



Contents lists available at ScienceDirect

Journal of Econometrics

journal homepage: www.elsevier.com/locate/jeconom

Nonparametric identification of a binary random factor in cross section data

Yingying Dong, Arthur Lewbel*

California State University Fullerton and Boston College, United States

ARTICLE INFO

Article history:

Received 22 January 2010

Received in revised form

13 March 2011

Accepted 16 March 2011

Available online 23 March 2011

JEL classification:

C25

C21

Keywords:

Mixture model

Random effects

Binary

Unobserved factor

Unobserved regressor

Nonparametric identification

Deconvolution

Treatment

ABSTRACT

Suppose V and U are two independent mean zero random variables, where V has an asymmetric distribution with two mass points and U has some zero odd moments (having a symmetric distribution suffices). We show that the distributions of V and U are nonparametrically identified just from observing the sum $V + U$, and provide a pointwise rate root n estimator. This can permit point identification of average treatment effects when the econometrician does not observe who was treated. We extend our results to include covariates X , showing that we can nonparametrically identify and estimate cross section regression models of the form $Y = g(X, D^*) + U$, where D^* is an unobserved binary regressor.

© 2011 Elsevier B.V. All rights reserved.

1. Introduction

We propose a method of nonparametrically identifying and estimating cross section regression models that contain an unobserved binary regressor or treatment, or equivalently an unobserved random effect that can take on two values. For example, suppose an experiment (natural or otherwise) with random or exogenous assignment to treatment was performed on some population, but we only have survey data collected in the region where the experiment occurred, and this survey does not report which (or even how many) individuals were treated. Then, given our assumptions, we can point identify the average treatment effect in this population and the probability of treatment, despite not observing who was treated.

No instruments or proxies for the unobserved binary regressor or treatment need to be observed. Identification is obtained by assuming that the unobserved exogenously assigned treatment or binary regressor effect is a location shift of the observed outcome, and that the regression or conditional outcome errors have zero

low order odd moments (a sufficient condition for which is symmetrically distributed errors). These identifying assumptions provide moment conditions that can be used to construct either an ordinary generalized method of moments (GMM) estimator or, in the presence of covariates, a nonparametric local GMM estimator for the model.

The zero low order odd moments used for identification here can arise in a number of contexts. Normal errors are of course symmetric and so have all odd moments equal to zero, and normality arises in many models such as those involving central limit theorems, e.g., Gibrat's law for wage or income distributions. Differences of independently, identically distributed errors, or more generally of exchangeable errors such as those following ARMA processes, are also symmetrically distributed (see, e.g., Proposition 1 of [Honore, 1992](#)). So, e.g., two period panel models with fixed effects and ARMA errors will generally have errors that are symmetric after time differencing. Our results could therefore be applied in a two period panel where individuals can have an unobserved mean shift at any time (corresponding to the unobserved binary regressor), fixed effects (which are differenced away) and exchangeable remaining errors (which yield symmetric errors after differencing). Below we give other more specific examples of models with the required odd moments being zero.

Ignoring covariates for the moment, suppose $Y = h + V + U$, where V and U are independent mean zero random variables and

* Corresponding address: Department of Economics, Boston College, 140 Commonwealth Avenue, Chestnut Hill, MA, 02467, United States. Tel.: +1 617 552 3678; fax: +1 617 552 2308.

E-mail address: lewbel@bc.edu (A. Lewbel).

URL: <http://www2.bc.edu/~lewbel/> (A. Lewbel).

h is a constant. The random V equals either b_0 or b_1 with unknown probabilities p and $1 - p$ respectively, where p does not equal a half, i.e., V is asymmetrically distributed. U is assumed to have its first few odd moments equal to zero. We observe a sample of observations of the random variable Y , and so can identify the marginal distribution of Y , but we do not observe h , V , or U .

We first show that the constant h and the distributions of V and U are nonparametrically identified just from observing Y . The only regularity assumption required is that some higher moments of Y exist. More precisely, the first three odd moments of U must be zero (and so also exist for Y) for local identification, while having the first five odd moments of U equal to zero suffices for global identification.

We also provide estimators for the distributions of V and U . We show that the constant h , the probability mass function of V , moments of the distribution of U , and points of the distribution function of U can all be estimated using GMM. Unlike common deconvolution estimators that can converge at slow rates, we estimate the distributions of V and U , and the density of U (if it is continuous) at the same rates of convergence as if V and U were separately observed, instead of just observing their sum.

We do not assume that the supports of V or U are known, so estimation of the distribution of V means identifying and estimating both of its support points b_0 and b_1 , as well as the probabilities p and $1 - p$, respectively, of V equaling b_0 or b_1 .

One can write V as $V = b_1 D^* + b_0(1 - D^*)$, where D^* is an unobserved binary indicator. For example, if D^* is the unobserved indicator of exogenously assigned treatment, then $b_1 - b_0$ is the average treatment effect, p is the probability of treatment, and U describes the remaining heterogeneity of outcomes Y (here treatment is assumed to only cause a shift in outcome means).

We also show how these results can be extended to allow for covariates. If h depends on a vector of covariates X while V and U are independent of X , then we obtain the random effects regression model $Y = h(X) + V + U$, which is popular for panel data, but which we identify and estimate just from cross section data.

More generally, we allow both h and the distributions of V and U to depend in unknown ways on X . This is equivalent to nonparametric identification and estimation of a regression model containing an unobserved binary regressor. The regression model is $Y = g(X, D^*) + U$, where g is an unknown function, D^* is an unobserved binary regressor (or unobserved indicator of treatment) that equals zero with unknown probability $p(X)$ and one with probability $1 - p(X)$, while U is a random error with an unknown conditional distribution $F_U(U | X)$ having its first few odd moments equal to zero (conditional symmetry conditioning on X suffices). The unobserved random variables U and D^* are assumed to be conditionally independent, conditioning upon X . By defining $h(x) = E(Y | X = x) = E[g(X, D^*) | X = x]$, $V = g(X, D^*) - h(X)$ and $U = Y - h(X) - V$, this regression model can then be rewritten as $Y = h(X) + V + U$, where $h(x)$ is a nonparametric regression function of Y on X , and the two support points of V conditional on $X = x$ are then $b_d(x) = g(x, d) - h(x)$ for $d = 0, 1$. Kitamura (unpublished manuscript) provides some nonparametric identification results for this model by placing constraints on how the distributions can depend upon X , while we place no such restrictions on the distribution of X and instead restrict the shape of the distribution of U .

The assumptions we impose on U in $Y = g(X, D^*) + U$ are common assumptions in regression models, e.g., they allow the error U to be heteroskedastic with respect to X , and they hold, e.g., if U given X is normal (though normality is not required). Also, regression model errors U are sometimes interpreted as measurement error in Y , and measurement errors are often assumed to be symmetric.

If D^* is an unobserved treatment indicator, then $g(X, 1) - g(X, 0)$ is the conditional average treatment effect, which may

be averaged over X to obtain an unconditional average treatment effect. Symmetry of errors is not usually assumed for treatment models, but suppose we have panel data (two periods of observations) and all treatments occur in one of the two periods. Then, as noted above, the required symmetry of U errors would result automatically from time differencing the data, given the standard panel model assumption of individual specific fixed effects plus independently, identically distributed (or more generally ARMA or other exchangeable) errors.

Another possible application of these extensions is a stochastic frontier model, where Y is the log of a firm's output, X are factors of production, and D^* indicates whether the firm operates efficiently at the frontier, or inefficiently. Existing stochastic frontier models obtain identification either by assuming parametric functional forms for both the distributions of V and U , or by using panel data and assuming that each firm's individual efficiency level is a fixed effect that is constant over time. See, e.g., Kumbhakar et al. (2007) and Simar and Wilson (2007). In contrast, our assumptions and associated estimators could be used to estimate a nonparametric stochastic frontier model using cross section data, given the restriction that unobserved efficiency is indexed by a binary D^* . Note that virtually all stochastic frontier models based on cross section data assume U given X is symmetrically distributed.

