject will receive the treatment. Rosenbaum and Rubin (1983b) show that matching treated and control subjects on their true propensity scores will, in expectation, result in balance on all covariates, both observed and unobserved. If this is the case, the difference between treatment and control groups after treatment will be an unbiased estimator of the average treatment effect. In practice, true propensity scores are not known and so must be estimated based on observed covariates, often with a logit or probit model. I return to the issue of unobserved covariates below in the section on sensitivity analysis.

The general method is relatively simple. *First*, a logit model is estimated with all covariates predicting whether an individual receives the treatment. My selection of covariates is based on the literature reviewed above.

Second, the predicted probabilities of receiving the treatment from the logit model are calculated. These predicted probabilities are the propensity scores.

Third, treated subjects are matched to controls based on their propensity scores. Many algorithms for matching exist, but I use a variant of the computationally simple algorithm called “nearest available pair matching,” which balances covariates well under most conditions (P. Rosenbaum 1995). I allow control cases to serve as a match for multiple treatment cases and constrain matches to have differences in treatment probabilities of less than two percentage points.¹⁰

Fourth, observed covariates are checked for balance. Because matching on the propensity score only creates covariate balance in expectation (i.e., over repeated samples), it is important to verify balance for the one sample used in any particular analysis. In addition, when the true functional form of the propensity score model is not clear, as is almost always the case, verifying balance ensures that the propensity score model is not grossly misspecified. If the covariate balance is not verified, the problem of functional form is simply shifted from the standard model of the outcome to the propensity score model. One commonly used measure of covariate

¹⁰ This is often referred to as “caliper” matching, where 2% is the caliper size.

balance is the standardized bias (e.g., Rosenbaum and Rubin 1985; Rubin 1991; D'Agostino and Rubin 2000), which is a standardized version of the difference in means of a single covariate, X , for the treatment and control groups:

$$B = \frac{|\bar{X}_T - \bar{X}_C|}{\sqrt{(S_T^2 + S_C^2)/2}},$$

where S^2 is the variance, and the subscripts T and C denote treatment and control groups. The standard bias basically measures the difference in means between the treatment and control group in terms of the number of standard deviations it is away from zero.¹¹ The standard deviation in the denominator is similar to an average of the standard deviation of the covariate in the treatment group and control group. Note that when the variances are small, as is often the case when observations are selected to be similar to one another, even a small difference in the covariate means can create a sizable standardized bias.

Fifth, treated and control groups are compared on the outcomes. In the analysis below, I use a version of the χ^2 statistic for a 2×2 table (treatment by dropout or treatment by pregnancy) that corrects for clustering in the data.¹² I also check statistically significant effects with logistic regression models on the matched sample predicting the outcome with the treatment variable and controlling matched covariates. This ensures that imperfect matching of covariates is not driving the estimated effects.

When all of the assumptions of regression analysis are met, particularly correct functional form and well-supported data, propensity score matching alone offers little advantage over traditional regression methods.¹³ However, when one cannot be sure that these assumptions are met, propensity score matching has several advantages over traditional regression methods.

First, because it makes no assumptions about the functional form of the relationship between covariates and the outcome, a matching estimator is nonparametric. While decisions must be made about the functional form of the model predicting treatment, one can also verify that the covariates are actually balanced by the matching procedure. When

¹¹ Note that the standardized bias is not a formal statistical test for covariate balance but rather simply a measure of covariate balance.

¹² I use Stata's `svytab` command.

¹³ One could imagine a model specification in which all variables were entered into the model as a series of dummy variables and all variables were fully interacted with all others, effectively making the estimation nonparametric and making functional form irrelevant. With the exception of the ease of employing the sensitivity analysis discussed below, matching would provide few if any advantages over such a model.

covariate balance is not achieved, the propensity score model can be respecified and balance reevaluated.

Second, matching ensures that comparisons between treatment and control groups are reasonable (Winship and Sobel, in press). The matching procedure itself provides important information about the comparability of treated and untreated cases by examining how well treated cases can be matched. The noncomparability of cases is often masked in estimates from regression models. In investigating neighborhood effects, it makes little sense to compare a young black person in a poor family in the inner city to a white young person in an upper-middle-class family in the suburbs, because the probability of either individual growing up in the other's neighborhood is extremely close to zero. This is the same as saying that the "supports in the data do not overlap"; there are certain combinations of independent variables for which there are no observations, which can bias estimates of causal effects (Heckman et al. 1998).

Third, matching estimators are generally more efficient (smaller standard errors) because fewer parameters are estimated (Winship and Sobel, in press). However, the efficiency gain is sometimes lessened by the need to drop cases that cannot be matched.

Fourth, one can estimate the propensity score without worrying about collinearity among the covariates because only the predicted values from the propensity score model are needed, not the coefficients. This avoids the problem of having too many controls in a regression model (Lieberman 1985).

One potential disadvantage of the propensity score approach is that it requires that treatment variables be binary. A propensity score is by definition a probability of being in either one state or another. The loss of information caused by this requirement will weaken statistical power, but it cannot be avoided here.¹⁴ Dichotomization also obscures nonlinear effects, which have been found previously in some studies of neighborhood effects (Crane 1991; South and Crowder 1999). In the analysis below, I group individuals into categories based on the mean poverty rate of the census tracts in which they lived during adolescence and make comparisons across categories. Propensity score estimation and matching are done separately for each of these comparisons, for each outcome, and for each racial group.

Propensity score matching by itself does not solve the problem of se-

¹⁴ To avoid this problem, one might imagine calculating a predicted neighborhood poverty rate based on an OLS regression and then matching cases with sufficiently different actual neighborhood poverty rates, and then comparing matched cases. Unfortunately, the statistical properties of such an estimator have not been investigated, so I do not use this approach here.

lection bias. Any unobserved covariate that is not highly correlated with the observed covariates could differ between treatment and control groups, biasing the estimates of the treatment effect. This amounts to omitted variable bias in the model of the propensity score. Sensitivity analysis can be used to test the robustness of results to selection bias due to an omitted covariate. One of the main advantages of propensity score matching and the counterfactual causal framework is the ease with which such a sensitivity analysis can be employed.

Sensitivity Analysis

The goal of the sensitivity analysis here is to assess how an unobserved covariate that affects both neighborhood choice and the outcome (either high school dropout or teenage pregnancy) would alter our conclusions about the neighborhood effect. Such hypothetical and unobserved covariates might include “parent’s commitment to children,” “future orientation,” “wealth,” or any similar factor that affects both the family’s decision of where to live and the probability of the child dropping out of high school or having a teenage pregnancy. For example, if adolescents living in high-poverty neighborhoods had parents who were less committed to their role as parents than otherwise similar adolescents in low-poverty neighborhoods and parental commitment affected high school dropout, then failing to account for differences in parental commitment would bias our estimates of effect of living in a high-poverty neighborhood. Sensitivity analysis asks, How do our inferences change given various hypothetical unobserved covariates?¹⁵

One approach might assess how the point estimate of the effect changed with the inclusion of hypothetical unobserved covariates of varying strengths. However, this approach would ignore sampling error. A second approach might examine how the statistical significance of a point estimate changed and therefore incorporates sampling error (P. Rosenbaum 1995). A third approach, that taken here, does both. It examines how the point estimate and its confidence interval change under the presence of an

¹⁵ Note that for an unobserved variable to be a potential source of selection bias, it must affect whether an individual receives the treatment and must affect the outcome. In particular, an unobserved variable that distinguishes between different subgroups in the treatment group, but does not have a residual effect on the outcome net of the variables already controlled, is not a threat to the robustness of the results. Such a variable is better understood in terms of “treatment effect heterogeneity,” or differences in the treatment effect between subgroups receiving treatment. The analyses in this article attempt to estimate the average total treatment effect among the treated and are not concerned with treatment effect heterogeneity, though treatment effect heterogeneity is an important theoretical issue in neighborhood effects.

unobserved covariate. The sensitivity analysis presented here is based on Rosenbaum and Rubin (1983a).¹⁶

Propensity score matching on data with a binary treatment and a binary outcome leads to a 2×2 table with the treatment (X) on one axis and the outcome (Y) on the other. A measure of the effect of treatment on the outcome with desirable properties is the odds ratio, or the odds of the outcome under the treatment condition ($X = 1$) divided by the odds of the outcome under the control condition ($X = 0$; see Fienberg 1980). Following P. Rosenbaum (1995) and Gastwirth, Krieger, and Rosenbaum (1998), I define an unobserved binary covariate, U , that affects both treatment assignment and outcome, denoting the effect of U on treatment assignment as Γ and the effect of U on the outcome as Δ , both expressed as odds ratios. The sensitivity analysis will involve calculating new estimates of the treatment effect for a series of combinations of specific values of these two sensitivity parameters.

Figure 2 shows the observed table and the latent table, a $2 \times 2 \times 2$ table that we would observe if we could observe U . Capital letters A through H indicate counts. (Note that the observed table is simply the latent table collapsed over U .) If we knew the counts in the latent table, we could estimate the effect of X on Y controlling for U , the “pure” treatment effect net of the unobserved covariate.¹⁷ Hence, the goal here is to determine these cell counts.

The $2 \times 2 \times 2$ latent table is uniquely defined by eight parameters: the grand mean, one-way marginals for each of the three variables, the three two-way associations between the variables (XY, UX, UY), and the three-way interaction between the three variables (UXY). By specifying these eight parameters, we can produce the latent table. Four of these parameters are provided by the observed table: the grand mean, the one-way marginals for X and Y , and the two-way association between X and Y . Two more are provided by Γ and Δ , which are the associations between U and X and between U and Y , respectively. We must therefore make assumptions about the one-way marginal for U and about the interaction term. Here I assume that both are zero.

When the one-way marginal for U is zero, cases are approximately evenly distributed between the top and bottom subtables in the latent table. This is a trivial assumption because changing this assumption would

¹⁶ Frank (2000) develops a similar sensitivity analysis for selection bias for analyses using ordinary least squares regression.

¹⁷ We can, e.g., create a data matrix with eight observations, one for each combination of X , Y , and U , with a frequency weight equal to the corresponding cell count, and then estimate a logit model: $\text{logit}(Y) = b_0 + b_1X + b_2U + e$. The desired treatment effect odds ratio is then $\exp(b_1)$, and its confidence interval can be easily calculated.

