

Identification of unobservables in observations

Yingyao Hu

The Institute for Fiscal Studies
Department of Economics, UCL

cemmap working paper CWP25/22



Economic
and Social
Research Council

Identification of Unobservables in Observations

Yingyao Hu*
Johns Hopkins University

December 5, 2022

Abstract

In empirical studies, the data usually don't include all the variables of interest in an economic model. This paper shows the identification of unobserved variables in observations at the population level. When the observables are distinct in each observation, there exists a function mapping from the observables to the unobservables. Such a function guarantees the uniqueness of the latent value in each observation. The key lies in the identification of the joint distribution of observables and unobservables from the distribution of observables. The joint distribution of observables and unobservables then reveal the latent value in each observation. Three examples of this result are discussed.

JEL classification: *C01, C14*

Keywords: *identification in observations, unobserved heterogeneity, latent variable model, measurement error model*

*Contact information: Department of Economics, Johns Hopkins University, 3400 N. Charles Street, Baltimore, MD 21218. Email: yhu@jhu.edu.

1 Introduction

“No two leaves are alike.”

Thousand years of human history show that no two leaves are alike. Suppose that each leaf has observed traits x and unobserved heterogeneity x^* . If we observe that no two leaves are alike, then the mapping from x to x^* is a function, i.e., a set of ordered pairs (x, x^*) in which no two different ordered pairs have the same first coordinate. Therefore, such a function can uniquely determine unobserved heterogeneity x^* from observed traits x for each leaf.

In empirical studies, an observation in the data is similar to observed traits x of a leaf and unobserved heterogeneity x^* corresponds to some variables of interest not observed in the data. For example, researchers observe a patient’s insurance policy, but not their health status in the data. In general, we observe an agent’s choices, but not their types or unobserved heterogeneity. In macroeconomics, we are interested in a country’s true GDP when only rough measurements are available. This paper intends to provide a framework to identify the value of a latent variable of interest in observations.

For a variable with a distinct value in each observation in a sample, researchers usually consider it as a continuous variable in the population. Such continuity only exists in assumptions given the discrete nature of a sample. It is observationally equivalent to assume that the population is a collection of a large but finite number of elements. To avoid an uncountable amount of unknowns, we adopt the latter in this paper.

Let x_i and x_i^* be measurements of observed traits and unobserved heterogeneity of leaf i , respectively. We define the property of leaves as follows:

Definition 1 *A population $\mathcal{P}_{X,X^*}$ satisfies **the property of leaves** if it is a collection of ordered pairs (x_i, x_i^*) for $i = 1, 2, \dots, N$; $N < \infty$ such that $x_i \neq x_j$ for any $i \neq j$. That is*

$$\mathcal{P}_{X,X^*} = \{(x_i, x_i^*) : x_i \neq x_j \text{ for } i \neq j \text{ and } i, j = 1, 2, \dots, N.\} \quad (1)$$

Furthermore, let F_{X,X^*} denote the cumulative distribution function of random variables (X, X^*) randomly drawn from population $\mathcal{P}_{X,X^*}$ with probability

$$Pr(\{(X, X^*) = (x_i, x_i^*)\}) = p_i > 0 \quad (2)$$

with $\sum_{i=1}^N p_i = 1$. A typical example is $p_i = \frac{1}{N}$. Here we focus on the case where the population size N is large but finite. In this case, distribution function F_{X,X^*} uniquely determines population $\mathcal{P}_{X,X^*}$ because F_{X,X^*} is a step function and each step corresponds to an element in $\mathcal{P}_{X,X^*}$. If we consider the set $\mathcal{P}_{X,X^*}$ as a mapping from X to X^* , then this mapping is a function, i.e., a set of ordered pairs in which no two different ordered pairs

have the same first coordinate. The population of observed traits x is

$$\mathcal{P}_X = \{x_i : (x_i, x_i^*) \in \mathcal{P}_{X,X^*} \text{ for some } x_i^*\} \quad (3)$$

with a distribution function F_X . In fact, its probability function is

$$Pr(\{X = x_i\}) = p_i \quad (4)$$

because x_i is distinct for all $i = 1, 2, \dots, N$, i.e., in the whole population. Therefore, the probabilities p_i is known for each i from F_X , but the value of x_i^* in Equation (2) still needs to be identified and distribution function F_{X,X^*} is still unknown. Notice that x_i^* is not necessarily distinct in observations.