Dong (forthcoming) empirically estimates a model where $Y = h(X) + V + U$, based on symmetry of U and using moments similar to the exponentials we suggest in an extension section. Our results formally prove identification of Dong's model, and our estimator is more general in that it allows V and the distribution of U to depend in arbitrary ways on X . Hu and Lewbel (2008) also nonparametrically identify some features of a model containing an unobserved binary regressor, using either a type of instrumental variable or an assumption of conditional independence of low order moments.

Models that allocate individuals into various types, as D^* does, are common in the statistics and marketing literatures. Examples include cluster analysis, latent class analysis, and mixture models (see, e.g., Clogg, 1995 and Hagenaars and McCutcheon, 2002). Our model resembles a finite (two distribution) mixture model, but differs crucially in that, for identification, finite mixture models usually require the distributions being mixed to be parametrically specified, while in our model U is nonparametric. While general mixture models are more flexible than ours in allowing for more than two groups and permitting the U distribution to vary across groups, ours is more flexible in letting U be nonparametric, essentially allowing for an infinite number of parameters versus finitely parameterized mixtures.

Some mixture models can be nonparametrically identified by observing draws of vectors of data, where the number of elements of the observed vectors is larger than the number of distributions being mixed. Examples include Hall and Zhou (2003) and Kasahara and Shimotsu (2009). In contrast, we obtain identification with a scalar Y . As noted above, Kitamura (unpublished manuscript) also obtains nonparametric identification with a scalar Y , but does so by requiring observation of a covariate that affects the component distributions with some restrictions. Another closely related mixture model result is Bordes et al. (2006), who impose strictly stronger conditions than we do, including that U is symmetric and continuously distributed.

Also related is the literature on mismeasured binary regressors, where identification generally requires instruments. An exception is Chen et al. (2008). Like our Theorem 1, they exploit error symmetry for identification, but unlike this paper they assume that the binary regressor is observed, though with some measurement (classification) error, instead of being completely unobserved. A more closely related result is Heckman and Robb (1985), who like us use zero low order odd moments to identify a binary effect,

though theirs is a restricted effect that is strictly nested in our results. Error symmetry has also been used to obtain identification in a variety of other econometric contexts, e.g., [Powell \(1986\)](#).

There are a few common ways of identifying the distributions of random variables given just their sum. One method of identification assumes that the exact distribution of one of the two errors is known a priori (e.g., from a validation sample as is common in the statistics literature on measurement error; see, e.g., [Carroll et al., 2006](#)) and uses deconvolution to obtain the distribution of the other one. For example, if U were normal, one would need to know a priori its mean and variance to estimate the distribution of V . A second standard way to obtain identification is to parameterize both the distributions of V and U , as in most of the latent class literature or in the stochastic frontier literature (see, e.g., [Kumbhakar and Lovell, 2000](#)), where a typical parameterization is to have V be log normal and U be normal. Panel data models often have errors of the form $V + U$ that are identified either by imposing specific error structures or by assuming one of the errors is fixed over time (see, e.g., [Baltagi, 2008](#) for a survey of random effects and fixed effects panel data models). Past nonparametric stochastic frontier models have similarly required panel data for identification, as described above. In contrast to all these identification methods, in our model both U and V have unknown distributions, and no panel data are required.

The next section contains our main identification result. We then provide moment conditions for estimating the model, including the distribution of V (its support points and the associated probability mass function), using ordinary GMM. Next we give estimators for the distribution and density function of U . We provide a Monte Carlo analysis showing that our estimator performs reasonably well even compared to infeasible maximum likelihood estimation. This is followed by some extensions showing how our identification and estimation methods can be augmented to provide additional moments for estimation, and to allow for covariates. Proofs are given in the [Appendix](#).

2. Identification

In this section, we provide our general identification results. Later we extend these results to include covariates X .

Assumption A1. Let $Y = h + V + U$. Assume the distribution of V is mean zero, asymmetric, and has exactly two points of support. Assume $E(U^d | V) = E(U^d)$ exists for all positive integers $d \leq 9$, and $E(U^{2d-1}) = 0$ for all positive integers $d \leq 5$.

Theorem 1. Let [Assumption A1](#) hold, and assume the distribution of Y is identified. Then the constant h and the distributions of V and U are identified.

Let b_0 and b_1 denote the two support points of the distribution of V , where without loss of generality $b_0 < b_1$, and let p be the probability that $V = b_0$, so $1 - p$ is the probability that $V = b_1$. Identification of the distribution of V by [Theorem 1](#) means identification of b_0 , b_1 , and p . If Y is the outcome of a treatment and D^* denotes an unobserved treatment indicator, then we can define $V = b_1 D^* + b_0(1 - D^*)$, and we will have identification of the probability of treatment p and identification of the average treatment effect $b_1 - b_0$ (which by mean independence of V and U will also equal the average treatment effect on the treated).

[Assumption A1](#) says that the first nine moments of U conditional on V are the same as the moments that would arise if U were distributed symmetrically and independent of V . The only regularity condition required for identification of the distribution of V in [Theorem 1](#) is existence of $E(Y^9)$.

The proof of [Theorem 1](#) also shows identification up to a finite set, and hence local identification of the distribution of V , assuming

just $E(U) = E(U^3) = E(U^5) = 0$ and $E(U^d | V) = E(U^d)$ for positive integers $d \leq 5$, thereby only needing existence of $E(Y^5)$. Higher moments going up to the ninth moment are required only to distinguish amongst the elements in the finite identified set and thereby provide global identification.

Note that [Theorem 1](#) assumes asymmetry of V (since otherwise it would be indistinguishable from U) and more than one point of support (since otherwise it would be indistinguishable from h), which is equivalent to requiring that p not be exactly equal to zero, one, or one half. This suggests that the identification and associated estimation will be weak if the actual p is very close to any of these values. In practice, it would be easy to tell if this problem exists, because if it did then the observed Y would itself be close to symmetrically distributed. Applying a formal test of data symmetry such as [Ahmed and Li \(1997\)](#) to the Y data is equivalent in our model to testing if p equals zero, one, or one half. More simply, one might look for substantial asymmetry in a histogram or kernel density estimate of Y .

We next consider estimation of h , b_0 , b_1 , and p , and then later show how the rest of the model, i.e., the distribution function of U , can be estimated.

3. Estimation

Our estimator will take the form of the standard generalized method of moments (GMM, as in [Hansen, 1982](#)), since given data $Y_1, \dots, Y_n$, we will below construct a set of moments of the form $E[G(Y, \theta)] = 0$, where G is a set of known functions and the vector θ consists of the parameters of interest h , b_0 , p , as well as u_2 , u_4 , and u_6 , where $u_d = E(U^d)$. The parameters u_2 , u_4 , and u_6 are nuisance parameters for estimating the V distribution, but in some applications they may be of interest as summary measures of the distribution of U .