Observed Table:

	Y=0	Y=1
X=0	A+E	B+F
X=1	C+G	D+H

Latent Table:

		Y=0	Y=1
U=0	X=0	A	B
	X=1	C	D
U=1	X=0	E	F
	X=1	G	H

FIG. 2.—Observed and latent tables for sensitivity analysis

simply shift cases between these subtables without implication for the relationship between X and Y . Assuming that there is no interaction implies that we have assumed that the relationship between X and Y does not vary for different values of U , that the relationship between U and X does not vary by Y , and that the relationship between U and Y does not vary by X . In terms of substantive intuition, the most important part of this assumption is that the effect of the treatment does not depend on the value of the unobserved covariate. In other words, the effect of treatment on the outcome is the same in the top and bottom subtables of the latent table in figure 2. This assumption is routinely made when controlling for a variable in a regression model, but when the variable is observed, it can be relaxed in a regression model by adding an interaction term (e.g., $U \times X$). Finally, this assumption also means that Γ (the relationship between U and X) is the same for both values of Y and that Δ (the relationship between U and Y) is the same for both values of X .

Together these two assumptions correspond to five equations specifying the relationships between the cell counts and sensitivity parameters:¹⁸

$$\begin{aligned}(\ln A + \ln F) - (\ln B + \ln E) &= \ln \Delta, \\(\ln C + \ln H) - (\ln D + \ln G) &= \ln \Delta, \\(\ln A + \ln G) - (\ln C + \ln E) &= \ln \Gamma, \\(\ln B + \ln H) - (\ln F + \ln D) &= \ln \Gamma, \\A + B + C + D &= E + F + G + H.\end{aligned}$$

These equations, along with four equations relating the cell counts in the observed table (e.g., $A + E$) to the actual observed counts, provide more than enough equations to solve for the eight cell counts in the latent table. Using mathematical software such as Mathematica, we can easily solve the system of equations for the eight cell counts given specific values of the sensitivity parameters.¹⁹ By systematically varying the sensitivity parameters and resolving the system of equations to produce various latent tables and then reestimating the treatment effect for each latent table, we can examine how the results change under varying assumptions about the unobserved covariate.²⁰

¹⁸ Readers more familiar with thinking of the odds ratio as a cross-product ratio may find equations that do not contain logarithms to be more intuitive: $AF/BE = \Delta$, $CH/DG = \Delta$, $AG/CE = \Gamma$, $BH/FD = \Gamma$; $A + B + C + D = E + F + G + H$.

¹⁹ In fact, we have one more equation than we need, so one of the first four equations can be dropped.

²⁰ I calculate the treatment effect and its confidence interval from the latent table using a logit model as described in n. 17, above. However, because I have no information about which cases in the cells of the observed table end up in the corresponding cells in the latent table, I cannot easily correct for clustering in the data (as is necessary because the PSID data is clustered by family). Some assumption about how cases are split must be made. If we assume no relationship between the unobserved covariate and the clustering (i.e., that cases from each cell of the observed table are assigned at random to the two corresponding cells in the latent table), then controlling for the unobserved covariate should not reduce the standard errors. Another assumption would be a strong relationship between the unobserved covariate and the clustering, which would suggest that cases from a cell in the observed table that are in the same cluster would appear in the same cell in the latent table. Under this assumption, controlling for the unobserved covariate would reduce the SEs from those in the observed table. I use the first assumption because it is the most conservative assumption, in that it suggests the larger SEs and thus the wider confidence interval. To calculate SEs for the latent table that take into account clustering, I multiply the regular SEs by an inflation factor. This inflation factor is the ratio of the clustered SEs in the observed table to the SEs in the observed table when clustering is ignored. A simple simulation exercise showed this to be a reasonable procedure. (I use Stata's clustered SEs option on the logit command.)

DATA

The PSID is a nationally representative longitudinal annual survey begun in 1968 with 4,800 families including about 18,000 individuals. New families formed by survey respondents were also followed, resulting in 6,434 families and over 60,000 individuals by 1999. The PSID contains extensive data on economic and demographic variables.

The PSID is especially well suited to the investigation of neighborhood effects. First, the PSID Geocode Match Files allow the user to attach geographic census data, such as poverty rates for census tracts and zip codes, from the 1970, 1980, and 1990 U.S. Censuses of Population and Housing to family records to measure neighborhood contextual variables. Second, the PSID contains an oversample of low-income families, ensuring a sizable group of respondents living in high-poverty neighborhoods. Third, selective attrition in the PSID panel is low (Duncan, Juster, and Morgan 1984), which is especially relevant here because of the importance of longitudinal data to applying the counterfactual model to neighborhood effects.

I use PSID data from 1968 through 1997, when Geocode data are available. I select respondents born between 1958 and 1977 for whom I could determine values on one or both outcome variables and who could be matched to a valid census tract for at least one year between the ages of 11 and 20. These respondents are age 10 between 1968 and 1987 and age 20 between 1977 and 1997. Variables are listed in table 2.²¹ All analyses take into account clustering of respondents by family in the PSID data.

The two outcome variables are high school dropout and teenage pregnancy. High school dropout is coded “1” if the respondent does not graduate from high school by age 20, regardless of whether he or she is still enrolled at that point and regardless of whether he or she received a general equivalency diploma (GED), and “0” otherwise. For females, teenage pregnancy is coded “1” if the respondent has a nonmarital live birth before age 20 and “0” otherwise.²²

²¹ Though residential mobility is associated with both neighborhood poverty and with adolescent outcomes (Hagan, MacMillan, and Wheaton 1996; Priresh and Downey 1999), the results presented in this article do not account for residential mobility during adolescence by matching. Mobility during adolescence occurs during the treatment period and therefore, in the counterfactual framework, is not a potential confounder. However, even when the number of neighborhood moves a family experiences after childhood is matched, results are substantively similar to those presented here and continue to be robust to selection bias (results not shown).

²² Technically, my measure of teenage pregnancy is not a measure of pregnancy but of live birth. If rates of miscarriage or abortion among pregnant teenagers vary by neighborhood poverty rate, this may bias estimates of the effects of neighborhood poverty on teenage pregnancy. The bias will be upward if abortion and miscarriage are more common in low-poverty neighborhoods and downward if they are more

TABLE 2
VARIABLES

Variable	Description
Outcomes:	
High school dropout	Not graduated from high school by age 20
Teenage pregnancy	Live birth by age 20 (females only)
Main independent variable:	
Neighborhood poverty rate	Mean census tract poverty rate age 11–20
Matching variables:	
Black/nonblack	
Female/male	
Low birth weight	Less than 88 ounces at birth
Year of birth (1958–77)	
Mother married at birth	
Male household head, age 10	
Household head HS grad, age 10	
Mother's age at birth	
Family owns home, age 10	
Family welfare receipt, age 10	
Household head work hours	Mean annual work hours, pre-age 11
Family income	Mean family income in 1979 dollars, pre-age 11
SMSA mean family income	1979 dollars; mean in SMSA of family, pre-age 11
SMSA segregation poor/nonpoor	Mean index of dissimilarity in SMSA of family, pre-age 11
SMSA poverty rate	Mean in SMSA of family, pre-age 11
SMSA % new housing	Mean percentage of housing units built in last decade in SMSA of family, pre-age 11
SMSA % vacant housing	Mean percentage of housing units vacant in SMSA of family, pre-age 11
SMSA % black	Mean in SMSA of family, pre-age 11
SMSA segregation black/nonblack	Mean index of dissimilarity in SMSA of family, pre-age 11
% individual/family variables imputed	
% tract/SMSA variables imputed	
Number of years age 11–20 tract poverty rate missing	
1978 PSID weight	

I measure the main independent variable, neighborhood poverty rate during adolescence, between the ages 11 and 20. Neighborhood poverty

common in high-poverty neighborhoods. I am not particularly concerned about this form of bias because the primary reasons we believe teenage pregnancy is a negative outcome are the negative effects of being born to a teenage mother and of being a teenage mother, both of which require a live birth.

rate is the mean poverty rate of the census tract over the years when the subject is between age 11 and age 20 and in which data are not missing. For the teenage pregnancy analyses, I do not include in this measure the years after the birth of a child because doing so would conflate movement to a higher-poverty neighborhood after a birth with the neighborhood effect I am attempting to measure. Because the mechanisms through which neighborhood context is theorized to affect adolescent outcomes relate to long-term developmental processes (e.g., socialization, peer relations, and interactions with neighborhood adults and external social institutions [Jencks and Mayer 1990]), a long-term measure of context better captures these processes than does measuring neighborhood context in a single year.²³ In addition, averaging over multiple years has the advantage of reducing measurement error, better capturing the true neighborhood conditions of an individual during adolescence.

I use neighborhood poverty rate as the measure of neighborhood disadvantage because of the role of structural disadvantages in theoretical accounts of the importance of neighborhoods (e.g., Wilson 1987, 1996). However, neighborhood disadvantage can also be thought of as a bundle of correlated variables. The methods used here cannot distinguish between neighborhood poverty and other highly correlated neighborhood characteristics—single-parent families, percentage of youth, or joblessness—as the causal factor in higher rates of teenage pregnancy or high school dropout in disadvantaged neighborhoods.

I measure matching covariates when the respondent is age 10 or under. For family and SMSA variables, I measure covariates when the respondent is age 10 or under or before he or she is born.²⁴ Gender and race covariates are always matched exactly, while other covariates are matched using propensity scores. Note that neighborhood poverty before age 11 is not among the matching variables, meaning that effects attributed to neighborhood during adolescence could potentially be due to neighborhood conditions before adolescence. However, though there has been little

²³ One potential disadvantage of averaging over time is that “treatments” with different theoretical meanings may be lumped together. For example, an adolescent who lived for five years in a low-poverty neighborhood and five years in a high-poverty neighborhood would be considered the same as one who lived for 10 years in a moderate-poverty neighborhood. A second potential disadvantage occurs because certain time periods in adolescence, such as the transition from middle school to high school, may be more important than others for the outcomes studied. Both of these problems lead to measurement error, which will bias results downward. However, downwardly biased effect estimates do not threaten the primary substantive goal of this article, that is, to show that neighborhood effects are robust to selection bias. If estimates that are biased downward are robust to selection bias, then unbiased estimates will be as well.