Here it is necessary to clarify that this identification analysis is at the population level, instead of at the sample level. In estimation, we usually start with a sample of X and use its sample statistics to estimate its population counterparts.¹ For example, we use empirical CDF $\hat{F}_X$ to consistently estimate the population CDF F_X . In the identification analysis, we start with population $\mathcal{P}_X$ and its CDF F_X when X is observed in a sample because $\mathcal{P}_X$ and F_X are identified as the limit of the sample and the empirical CDF, respectively. This paper presents sufficient conditions, under which one can uniquely determine the unobserved x^* from the observed x in each observation in the population. We summarize the immediate conditions as follows:

Proposition 1 *Suppose that Conditions 1 and 2 hold as follows:*

1. *Population $\mathcal{P}_{X,X^*}$, with distribution function F_{X,X^*} , satisfies the property of leaves in Equations (1) ;*
2. *F_X uniquely determines F_{X,X^*} , where distribution function F_X corresponds to population $\mathcal{P}_X$ in Equations (3).*

Then, $\mathcal{P}_X$ and F_X uniquely determine $\mathcal{P}_{X,X^}$ and F_{X,X^*} , i.e., each x_i in $\mathcal{P}_X$ uniquely determines its corresponding x_i^* through $\mathcal{P}_{X,X^*}$.*

Proof: Distribution function F_X corresponds to population $\mathcal{P}_X$. Condition 2) requires that F_X uniquely determines F_{X,X^*} . Given that population $\mathcal{P}_{X,X^*}$ contains a large but finite number N of different elements, distribution function F_{X,X^*} is a step function and each step corresponds to one element in population $\mathcal{P}_{X,X^*}$, and therefore, F_{X,X^*} uniquely determines population $\mathcal{P}_{X,X^*}$. In summary, $\mathcal{P}_X$ with F_X uniquely determines $\mathcal{P}_{X,X^*}$ with F_{X,X^*} . Q.E.D.

¹If we draw from such a population of X with placement, it is possible to have two draws with the same value x . The property of leaves guarantees that the value x corresponds to a unique x^* and that the two draws are from the same leaf, which can be represented in the sample proportions of each distinct value of x . Therefore, such a generated sample of X is representative of the population of leaves and its distribution.

Although it is for a large but finite N , this result can be extended to the case with $N \rightarrow \infty$ as long as population $\mathcal{P}_{X,X^*}$ contains a countably number of elements. Note that the probability in Equation (2) can't be uniform, i.e., $p_i = p$, in this case. The discreteness of $\mathcal{P}_{X,X^*}$ implies that population distribution F_{X,X^*} is a step function and each step still corresponds one element in population $\mathcal{P}_{X,X^*}$. Such a property is lost when there are uncountably many elements in the population, e.g., the unit interval.²

Conditions 1) in Proposition 1 requires that the observed X should satisfy the property of leaves so that there exists a function mapping from the observed to the unobserved. Condition 2 is the key to achieve the identification of unobservables in observations. In order to make Proposition 1 useful, it is important to provide sufficient conditions to identify F_{X,X^*} from F_X .

Proposition 1 can be adapted to the case where additional variables are observed. If the data include (X, Z) instead of X only, and if $F_{X,Z}$ uniquely determines F_{X,X^*} , then there is no need to identify F_{X,Z,X^*} , which may require more assumptions than those for the identification of F_{X,X^*} . In that case, the result in Proposition 1 remains with $F_{X,Z}$ uniquely determining F_{X,X^*} in Condition 2).