Let $v_d = E(V^d)$. Then $v_1 = E(V) = b_0 p + b_1(1 - p) = 0$, so

$$b_1 = b_0 p / (p - 1) \quad (1)$$

and, therefore,

$$v_d = E(V^d) = b_0^d p + \left(\frac{b_0 p}{p - 1} \right)^d (1 - p). \quad (2)$$

Now expand the expression $E[(Y - h)^d - (V + U)^d] = 0$ for integers d , noting by [Assumption A1](#) that the first five odd moments of U are zero. The results are

$$E(Y - h) = 0 \quad (3)$$

$$E((Y - h)^2 - (v_2 + u_2)) = 0 \quad (4)$$

$$E((Y - h)^3 - v_3) = 0 \quad (5)$$

$$E((Y - h)^4 - (v_4 + 6v_2 u_2 + u_4)) = 0 \quad (6)$$

$$E((Y - h)^5 - (v_5 + 10v_3 u_2)) = 0 \quad (7)$$

$$E((Y - h)^6 - (v_6 + 15v_4 u_2 + 15v_2 u_4 + u_6)) = 0 \quad (8)$$

$$E((Y - h)^7 - (v_7 + 21v_5 u_2 + 35v_3 u_4)) = 0 \quad (9)$$

$$E((Y - h)^9 - (v_9 + 36v_7 u_2 + 126v_5 u_4 + 84v_3 u_6)) = 0. \quad (10)$$

Substituting Eq. (2) into Eqs. (3)–(10) gives eight moments we can write as $E[G(Y, \theta)] = 0$ in the six unknown parameters $\theta = (h, b_0, p, u_2, u_4, u_6)$, which we use for estimation via GMM. We use these moments because the proof of [Theorem 1](#) shows that these particular eight equations are exactly those required to point identify the parameters defining the distribution of V . As shown in the proof, more equations than unknowns are required for global identification because of the nonlinearity of these equations, and in particular the presence of multiple roots. Given an estimate of θ , the estimate of b_1 is then obtained by Eq. (1).

Based on [Theorem 1](#), for this estimator we assume that $Y_1, \dots, Y_n$ are identically distributed (or more precisely, have identical first nine moments); however, the Y observations do not need to be independent, since GMM estimation theory permits some serial dependence in the data. Standard GMM limiting distribution theory applied to our moments provides root n consistent, asymptotically normal estimates of θ and hence of h and of the distribution of V (i.e., the support points b_0 and b_1 and the probability p , where $\hat{b}_1$ is obtained by $\hat{b}_1 = \hat{b}_0 \hat{p} / (\hat{p} - 1)$ from [Eq. \(11\)](#)). Without loss of generality we have imposed $b_0 < b_1$ (if this is violated then the definitions of these two parameters can be switched to make the inequality hold), and this along with $E(V) = 0$ implies that $\hat{b}_0$ is negative and $\hat{b}_1$ is positive, which may be imposed on estimation. One could also impose that $\hat{p}$ lie between zero and one, and that $\hat{u}_2, \hat{u}_4$, and $\hat{u}_6$ be positive.

One might anticipate poor empirical results, and great sensitivity to outliers or extreme observations of Y , given the use of such high order moments for estimation. For example, [Altonji and Segal \(1996\)](#) document bias in GMM estimates associated with just second order moments. However, we found that these problems rarely arose in our Monte Carlo simulations (in applications one would want to carefully scale Y to avoid the effects of computer rounding errors based on inverting matrix entries of varying orders of magnitude). We believe the reason the estimator performs reasonably well is that only lower order moments are required for local identification, up to a small finite set of values. The higher order moments (specifically those above the fifth) are only needed for global identification to distinguish between these few possible multiple solutions of the low order polynomials. For example, if by the low order moments b_1 was identified up to a value in the neighborhood of either 1 or 3, then even poorly estimated higher moments could succeed in distinguishing between these two neighborhoods, by having a sample GMM moment mean be substantially closer to zero in the neighborhood of one value rather than the other.

In an extension section we describe how additional moments could be constructed for estimation based on symmetry of U . These alternative moments might be employed in applications where the polynomial based moments are found to be problematic. Another possibility would be to Winsorize extreme observations of the Y data prior to estimation to robustify the higher moment estimates.

4. The distribution of U

For any random variable Z , let F_Z denote the marginal cumulative distribution function of Z . Define $\varepsilon = V + U$. Define $F_{KU}(u)$ by

$$F_{KU}(u) = \begin{cases} \sum_{k=0}^{K-1} \left(\frac{1-p}{p} \right)^k \frac{1}{p} F_\varepsilon(u + b_0 + (b_0 - b_1)k) & \text{if } p > 1/2 \\ \sum_{k=0}^{K-1} \left(\frac{p}{1-p} \right)^k \frac{1}{1-p} F_\varepsilon(u + b_1 + (b_1 - b_0)k) & \text{if } p < 1/2. \end{cases}$$

The proof of [Theorem 1](#) shows that $F_U(u) = F_{KU}(u) + R_K$, where $0 \leq R_K \leq \min \left(\left(\frac{p}{1-p} \right)^K, \left(\frac{1-p}{p} \right)^K \right)$, so the remainder term $R_K \rightarrow 0$ as $K \rightarrow \infty$ (since $p = 1/2$ is ruled out). This suggests that F_U could be estimated by $F_{KU}(u)$ after replacing F_ε with the empirical distribution of $Y - \hat{h}$, replacing b_0, b_1 , and p with their estimates, and letting $K \rightarrow \infty$ as $N \rightarrow \infty$.

However, under the assumption that U is symmetrically distributed, the following theorem provides a more convenient way to estimate the distribution function of U . Define

$$\Psi(u) = \frac{[F_\varepsilon(-u + b_0) - 1]p + F_\varepsilon(u + b_1)(1-p)}{1-2p}. \quad (11)$$

Theorem 2. Let [Assumption A1](#) hold. Assume U is symmetrically distributed. Then

$$F_U(u) = \frac{\Psi(u) - \Psi(-u) + 1}{2}. \quad (12)$$

[Theorem 2](#) provides a direct expression for the distribution of U in terms of b_0, b_1, p and the distribution of ε , all of which are previously identified. This can be used to construct an estimator for $F_U(u)$ as follows.

Let $I(\cdot)$ denote the indicator function that equals one if $\cdot$ is true and zero otherwise, and let θ be a vector containing h, b_0, b_1 , and p . Define the function $\omega(Y, u, \theta)$ by

$$\omega(Y, u, \theta) = \frac{[I(Y \leq h - u + b_0) - 1]p + I(Y \leq h + u + b_1)(1-p)}{1-2p}. \quad (13)$$

Then using $Y = h + \varepsilon$ it follows immediately from [Eq. \(11\)](#) that

$$\Psi(u) = E(\omega(Y, u, \theta)). \quad (14)$$

An estimator for $F_U(u)$ can now be constructed by replacing the parameters in [Eq. \(14\)](#) with estimates, replacing the expectation with a sample average, and plugging the result into [Eq. \(12\)](#). The resulting estimator is

$$\hat{F}_U(u) = \frac{1}{n} \sum_{i=1}^n \frac{\omega(Y_i, u, \hat{\theta}) - \omega(Y_i, -u, \hat{\theta}) + 1}{2}. \quad (15)$$

Alternatively, $F_U(u)$ for a finite number of values of u , say $u_1, \dots, u_J$, can be estimated as follows. Recall that $E[G(Y, \theta)] = 0$ was used to estimate the parameters h, b_0, b_1, p by GMM. For notational convenience, let $\eta_j = F_U(u_j)$ for each u_j . Then by [Eqs. \(12\) and \(14\)](#)

$$E \left[\eta_j - \frac{\omega(Y, u_j, \theta) - \omega(Y, u_j, \theta) + 1}{2} \right] = 0. \quad (16)$$

Adding [Eq. \(16\)](#) for $j = 1, \dots, J$ to the set of functions defining G , including $\eta_1, \dots, \eta_J$ in the vector θ , and then applying GMM to this augmented set of moment conditions $E[G(Y, \theta)] = 0$ simultaneously yields root n consistent, asymptotically normal estimates of h, b_0, b_1, p and $\eta_j = F_U(u_j)$ for $j = 1, \dots, J$. An advantage of this approach versus [Eq. \(15\)](#) is that GMM limiting distribution theory then provides standard error estimates for each $\hat{F}_U(u_j)$.

While p is the unconditional probability that $V = b_0$, given $\hat{F}_U$ it is straightforward to estimate conditional probabilities as well. In particular,

$$\begin{aligned} \Pr(V = b_0 | Y \leq y) &= \Pr(V = b_0, Y \leq y) / \Pr(Y \leq y) \\ &= F_U(y - h - b_0) / F_Y(y), \end{aligned}$$

which could be estimated as $\hat{F}_U(y - \hat{h} - \hat{b}_0) / \hat{F}_Y(y)$, where $\hat{F}_Y$ is the empirical distribution of Y .

