

Zheng, Buhong

Working Paper

Statistical Inferences for Poverty Measures with Relative Poverty Rates

LIS Working Paper Series, No. 167

Provided in Cooperation with:

Luxembourg Income Study (LIS)

Suggested Citation: Zheng, Buhong (1997) : Statistical Inferences for Poverty Measures with Relative Poverty Rates, LIS Working Paper Series, No. 167, Luxembourg Income Study (LIS), Luxembourg

This Version is available at:

<https://hdl.handle.net/10419/160839>

Standard-Nutzungsbedingungen:

Die Dokumente auf EconStor dürfen zu eigenen wissenschaftlichen Zwecken und zum Privatgebrauch gespeichert und kopiert werden.

Sie dürfen die Dokumente nicht für öffentliche oder kommerzielle Zwecke vervielfältigen, öffentlich ausstellen, öffentlich zugänglich machen, vertreiben oder anderweitig nutzen.

Sofern die Verfasser die Dokumente unter Open-Content-Lizenzen (insbesondere CC-Lizenzen) zur Verfügung gestellt haben sollten, gelten abweichend von diesen Nutzungsbedingungen die in der dort genannten Lizenz gewährten Nutzungsrechte.

Terms of use:

Documents in EconStor may be saved and copied for your personal and scholarly purposes.

You are not to copy documents for public or commercial purposes, to exhibit the documents publicly, to make them publicly available on the internet, or to distribute or otherwise use the documents in public.

If the documents have been made available under an Open Content Licence (especially Creative Commons Licences), you may exercise further usage rights as specified in the indicated licence.

Luxembourg Income Study Working Paper Series

Working Paper No. 167

**Statistical Inferences for Poverty Measures
with Relative Poverty Rates**

Buhong Zheng

August 1997



Luxembourg Income Study (LIS), asbl

Statistical Inferences for Poverty Measures with Relative Poverty Lines

Buhong Zheng

Department of Economics
University of Colorado
Denver, CO, USA
Phone: (303) 556-8375
Fax: (303) 556-3547
bzheng@carbon.cudenver.edu

Abstract

Relative poverty lines such as one-half of median income have been widely used in poverty comparisons. The U.S. government has also been urged to adopt the notion of relative poverty lines. This paper contributes to the literature by developing statistical inferences for testing poverty measures with relative poverty lines. The poverty measures we consider are the decomposable class and the measure proposed by Sen. The poverty lines we specify are percentages of mean income and percentages of quantiles. We show that poverty indices can be consistently estimated and the sample estimates are asymptotically normally distributed. As a consequence, distribution-free statistical inferences can be established in a straightforward manner. We illustrate the inference procedures by comparing poverty across ten countries and over two time periods.

JEL Classifications: C40, I32

Key Words: Relative Poverty Line, Statistical Inference, Asymptotic Distribution, Decomposable Poverty Measures, Sen Poverty Measure, Luxembourg Income Study

Statistical Inferences for Poverty Measures with Relative Poverty Lines

Buhong Zheng

I. Introduction

Professor Sen's (1976) ground-breaking work on poverty measurement has fundamentally changed the way poverty is viewed and measured. It is now well recognized that a poverty measure needs to consider not only the incidence of poverty (the proportion of the people living below the poverty line) but also the income distribution within the poor. In the past two decades, a growing body of literature has been devoted to the way poverty should be measured and many new distribution-sensitive poverty measures, in addition to the one proposed by Sen (1976), have been introduced¹. The recent literature is replete with numerous empirical studies using these new poverty measures to address distributional issues.

The application of a poverty measure requires the specification of a poverty line which separates population into poor and non-poor. In the literature, there are three distinct ways to specify a poverty line, namely, absolute, relative and subjective methods and the defined poverty lines are referred to as absolute, relative and subjective poverty lines. The absolute method sets the poverty line as a minimum amount of resources at a point in time and updates the line only for price changes over time. The poverty line used in the U.S. official poverty statistics is an example of the absolute poverty line. The relative method specifies the poverty line as a point in the distribution of income or expenditure and, hence, the line can be updated automatically over time for changes in living standards. In practice, it is often to specify the relative poverty line as a percentage of mean income or expenditure (*e.g.*, Expert Committee on Family Budget Revisions (1980), O'Higgins and Jenkins (1990), Johnson and Webb (1992), and Wolfson and Evans (1989)), as a percentage of median income or expenditure (*e.g.*, Fuchs (1967), Blackburn (1990, 1994), and Smeeding (1991)), or simply as a quantile (*e.g.*, OECD (1982)). The subjective method derives the poverty line based on public opinion on minimum income or expenditure levels that can "get along" or "make ends meet." Compared with the first two approaches, the subjective method is relatively less popular and has been rarely used.