It is useful to understand the result in Proposition 1 in the case of the widely-used linear regression model, i.e.,

$$Y = W\beta + \eta$$

with $E[\eta|W] = 0$. A researcher observes $X = (Y, W)$ but not $X^* = \eta$. Suppose that we never observe repeated values of (Y, W) in the population. Therefore, Condition 1) is satisfied. In fact, the function implied by Condition 1), which maps from $X = (Y, W)$ to $X^* = \eta$, is given by the model, i.e., $\eta = Y - W\beta$. Given that we can identify and estimate β using distribution $F_{Y,W}$ through the moment equation $E[Y - W\beta|W] = 0$, parameter β can be considered as known from the population. Then, it can be shown that F_X , i.e., the distribution of (Y, W) , uniquely determines F_{X,X^*} , i.e., the distribution of (Y, W, η) , because $\eta = Y - W\beta$, and therefore, $\mathcal{P}_X$, the population of (Y, W) , uniquely determines $\mathcal{P}_{X,X^*}$, the population of (Y, W, η) . That means, the regression error η_i , although unobserved, is uniquely determined by (y_i, w_i) in each observation through $\mathcal{P}_{X,X^*}$ as $\eta_i = y_i - w_i\beta$. The estimation of residuals in the linear regression model is the sample counterpart of this procedure.

Condition 1) in Proposition 1 holds as long as there is a traditionally-defined continuous variable in the sample. The challenging part of Proposition 1 is to show that the joint distribution of observables and unobservables is uniquely determined by that of the observables. The next two sections present examples of sufficient conditions for the identification of F_{X,X^*} from F_X . We adopt the framework in Hu (2017) to consider cases with a difference number of measurements of X^* in X .

²A recent working paper Hu et al. (2022) provides some identification arguments in that case.

2 A 2-measurement case

In this section, we consider a 2-measurement setting, where $X = (X_1, X_2)$. Assume that X_1 , X_2 , and X^* are a scalar random variable satisfying

$$\begin{aligned} X_1 &= X^* + \epsilon_1 \\ X_2 &= X^* + \epsilon_2 \end{aligned} \tag{5}$$

where i) ϵ_1 is independent of (X^*, ϵ_2) , ii) the characteristic function of X_1 is absolutely integrable and does not vanish on the real line, and iii) $E[\epsilon_2|X^*] = 0$.

Before presenting the technical results, it is useful to illustrate the idea of identification in observations with a simple example. Suppose $X^* \in \{0, 1\}$, $\epsilon_1 \in \{-1, 2\}$ and $\epsilon_2 \in \{-1, 0, 1\}$ satisfying Equation (5) with distribution functions f_{X^*, ϵ_2} and f_{ϵ_1} . Notice that ϵ_2 should have a zero mean conditional on X^* . The population and its distribution can be presented as in Table 1. Given the population of 12 observations of (X_1, X_2) with distribution function f_{X_1, X_2} , the goal is to show that the value of X^* is uniquely determined in each observation.

Table 1: An illustration of identification in observations

observation i	observables		unobservables			probability
	$X_1 = X^* + \epsilon_1$	$X_2 = X^* + \epsilon_2$	ϵ_1	X^*	ϵ_2	p_i
1	0	0	-1	1	-1	$f_{X_1, X_2}(0, 0) = f_{\epsilon_1}(-1)f_{X^*, \epsilon_2}(1, -1)$
2	0	1	-1	1	0	$f_{X_1, X_2}(0, 1) = f_{\epsilon_1}(-1)f_{X^*, \epsilon_2}(1, 0)$
3	0	2	-1	1	1	...
4	-1	-1	-1	0	-1	...
5	-1	0	-1	0	0	...
6	-1	1	-1	0	1	...
7	3	0	2	1	-1	...
8	3	1	2	1	0	...
9	3	2	2	1	1	...
10	2	-1	2	0	-1	...
11	2	0	2	0	0	...
12	2	1	2	0	1	...

Note: In each group, mean of X_2 reveals X^* .

In this example, the observed (X_1, X_2) are distinct. If we group the 12 observations by X_1 , then the mean of X_2 within each group is equal to the value of latent X^* . In general, the property of leaves guarantees the uniqueness of X^* in each observation. The restrictions on the distribution, i.e., ϵ_2 should have a zero mean conditional on X^* , reveal the value of X^* in each observation. Therefore, the unobserved is uniquely determined by the observed in observations. Notice that the four groups are actually categorized by the values of X^*

and ϵ_1 .