Let f_Z denote the probability density function of any continuously distributed random variable Z . So far no assumption has been made about whether U is continuous or discrete. However, if U is

continuous, then ε and Y are also continuous, and then taking the derivative of Eqs. (11) and (12) with respect to u gives

$$\psi(u) = \frac{-f_\varepsilon(-u + b_0)p + f_\varepsilon(u + b_1)(1 - p)}{1 - 2p}, \quad (17)$$

$$f_U(u) = \frac{\psi(u) + \psi(-u)}{2},$$

which suggests the estimators

$$\hat{\psi}(u) = \frac{-\hat{f}_\varepsilon(-u + \hat{b}_0)\hat{p} + \hat{f}_\varepsilon(u + \hat{b}_1)(1 - \hat{p})}{1 - 2\hat{p}}, \quad (18)$$

$$\hat{f}_U(u) = \frac{\hat{\psi}(u) + \hat{\psi}(-u)}{2}, \quad (19)$$

where $\hat{f}_\varepsilon(\varepsilon)$ is a kernel density or other estimator of $f_\varepsilon(\varepsilon)$, constructed using data $\hat{\varepsilon}_i = Y_i - \hat{h}$ for $i = 1, \dots, n$. Since densities converge at slower than rate root n , the limiting distribution of this estimator will generally be the same as if $\hat{h}$, $\hat{b}_0$, $\hat{b}_1$, and $\hat{p}$ were evaluated at their true values (e.g., this holds if f is differentiable by a mean value expansion of $\hat{\psi}$ around the true values of h , b_0 , b_1 , and p). The above $\hat{f}_U(u)$ is just the weighted sum of two density estimators, each one dimensional, and so will converge at the same rate as a one dimensional density estimator. For example, this will be the pointwise rate $n^{-2/5}$ using a kernel density estimator of $\hat{f}_\varepsilon$ under standard assumptions as in Silverman (1986) (independent observations, f_ε twice differentiable, evaluated at points not on the boundary of the support of ε , bandwidth proportional to $n^{-1/5}$, and a second order kernel function) with p bounded away from 1/2. It is possible for $\hat{f}_U(u)$ to be negative in finite samples, so if desired one could replace negative values of $\hat{f}_U(u)$ with zero.

A potential numerical problem is that Eq. (18) may require evaluating $\hat{f}_\varepsilon$ at a value that is outside the range of observed values of $\hat{\varepsilon}_i$. Since both $\hat{\psi}(u)$ and $\hat{\psi}(-u)$ are consistent estimators of $\hat{f}_U(u)$ (though generally less precise than Eq. (19) because they individually ignore the symmetry constraint), one could use either $\hat{\psi}(u)$ or $\hat{\psi}(-u)$ instead of their average to estimate $\hat{f}_U(u)$ whenever $\hat{\psi}(-u)$ or $\hat{\psi}(u)$, respectively, requires evaluating $\hat{f}_\varepsilon$ at a point outside the range of observed values of $\hat{\varepsilon}_i$.

This construction also suggests a specification test for the model. Since symmetry of U implies that $\hat{\psi}(u) = \hat{\psi}(-u)$ one could base a test on whether $\int_0^L [\hat{\psi}(u) - \hat{\psi}(-u)]^2 w(u) du = 0$, where $w(u)$ is a weighting function that integrates to one, and L is in the range of values for which neither $\hat{\psi}(-u)$ nor $\hat{\psi}(u)$ requires evaluating $\hat{f}_\varepsilon$ at a point outside the range of observed values of $\hat{\varepsilon}_i$. The limiting distribution theory for this type of test statistic (a degenerate U statistic under the null) based on functions of kernel densities is standard, and in this case would closely resemble Ahmed and Li (1997).

5. Monte Carlo analysis

Our Monte Carlo design takes $h = 0$, U standard normal, and $-b_0p = b_1(1 - p) = 1$. We consider three different values of p between one half and one, specifically 0.6, 0.8, and 0.95. By symmetry of this design, we should obtain the same results in terms of accuracy if we took p equal to 0.4, 0.2, and 0.05, respectively. The nuisance parameters, derived from the distribution of U , are then $u_2 = 1$, $u_4 = 3$, and $u_6 = 15$. Our sample size is $n = 1000$, and for each design we perform 1000 Monte Carlo replications. Each draw of Y in each replication is constructed by drawing an observation of V and one of U from their above described distributions and then summing the two.

In each simulated data set we first estimated h as the sample mean of Y , then performed standard two step GMM (using the identity matrix as the weighting matrix in the first step) with the

Table 1
 $p = 0.6$.

Parameter	b_1	b_0	p	u_2	u_4	u_6
True	2.500	-1.667	0.600	1.000	3.000	15.000
Median	2.491	-1.657	0.601	0.983	2.786	12.309
Mean	2.374	-1.578	0.586	1.234	4.846	34.142
Stddev	0.506	0.347	0.099	0.925	7.946	90.287
Root MSE	0.522	0.358	0.100	0.954	8.154	92.250
Median ABS ERR	0.068	0.057	0.013	0.049	0.380	3.927
Mean ABS ERR	0.201	0.146	0.034	0.302	2.399	25.013
25% quantile	2.422	-1.712	0.588	0.942	2.518	9.973
75% quantile	2.561	-1.599	0.614	1.034	3.106	15.210

Table 2
 $p = 0.8$.

Parameter	b_1	b_0	p	u_2	u_4	u_6
True	5.000	-1.250	0.800	1.000	3.000	15.000
Median	4.955	-1.242	0.799	1.007	2.858	13.222
Mean	4.601	-1.267	0.756	1.018	3.228	17.151
Stddev	1.235	0.402	0.153	0.368	2.657	26.392
Root MSE	1.298	0.403	0.159	0.368	2.665	26.466
Median ABS ERR	0.094	0.065	0.009	0.065	0.337	3.106
Mean ABS ERR	0.459	0.175	0.054	0.164	0.884	7.905
25% quantile	4.871	-1.305	0.789	0.946	2.548	10.549
75% quantile	5.036	-1.178	0.808	1.078	3.169	15.951

Table 3
 $p = 0.95$.

Parameter	b_1	b_0	p	u_2	u_4	u_6
True	20.000	-1.053	0.950	1.000	3.000	15.000
Median	19.948	-1.051	0.950	0.988	2.890	14.997
Mean	19.852	-1.051	0.947	0.987	2.852	14.501
Stddev	1.404	0.161	0.043	0.122	0.529	16.208
Root MSE	1.411	0.161	0.043	0.123	0.549	16.208
Median ABS ERR	0.142	0.103	0.005	0.083	0.315	0.010
Mean ABS ERR	0.265	0.123	0.008	0.096	0.397	4.348
25% quantile	19.816	-1.151	0.945	0.903	2.527	14.946
75% quantile	20.083	-0.943	0.955	1.064	3.121	15.002

moments Eqs. (4)–(10). We imposed the inequality constraints on estimation that p lie between zero and one and that b_1 , u_2 , u_4 , and u_6 are positive. Rarely, this GMM either failed to converge after many iterations, or iterated towards one of these boundary points. When this happened, we applied two step GMM to just the low order (sufficient for local identification) moments (4), (5) and (7), and then used the results as starting values for two step GMM using all the moments (4)–(10). We could have alternatively performed a more time consuming grid search, but this procedure led to estimates that converged to interior points in all but a handful of simulations. Specifically, in fewer than one half of one percent of replications this procedure either failed to converge or produced estimates of p that approached the boundaries of zero or one. We drop these few failed replications from our reported results.