²⁴ For respondents born into PSID families, family information is available in survey years prior to birth.

research on the timing of neighborhood effects, theory suggests that context during adolescence will have a far stronger impact on teenage pregnancy and high school dropout than context during childhood. Jencks and Mayer (1990) describe three mechanisms by which neighborhood context may influence adolescent outcomes: peer relations, socialization by neighborhood adults, and interaction with nonneighborhood adults through social institutions. Adolescence is the developmental period in which social interaction shifts away from the family and toward peers, other adults, and institutions. It is also the period in which attitudes and values regarding education and fertility are formed and decisions are made.²⁵

With the exception of race and gender, missing values on covariates are imputed using best subset regression, and I also match on the proportion of individual or family variables imputed and the proportion of tract or SMSA variables imputed. Finally, I match on the 1978 PSID individual weight to ensure that particularly unusual individuals are not concentrated in either the treatment or control groups.²⁶ The only single survey year in which all cases selected for this analysis were alive and lived in PSID families is 1978.

Following most previous research on neighborhood effects, I use census tracts to approximate neighborhoods.²⁷ The typical census tract is geographically contiguous and contains between 4,000 and 8,000 people. Tract boundaries are selected by resident committees to reflect local opinions about neighborhood borders. While it is widely recognized that char-

²⁵ Matching on neighborhood poverty before adolescence would ensure that neighborhood effects are only those of neighborhood poverty during adolescence. However, doing so would also restrict the variation in neighborhood poverty level during adolescence to movement of families during adolescence and changing neighborhood conditions among those that do not move, confounding neighborhood contextual effects with the effects of residential mobility *per se*. (I thank an anonymous reviewer for pointing this out.) In fact, results are not very sensitive to the decision of whether to match on neighborhood poverty level before adolescence, with effects for whites being the same and effects for blacks being slightly smaller when this variable is matched (results not shown).

²⁶ For those with missing weights, I first impute their weight by assigning them their weight in the nearest year in which their weight is not missing. Weights in nearby years are correlated .97–.99 in the PSID. For those with missing weights after this procedure, I impute using nearest available subset regression.

²⁷ Because many parts of the country were not divided into census tracts in 1970, PSID subjects who were adolescents earlier in the PSID panel are more likely to have missing data on neighborhood characteristics and therefore to have been dropped from the analysis. In particular, many smaller metropolitan areas were not “traced.” As such, results should only be considered representative of young people in larger metropolitan areas.

acteristics of census tracts are not the ideal measures of neighborhood context, they are the best available.²⁸

Census data comes from the Urban Institute's Underclass Database (UDB), which contains an extensive set of census tract and metropolitan area characteristics from the 1970, 1980, and 1990 censuses. The Urban Institute staff has created 1970 and 1990 data sets that use 1980 tract boundaries, so that researchers can study tracts and metropolitan areas with consistent geographic boundaries over the 20-year period. The PSID Geocode data matches PSID families to 1980 census tracts for the 1968–85 survey years and PSID families to 1990 census tracts for 1986–97. For the years 1971–79 and 1981–85, I create census tract characteristics by linear interpolation between decennial census years. Census tract characteristics for 1968 and 1969 come from the 1970 census, and census tract characteristics for 1986–97 come from the 1990 census.

Table 1 shows unadjusted dropout and teen pregnancy rates by race and neighborhood poverty rate among PSID respondents born during 1958–77. Since there are few respondents who are neither black nor white, I group all nonblacks together.²⁹ The base rates of dropout and teenage pregnancy increase monotonically but not linearly as the neighborhood poverty rate increases. The accepted wisdom on neighborhood effects suggests that high-poverty neighborhoods are especially destructive, implying that rates of dropout and pregnancy will show the largest differences between high-poverty and extreme-poverty neighborhoods. In contrast, the unadjusted association between neighborhood poverty and outcomes is stronger when comparing low-poverty to moderate-poverty neighborhoods than when comparing moderate-poverty and high- or extreme-poverty neighborhoods. This is the case across races and for both outcomes considered here.

Blacks in this sample have higher dropout and teen pregnancy rates than nonblacks, as they do nationally. In addition, while nonblacks are concentrated in neighborhoods with the lowest poverty rates, higher-

²⁸ As Tienda (1991) notes, using census-defined boundaries defines neighborhoods spatially but not socially. This is dangerous because social interaction that occurs in neighborhoods presumably contributes to neighborhood effects. In addition, census tracts probably do not represent the right concept of neighborhood for individuals of all ages and for all outcomes. The reader should recognize that using census tracts to approximate neighborhoods measures true neighborhood characteristics with error, which will bias estimates of neighborhood effects toward zero. An alternative view is that the term “neighborhood” is solely a spatial designation and is not intended to refer to social interaction. If one accepts this use of the term, then using physical boundaries will lead to measurement error only if inappropriate spatial boundaries are used since the congruence of spatial and social boundaries is irrelevant by definition.

²⁹ Unfortunately, PSID data do not allow the identification of Latino ethnicity prior to 1985. Latinos may be either black or nonblack.

poverty neighborhoods are the norm for blacks. For nonblacks, I use less than 10% poverty as the control group and define the treatments as living in 10%–20% poor neighborhoods (moderate) and living in 20% or higher poor neighborhoods (high). For blacks, I use greater than 20% poor as the control group because high-poverty neighborhoods are the most common context experienced by blacks in this sample. I define treatment as living in a less than 10% poor neighborhood (low) or as living in a 10%–20% poor neighborhood (moderate).³⁰

RESULTS

High School Dropout

Table 3 shows the comparison between matched groups of nonblacks in low-poverty (less than 10%) and moderate-poverty neighborhoods (10%–20% poor). While the base unmatched (and unadjusted) dropout rates differ significantly between the two groups, the treatment effects odds ratio is only 1.36, an effect that is not statistically significant. The matching algorithm found matches for 663 (99.4%) of the 667 treatment cases, and the groups are well matched, showing good covariate balance. The bottom panel of the table shows the “sensitivity matrix,” that is, the estimated treatment effect and its confidence interval given various values of sensitivity parameters Γ and Δ . Note that when either $\Gamma = 1$ or $\Delta = 1$, the treatment effect is not reduced by the unobserved covariate because the unobserved covariate either has no effect on treatment assignment or no effect on the outcome. As Γ and Δ increase, the estimated treatment effect decreases. However, for all values of Γ and Δ , the confidence interval includes one, meaning the effect is not statistically significant.³¹

Table 4 shows the comparison between matched groups of nonblacks in low-poverty and high-poverty neighborhoods (greater than 20% poor). The matching algorithm found matches for 276 (83.6%) of the 330 treatment cases, and the groups are well matched. In this comparison, the matched treatment and control groups differ by 12.3 percentage points.

³⁰ For practical reasons associated with the matching software, it is desirable to have more control cases than treatment cases.

³¹ Throughout the presentation of results, I assess sensitivity to an unobserved covariate that, if controlled for, would reduce the estimated effect because the usual concern is that selection bias leads to estimated neighborhood effects that are too large rather than too small. Thus, I am considering unobserved covariates that are either positively related to both treatment assignment and outcome (e.g., single-parent family) or negatively related to both treatment assignment and outcome (e.g., family income). However, using these methods, one could easily consider an unobserved covariate that, when controlled for, would increase the neighborhood effect by considering values of either Γ or Δ that are between zero and one.

TABLE 3
EFFECT OF NEIGHBORHOOD POVERTY ON DROPOUT AMONG NONBLACKS,
LOW (< 10%) VS. MODERATE (10%–20%) POVERTY

	BASE DROPOUT RATE	MATCHED SAMPLE DROPOUT RATE	MATCHED SAMPLE COVAR- IATE BALANCE		
			Standard Bias*	Covariates	
Control (low)	11.8	17.8	> .50	0	
Treatment (moderate) ...	23.1	22.8	.25–.50	0	
Difference					
($T - C$)	11.3	5.0	.15–.25	0	
Odds ratio (T/C)	2.24	1.36	.10–.15	1	
			.05–.10	6	
χ^2 (1 df)	36.6	2.91	< .05	15	
	$p < .001$	$p = .089$			
Treatment cases	667	663			
% matched		99.4			
Unique control					
cases	1,574	407			
SENSITIVITY TO SELECTION BIAS					
(Odds Ratio Effect Estimate and 95% Confidence Interval)					
	$\Delta = 1.0$	$\Delta = 1.2$	$\Delta = 1.4$	$\Delta = 1.6$	$\Delta = 1.8$
$\Gamma = 1.0$	1.36	1.36	1.36	1.36	1.36
	(.95–1.94)	(.95–1.94)	(.95–1.95)	(.95–1.95)	(.95–1.95)
$\Gamma = 1.2$	1.36	1.35	1.34	1.33	1.33
	(.95–1.94)	(.95–1.93)	(.94–1.92)	(.93–1.91)	(.93–1.90)
$\Gamma = 1.4$	1.36	1.34	1.32	1.31	1.3
	(.95–1.95)	(.94–1.92)	(.93–1.89)	(.92–1.87)	(.91–1.86)
$\Gamma = 1.6$	1.36	1.33	1.31	1.29	1.27
	(.95–1.95)	(.93–1.91)	(.91–1.87)	(.90–1.85)	(.89–1.82)
$\Gamma = 1.8$	1.36	1.33	1.3	1.27	1.25
	(.95–1.95)	(.93–1.90)	(.90–1.86)	(.89–1.82)	(.87–1.80)

NOTE.— Γ = effect of hypothetical unobserved covariate on treatment (odds ratio). Δ = effect of hypothetical unobserved covariate on outcome (odds ratio).

* The standard bias of the propensity score is .0005.