Although we use X_2 to put the 12 observations into 4 groups in this particular example in Table 1, it really is the combination of the 12 observations of (X_1, X_2) and the distribution function of f_{X_1, X_2, X^*} or $f_{\epsilon_1} f_{X^*, \epsilon_2}$ that identifies the 4 groups out of the 12 observations. It is possible that X_1 is not enough to distinguish all the groups with the same values of X^* and ϵ_1 . For example, we may change the support of ϵ_1 and ϵ_2 to be $\epsilon_1 \in \{-1, 0\}$ and $\epsilon_2 \in \{-1.5, 0.5, 1\}$, still assuming ϵ_2 should have a zero mean. Table 2 shows the population and the probabilities in this case, where X_1 is no longer enough to identify the four groups. From the observed 12 probabilities, i.e., p_i , in f_{X_1, X_2} , however, we are able to identify the 8 unknown probabilities in f_{ϵ_1} and f_{X^*, ϵ_2} , which will be shown below in a more general setup. The identified distribution function, i.e., the probabilities in f_{ϵ_1} and f_{X^*, ϵ_2} can determine how to put the observations into four groups with the same X^* and ϵ_1 as in the last column in Table 2. Then, mean of X_2 reveals X^* in each group.

Table 2: A second example

observation i	observables		unobservables			probability p_i
	$X_1 = X^* + \epsilon_1$	$X_2 = X^* + \epsilon_2$	ϵ_1	X^*	ϵ_2	
1	0	-0.5	-1	1	-1.5	$f_{X_1, X_2}(0, -0.5) = f_{\epsilon_1}(-1)f_{X^*, \epsilon_2}(1, -1.5)$
2	0	1.5	-1	1	0.5	$f_{X_1, X_2}(0, 1.5) = f_{\epsilon_1}(-1)f_{X^*, \epsilon_2}(1, 0.5)$
3	0	2	-1	1	1	$f_{X_1, X_2}(0, 2) = f_{\epsilon_1}(-1)f_{X^*, \epsilon_2}(1, 1)$
4	-1	-1.5	-1	0	-1.5	...
5	-1	0.5	-1	0	0.5	...
6	-1	1	-1	0	1	...
7	1	-0.5	0	1	-1.5	...
8	1	1.5	0	1	0.5	...
9	1	2	0	1	1	...
10	0	-1.5	0	0	-1.5	$f_{X_1, X_2}(0, -1.5) = f_{\epsilon_1}(0)f_{X^*, \epsilon_2}(0, -1.5)$
11	0	0.5	0	0	0.5	$f_{X_1, X_2}(0, 0.5) = f_{\epsilon_1}(0)f_{X^*, \epsilon_2}(0, 0.5)$
12	0	1	0	0	1	$f_{X_1, X_2}(0, 1) = f_{\epsilon_1}(0)f_{X^*, \epsilon_2}(0, 1)$

Note: In each group, mean of X_2 reveals X^* .

The setup in Equation (5) is well known because the distribution of the latent variable X^* can be written as a closed-form function of the observed distribution f_{X_1, X_2} . The characteristic function of X^* is defined as $\phi_{X^*}(t) = E[e^{itX^*}]$ with $i = \sqrt{-1}$. One can show

that

$$\begin{aligned} f_{X^*}(x^*) &= \frac{1}{2\pi} \int_{-\infty}^{\infty} e^{-ix^*t} \phi_{X^*}(t) dt \\ \phi_{X^*}(t) &= \exp \left[\int_0^t \frac{iE[X_2 e^{isX_1}]}{E[e^{isX_1}]} ds \right]. \end{aligned} \quad (6)$$

This is the so-called Kotlarski's identity (Kotlarski (1966) and Rao (1992)). Notice that

$$\begin{aligned} \phi_{X_1, X_2}(s, t) &= \exp[isX_1 + itX_2] \\ &= \exp[isX^* + itX_2] \exp[is\epsilon_1] \\ &= \phi_{X^*, X_2}(s, t) \frac{\phi_{X_1}(s)}{\phi_{X^*}(s)} \end{aligned}$$