¹ For a survey on this literature, see Foster (1984), Chakravarty (1990) or Zheng (1997).

While absolute poverty lines have been used in most government poverty statistics, there is an increasing use of relative poverty lines in both international comparisons and intra-national cross-time analyses of poverty. In a recent important report on measuring poverty threshold by the Panel on Poverty and Family Assistance (1995), a group of leading scholars strongly urge the U.S. government to abandon the absolute approach that has been used since 1963, and to adopt the relative approach. A relative approach, argued the panel, "recognizes the social nature of economic deprivation and provides a way to keep the poverty line up to date with overall economic changes in a society" (p. 125). Compared with absolute poverty lines, relative poverty lines such as one-half of median income are easy to understand, easy to calculate and easy to update; they avoid the difficulty of periodic reassessments needed for absolute poverty lines. Besides, an absolute approach such as "expert budgets" also contains large elements of relativity. In fact, the initial U.S. official poverty line constructed in 1963 using the Orshansky method was just about one-half of median after-tax income. If the same method were applied to other years directly, as shown by the panel, it would yield poverty lines that are quite different from the official lines and are rather close to one-half of median after-tax incomes.

The purpose of this paper is to develop appropriate statistical inferences for poverty measures with relative poverty lines.² Specifically, we consider two popular types of poverty measures, the decomposable class and the Sen measure, and two types of relative poverty lines, percentages of mean income and percentages of quantiles (which include median income as a special case). We show that both the decomposable and the Sen poverty indices can be consistently estimated and the sample estimates are asymptotically normally distributed. Furthermore, we derive the variance-covariance structure and show that the structure can be consistently estimated and, hence, asymptotic nonparametric distribution-free statistical inferences can be established in a straightforward manner. Finally, using data from the Luxembourg Income Study database, we apply the statistical inferences to compare poverty levels across ten countries and over two time periods.

² The statistical inferences for poverty measures with absolute poverty lines have been well established in the literature (Jäntti (1992), Kakwani (1994) and Bishop et al. (1995) for decomposable poverty measures, and Bishop et al. (1997) for the Sen poverty measure). However, the task we are pursuing here is quite different since poverty lines have to be estimated from the sample so one needs to take into account of the sampling variability of poverty lines.

The rest of the paper is organized as follows. Section II develops large sample properties of the estimates of decomposable poverty measures with relative poverty lines. Section III develops large sample property of the estimates of the Sen poverty measure. Section IV shows that the variance-covariances of decomposable poverty measures and the Sen measure can be consistently estimated. Section V presents an empirical illustration of the statistical inferences. Section VI concludes the paper.

II. Large Sample Properties of Decomposable Poverty Measures

Consider a continuous (after-tax) income distribution with *d.f.* $F(x)$ where x is defined over $(0, \infty)$.

We further assume that $F(x)$ is strictly monotonic and the first two moments of x exist and are finite. Let z be a relative poverty line, *i.e.*, z is determined by some parameter of the income distribution, a decomposable (additively separable) poverty measure in its continuous form can be defined as

$$(2.1) \quad P(F; z) = \int_0^z p(x, z) dF(x),$$

where $p(x, z)$ is a positive poverty-deprivation function which is continuous in both x and z with

$\partial p(x, z) / \partial x \leq 0$ and $\partial^2 p(x, z) / \partial x^2 \geq 0$.³ In addition, we further assume that $p_z \equiv \partial p(x, z) / \partial z$

exists and is uniformly continuous over $(0, \infty)$ and that $a = E[p_z | x < z]$ exists and is finite. It is useful to point out that the last condition is not a standard one in the literature of poverty measurement, it is assumed here for us to derive the large sample properties of decomposable poverty measures. Fortunately, this requirement is satisfied by all perceivable decomposable poverty measures.

The class of poverty measures defined in (2.1) includes many commonly used decomposable poverty measures. For example, if $p(x, z) = (1 - x/z)^k$ with $k \geq 2$, then the measure $P(F; z)$ is the class of measures proposed by Foster et al. (1984). If the parameter k were allowed to take values 0 and 1, this class contains two well-known measures of poverty: the headcount ratio and the poverty gap ratio. The headcount ratio has been the official poverty measure of many countries including the United States. If

³ These two conditions ensure the weak-form monotonicity axiom and the weak-form transfer axiom to be satisfied. For a complete discussion of the poverty axioms, see, for instance, Zheng (1997).