These estimates suggest that a high-poverty neighborhood almost doubles the odds of high school dropout among nonblacks. This effect is statistically significant. I checked these results with a logit model controlling for the matching variables in the matched sample, and the estimates are of a similar magnitude (estimates not shown). The bottom panel of the table shows that this effect is fairly robust to selection bias. For example, if there were an unobserved covariate that doubled the odds of growing up in a high-poverty neighborhood and doubled the odds of high school dropout, the treatment would still multiply the odds of dropout by 1.69, an effect that would still be statistically significant at the 0.05 level. Note

TABLE 4
EFFECT OF NEIGHBORHOOD POVERTY ON DROPOUT AMONG NONBLACKS,
LOW (< 10%) vs. HIGH (> 20%) POVERTY

	BASE DROPOUT RATE	MATCHED SAMPLE DROPOUT RATE	MATCHED SAMPLE COVARI- ATE BALANCE		
			Standard Bias*	Covariates	
Control (low)	11.8	20.3	> .50	0	
Treatment (high)	34.2	32.6	.25-.50	0	
Difference					
($T - C$)	22.4	12.3	.15-.25	1	
Odds Ratio (T/C) ...	3.88	1.90	.10-.15	1	
			.05-.10	5	
χ^2 (1 df)	61.5	6.55	< .05	15	
	$p < .001$	$p = .011$			
Treatment cases	330	276			
% matched		83.6			
Unique control					
cases	1,574	207			
SENSITIVITY TO SELECTION BIAS					
(Odds Ratio Effect Estimate and 95% Confidence Interval)					
	$\Delta = 1.0$	$\Delta = 1.5$	$\Delta = 2.0$	$\Delta = 2.5$	$\Delta = 3.0$
$\Gamma = 1.0$	1.90	1.90	1.90	1.90	1.90
	(1.16-3.12)	(1.16-3.12)	(1.15-3.14)	(1.15-3.15)	(1.14-3.16)
$\Gamma = 1.5$	1.90	1.82	1.77	1.74	1.71
	(1.16-3.13)	(1.11-3.00)	(1.07-2.93)	(1.05-2.88)	(1.02-2.84)
$\Gamma = 2.0$	1.90	1.77	1.69	1.63	1.58
	(1.15-3.14)	(1.07-2.93)	(1.02-2.80)	(.98-2.71)	(.95-2.64)
$\Gamma = 2.5$	1.90	1.74	1.63	1.55	1.49
	(1.14-3.16)	(1.04-2.89)	(.98-2.72)	(.93-2.60)	(.89-2.51)
$\Gamma = 3.0$	1.90	1.71	1.58	1.49	1.43
	(1.14-3.18)	(1.02-2.85)	(.94-2.65)	(.89-2.51)	(.84-2.41)

NOTE.— Γ = effect of hypothetical unobserved covariate on treatment (odds ratio). Δ = effect of hypothetical unobserved covariate on outcome (odds ratio).

* The standard bias of the propensity score is .0013.

that such an unobserved covariate would need to have such relationships with the treatment and outcome net of the observed covariates in table 2 since the treatment and control groups are already balanced on these observed covariates.

Another way to read the sensitivity matrix is to look for the smallest values of Γ and Δ for which the effect of neighborhood poverty is not statistically significant (i.e., for which the value "1" is included in the confidence interval). These values correspond to the "weakest" unobserved covariate that would lead us to conclude that the neighborhood effect is not statistically different from no effect. For this table, this occurs when

$\Gamma = 2.0$ and $\Delta = 2.5$ (or vice versa since the matrix is symmetric about the diagonal). In the discussion section below, I return to the issue of how to substantively interpret the sensitivity parameters.

Tables 5 and 6 examine the effect of neighborhood poverty among blacks. Recall that because high-poverty neighborhoods are the most common condition for blacks (see table 1), this category is the control group and moderate-poverty and low-poverty are considered treatments. However, to ease comparison with results for nonblacks, treatment effects for blacks are displayed as the inverse of the effect of treatment.

Table 5 compares blacks in high-poverty neighborhoods to those in low-poverty neighborhoods. The matching algorithm found matches for 195 (94.7%) of the 206 observations in the low-poverty neighborhoods, and the matched samples show good covariate balance. In the matched sample, living in a low-poverty neighborhood reduces the rate of high school dropout by about 12 percentage points, almost halving the dropout rate. According to this analysis, growing up in a high-poverty neighborhood slightly more than doubles the odds of high school dropout, an effect slightly larger than that for nonblacks, though the difference is not statistically significant. I checked the treatment effect results with a logit model controlling for the matching variables in the matched sample, and the estimates are comparable (estimates not shown). Again, the effect is robust to selection bias. Even in the presence of an unobserved covariate that doubles the odds of living in a high-poverty neighborhood and doubles the odds of dropout, the high-poverty neighborhood multiplies the odds of dropout by 1.91. An unobserved covariate would have to double the odds of living in a high-poverty neighborhood and multiply by 2.5 the odds of dropout (or vice versa) to render the effect statistically insignificant.

Table 6 compares matched samples of blacks in high-poverty neighborhoods with those in moderate-poverty neighborhoods. The sample is well-matched, but with an odds ratio of only 1.18, the effect is small, statistically insignificant, and not surprisingly, not robust to selection bias.

Teenage Pregnancy

Table 7 compares teenage pregnancy rates among the matched sample of nonblack youth in low-poverty and moderate-poverty neighborhoods. The algorithm found matches for 432 (99.5%) of the 434 treatment cases, and the matched sample contains good covariate balance. The estimates show that a moderate-poverty neighborhood increases the pregnancy rate by 6.7 percentage points, more than doubling the odds of teenage pregnancy. I checked these results with a logit model controlling for the matching variables in the matched sample, and again the estimates are comparable

TABLE 5
EFFECT OF NEIGHBORHOOD POVERTY ON DROPOUT AMONG BLACKS,
HIGH (> 20%) VS. LOW (< 10%) POVERTY

	BASE DROPOUT RATE	MATCHED SAMPLE DROPOUT RATE	MATCHED SAMPLE COVARI- ATE BALANCE		
			Standard Bias*	Covariates	
Control (high)	34.2	25.6	> .50	0	
Treatment (low)	13.1	13.6	.25-.50	0	
Difference					
(C - T)	21.1	12.0	.15-.25	0	
Odds ratio (C/T)	3.45	2.15	.10-.15	4	
			.05-.10	7	
χ^2 (1 df)	30.3	6.35	< .05	11	
	$p < .001$	$p = .013$			
Treatment cases	206	195			
% matched		94.7			
Unique control					
cases	1,610	147			
SENSITIVITY TO SELECTION BIAS					
(Odds Ratio Effect Estimate and 95% Confidence Interval)					
	$\Delta = 1.0$	$\Delta = 1.5$	$\Delta = 2.0$	$\Delta = 2.5$	$\Delta = 3.0$
$\Gamma = 1.0$	2.15 (1.18-3.91)	2.15 (1.17-3.92)	2.15 (1.17-3.93)	2.15 (1.17-3.95)	2.15 (1.16-3.97)
$\Gamma = 1.5$	2.15 (1.17-3.93)	2.06 (1.12-3.77)	2.00 (1.09-3.68)	1.96 (1.06-3.61)	1.93 (1.04-3.56)
$\Gamma = 2.0$	2.15 (1.17-3.95)	2.00 (1.09-3.69)	1.91 (1.03-3.52)	1.84 (.99-3.41)	1.78 (.96-3.32)
$\Gamma = 2.5$	2.15 (1.16-3.97)	1.96 (1.06-3.63)	1.84 (.99-3.41)	1.75 (.94-3.26)	1.69 (.90-3.15)
$\Gamma = 3.0$	2.15 (1.15-4.00)	1.93 (1.03-3.59)	1.78 (.95-3.34)	1.69 (.90-3.16)	1.61 (.85-3.03)

NOTE.— Γ = effect of hypothetical unobserved covariate on treatment (odds ratio). Δ = effect of hypothetical unobserved covariate on outcome (odds ratio).

* The standard bias of the propensity score is .0019.

(estimates not shown). The statistical significance of the effect is fairly robust to selection bias. An unobserved covariate that multiplied the odds of treatment by 2.2 and the odds of teenage pregnancy by 2.2 would lower the treatment effect such that its confidence interval would include one.

Table 8 shows the comparison between pregnancy rates in the matched sample of nonblack adolescents in low-poverty and high-poverty neighborhoods. The algorithm found matches for 166 (82.2%) of the 202 treatment cases, and the matched sample contains good covariate balance. The estimates show that a high-poverty neighborhood increases the teenage pregnancy rate by 11.5 percentage points, almost tripling the odds of

TABLE 6
EFFECT OF NEIGHBORHOOD POVERTY ON DROPOUT AMONG BLACKS,
HIGH (> 20%) VS. MODERATE (10%–20%) POVERTY

	BASE DROPOUT RATE	MATCHED SAMPLE DROPOUT RATE	MATCHED SAMPLE COVAR- IATE BALANCE		
			Standard Bias*	Covariates	
Control (high)	34.2	30.9	> .50	0	
Treatment (moderate) ...	26.9	27.4	.25-.50	0	
Difference					
($C - T$)	7.3	3.5	.15-.25	0	
Odds ratio (C/T)	1.41	1.18	.10-.15	0	
			.05-.10	2	
χ^2 (1 df)	7.05	1.12	< .05	20	
	$p = .008$	$p = .291$			
Treatment cases	587	577			
% matched		98.3			
Unique control					
cases	1,610	434			
SENSITIVITY TO SELECTION BIAS					
(Odds Ratio Effect Estimate and 95% Confidence Interval)					
	$\Delta = 1.0$	$\Delta = 1.2$	$\Delta = 1.4$	$\Delta = 1.6$	$\Delta = 1.8$
$\Gamma = 1.0$	1.18	1.18	1.18	1.18	1.18
	(.87-1.61)	(.87-1.61)	(.87-1.61)	(.87-1.62)	(.87-1.62)
$\Gamma = 1.2$	1.18	1.17	1.17	1.16	1.15
	(.87-1.61)	(.86-1.60)	(.85-1.59)	(.85-1.58)	(.84-1.58)
$\Gamma = 1.4$	1.18	1.17	1.15	1.14	1.13
	(.87-1.62)	(.85-1.59)	(.84-1.57)	(.83-1.56)	(.82-1.54)
$\Gamma = 1.6$	1.18	1.16	1.14	1.12	1.10
	(.87-1.62)	(.85-1.58)	(.83-1.56)	(.82-1.53)	(.81-1.51)
$\Gamma = 1.8$	1.18	1.15	1.13	1.10	1.09
	(.86-1.62)	(.84-1.58)	(.82-1.54)	(.81-1.51)	(.79-1.49)

NOTE.— Γ = effect of hypothetical unobserved covariate on treatment (odds ratio). Δ = effect of hypothetical unobserved covariate on outcome (odds ratio).