The results are reported in Tables 1–3. The parameters of the distribution of V (b_0 , b_1 , and p) are estimated with reasonable accuracy, having relatively small root mean squared errors and interquartile ranges. These parameters are very close to median unbiased, but have mean bias of a few percent, with p always mean biased downwards. This is likely because estimates of $\hat{p}$ are more or less equally likely to be above or below the true (yielding very small median bias), but when they are below the true they can be much further from the truth than when they are too high, e.g., when $p = 0.8$ an estimate that is too high can be biased by at most 0.2, while the downward bias can be as large as -0.8. Some fraction of

these replications may be centering around incorrect roots of the polynomial moments, which can be quite distant from the correct roots.

The high order nuisance parameters u_4 and u_6 are sometimes estimated very poorly. In particular, for $p = 0.6$ the median bias of u_6 is almost -20% while the mean bias is over five times larger and has the opposite sign of the median bias. The fact that the high order moment nuisance parameters are generally much more poorly estimated than the parameters of interest supports our claim that low order moments are providing most of the parameter estimation precision, while higher order moments mainly serve to distinguish among discretely separated local alternatives.

We performed limited experiments with alternative sample sizes, which are not reported to save space. Precision increases with sample size pretty much as one would expect. More substantial is the frequency with which numerical problems were encountered, e.g., at $n = 500$ we encountered convergence or boundary problems in 1.4% of replications while these problems were almost nonexistent at $n = 5000$.

6. Extension 1: additional moments

Here we provide additional moments that might be used for estimating the parameters h , b_0 , b_1 and p .

Proposition 1. *Let $Y = h + V + U$. Assume the distribution of V is mean zero, asymmetric, and has exactly two points of support. Assume U is symmetrically distributed around zero and is independent of V . Assume $E[\exp(TU)]$ exists for some positive constant T . Then for any positive $\tau \leq T$ there exists a constant α_τ such that the following two equations hold with $r = p/(1-p)$:*

$$E[\exp(\tau(Y-h)) - (r \exp(\tau b_0) + \exp(-\tau r))\alpha_\tau] = 0 \quad (20)$$

$$E[\exp(-\tau(Y-h)) - (r \exp(-\tau b_0) + \exp(\tau b_0))\alpha_\tau] = 0. \quad (21)$$

Given a set of L positive values for τ , i.e., constants $\tau_1, \dots, \tau_L$, each of which is less than T , Eqs. (20) and (21) provide $2L$ moment conditions satisfied by the set of $L + 3$ parameters $\alpha_{\tau_1}, \dots, \alpha_{\tau_L}$, h , p , and b_0 . Although the order condition for identification is therefore satisfied with $L \geq 3$, we do not have a proof analogous to Theorem 1 showing that the parameters are actually globally identified based on any number of these moments. Also, Proposition 1 is based on means of exponents, and so requires Y to have a thinner tailed distribution than estimation based on the polynomial Eqs. (3)–(10). Still, if global identification holds with these parameters, then they could be used by themselves for estimation, otherwise they could be combined with the polynomial moments to possibly increase estimation efficiency.

One could also construct moments of complex exponentials based on the characteristic function of $Y - h$ instead of those based on the moment generating function as in Proposition 1, which avoids the requirement for thin tailed distributions. However, such moments could sometimes vanish and thereby be uninformative, as when U is uniform.

Proposition 1 actually provides a continuum of moments, so rather than just choose a finite number of values for τ , it would also be possible to efficiently combine all the moments given by an interval of values of τ using, e.g., Carrasco and Florens (2000).

7. Extension 2: h depends on covariates

We now extend our results by permitting h to depend on covariates X . Estimators associated with this extension will take the form of standard two step estimators with a uniformly consistent first step.

Corollary 1. *Assume the conditional distribution of Y given X is identified and its mean exists. Let $Y = h(X) + V + U$. Let Assumption A1 hold. Assume V and U are independent of X . Then the function $h(X)$ and distributions of U and V are identified.*

Corollary 1 extends Theorem 1 by allowing the conditional mean of Y to nonparametrically depend on X . Given the assumptions of Corollary 1, it follows immediately that Eqs. (3)–(10) hold replacing h with $h(X)$, and if U is symmetrically distributed and independent of V and X then Eqs. (20) and (21) also hold replacing h with $h(X)$. This suggests a couple of ways of extending the GMM estimators of the previous section. One method is to first estimate $h(X)$ by a uniformly consistent nonparametric mean regression of Y on X (e.g., a kernel regression over a compact set of X values on the interior of its support), then replace $Y - h$ in Eqs. (3)–(10) and/or Eqs. (20) and (21) with $\varepsilon = Y - h(X)$, and apply ordinary GMM to the resulting moment conditions (using as data $\hat{\varepsilon}_i = Y_i - \hat{h}(X_i)$ for $i = 1, \dots, n$) to estimate the parameters b_0 , b_1 , p , u_2 , u_4 , and u_6 . Consistency of this estimator follows immediately from the uniform consistency of $\hat{h}$ and ordinary consistency of GMM. This estimator is easy to implement because it only depends on ordinary nonparametric regression and ordinary GMM. Root n limiting distribution theory may be immediately obtained by applying generic two step estimation theorems as in Newey and McFadden (1994).

After replacing $\hat{h}$ with $\hat{h}(X_i)$, Eq. (15) can be used to estimate the distribution of U , or alternatively Eq. (16) for $j = 1, \dots, J$, replacing h with $h(X)$, can be included in the set of functions defining G in the estimator described above. Since ε has the same properties here as before, given uniform consistency of $\hat{h}(X)$, the estimator (19) will still consistently estimate the density of U if it is continuous, using as data $\hat{\varepsilon}_i = Y_i - \hat{h}(X_i)$ for $i = 1, \dots, n$ to estimate the density function f_ε .

8. Extension 3: nonparametric regression with an unobserved binary regressor

This section extends previous results to a more general nonparametric regression model of the form $Y = g(X, D^*) + U$. Specifically, we have the following corollary.

Corollary 2. *Assume the joint distribution of Y, X is identified and that $g(X, D^*) = E(Y | X, D^*)$ exists, where D^* is an unobserved variable with support $\{0, 1\}$. Assume that the distribution of $g(X, D^*)$ conditional upon X is asymmetric for all X on its support. Define $p(X) = E(1 - D^* | X)$ and define $U = Y - g(X, D^*)$. Assume $E(U^d | X, D^*) = E(U^d | X)$ exists for all integers $d \leq 9$ and $E(U^{2d-1} | X) = 0$ for all positive integers $d \leq 5$. Then the functions $g(X, D^*)$, $p(X)$, and the distribution of U are identified.*

Corollary 2 permits all of the parameters of the model to vary nonparametrically with X . It provides identification of the regression model $Y = g(X, D^*) + U$, allowing the unobserved model error U to be heteroskedastic (and have nonconstant higher moments as well), though the variance and other low order even moments of U can only depend on X and not on the unobserved regressor D^* . As noted in the introduction and in the proof of this corollary, $Y = g(X, D^*) + U$ is equivalent to $Y = h(X) + V + U$. However, unlike Corollary 1, now V and U have distributions that can depend on X . As with Theorem 1, symmetry of U (now

conditional on X) suffices to make the required low order odd moments of U be zero.

Given the assumptions of [Corollary 2](#), Eqs. (3)–(10), and given symmetry of U , Eqs. (20) and (21), will all hold after replacing the parameters h, b_0, b_1, p, u_j , and τ_ℓ , and with functions $h(X), b_0(X), b_1(X), p(X), u_j(X)$, and $\tau_\ell(X)$ and replacing the unconditional expectations in these equations with conditional expectations, conditioning on $X = x$. If desired, we can further replace $b_0(X)$ and $b_1(X)$ with $g(x, 0) - h(x)$ and $g(x, 1) - h(x)$, respectively, to directly obtain estimates of the function $g(X, D^*)$ instead of $b_0(X)$ and $b_1(X)$.