$p(x, z) = \ln(z/x)$, then the measure $P(F; z)$ becomes the poverty measure introduced by Watts (1968). If $p(x, z) = 1 - (x/z)^b$ with $0 < b < 1$, then $P(F, z)$ is the measure proposed by Chakravarty (1983) and is also a transformation of a measure introduced in Clark et al. (1981). Of these measures, the headcount ratio and the poverty gap ratio do not satisfy the transfer axiom which requires that a poor-to-rich transfer of income increase poverty. The measures of Foster et al. for $k \geq 2$, Watts and Clark et al. satisfy the transfer axiom and are usually referred to as distribution-sensitive poverty measures.

Clearly all these poverty measures satisfy the conditions specified for $p(x, z)$. In particular, all measures satisfy the condition that p_z is uniformly continuous and $a = E[p_z | x < z]$ exists and is finite. For example, $a = 0$ for the headcount ratio, $a = E[x | x < z] / z^2$ for the poverty gap ratio, and $a = F(z) / z$ for the Watts poverty measure.

Assume a random sample of size n , $x_1, x_2, \dots, x_n$, is identically and independently drawn from a population with c.d.f. $F(x)$, then the decomposable poverty measure defined in (2.1) can be estimated as

$$(2.2) \quad \hat{P} = \frac{1}{n} \sum_{i=1}^n p(x_i, \hat{z}) I(x_i < \hat{z}),$$

where $\hat{z}$ is the sample estimate of z and the indicator variable $I(\Omega)$ for a condition Ω is defined as

$$(2.3) \quad I(\Omega) = \begin{cases} 1, & \text{if } \Omega \text{ is satisfied} \\ 0, & \text{if } \Omega \text{ is not satisfied} \end{cases}.$$

Hence $I(x_i < \hat{z})$ is one if $x_i < \hat{z}$ and is zero otherwise.

In this paper we consider two types of relative poverty lines--mean poverty lines and quantile poverty lines. The following definition formally specifies these two types of poverty lines and their sample estimates.

Definition 2.1. A poverty line z is a mean poverty line if $z = a\bar{m}$ where $\bar{m}$ is mean income and $0 < a \leq 1$;

a poverty line z is a quantile poverty line if $z = ax_q$ where x_q is a quantile of order q , i.e., $F(x_q) = q$.

The sample estimate of $z = a\bar{m}$ is $\hat{z} = a\bar{x}$ with $\bar{x} = n^{-1} \sum_{i=1}^n x_i$; the sample estimate of $z = ax_q$ is

⁴ Note that we adopt the weak definition of poverty, a person is poor if her income is strictly less than the poverty line. The strong definition, on the other hands, regards the person at the poverty line as being poor. See Donaldson and Weymark (1986) for a detailed discussion on these two definitions of poverty.

$\hat{z} = \mathbf{a}x_{(r)}$ where $r = [nq]$ and $x_{(r)}$ is the r th order statistic of $(x_1, x_2, \dots, x_n)$.

It is well known that $\bar{x}$ converges almost surely to $\mathbf{m}$ and $x_{(r)}$ converges almost surely to $\mathbf{X}_q$ (see, *e.g.*, Theorems 2.2.1A and 2.3.1 of Serfling (1980)). Hence $\hat{z}$ converges almost surely to $\mathbf{z}$ for both types of poverty lines. In what follows we derive the large sample properties of $\hat{P}$ for these two types of poverty lines.

First note that $\hat{P}$ can be expressed as

$$(2.4) \quad \hat{P} = \frac{1}{n} \sum_{i=1}^n p(x_i, z) I(x_i < z) + \frac{1}{n} \sum_{i=1}^n \{p(x_i, \hat{z}) - p(x_i, z)\} I(x_i < z) \\ + \frac{1}{n} \sum_{i=1}^n p(x_i, \hat{z}) \{I(x_i < \hat{z}) - I(x_i < z)\}.$$

Applying the mean-value theorem to $\{p(x_i, \hat{z}) - p(x_i, z)\}$, we may write the second term of the right hand side of (2.4) as

$$(2.5) \quad \frac{1}{n} \sum_{i=1}^n \{p(x_i, \hat{z}) - p(x_i, z)\} I(x_i < z) = \frac{1}{n} \sum_{i=1}^n p_z(x_i, \tilde{z})(\hat{z} - z) I(x_i < z) \\ = (\hat{z} - z) \left\{ \frac{1}{n} \sum_{i=1}^n p_z(x_i, \tilde{z}) I(x_i < z) \right\},$$

where $\tilde{z}$ is a value between $\hat{z}$ and z .