* The standard bias of the propensity score is .0002.

teenage pregnancy. I checked these results with a logit model controlling for the matching variables in the matched sample, and again the estimates are comparable (estimates not shown). This effect is also robust to selection bias. Even if an unobserved covariate almost doubled the odds of growing up in a high-poverty neighborhood and more than doubled the odds of teenage pregnancy, growing up in a high-poverty neighborhood would still multiply the odds of teenage pregnancy by 2.63, an effect which would still be statistically significant.

Table 9 shows the comparison between pregnancy rates in the matched sample of black adolescents in high-poverty and low-poverty neighbor-

TABLE 7
EFFECT OF NEIGHBORHOOD POVERTY ON TEENAGE PREGNANCY AMONG NONBLACKS,
LOW (< 10%) VS. MODERATE (10%–20%) POVERTY

	BASE PREGNANCY RATE	MATCHED SAMPLE PREGNANCY RATE	MATCHED SAMPLE COVARIATE BALANCE		
			Standard Bias*	Covariates	
Control (low)	4.6	6.0	> .50	0	
Treatment (moderate)	12.7	12.7	.25–.50	0	
Difference					
($T - C$)	8.1	6.7	.15–.25	0	
Odds ratio (T/C)	3.02	2.28	.10–.15	2	
			.05–.10	6	
χ^2 (1 df)	30.0	6.14	< .05	13	
	$p < .001$	$p = .014$			
Treatment cases	434	432			
% matched		99.5			
Unique control cases	911	254			
SENSITIVITY TO SELECTION BIAS					
(Odds Ratio Effect Estimate and 95% Confidence Interval)					
	$\Delta = 1.0$	$\Delta = 1.3$	$\Delta = 1.6$	$\Delta = 1.9$	$\Delta = 2.2$
$\Gamma = 1.0$	2.28 (1.17–4.44)	2.28 (1.17–4.44)	2.28 (1.17–4.45)	2.28 (1.17–4.45)	2.28 (1.16–4.46)
$\Gamma = 1.3$	2.28 (1.17–4.45)	2.24 (1.15–4.37)	2.21 (1.13–4.32)	2.19 (1.12–4.27)	2.16 (1.11–4.24)
$\Gamma = 1.6$	2.28 (1.16–4.46)	2.21 (1.13–4.32)	2.16 (1.10–4.22)	2.11 (1.08–4.15)	2.08 (1.06–4.08)
$\Gamma = 1.9$	2.28 (1.16–4.47)	2.19 (1.11–4.29)	2.11 (1.08–4.16)	2.06 (1.05–4.05)	2.01 (1.02–3.96)
$\Gamma = 2.2$	2.28 (1.16–4.49)	2.16 (1.10–4.27)	2.08 (1.05–4.10)	2.01 (1.02–3.97)	1.96 (.99–3.87)

NOTE.— Γ = effect of hypothetical unobserved covariate on treatment (odds ratio). Δ = effect of hypothetical unobserved covariate on outcome (odds ratio).

* The standard bias of the propensity score is .0002.

hoods. The algorithm found matches for 121 (89.6%) of the 135 treatment cases, and the matched sample contains good covariate balance. The estimates show that a high-poverty neighborhood increases the teenage pregnancy rate by almost 16 percentage points, multiplying the odds of teenage pregnancy by 2.59. I checked these results with a logit model controlling for the matching variables in the matched sample, and again the estimates are comparable (estimates not shown). The effect is robust to selection bias. Even with an unobserved covariate that doubled both the odds of growing up in a high-poverty neighborhood and the odds of

TABLE 8
EFFECT OF NEIGHBORHOOD POVERTY ON TEENAGE PREGNANCY AMONG NONBLACKS,
LOW (< 10%) vs. HIGH (> 20%) POVERTY

	BASE PREGNANCY RATE	MATCHED SAMPLE PREGNANCY RATE	MATCHED SAMPLE CO- VARIATE BALANCE	
			Standard Bias*	Covariates
Control (low)	4.6	7.2	> .50	0
Treatment (high)	18.8	18.7	.25-.50	0
Difference				
($T - C$)	14.2	11.5	.15-.25	0
Odds ratio (T/C)	4.80	2.95	.10-.15	4
			.05-.10	5
χ^2 (1 df)	38.6	5.55	< .05	12
	$p < .001$	$p = .019$		
Treatment cases	202	166		
% matched		82.2		
Unique control				
cases	911	122		

SENSITIVITY TO SELECTION BIAS (Odds Ratio Effect Estimate and 95% Confidence Interval)					
	$\Delta = 1.0$	$\Delta = 1.4$	$\Delta = 1.8$	$\Delta = 2.2$	$\Delta = 2.6$
$\Gamma = 1.0$	2.95 (1.16-7.48)	2.95 (1.16-7.49)	2.95 (1.16-7.51)	2.95 (1.15-7.53)	2.95 (1.15-7.55)
$\Gamma = 1.4$	2.95 (1.16-7.50)	2.87 (1.13-7.29)	2.81 (1.10-7.16)	2.76 (1.08-7.06)	2.72 (1.06-6.99)
$\Gamma = 1.8$	2.95 (1.15-7.54)	2.81 (1.10-7.18)	2.71 (1.06-6.94)	2.63 (1.02-6.76)	2.57 (1.00-6.62)
$\Gamma = 2.2$	2.95 (1.14-7.59)	2.76 (1.07-7.11)	2.63 (1.02-6.79)	2.53 (.98-6.55)	2.46 (.95-6.37)
$\Gamma = 2.6$	2.95 (1.14-7.65)	2.72 (1.05-7.07)	2.57 (.99-6.68)	2.46 (.94-6.39)	2.36 (.91-6.17)

NOTE.— Γ = effect of hypothetical unobserved covariate on treatment (odds ratio). Δ = effect of hypothetical unobserved covariate on outcome (odds ratio).

* The standard bias of the propensity score is .0010.

teenage pregnancy, the effect would still be 2.30. A covariate that multiplied both the odds of treatment by 2.5 and the odds of teenage pregnancy by 3.0 would be required to drive the statistical significance of the effect to nonsignificance.

Table 10 compares pregnancy rates in the matched sample of black adolescents in high-poverty and moderate-poverty neighborhoods. While the algorithm found matches for 385 (99.7%) of 386 treatment cases and covariates are well-balanced, the estimated effect is essentially zero. This is not surprising given the small effect in the unmatched sample, a dif-

TABLE 9
EFFECT OF NEIGHBORHOOD POVERTY ON TEENAGE PREGNANCY AMONG BLACKS,
HIGH (> 20%) VS. LOW (< 10%) POVERTY

	BASE PREGNANCY RATE	MATCHED SAMPLE PREGNANCY RATE	MATCHED SAMPLE CO- VARIATE BALANCE		
			Standard Bias*	Covariates	
Control (high)	30.1	29.8	> .50	0	
Treatment (low)	12.6	14.1	.25-.50	0	
Difference					
$(C - T)$	17.5	15.7	.15-.25	1	
Odds ratio (C/T)	2.99	2.59	.10-.15	1	
			.05-.10	8	
χ^2 (1 df)	17.5	7.51	< .05	11	
	$p < .001$	$p = .007$			
Treatment cases	135	121			
% matched		89.6			
Unique control					
cases	1,034	100			
SENSITIVITY TO SELECTION BIAS					
(Odds Ratio Effect Estimate and 95% Confidence Interval)					
	$\Delta = 1.0$	$\Delta = 1.5$	$\Delta = 2.0$	$\Delta = 2.5$	$\Delta = 3.0$
$\Gamma = 1.0$	2.59 (1.29-5.19)	2.59 (1.29-5.21)	2.59 (1.28-5.23)	2.59 (1.28-5.26)	2.59 (1.27-5.28)
$\Gamma = 1.5$	2.59 (1.29-5.21)	2.49 (1.23-5.01)	2.42 (1.20-4.89)	2.37 (1.16-4.80)	2.33 (1.14-4.74)
$\Gamma = 2.0$	2.59 (1.28-5.25)	2.42 (1.19-4.90)	2.30 (1.13-4.68)	2.22 (1.09-4.53)	2.16 (1.05-4.42)
$\Gamma = 2.5$	2.59 (1.27-5.28)	2.37 (1.16-4.82)	2.22 (1.09-4.54)	2.12 (1.03-4.34)	2.03 (.99-4.20)
$\Gamma = 3.0$	2.59 (1.26-5.32)	2.33 (1.13-4.78)	2.16 (1.05-4.44)	2.03 (.98-4.21)	1.94 (.94-4.04)

NOTE.— Γ = effect of hypothetical unobserved covariate on treatment (odds ratio). Δ = effect of hypothetical unobserved covariate on outcome (odds ratio).

* The standard bias of the propensity score is .0008.

ference between unmatched treatment and control groups of less than 4 percentage points.

Discussion

Two remaining issues in the interpretation of the results warrant further discussion. First is the issue of the unmatched treatment cases. In the analyses in which robust effects are found, between 82% and 99% of the treatment cases are actually matched. As noted above, there is a trade-off between estimating the effect for the entire treatment group but having

TABLE 10
EFFECT OF NEIGHBORHOOD POVERTY ON TEENAGE PREGNANCY AMONG BLACKS,
HIGH (> 20%) VS. MODERATE (10%–20%) POVERTY

	BASE PREGNANCY RATE	MATCHED SAMPLE PREGNANCY RATE	MATCHED SAMPLE COVAR- IATE BALANCE		
			Standard Bias*	Covariates	
Control (high)	30.1	30.4	> .50	0	
Treatment (moder- ate)	26.9	27.1	.25-.50	0	
Difference ($C - T$)	3.2	3.3	.15-.25	0	
Odds ratio (C/T)	1.17	1.17	.10-.15	3	
			.05-.10	7	
χ^2 (1 df)	1.19	.675	< .05	11	
	$p = .275$	$p = .412$			
Treatment cases	386	385			
% matched		99.7			
Unique control cases	1,057	274			
SENSITIVITY TO SELECTION BIAS (Odds Ratio Effect Estimate and 95% Confidence Interval)					
	$\Delta = 1.0$	$\Delta = 1.2$	$\Delta = 1.4$	$\Delta = 1.6$	$\Delta = 1.8$
$\Gamma = 1.0$	2.08 (.89-4.85)	2.08 (.89-4.85)	2.08 (.89-4.85)	2.08 (.89-4.86)	2.08 (.89-4.86)
$\Gamma = 1.2$	2.08 (.89-4.85)	2.07 (.89-4.84)	2.07 (.88-4.83)	2.06 (.88-4.83)	2.06 (.88-4.82)
$\Gamma = 1.4$	2.08 (.89-4.86)	2.07 (.88-4.84)	2.06 (.88-4.82)	2.05 (.88-4.80)	2.05 (.87-4.79)
$\Gamma = 1.6$	2.08 (.89-4.86)	2.06 (.88-4.83)	2.05 (.88-4.80)	2.04 (.87-4.78)	2.03 (.87-4.76)
$\Gamma = 1.8$	2.08 (.89-4.87)	2.06 (.88-4.83)	2.05 (.87-4.79)	2.03 (.87-4.76)	2.02 (.86-4.74)

NOTE.— Γ = effect of hypothetical unobserved covariate on treatment (odds ratio). Δ = effect of hypothetical unobserved covariate on outcome (odds ratio).