Therefore, we may identify the joint distribution F_{X^*, X_2} as follows:

$$\phi_{X^*, X_2}(s, t) = \phi_{X_1, X_2}(s, t) \frac{\phi_{X^*}(s)}{\phi_{X_1}(s)}$$

Finally, the distribution of (X_1, X_2, X^*) can be uniquely determined by the distribution of (X_1, X_2) as follows:

$$\begin{aligned} \phi_{X_1, X_2, X^*}(s, t, v) &= \exp[isX_1 + itX_2 + ivX^*] \\ &= \exp[i(s+v)X^* + itX_2] \exp[is\epsilon_1] \\ &= \phi_{X^*, X_2}(s+v, t) \frac{\phi_{X_1}(s)}{\phi_{X^*}(s)} \end{aligned} \quad (7)$$

That means F_{X_1, X_2} uniquely determines F_{X_1, X_2, X^*} . Under the assumption that observations of (X_1, X_2) are distinct, Condition 1) of Proposition 1 holds. Then, the identification of F_{X_1, X_2, X^*} implies that the value of X^* can be uniquely determined by that of (X_1, X_2) . We summarize the results as follows:

Lemma 1 *Suppose that $X = (X_1, X_2)$ satisfies Equation (5) and that observations of (X_1, X_2) are distinct in population $\mathcal{P}_{X_1, X_2, X^*}$. Then, $\mathcal{P}_{X_1, X_2}$ and F_{X_1, X_2} uniquely determine $\mathcal{P}_{X_1, X_2, X^*}$ and F_{X_1, X_2, X^*} , i.e., each $(x_{1,i}, x_{2,i})$ in $\mathcal{P}_{X_1, X_2}$ uniquely determines its corresponding x_i^* through $\mathcal{P}_{X_1, X_2, X^*}$.*

3 A 3-measurement case

It is possible to avoid the additivity and linearity in Equation (5), when there are more observables. In the case where $X = (X_1, X_2, X_3)$, Hu (2008) provides sufficient conditions

to identify the distribution of (X_1, X_2, X_3, X^*) from that of (X_1, X_2, X_3) . A version of the conditions is presented here:

Assumption 1 *The two measurements X_1 and X_2 and the latent variable X^* share the same support $\mathcal{X} = \{v_1, v_2, \dots, v_K\}$ with $K < N$.*

This condition is not restrictive because the results can be straightforwardly extended to the case where supports of measurements X_1 and X_2 are larger than that of X^* .

Assumption 2 *The observables satisfy conditional independence as follows:*

$$F[X_1, X_2, X_3 | X^*] = F[X_1 | X^*] F[X_2 | X^*] F[X_3 | X^*] \quad (8)$$

where $F[X | X^*]$ is the conditional CDF of X on X^* .

Let f_{X_1, X_2} be the probability function of (X_1, X_2) . Define a matrix representation of the joint distribution as follows:

$$M_{X_1, X_2} = [f_{X_1, X_2}(v_i, v_j)]_{i=1,2,\dots,K; j=1,2,\dots,K} \quad (9)$$

We assume

Assumption 3 *Matrix M_{X_1, X_2} has rank K .*

Assumption 4 *There exists a function $g(\cdot)$ such that $E[g(X_3) | X^* = \bar{v}] \neq E[g(X_3) | X^* = \tilde{v}]$ for any $\bar{v} \neq \tilde{v}$ in $\mathcal{X}$.*

Assumption 5 *$f_{X_1 | X^*}(v | v) > f_{X_1 | X^*}(\tilde{v} | v)$ for any $\tilde{v} \neq v \in \mathcal{X}$, i.e., v is the mode of distribution $f_{X_1 | X^*}(\cdot | v)$.*

Under assumptions 1, 2, 3, 4, and 5, Hu (2008) shows that the joint distribution of the observables (X_1, X_2, X_3) uniquely determines the joint distribution of the observed variables and the unobservable (X_1, X_2, X_3, X^*) . Furthermore, if (X_1, X_2, X_3) satisfies Condition 1) in Proposition 1, then we can apply Proposition 1 as follows:

Lemma 2 *Suppose that Assumptions 1, 2, 3, 4, and 5 hold, and that observations of $X = (X_1, X_2, X_3)$ are distinct in population $\mathcal{P}_{X_1, X_2, X_3, X^*}$. Then, $\mathcal{P}_{X_1, X_2, X_3}$ and F_{X_1, X_2, X_3} uniquely determine $\mathcal{P}_{X_1, X_2, X_3, X^*}$ and F_{X_1, X_2, X_3, X^*} , i.e., each $(x_{1,i}, x_{2,i}, x_{3,i})$ in $\mathcal{P}_{X_1, X_2, X_3}$ uniquely determines its corresponding x_i^* through $\mathcal{P}_{X_1, X_2, X_3, X^*}$.*

Because X_1 and X_2 have a small discrete support, i.e., $K < N$, we need X_3 to have a large support to make $X = (X_1, X_2, X_3)$ distinct in the population. For the identification of distributions, X_3 can be a measurement as little informative as a binary indicator. But for identification in observations in this paper, X_3 needs to have a large support so that $X = (X_1, X_2, X_3)$ is distinct in each observation.

Given the general result above, it is still useful to provide a simple example to illustrate the idea of identification in observations in this case. Suppose X_1, X_2, X^* share the same support $\{0, 1\}$ with non-degenerated misclassification probabilities $f_{X_1|X^*}(1|0) > 0$, $f_{X_1|X^*}(0|1) > 0$, $f_{X_2|X^*}(1|0) > 0$, $f_{X_2|X^*}(0|1) > 0$, and $X_3 \in \{1, 2, 3, 4\}$ with $f_{X_3|X^*}$ satisfying

$$f_{X_3|X^*}(1|1) = f_{X_3|X^*}(2|0) = f_{X_3|X^*}(3|0) = f_{X_3|X^*}(4|1) = 0.$$

This population is presented in Table 3 and satisfies the property of leaves. The goal is to show that the observations of (X_1, X_2, X_3) and the distribution of (X_1, X_2, X_3) can uniquely determine the value of X^* in each observation.

Assumptions in Lemma 2 hold for the example in Table 3. The conditional independence, together with other assumptions, identifies the probability functions, including $f_{X_1|X^*}$. Then, for each given value X_3 , X^* takes a unique value x^* , which is equal to the mode of $f_{X_1|X^*}(\cdot|x^*)$ under Assumption 5. Therefore, the unobserved x^* is uniquely determined by the observed variables in each observation.

Because every observation of the observables (X_1, X_2, X_3) has to be distinct in the population of (X_1, X_2, X_3, X^*) , each value of (X_1, X_2, X_3) can only map to one unique value of X^* . It is also useful to present a case where the property of leaves fails. Table 4 shows a violation of the property of leaves with $f_{X_3|X^*}$ satisfying

$$\begin{aligned} f_{X_3|X^*}(1|1) &= f_{X_3|X^*}(2|0) = f_{X_3|X^*}(3|0) = 0 \\ f_{X_3|X^*}(4|0) &> 0, f_{X_3|X^*}(4|1) > 0 \end{aligned}$$

in the example above because observations 13, 14, 15, 16 are the same as observations 17, 18, 19, 20, respectively, if we only observe (X_1, X_2, X_3) . In other words, $X_3 = 4$ corresponds to $X^* = 0$ and $X^* = 1$. In particular, leaves (or observations) 13 and 17 are different in population, but they are the same from a researcher's view because they don't observe X^* . That is the case we rule out here because it is not consistent with the common knowledge that no two leaves are alike.

4 Summary

This paper provides sufficient conditions for the identification of unobserved variables in observations at the population level. Based on an observed feature of the data – the property

Table 3: An illustration of identification in observations

observation i	observables $X_1 \quad X_2 \quad X_3$			unobservables X^*	probability p_i
1	0	0	1	0	$f_{X_1, X_2, X_3}(0, 0, 1) = f_{X_1 X^*}(0 0)f_{X_2 X^*}(0 0)f_{X_3 X^*}(1 0)f_{X^*}(0)$ $f_{X_1, X_2, X_3}(1, 0, 1) = f_{X_1 X^*}(1 0)f_{X_2 X^*}(0 0)f_{X_3 X^*}(1 0)f_{X^*}(0)$
2	1	0	1	0	
3	0	1	1	0	
4	1	1	1	0	
5	0	0	2	1	...
6	1	0	2	1	...
7	0	1	2	1	...
8	1	1	2	1	...
9	0	0	3	1	...
10	1	0	3	1	...
11	0	1	3	1	...
12	1	1	3	1	...
13	0	0	4	0	...
14	1	0	4	0	...
15	0	1	4	0	...
16	1	1	4	0	...