Let $q(x)$ be the vector of all of the above listed unknown functions. Then these conditional expectations can be written as

$$E[G(q(x), Y) | X = x] = 0 \quad (22)$$

for a vector of known functions G . Eq. (22) is in the form of conditional GMM which could be estimated using [Ai and Chen \(2003\)](#), replacing all of the unknown functions $q(x)$ with sieves (related estimators are [Carrasco and Florens, 2000](#) and [Newey and Powell, 2003](#)). However, given independent, identically distributed draws of X, Y , the local GMM estimator of [Lewbel \(2007\)](#) may be easier to use because it exploits the special structure we have here, where all the functions $q(x)$ to be estimated depend on the same variables that the moments are conditioned upon, that is, $X = x$. We summarize here how this local GMM estimator would be implemented. See the online supplement [Appendix B](#) to this paper or [Lewbel \(2007\)](#) for details regarding the associated limiting distribution theory.

1. For any value of x , construct data $Z_i = K((x - X_i)/b)$ for $i = 1, \dots, n$, where K is an ordinary kernel function (e.g., the standard normal density function) and b is a bandwidth parameter. As is common practice when using kernel functions, it is a good idea to first standardize the data by scaling each continuous element of X by its sample standard deviation.
2. Obtain $\hat{\theta}$ by applying standard two step GMM based on the moment conditions $E(G(\theta, Y)Z) = 0$ for G from Eq. (22).
3. For the given value of x , let $\hat{q}(x) = \hat{\theta}$.
4. Repeat these steps using every value of x for which one wishes to estimate the vector of functions $q(x)$. For example, one may repeat these steps for a fine grid of x points on the support of X , or repeat these steps for x equal to each data point X_i to just estimate the functions $q(x)$ at the observed data points.

Note that this local GMM estimator can be used when X contains both continuous and discretely distributed elements. If all elements of X are discrete, then the estimator simplifies back to [Hansen's \(1982\)](#) original GMM.

9. Discrete V with more than two support points

A simple counting argument suggests that it may be possible to extend this paper's identification and associated estimators to applications where V is discrete with more than two points of support, as follows. Suppose V takes on the values $b_0, b_1, \dots, b_H$ with probabilities $p_0, p_1, \dots, p_H$. Let $u_j = E(U^j)$ for integers j as before. Then for any positive odd integer S , the moments $E(Y^s)$ for $s = 1, \dots, S$ equal known functions of the $2H + (S + 1)/2$ parameters $b_1, b_2, \dots, b_H, p_1, p_2, \dots, p_H, u_2, u_4, \dots, u_{S-1}, h$. Note that p_0 and b_0 can be expressed as functions of the other parameters by probabilities summing to one and V having mean zero, and we assume u_s for odd values of $s \leq S$ are zero. Therefore, with any odd $S \geq 4H + 1$, $E(Y^s)$ for $s = 1, \dots, S$ provides at least as many moment equations as unknowns, which could be used to estimate these parameters by GMM and will generally suffice for local identification. These moments include polynomials with up to $S - 1$ roots, so having S much larger than $4H + 1$ may be necessary for

global identification, just as the proof of [Theorem 1](#) requires $S = 9$ even though in that theorem $H = 1$. Still, as long as U has sufficiently thin tails, $E(Y^s)$ can exist for arbitrarily high integers s , thereby providing far more identifying equations than unknowns.

The above analysis is only suggestive. We do not have a proof of global identification with more than two points of support, though local identification up to a finite set should hold, given that the moments are polynomials, which must have a finite number of roots.

Assuming that a given model where V takes on more than two values is identified, moment conditions for estimation analogous to those we provided earlier are available. For example, as in the proof of [Proposition 1](#), it follows from symmetry of U that

$$E[\exp(\tau(Y - h))] = E[\exp(\tau V)]\alpha_\tau$$

with $\alpha_\tau = \alpha_{-\tau}$ for any τ for which these expectations exist, and therefore by choosing constants $\tau_1, \dots, \tau_L$, GMM estimation could be based on the $2L$ moments

$$E \left[\sum_{k=0}^H [\exp(\tau_\ell(Y - h))] - \exp(\tau_\ell b_k) \alpha_{\tau_\ell} p_k \right] = 0$$

$$E \left[\sum_{k=0}^H [\exp(-\tau_\ell(Y - h))] - \exp(-\tau_\ell b_k) \alpha_{\tau_\ell} p_k \right] = 0$$

for $\ell = 1, \dots, L$. The number of parameters b_k, p_k and α_{τ_ℓ} to be estimated would be $2H + L$, so taking $L > 2H$ provides more moments than unknowns.

10. Conclusions

We have proved global point identification and provided estimators for the models $Y = h + V + U$ or $Y = h(X) + V + U$, and more generally for $Y = g(X, D^*) + U$. In these models, D^* or V are unobserved regressors with two points of support, and the unobserved U is drawn from an unknown distribution having some odd central moments equal to zero, as would be the case if U is symmetrically distributed. No instruments, measures, or proxies for D^* or V are observed. A small Monte Carlo analysis shows that our estimator works reasonably well with a moderate sample size, despite involving high order data moments.

To further illustrate the estimator, in an online supplemental [Appendix B](#) to this paper we provide a small empirical application involving distribution of income across countries.

Interesting work for the future could include derivation of semiparametric efficiency bounds for the model, and obtaining conditions for global identification when V can take on more than two values.

Appendix A. Proofs

Proof of Theorem 1. To save space, a great deal of tedious but straightforward algebra is omitted. These details are available in an online supplemental [Appendix B](#).

First identify h by $h = E(Y)$, since V and U are mean zero. Then the distribution of ε defined by $\varepsilon = Y - h$ is identified, and $\varepsilon = U + V$. Define $e_d = E(\varepsilon^d)$, $u_d = E(U^d)$, and $v_d = E(V^d)$. Now evaluate e_d for integers $d \leq 9$. These e_d exist by assumption, and are identified because the distribution of ε is identified. Using independence of V and U , $v_1 = 0$, and $u_d = 0$ for odd values of d up to nine, evaluate $e_d = E((U + V)^d)$ to obtain $e_2 = v_2 + u_2$, $e_3 = v_3$, $e_4 = v_4 + 6v_2u_2 + u_4$, so $u_4 = e_4 - v_4 - 6v_2e_2 + 6v_2^2$, and $e_5 = v_5 + 10v_3u_2 = v_5 + 10v_3(e_2 - v_2)$. Define $s = e_5 - 10e_3e_2$, and note that s depends only on identified objects and so is identified. Then $s = v_5 - 10e_3v_2$.

Similarly, $e_6 = v_6 + 15v_4u_2 + 15v_2u_4 + u_6$, which can be solved for u_6 , $e_7 = v_7 + 21v_5u_2 + 35v_3u_4$, and $e_9 = v_9 + 36v_7u_2 +$

$126v_5u_4 + 84v_3u_6$. Substituting out the earlier expressions for u_2, u_4 , and u_6 in the e_7 and e_9 equations gives results that can be written as $q = v_7 - 35e_3v_4 - 21sv_2$ and $w = v_9 - 36qv_2 - 126sv_4 - 84e_3v_6$, where q and w are identified by $q = e_7 - 21se_2 - 35e_3e_4 = e_7 - 21e_5e_2 + e_3(210e_2^2 - 35e_4)$, and $w = e_9 - 36qe_2 - 126se_4 - 84e_3e_6 = e_9 - 36e_7e_2 + e_5(756e_2^2 - 126e_4) + e_3(2520e_2e_4 - 84e_6 - 7560e_2^3)$.

Summarizing, we have w, s, q, e_3 are all identified and $e_3 = v_3$, $s = v_5 - 10e_3v_2$, $q = v_7 - 35e_3v_4 - 21sv_2$, and $w = v_9 - 84e_3v_6 - 126sv_4 - 36qv_2$.