* The standard bias of the propensity score is .0002.

little knowledge about the robustness of the results and estimating the effect for only some of the treatment group but understanding the robustness of the results. However, it is still valuable to consider the unmatched treatment cases. What are the consequences of not finding control cases for a nontrivial portion of the treatment cases?

The matching estimator can only estimate the average treatment effect on the treated for the group of cases for which matches are found. The relationship between the estimate from the matched group presented here and the hypothetical estimate one would get if it were possible to match

all treatment cases depends on the difference between the matched and unmatched treatment cases. The treatment cases that are not matched are those with high estimated probabilities of living in high-poverty neighborhoods. Those individuals are more likely to grow up in female-headed households, experience lower family incomes, have household heads with lower education, and so on.

The key question is whether we would expect the unmatched to experience a larger neighborhood effect or a smaller neighborhood effect than the matched. On the one hand, since those who are unmatched appear to be worse off, we might expect them to be more affected by neighborhood context because they have fewer resources in the home to protect them from neighborhood risk factors (Furstenberg et al. 1999). This reasoning would suggest that the estimates presented above are lower bounds. On the other hand, since those who are unmatched are already worse off, we might expect them to be less affected by neighborhood context. Because they are already at such high risk of poor outcomes, the additional risks associated with a high-poverty neighborhood may not make much difference. This reasoning would suggest that the estimates presented above are upper bounds. I would argue that the former is correct, implying that the above estimates are lower bounds. Table 1 shows that while even in the most impoverished neighborhoods, teenage pregnancy and high school dropout are not uncommon, they are also not the norm. There is still plenty of room for neighborhood effects to operate on those who are otherwise most disadvantaged.

The second issue is how one might arrive at a more substantively grounded interpretation of the Γ and Δ values produced by the sensitivity analyses. To review, the sensitivity matrix presents the estimated effect of the treatment on the outcome and its 95% confidence interval, taking into account the selection bias due to an unobserved covariate of varying relationships with both treatment assignment (Γ) and the outcome (Δ). The unobserved covariate would have to have such effects on treatment and outcome net of the effects of all of the other variables already matched. Since Γ and Δ are expressed as odds ratios, one can compare their values to odds ratios from logit models predicting treatment and predicting the outcome, respectively. Fortunately, models predicting treatment were estimated (but not shown) in the process of calculating the estimated propensity scores, and models predicting the outcome can be easily estimated from the data at hand. One can look at these models to perform a “back-of-the envelope” calculation to see which observed variables produce coefficients similar to Γ and Δ and then think about how these variables compare with potential unobserved covariates, allowing for a more substantive interpretation of the sensitivity analysis.

The most-often mentioned unobserved covariate in studies of neigh-

neighborhood effects is parental involvement with or commitment to their children. The story is that, among families at-risk of residing in a high-poverty neighborhood, more committed parents are more likely to make personal sacrifices to afford the higher rents or home prices in lower-poverty neighborhoods, or perhaps are more knowledgeable about neighborhood conditions and their potential effects on children. These same parents are more likely to be involved in their children's schooling, better able to shepherd at-risk adolescents through high school, and do a better job monitoring their adolescents' peer relations and free time and instilling values like hard work and abstinence. If data sets used to study neighborhood effects identified and measured such parental characteristics and behaviors, neighborhood effects might disappear. How big would the effects of parental involvement have to be for this to be the case?

Consider the analysis of the effect of a high-poverty neighborhood on dropout among nonblacks, in which an unobserved covariate with $\Gamma = 2.0$ and $\Delta = 2.5$ would drive the neighborhood effect to statistical insignificance. In this sample, having a household head with less than a high school education has $\Gamma = 2.5$ and $\Delta = 2.5$. In addition, growing up in a family with an average income of \$22,000 or lower per year (in 1999 dollars) for 10 years has $\Gamma = 2.0$ and $\Delta = 1.65$. Thus, to completely wipe out the neighborhood effect, parental involvement would have to be just as powerful as parental education or more powerful than a huge change in family income, net of both of those variables. In a similar way, the sensitivity analysis for dropout among black adolescents has the same Γ and Δ , but having a household head who is a high school dropout has $\Gamma = 2.0$ and $\Delta = 1.5$, and an increase in average annual family income of \$25,000 for 10 years has $\Gamma = 2.5$ and $\Delta = 1.4$. Therefore, to completely wipe out the effect of neighborhood poverty on dropout among blacks, parental involvement would have to be a more powerful predictor of both neighborhood and dropout than either parental education or family income, net of these variables.

For teenage pregnancy among nonblacks in high- vs. low-poverty neighborhoods, the sensitivity analysis produces approximately $\Gamma = 2.2$ and $\Delta = 2.2$ as the point at which neighborhood effects are not significant. In this sample, living in a family that has ever received welfare has $\Gamma = 2.2$ and $\Delta = 1.8$ and experiencing low weight at birth has $\Gamma = 1.5$ and $\Delta = 2.4$. Thus, for parental involvement to wipe out the effect of neighborhood poverty, it would have to have larger effects than welfare receipt or low birth weight, net of these variables. For the analysis of teenage pregnancy among blacks in low-poverty vs. high-poverty neighborhoods, $\Gamma = 3.0$ and $\Delta = 3.0$. In this sample, having an average family income that is \$31,000 per year or higher over 10 years has $\Gamma = 3.0$ and $\Delta = 1.4$, and having a household head with less than a high school ed-

ucation has $\Gamma = 1.6$ and $\Delta = 1.8$. Thus, to drive the neighborhood effect to nonsignificance, parental involvement would have to be more powerful than family income or having poorly educated parents, net of these variables.³²

CONCLUSION

Since the publication of Wilson's *The Truly Disadvantaged* (1987), social scientists have been concerned with estimating the effects of neighborhood poverty on individual outcomes, especially its effects on young people. Research on this issue using observational data has been particularly hampered by the problem of selection bias: individuals living in high-poverty neighborhoods are different from those living in low-poverty neighborhoods on both observed, unobserved, and perhaps even on unobservable factors. Studies that have attempted to statistically control for these factors using regression models have been unconvincing because it is impossible to control for unobserved variables. More sophisticated methods such as instrumental variables and sibling fixed effects have attempted to harness exogenous sources of variation in levels of neighborhood poverty with mixed success.

This study employs counterfactual models, propensity score matching, and sensitivity analysis to estimate the effects of neighborhood poverty on high school dropout and teenage pregnancy. It uses a new method to deal with selection bias. Instead of attempting to remove or avoid selection bias caused by unobserved factors, these methods assess how the presence of varying levels of selection bias would alter conclusions about the effect. These methods are of course not without their own shortcomings, but they offer another angle of attack on a particularly difficult but extremely important social science problem.

This study shows that when two groups of children who are identical at age 10 on observed variables (including but not limited to family income, parents' education, welfare receipt, and family structure) experience different neighborhood contexts during adolescence, those who grow up in high-poverty neighborhoods are more likely to drop out of high school and have a teenage pregnancy than those who grow up in low-poverty

³² An alternative (and perhaps better) method for examining the potential for an unobserved covariate such as parental involvement to drive effects to nonsignificance would be to compare Γ and Δ with effect estimates from previously published studies that controlled for a similar group of observed covariates. Unfortunately, I was not able to locate any previous studies that examine the effect of parental involvement on both neighborhood choice and teenage pregnancy or high school dropout in a similar sample.

neighborhoods. The presence of these effects is robust to hypothetical unobserved factors that affect both neighborhood and the outcomes.

The neighborhood effects presented here are considerably larger than those presented in earlier analyses using observational data. For example, Brooks-Gunn et al. (1993) also used PSID data and examine the effect of neighborhood context on high school dropout and teenage pregnancy. In models that control for family-level factors (but not metropolitan-level factors), moving from a neighborhood in which none of the families are affluent to one in which all the families are affluent reduces only threefold the probability of high school dropout and teenage pregnancy. Two explanations may account for the difference in magnitude. First, previous estimates have measured neighborhood characteristics at a single point in time, whereas this analysis averages neighborhood poverty rates over adolescence, thereby reducing measurement error, which can bias estimates toward zero (for an exception, see South and Crowder [1999]). In the PSID sample used here, the correlation between neighborhood poverty rate at age 11 and neighborhood poverty at age 18 is only 0.75, suggesting that there is meaningful variation in neighborhood context over time. Second, this study defines neighborhood effects counterfactually as the difference in outcomes between groups of youth who were, on average, identical prior to adolescence but experience different neighborhood contexts during adolescence. Previous studies have controlled for variables such as school or family characteristics measured during adolescence that were themselves likely affected by neighborhood context (again, for an exception, see South and Crowder [1999]). These studies bias toward zero estimates of neighborhood effects by removing the portion of the neighborhood effect that operates through these variables.

In contrast, results from this analysis are smaller than comparable results for the quasi-experimental Gatreux study, in which the odds of failing to graduate from high school among black children who stayed in the city were almost four and a half times those who moved to the suburbs (Rubinowitz and Rosenbaum 2000, p. 163). In this analysis, growing up in a high-poverty neighborhood slightly more than doubled the odds of high school dropout among blacks. This larger effect in Gatreux is expected, however, because people who thought they would benefit from the Gatreux program selected into it.

This study has attempted to measure the total effect of growing up in neighborhoods characterized by high poverty rates, focusing on high school dropout and teenage pregnancy. Further research is needed to identify the causal mechanisms through which these effects operate and to measure their relative magnitudes. An important next step is to identify the major domains through which neighborhood effects on adolescents are mediated. Families, schools, and peer groups are the obvious candi-

dates, but more qualitative research is needed to identify and to understand the social processes through which neighborhood effects occur. Further quantitative research is needed to estimate the relative importance of each of these domains in the total effect of neighborhood poverty.