Note: For a given X_3 , X^* takes a unique value equal to the mode of $f_{X_1|X^*}(\cdot|x^*)$.

of leaves, the results in this paper imply that when the joint distribution of the observable X and the unobservable X^* satisfies certain conditions, it is not only possible to identify their joint distribution F_{X, X^*} from F_X , but also possible to identify the value of the unobservable X^* in observations. The distinctness of observed variables in observations implies there exists a function mapping from the observables to the unobservables. Such a function guarantees the uniqueness of the latent value in each observation. The joint distribution can then reveal the latent value in each observation.

In the identification analysis, we consider distribution function F_X to be identified as the limit of the empirical distribution of X from a sample. The results in this paper suggests that it is possible to use the sample counterpart of this argument to estimate unobservables at the observation level. A simple existing example is OLS residuals in a linear regression model. The identification results here imply that researchers may tackle unobservables by directly estimating them in a broad range of models.

References

Hu, Yingyao, “Identification and Estimation of Nonlinear Models with Misclassification Error Using Instrumental Variables: A General Solution,” *Journal of Econometrics*, 2008,

Table 4: A violation of the property of leaves in Equation 1

observation i	observables $X_1 \quad X_2 \quad X_3$			unobservables X^*	probability p_i
1	0	0	1	0	$f_{X_1, X_2, X_3}(0, 0, 1) = f_{X_1 X^*}(0 0)f_{X_2 X^*}(0 0)f_{X_3 X^*}(1 0)f_{X^*}(0)$ $f_{X_1, X_2, X_3}(1, 0, 1) = f_{X_1 X^*}(1 0)f_{X_2 X^*}(0 0)f_{X_3 X^*}(1 0)f_{X^*}(0)$
2	1	0	1	0	
3	0	1	1	0	
4	1	1	1	0	
5	0	0	2	1	...
6	1	0	2	1	...
7	0	1	2	1	...
8	1	1	2	1	...
9	0	0	3	1	...
10	1	0	3	1	...
11	0	1	3	1	...
12	1	1	3	1	...
13	0	0	4	0	$f_{X_1 X^*}(0 0)f_{X_2 X^*}(0 0)f_{X_3 X^*}(4 0)f_{X^*}(0)$
14	1	0	4	0	
15	0	1	4	0	
16	1	1	4	0	
17	0	0	4	1	$f_{X_1 X^*}(0 1)f_{X_2 X^*}(0 1)f_{X_3 X^*}(4 1)f_{X^*}(1)$
18	1	0	4	1	
19	0	1	4	1	
20	1	1	4	1	

Note: For $X_3 = 4$, X^* is not unique.

$$f_{X_1, X_2, X_3}(0, 0, 4) = \sum_{x^* \in \{0, 1\}} f_{X_1|X^*}(0|x^*)f_{X_2|X^*}(0|x^*)f_{X_3|X^*}(4|x^*)f_{X^*}(x^*)$$

144, 27–61.

– , “The econometrics of unobservables: Applications of measurement error models in empirical industrial organization and labor economics,” *Journal of econometrics*, 2017, 200(2), 154–168.

– , **Yang Liu**, and **Jiaxiong Yao**, “Revealing Unobservables by Deep Learning: Generative Element Extraction Networks (GEEN),” *working paper*, 2022.

Kotlarski, Ignacy, “On Some Characterizations of Probability Distributions in Hilbert Spaces,” *Annali di Matematica Pura et Applicata*, 1966, 74, 129–134.

Rao, B.L.S. Prakasa, *Identifiability in Stochastic Models: Characterization of Probability Distributions*, Academic Press, Inc., 1992.