Now V only takes on two values, so let V equal b_0 with probability p_0 and b_1 with probability p_1 . Let $r = p_0/p_1$. Using $p_1 = 1 - p_0$ and $E(V) = b_0p_0 + b_1p_1 = 0$ we have

$$p_0 = r/(1+r), \quad p_1 = 1/(1+r), \quad b_1 = -b_0r,$$

and for any integer d

$$\begin{aligned} v_d &= b_0^d p_0 + b_1^d p_1 = b_0^d (p_0 + (-r)^d p_1) \\ &= b_0^d [r + (-r)^d] / (1+r). \end{aligned}$$

Substituting this v_d into the expression for e_3, s, q , and w reduces to

$$\begin{aligned} e_3 &= b_0^3 r(1-r), \quad s = b_0^5 r(1-r)(r^2 - 10r + 1) \\ q &= b_0^7 r(1-r)(r^4 - 56r^3 + 246r^2 - 56r + 1) \\ w &= b_0^9 r(1-r)(r^6 - 246r^5 + 3487r^4 - 10452r^3 \\ &\quad + 3487r^2 - 246r + 1). \end{aligned}$$

These are four equations in the two unknowns b_0 and r . We require all four equations for point identification, because these are polynomials in r and so have multiple roots. However, from just the e_3 and s equations and $b_0 \neq 0$ we have the identified polynomial in r

$$e_3^5 (r^2 - 10r + 1)^3 - s^3 r^2 (1-r)^2 = 0$$

which has at most six roots. Associated with each possible root r is a corresponding identified distribution for V and U as described at the end of this proof. This shows set identification of the model up to a finite set, and hence local identification, using just the first five moments of Y .

To show global identification, we will first show that the four equations for e_3, s, q , and w imply that $r^2 - \gamma r + 1 = 0$, where γ is finite and identified.

First we have $e_3 = v_3 \neq 0$ and $r \neq 1$ by asymmetry of V . Also $r \neq 0$ and $b_0 \neq 0$ because then V would only have one point of support instead of two. Applying these results to the s equation shows that if s (which is identified) is zero then $r^2 - 10r + 1 = 0$, and so in that case γ is identified. So now consider the case where $s \neq 0$.

Define $R = qe_3/s^2$, which is identified because its components are identified. Then

$$\begin{aligned} R &= (r^4 - 56r^3 + 246r^2 - 56r + 1)(r^2 - 10r + 1)^{-2} \quad \text{so} \\ 0 &= (1-R)r^4 + (-56+20R)r^3 + (246-102R)r^2 \\ &\quad + (-56+20R)r + (1-R). \end{aligned}$$

If $R = 1$, then (using $r \neq 0$) this polynomial reduces to the quadratic $0 = r^2 - 4r + 1$, so in this case $\gamma = -4$ is identified. Now consider the case where $R \neq 1$.

Define $Q = s^3/e_3^5$ and $S = w/e_3^9$. Both Q and S exist because $e_3 \neq 0$, and they are identified because their components are identified. Then

$$\begin{aligned} Q &= (r^2 - 10r + 1)^3 (r(1-r))^{-2} \quad \text{so} \\ 0 &= r^6 - 30r^5 + (303-Q)r^4 + (2Q-1060)r^3 \\ &\quad + (303-Q)r^2 - 30r + 1. \end{aligned}$$

Also

$$\begin{aligned} \frac{w}{e_3^9} &= S \\ &= \frac{b_0^9 r(1-r)(r^6 - 246r^5 + 3487r^4 - 10452r^3 + 3487r^2 - 246r + 1)}{(b_0^3 r(1-r))^3} \end{aligned}$$

$$0 = r^6 - 246r^5 + (3487-S)r^4 + (2S-10452)r^3 + (3487-S)r^2 - 246r + 1.$$

Subtracting the polynomial with Q from the polynomial with S gives

$$0 = 216r^4 + (S-Q-3184)r^3 + (9392+2Q-2S)r^2 + (S-Q-3184)r + 216.$$

Multiply this by $(1-R)$, multiply the polynomial based on R by 216, subtract one from the other and divide by r (which is nonzero) to obtain an expression that simplifies to

$$0 = Nr^2 - (2(1-R)(6320+S-Q) + 31104)r + N,$$

where $N = (1-R)(1136+S-Q) + 7776$, which after substituting in for R, S , and Q becomes

$$N = \frac{15552r(r+1)^4}{(r^2-10r+1)^2(1-r)^2}.$$

The denominator of this expression for N is not equal to zero, because that would imply $s = 0$, and we are currently specifically considering the case where $s \neq 0$ (having already analyzed the case where $s = 0$). Also $N \neq 0$ because $r \neq 0$, and $r \neq -1$. We therefore have $0 = r^2 - \gamma r + 1$, where $\gamma = (2(1-R)(6320+S-Q) + 31104)/N$, which is identified because all of its components are identified.

We have now shown that $0 = r^2 - \gamma r + 1$, where γ is identified. This equation says that $\gamma = r + r^{-1} = [p_0/(1-p_0)] + [(1-p_0)/p_0]$. Whatever value p_0 takes on between zero and one makes this expression for γ greater than or equal to two. The equation $0 = r^2 - \gamma r + 1$ has solutions

$$r = \frac{1}{2}\gamma + \frac{1}{2}\sqrt{\gamma^2 - 4} \quad \text{and} \quad r = \frac{1}{\frac{1}{2}\gamma + \frac{1}{2}\sqrt{\gamma^2 - 4}}$$

with $\gamma^2 \geq 4$, so one of these solutions must be the true value of r . Given r , we can then solve for b_0 by $b_0 = e_3^{1/3} (r(1-r))^{1/3}$. Recall that $r = p_0/p_1$. If we exchanged b_0 with b_1 and exchanged p_0 with p_1 everywhere, all of the above equations would still hold. It follows that one of the above two values of r must equal p_0/p_1 , and the other equals p_1/p_0 . The former when substituted into $e_3(r(1-r))$ will yield b_0^3 and the latter must yield b_1^3 . Without loss of generality imposing the constraint $b_0 < 0 < b_1$ shows that the correct solution for r will be the one that satisfies $e_3(r(1-r)) < 0$, and so r and b_0 is identified. The remainder of the distribution of V is then given by $p_0 = r/(1+r)$, $p_1 = 1/(1+r)$, and $b_1 = -b_0r$.

Finally, we show identification of the distribution of U . For any random variable Z , let F_Z denote the marginal cumulative distribution function of Z . By the probability mass function of the V distribution, $F_\varepsilon(\varepsilon) = (1-p)F_U(\varepsilon - b_1) + pF_U(\varepsilon - b_0)$. Letting $\varepsilon = u + (b_0 - b_1)k - b_0$ and rearranging gives

$$\begin{aligned} F_U(u + (b_0 - b_1)k) &= \frac{1}{p} F_\varepsilon(u + b_0 + (b_0 - b_1)k) \\ &\quad - \frac{1-p}{p} F_U(u + (b_0 - b_1)(k+1)), \end{aligned}$$

so for positive integers K , $F_U(u) = R_K + \sum_{k=0}^{K-1} \left(\frac{1-p}{p}\right)^k \frac{1}{p} F_\varepsilon(u + b_0 + (b_0 - b_1)k)$, where the remainder term $R_K = r^{-K} F_U(u +$

$(b_0 - b_1)K) \leq r^{-K}$. If $r > 1$ then $R_K \rightarrow 0$ as $K \rightarrow \infty$, so $F_U(u)$ is identified by

$$F_U(u) = \sum_{k=0}^{\infty} \left(\frac{1-p}{p} \right)^k \frac{1}{p} F_{\varepsilon}(u + b_0 + (b_0 - b_1)k) \quad (23)$$

since all the terms on the right of this expression are identified, given that the distributions of ε and of V are identified. If $r < 1$, then exchange the roles of b_0 and b_1 (e.g., start by letting $\varepsilon = u + (b_1 - b_0)k - b_1$) which will correspondingly exchange p and $1-p$ to obtain $F_U(u) = \sum_{k=0}^{\infty} \left(\frac{p}{1-p} \right)^k \frac{1}{1-p} F_{\varepsilon}(u + b_1 + (b_1 - b_0)k)$, where now the remainder term is $R_K = r^K F_U(u + (b_1 - b_0)K) \leq r^K \rightarrow 0$ as $K \rightarrow \infty$ since now $r < 1$. The case of $r = 0$ is ruled out, since that is equivalent to $p = 1/2$. $\square$