APPENDIX

Comparison with Traditional Logit Estimates

As a comparison for the matching estimates, the top portion of table A1 shows estimates from traditional logit regression models of high school dropout among blacks and nonblacks, using all observations in the data. Neighborhood poverty rate is entered into the model as a series of dummy variables that correspond to the comparisons made in the matching analyses, with the control group from the matching estimators serving as the omitted category. The models also control for the matching variables listed in table 2. For nonblacks, the logit model shows comparable estimates. Neither the logit model nor the matching estimator finds a significant effect of moderate poverty compared to low-poverty neighborhoods. The logit model indicates a slightly smaller effect of a high-poverty neighborhood in comparison to a low-poverty neighborhood (odds ratio of 1.66 for the logit vs. 1.96 for the matching estimator).

For blacks, the logit model and matching estimators also produce roughly comparable results. In neither case is the difference between high poverty and moderate poverty meaningful. For the comparison between low poverty and high poverty, the logit model indicates a slightly smaller effect (odds ratio of 1.79 for the logit and 2.15 for the matching estimator).

The bottom portion of table A1 shows similar estimates from traditional logit regression models of teenage pregnancy among blacks and nonblacks, using all observations in the data. Again, the estimates from teenage pregnancy models are very consistent with those from the matching estimators, though the estimates differ slightly in magnitude. For nonblacks, the estimated effect of a moderate-poverty neighborhood is larger in the logit estimates (odds ratio of 2.79 for the logit and 2.28 for the matching estimators). The estimated effect of a high-poverty neighborhood from the logit model is also considerably larger than that of the matching estimator (odds ratio of 3.90 for the logit and 2.95 for the matching estimator). For blacks, the estimated effect of a high-poverty neighborhood is considerably larger in the logit model (odds ratio of 2.59 for the matching estimator and 3.20 for the logit model). Neither the logit model nor the matching estimator finds an effect of moderate-poverty neighborhood on teenage pregnancy among blacks.

There are at least two differences between the matching and logit re-

TABLE A1
LOGIT REGRESSION ESTIMATES OF THE EFFECT OF NEIGHBORHOOD POVERTY ON HIGH SCHOOL DROPPING OUT AND TEENAGE PREGNANCY

NEIGHBORHOOD POVERTY RATE	NONBLACKS		BLACKS	
	Coefficient	SE	Coefficient	SE
Dropout:				
Low (0%–9.99%)	. . .	. . .	–.579 (.56)	.374
Moderate (10%–19.99%)	.349 (1.42)	.158	.057 (1.06)	.203
High (20%+)	.608 (1.84)	.206	. . .	. . .
<i>N</i>	2,571		2,403	
Pregnancy:				
Low (0–9.99%)	. . .	. . .	–1.158 (.31)	.403
Moderate (10–19.99%)	1.025 (2.79)	.253	.077 (1.08)	.207
High (20%+)	1.360 (3.90)	.335	. . .	. . .
<i>N</i>	1,547		1,578	

NOTE.—Models control for matching variables in table 2 and use the 1978 PSID weight. SEs corrected for clustering by family. Odds ratios are given in parentheses.

gression estimators that may account for the different estimates. First, the logit models are estimated over a larger, and different, sample. The logit model uses all cases, while the matching estimator selects those cases that, based on observed characteristics, appear to be comparable. As discussed above, when observations are not comparable (the data are not “supported”) and functional form is not correctly specified, biases can result. Since we have verified covariate balance, the matching estimator is nonparametric and therefore does not have to contend with specifying the correct functional form of all observed variables, as the logit model does. Second, slightly larger effects from the matching estimators may be the result of imperfect matching on the observed characteristics, which were controlled statistically in the logit models.

REFERENCES

- Aaronson, Daniel. 1997. “Sibling Estimates of Neighborhood Effects.” Pp. 80–93 in *Neighborhood Poverty: Vol. 2, Policy Implications in Studying Neighborhoods*, edited by Jeanne Brooks-Gunn, Greg J. Duncan, and J. Lawrence Aber. New York: Russell Sage.
- . 1998. “Using Sibling Data to Estimate the Impact of Neighborhoods on Children’s Economic Outcomes.” *Journal of Human Resources* 33:915–46.
- Angrist, Joshua, Guido Imbens, and Donald Rubin. 1996. “Identification of Causal

- Effects Using Instrumental Variables (with Discussion)." *Journal of the American Statistical Association* 91:444–72.
- Ashenfelter, Orley, and Cecelia Rouse. 2000. "Schooling, Intelligence and Income in America: Cracks in the Bell Curve Myth." In *Meritocracy in America*, edited by Kenneth Arrow, Samuel Bowles, and Steven Durlauf. Princeton N.J.: Princeton University Press.
- Brooks-Gunn, Jeanne, Greg J. Duncan, Pamela K. Klebanov, and Naomi Sealand. 1993. "Do Neighborhoods Influence Child and Adolescent Development?" *American Journal of Sociology* 99:353–95.
- Clark, Kenneth B. 1965. *Dark Ghetto: Dilemmas of Social Power*. New York: Harper Torchbooks.
- Cook, Christine C., and Marilyn J. Bruin. 1993. "Housing and Neighborhood Assessment Criteria among Black Urban Households." *Urban Affairs Quarterly* 29: 328–39.
- Crane, Jonathan. 1991. "The Epidemic Theory of Ghettos and Neighborhood Effects on Dropping Out and Teenage Childbearing." *American Journal of Sociology* 96: 1226–59.
- Crowder, Kyle. 2000. "The Racial Context of White Mobility: An Individual-Level Assessment of the White Flight Hypothesis." *Social Science Research* 29:223–57.
- Cullen, Julie Berry, and Steven D. Levitt. 1996. "Crime, Urban Flight, and the Consequences for Cities." NBER Working Paper no. 5737.
- D'Agostino, Ralph B., and Donald B. Rubin. 2000. "Estimating and Using Propensity Scores with Partially Missing Data." *Journal of the American Statistical Association* 95:749–59.
- Darling, Nancy, and Laurence Steinberg. 1997. "Community Influences on Adolescent Achievement and Deviance." Pp. 120–31 in *Neighborhood Poverty: Vol. 2, Policy Implications in Studying Neighborhoods*, edited by Jeanne Brooks-Gunn, Greg J. Duncan, and J. Lawrence Aber. New York: Russell Sage.
- Denton, Nancy A., and Douglas S. Massey. 1991. "Patterns of Neighborhood Transition in a Multiethnic World: U.S. Metropolitan Areas, 1970–1980." *Demography* 28:41–63.
- Dowding, Keith, Peter John, and Stephen Biggs. 1994. "Tiebout: A Survey of the Empirical Literature." *Urban Studies* 31:767–97.
- Dugan, Laura. 1999. "The Effect of Criminal Victimization on a Household's Moving Decision." *Criminology* 37:903–30.
- Duncan, Greg J., James P. Connell, and Pamela K. Klebanov. 1997. "Conceptual and Methodological Issues in Estimating Causal Effects of Neighborhoods and Family Conditions on Individual Development." Pp. 219–50 in *Neighborhood Poverty: Vol. 1, Context and Consequences for Children*, edited by Jeanne Brooks-Gunn, Greg J. Duncan, and J. Lawrence Aber. New York: Russell Sage.
- Duncan, Greg J., Thomas F. Juster, and James N. Morgan. 1984. "The Role of Panel Studies in the World of Scarce Research Resources." Pp. 301–28 in *The Collection and Analysis of Economic and Consumer Data: In Memory of Robert Ferber*, edited by Seymour Sudman and Mary A. Spaeth. Urbana: University of Illinois, Bureau of Economic and Business Research and Survey Research Laboratory.
- Evans, William N., Wallace E. Oates, and Robert M. Schwab. 1992. "Measuring Peer Group Effects: A Study of Teenage Behavior." *Journal of Political Economy* 100: 966–91.
- Farley, Reynolds, and William H. Frey. 1994. "Changes in the Segregation of Whites from Blacks during the 1980s: Small Steps toward a More Integrated Society." *American Sociological Review* 59:23–45.
- Fienberg, Stephen E. 1980. *The Analysis of Cross-Classified Categorical Data*, 2d ed. Cambridge, Mass.: MIT Press.
- Figlio, David N., and Maurice E. Lucas. 2000. "What's in a Grade? School Report Cards and House Prices." NBER Working Paper no. 8019.

- Frank, Kenneth A. 2000. "Impact of a Confounding Variable on a Regression Coefficient." *Sociological Methods and Research* 29:147–94.
- Freeman, Richard. 1997. *When Earnings Diverge: Causes, Consequences, and Cures for the New Inequality in the U.S.* Washington, D.C.: National Policy Association.
- Frey, William H., and Reynolds Farley. 1996. "Latino, Asian, and Black Segregation in U.S. Metropolitan Areas: Are Multiethnic Metros Different?" *Demography* 33: 35–50.
- Furstenberg, Frank F., Jr., Thomas D. Cook, Jacquelynne Eccles, Glen H. Elder, Jr., and Arnold Sameroff. 1999. *Managing to Make It: Urban Families and Adolescent Success*. Chicago: University of Chicago Press.
- Furstenberg, Frank, Jr., and Mary Elizabeth Hughes. 1997. "The Influence of Neighborhoods on Children's Development: A Theoretical Perspective and Research Agenda." Pp. 23–47 in *Neighborhood Poverty: Vol. 2, Policy Implications in Studying Neighborhoods*, edited by Jeanne Brooks-Gunn, Greg J. Duncan, and J. Lawrence Aber. New York: Russell Sage.
- Gastwirth, Joseph L., Abba M. Krieger, and Paul R. Rosenbaum. 1998. "Dual and Simultaneous Sensitivity Analysis for Matched Pairs." *Biometrika* 85:907–20.
- Ginther, Donna, Robert Haveman, and Barbara Wolfe. 2000. "Neighborhood Attributes as Determinants of Children's Outcomes: How Robust Are the Relationships?" *Journal of Human Resources* 35:603–42.
- Gramlich, Edward, Deborah Laren, and Naomi Sealand. 1992. "Moving into and out of Poor Areas." *Journal of Policy Analysis and Management* 11:273–87.
- Hagan, John, Ross MacMillan, and Blair Wheaton. 1996. "New Kid in Town: Social Capital and the Life Course Effects of Family Migration on Children." *American Sociological Review* 61:368–85.
- Hanratty, Maria, Sara S. McLanahan, and Becky Pettit. 1998. "The Impact of the Los Angeles Moving to Opportunity Program on Residential Mobility, Neighborhood Characteristics, and Early Child and Parent Outcomes." Working Paper no. 98-18. Princeton University, Bendheim-Thoman Center for Research on Child Wellbeing.
- Harris, David R. 1999. "'Property Values Drop When Blacks Move In, Because . . .': Racial and Socioeconomic Determinants of Neighborhood Desirability." *American Sociological Review* 64:461–79.
- Heckman, James, Hidehiko Ichimura, Jeffrey Smith, and Petra Todd. 1998. "Characterizing Selection Bias Using Experimental Data." *Econometrica* 66: 1017–98.
- Holland, Paul W. 1986. "Statistics and Causal Inference (with comments)." *Journal of the American Statistical Association* 81:945–70.
- Imbens, Guido W., and Joshua D. Angrist. 1994. "Identification and Estimation of Local Average Treatment Effects." *Econometrica* 62:467–75.
- Jargowsky, Paul A. 1997. *Poverty and Place: Ghettos, Barrios, and the American City*. New York: Russell Sage.
- Jencks, Christopher, and Susan E. Mayer. 1990. "The Social Consequences of Growing Up in a Poor Neighborhood," Pp. 111–86 in *Inner-City Poverty in the United States*, edited by Lawrence E. Lynn, Jr. and Michael G. H. McGreary. Washington, D.C.: National Academy Press.
- Katz, Lawrence F., Jeffrey R. Kling, and Jeffrey B. Liebman. 2001. "Moving to Opportunity in Boston: Early Results of a Randomized Mobility Experiment." *Quarterly Journal of Economics* 116:607–54.
- Ladd, Helen F., and Jens Ludwig. 1997. "Federal Housing Assistance: Residential Relocation, and Educational Opportunities: Evidence from Baltimore." *American Economic Review* 87 (AEA Papers and Proceedings): 272–77.
- Lee, Barrett A., R. S. Oropesa, and James W. Kanan. 1994. "Neighborhood Context and Residential Mobility." *Demography* 31:249–70.

- Lieberson, Stanley. 1985. *Making It Count: The Improvement of Social Research and Theory*. Berkeley and Los Angeles: University of California Press.
- Ludwig, Jens, Greg J. Duncan, and Paul Hirschfield. 2001. "Urban Poverty and Juvenile Crime: Evidence from a Randomized Housing Mobility Experiment." *Quarterly Journal of Economics* 116:655–80.
- Ludwig, Jens, Helen F. Ladd, and Greg J. Duncan. 2001. "Urban Poverty and Educational Outcomes. Pp. 147–201 in *Brookings-Wharton Papers on Urban Affairs*, edited by William Gale and Judith Rothenberg Pack. Washington, D.C.: Brookings Institution.
- Massey, Douglass S., and Nancy A. Denton. 1993. *American Apartheid: Segregation and the Making of the Underclass*. Cambridge, Mass.: Harvard University Press.
- Maynard, Rebecca A., ed. 1997. *Kids Having Kids: Economic Costs and Social Consequences of Teen Pregnancy*. Washington, D.C.: Urban Institute Press.
- Morgan, Stephen L. 2001. "Counterfactuals, Causal Effect Heterogeneity, and the Catholic School Effect on Learning." *Sociology of Education* 74:341–74.
- Pribesh, Shana, and Douglas B. Downey. 1999. "Why Are Residential and School Moves Associated with Poor School Performance?" *Demography* 36:521–34.
- Quillian, Lincoln. 1999. "Migration Patterns and the Growth of High-Poverty Neighborhoods, 1970–1990." *American Journal of Sociology* 105:1–37.
- Rhode, Paul W., and Koleman S. Strumpf. 2000. "A Historical Test of the Tiebout Hypothesis: Local Heterogeneity from 1850–1990." NBER Working Paper no. 7946.
- Rosenbaum, James. 1995. "Changing the Geography of Opportunity by Expanding Residential Choice: Lessons from the *Gautreaux* Program." *Housing Policy Debate* 6:231–69.
- Rosenbaum, James E., Nancy Fishman, Allison Brett, and Patricia Meaden. 1993. "Can the Kerner Commission's Housing Strategy Improve Employment, Education, and Social Integration for Low-Income Blacks?" *North Carolina Law Review* 71: 1519–56.
- Rosenbaum, James E., and Susan J. Popkin. 1991. "Employment and Earnings of Low-Income Blacks Who Move to Middle-Class Suburbs." Pp. 342–56 in *The Urban Underclass*, edited by Christopher Jencks and Paul E. Peterson. Washington, D.C.: Brookings.
- Rosenbaum, Paul R. 1984a. "The Consequences of Adjustment for a Concomitant Covariate That Has Been Affected by the Treatment." *Journal of the Royal Statistical Society*, ser. A, 147:656–66.
- . 1984b. "From Association to Causation in Observational Studies: The Role of Tests of Strongly Ignorable Treatment Assignment." *Journal of the American Statistical Association* 79:41–48.
- . 1995. *Observational Studies*. New York: Springer-Verlag.
- Rosenbaum, Paul R., and Donald B. Rubin. 1983a. "Assessing Sensitivity to an Unobserved Covariate in an Observational Study with Binary Outcome." *Journal of the Royal Statistical Society* 45:212–18.
- . 1983b. "The Central Role of the Propensity Score in Observational Studies for Causal Effects." *Biometrika* 70:41–55.
- . 1984. "Reducing Bias in Observational Studies Using Subclassification on the Propensity Score." *Journal of the American Statistical Association* 79:516–24.
- . 1985. "Constructing a Control Group Using Multivariate Matched Sampling Methods That Incorporate the Propensity Score." *American Statistician* 39:33–38.
- Rossi, Peter H. 1980. *Why Families Move*, 2d ed. London: Sage.
- Rubin, Donald B. 1974. "Estimating the Causal Effects of Treatment in Randomized and Nonrandomized Studies." *Journal of Educational Psychology* 66:688–701.
- . 1977. "Assignment to Treatment Group on the Basis of a Covariate." *Journal of Educational Statistics* 2:1–26.

- . 1991. "Practical Implications of the Modes of Statistical Inference for Causal Effects and the Critical Role of the Assignment Mechanism." *Biometrics* 47:1213–34.
- Rubinowitz, Leonard S., and James E. Rosenbaum. 2000. *Crossing the Class and Color Lines: From Public Housing to White Suburbia*. Chicago: University of Chicago Press.
- Sampson, Robert J., and Jeffrey D. Morenoff. 1997. "Ecological Perspectives on the Neighborhood Context of Urban Poverty: Past and Present." Pp. 1–22 in *Neighborhood Poverty: Vol. 2, Policy Implications in Studying Neighborhoods*, edited by Jeanne Brooks-Gunn, Greg J. Duncan, and J. Lawrence Aber. New York: Russell Sage.
- Small, Mario L., and Katherine Newman. 2001. "Urban Poverty after *The Truly Disadvantaged*: The Rediscovery of the Family, the Neighborhood, and Culture." *Annual Review of Sociology* 27:23–45.
- Smith, Herbert L. 1997. "Matching with Multiple Controls to Estimate Treatment Effects in Observational Studies." *Sociological Methodology* 27:325–53.
- South, Scott J., and Kyle D. Crowder. 1997a. "Escaping Distressed Neighborhoods: Individual, Community, and Metropolitan Influences." *American Journal of Sociology* 102:1040–84.
- . 1997b. "Residential Mobility between Cities and Suburbs: Race, Suburbanization, and Back-to-the-City Moves." *Demography* 34:525–38.
- . 1998a. "Avenues and Barriers to Residential Mobility among Single Mothers." *Journal of Marriage and the Family* 60:866–77.
- . 1998b. "Housing Discrimination and Residential Mobility: Impacts for Blacks and Whites." *Population Research and Policy Review* 17:369–87.
- . 1999. "Neighborhood Effects on Family Formation: Concentrated Poverty and Beyond." *American Sociological Review* 64:113–32.
- South, Scott J., and Glenn D. Deane. 1993. "Race and Residential Mobility: Individual Determinants and Structural Constraints." *Social Forces* 72:147–67.
- Speare, Alden Jr., and Frances K. Goldscheider. 1987. "Effects of Marital Status Change on Residential Mobility." *Journal of Marriage and the Family* 49:455–64.
- St. John, Craig, Mark Edwards, and Deeann Wenk. 1995. "Racial Differences in Intraurban Residential Mobility." *Urban Affairs Quarterly* 30:709–29.
- Tiebout, Charles. 1956. "A Pure Theory of Local Expenditures." *Journal of Political Economy* 64:416–24.
- Tienda, Marta. 1991. "Poor People and Poor Places: Deciphering Neighborhood Effects on Poverty Outcomes." Pp. 244–62 in *Macro-Micro Linkages in Sociology*, edited by J. Haber. Newbury Park, Calif.: Sage.
- Tolnay, Stewart E., Kyle D. Crowder, and Robert M. Adelman. 2000. "'Narrow and Filthy Alleys of the City'? The Residential Settlement Patterns of Black Southern Migrants to the North." *Social Forces* 78:989–1015.
- Wilson, William Julius. 1987. *The Truly Disadvantaged: The Inner-City, the Underclass, and Public Policy*. Chicago: University of Chicago Press.
- . 1996. *When Work Disappears: The World of the New Urban Poor*. New York: Knopf.
- Winship, Christopher, and Robert D. Mare. 1992. "Models for Sample Selection Bias." *Annual Review of Sociology* 18:327–50.
- Winship, Christopher, and Stephen L. Morgan. 1999. "The Estimation of Causal Effects from Observational Data." *Annual Review of Sociology* 25:659–707.
- Winship, Christopher, and Michael Sobel. In press. "Causal Inference in Sociological Studies." In *Handbook of Data Analysis*, edited by Melissa Hardy and Alan Bryman. Newbury Park, Calif.: Sage.
- Yinger, John. 1995. *Closed Doors, Opportunities Lost: The Continuing Costs of Housing Discrimination*. New York: Russell Sage.