Proof of Proposition 1. $Y = h + V + U$ and independence of U and V implies that

$$E[\exp(\tau(Y - h))] = E[\exp(\tau V)] E[\exp(\tau U)].$$

Now $E[\exp(\tau V)] = p \exp(\tau b_0) + (1-p) \exp(\tau b_1)$. Define $\alpha_{\tau} = (1-p)E(e^{\tau U})$. By symmetry of U , $\alpha_{\tau} = \alpha_{-\tau}$. These equations with $r = p/(1-p)$ and $b_1 = b_0 p/(p-1)$ give Eqs. (20) and (21). $\square$

Proof of Theorem 2. By the probability mass function of the V distribution, $F_{\varepsilon}(\varepsilon) = (1-p)F_U(\varepsilon - b_1) + pF_U(\varepsilon - b_0)$. Evaluating this expression at $\varepsilon = u + b_1$ gives

$$F_{\varepsilon}(u + b_1) = (1-p)F_U(u) + pF_U(u + b_1 - b_0) \quad (24)$$

and evaluating at $\varepsilon = -u + b_0$ gives $F_{\varepsilon}(-u + b_0) = (1-p)F_U(-u - b_1 + b_0) + pF_U(-u)$. Apply symmetry of U which implies $F_U(u) = 1 - F_U(-u)$ to this last equation to obtain

$$F_{\varepsilon}(-u + b_0) = (1-p)[1 - F_U(u + b_1 - b_0)] + p[1 - F_U(u)]. \quad (25)$$

Eqs. (24) and (25) are two equations in the two unknowns $F_U(u + b_1 - b_0)$ and $F_U(u)$. Solving for $F_U(u)$ gives $F_U(u) = \Psi(u)$ with $\Psi(u)$ given by Eq. (11). It follows from symmetry of U that $F_U(u)$ must also equal $1 - \Psi(-u)$, which gives Eq. (12). $\square$

Proof of Corollary 1. First identify $h(x)$ by $h(x) = E(Y | X = x)$, since $E(Y - h(X) | X = x) = E(V + U | X = x) = E(V + U) = 0$. Next define $\varepsilon = Y - h(X)$, and then the rest of the proof is identical to the proof of Theorem 1. $\square$

Proof of Corollary 2. Define $h(x) = E(Y | X)$ and $\varepsilon = Y - h(X)$. Then $h(x)$ and the distribution of ε conditional upon X is identified and $E(\varepsilon | X) = 0$. Define $V = g(X, D^*) - h(X)$ and let $b_d(X) = g(X, d) - h(X)$ for $d = 0, 1$. Then $\varepsilon = V + U$, where V (given X) has the distribution with support equal to the two values $b_0(X)$ and $b_1(X)$ with probabilities $p(X)$ and $1 - p(X)$, respectively. Also U and ε have mean zero given X so $E(V | X) = 0$. Applying Theorem 1 separately for each value x on the support of X shows that $b_0(x)$, $b_1(x)$, $p(x)$, and the conditional distribution of U given $X = x$ is identified for each such x , and it follows that the function $g(x, d)$ is identified by $g(x, d) = b_d(x) + h(x)$. $\square$

Appendix B. Supplementary data

Supplementary material related to this article can be found online at doi:10.1016/j.jeconom.2011.03.003.

References

- Ahmed, I.A., Li, Q., 1997. Testing symmetry of an unknown density by kernel method. *Nonparametric Statistics* 7, 279–293.
- Ai, C., Chen, X., 2003. Efficient estimation of models with conditional moment restrictions containing unknown functions. *Econometrica* 71, 1795–1844.
- Altonji, J., Segal, L., 1996. Small-sample bias in GMM estimation of covariance structures. *Journal of Business & Economic Statistics* 14, 353–366.
- Baltagi, B.H., 2008. *Econometric Analysis of Panel Data*, 4th ed. Wiley.
- Bordes, L., Mottelet, S., Vandekerckhove, P., 2006. Semiparametric estimation of a two-component mixture model. *Annals of Statistics* 34, 1204–1232.
- Carrasco, M., Florens, J.P., 2000. Generalization of GMM to a continuum of moment conditions. *Econometric Theory* 16, 797–834.
- Carroll, R.J., Ruppert, D., Stefanski, L.A., Crainiceanu, C.M., 2006. *Measurement Error in Nonlinear Models: A Modern Perspective*, 2nd ed. Chapman & Hall/CRC.
- Chen, X., Hu, Y., Lewbel, A., 2008. Nonparametric identification of regression models containing a misclassified dichotomous regressor without instruments. *Economics Letters* 100, 381–384.
- Clogg, C.C., 1995. Latent class models. In: Arminger, G., Clogg, C.C., Sobel, M.E. (Eds.), *Handbook of Statistical Modeling for the Social and Behavioral Sciences*. Plenum, New York, pp. 311–359 (Chapter 6).
- Dong, Y., 2011. Semiparametric binary random effects models: estimating two types of drinking behavior. *Economics Letters* (forthcoming).
- Hagenaars, J.A., McCutcheon, A.L., 2002. *Applied Latent Class Analysis Models*. Cambridge University Press, Cambridge.
- Hall, P., Zhou, X.-H., 2003. Nonparametric estimation of component distributions in a multivariate mixture. *Annals of Statistics* 31, 201–224.
- Hansen, L., 1982. Large sample properties of generalized method of moments estimators. *Econometrica* 50, 1029–1054.
- Heckman, J.J., Robb, R., 1985. Alternative methods for evaluating the impact of interventions. In: Heckman, James J., Singer, B. (Eds.), *Longitudinal Analysis of Labor Market Data*. Cambridge University Press, New York, pp. 156–245.
- Honore, B., 1992. Trimmed lad and least squares estimation of truncated and censored regression models with fixed effects. *Econometrica* 60, 533–565.
- Hu, Y., Lewbel, A., 2008. Identifying the returns to lying when the truth is unobserved. Boston College Working Paper.
- Kasahara, H., Shimotsu, K., 2009. Nonparametric identification of finite mixture models of dynamic discrete choices. *Econometrica* 77, 135–175.
- Kitamura, Y., 2004. Nonparametric identifiability of finite mixtures. Unpublished Manuscript, Yale University.
- Kumbhakar, S.C., Lovell, C.A.K., 2000. *Stochastic Frontier Analysis*. Cambridge University Press.
- Kumbhakar, S.C., Park, B.U., Simar, L., Tsionas, E.G., 2007. Nonparametric stochastic frontiers: a local maximum likelihood approach. *Journal of Econometrics* 137, 1–27.
- Lewbel, A., 2007. A local generalized method of moments estimator. *Economics Letters* 94, 124–128.
- Newey, W.K., McFadden, D., 1994. Large sample estimation and hypothesis testing. In: Engle, R.F., McFadden, D.L. (Eds.), *Handbook of Econometrics*, vol. IV. Elsevier, Amsterdam, pp. 2111–2245.
- Newey, W.K., Powell, J.L., 2003. Instrumental variable estimation of nonparametric models. *Econometrica* 71, 1565–1578.
- Powell, J.L., 1986. Symmetrically trimmed least squares estimation of Tobit models. *Econometrica* 54, 1435–1460.
- Silverman, B.W., 1986. *Density Estimation for Statistics and Data Analysis*. Chapman and Hall, London.
- Simar, L., Wilson, P.W., 2007. Statistical inference in nonparametric frontier models: recent developments and perspectives. In: Fried, H., Lovell, C.A.K., Schmidt, S.S. (Eds.), *The Measurement of Productive Efficiency*, 2nd ed. Oxford University Press, Oxford (Chapter 4).