For a large n , the third term in the right hand side of (2.4) can be approximated as follows⁵:

$$(2.6) \quad \frac{1}{n} \sum_{i=1}^n p(x_i, \hat{z}) \{I(x_i < \hat{z}) - I(x_i < z)\} \sim p(z, \hat{z}) \left\{ \frac{1}{n} \sum_{i=1}^n [I(x_i < \hat{z}) - I(x_i < z)] \right\},$$

where $u_n(x) \sim v_n(x)$ denotes that $u_n(x) - v_n(x)$ converges in probability to zero. It follows from Slutsky's theorem (Theorem 1.5.4 of Serfling (1980)) that both sides of (2.6) have the same limiting distribution.

Since $\sum_{i=1}^n [I(x_i < \hat{z}) - I(x_i < z)]$ in (2.6) is the signed number of observations between $\hat{z}$ and z , it

can be approximated by $n(F(\hat{z}) - F(z))$ at a rate of convergence $o(n^{\frac{1}{2}})$, *i.e.*,

$$(2.7) \quad \sum_{i=1}^n [I(x_i < \hat{z}) - I(x_i < z)] = n(F(\hat{z}) - F(z)) + o(n^{\frac{1}{2}}).$$

⁵ This approximation can be obtained by using Young's form Taylor's theorem (Serfling (1980), p. 45) and the fact that $\hat{z}$ converges almost surely to $\mathbf{z}$.

Further applying the one-term Taylor expansion, we have

$$(2.8) \quad F(\hat{z}) - F(z) = f(z)(\hat{z} - z) + o(n^{-\frac{1}{2}}),$$

where $f(x)$ is the density function of $F(x)$. It follows that (2.6) can be rewritten as

$$(2.9) \quad \frac{1}{n} \sum_{i=1}^n p(x_i, \hat{z}) \{I(x_i < \hat{z}) - I(x_i < z)\} \sim p(z, \hat{z}) f(z)(\hat{z} - z) + o(n^{-\frac{1}{2}}).$$

Substituting (2.5) and (2.9) into (2.4), we have

$$(2.10) \quad \hat{P} \sim \frac{1}{n} \sum_{i=1}^n p(x_i, z) I(x_i < z) + (\hat{z} - z) \left\{ \frac{1}{n} \sum_{i=1}^n p_z(x_i, \tilde{z}) I(x_i < z) \right\} \\ + p(z, \hat{z}) f(z)(\hat{z} - z) + o(n^{-\frac{1}{2}}).$$

By assumption $p_z(x, z)$ is uniformly continuous in x and z , thus $n^{-1} \sum_{i=1}^n p_z(x_i, \tilde{z}) I(x_i < z)$ converges almost surely to $a = E[p_z | x < z]$. This result, together with the fact that $p(z, \hat{z})$ also converges almost surely to $p(z, z)$, yields the following approximation of $\hat{P}$:

$$(2.11) \quad \hat{P} \sim \frac{1}{n} \sum_{i=1}^n p(x_i, z) I(x_i < z) + (\hat{z} - z) \{a + p(z, z) f(z)\} + o(n^{-\frac{1}{2}}).$$

If $z = \mathbf{am}$, then (2.11) becomes

$$(2.11a) \quad \hat{P} \sim \frac{1}{n} \sum_{i=1}^n p(x_i, \mathbf{am}) I(x_i < \mathbf{am}) + \mathbf{a}(\bar{x} - \mathbf{m}) \{a + p(\mathbf{am}, \mathbf{am}) f(\mathbf{am})\} + o(n^{-\frac{1}{2}}).$$

It can be easily verified that $\lim_{n \rightarrow \infty} E(\hat{P}) = E[p(x, \mathbf{am}) I(x < \mathbf{am})] = P(F; \mathbf{am})$ which establishes the asymptotic mean of $\hat{P}$. The asymptotic normality of $\hat{P}$ can be directly verified by applying the Kolmogorov (strong) law of large numbers and the Lindeberg-Lévy central limit theorem. The variance of $n^{\frac{1}{2}}(\hat{P} - P)$

can also be readily calculated as

$$(2.12) \quad \mathbf{e}^2 = \left\{ \int_0^z p^2(x, z) dF(x) - P^2 \right\} + 2\mathbf{a}[a + p(z, z) f(z)] \left\{ \int_0^z x p(x, z) dF(x) - \mathbf{m}P \right\} \\ + \mathbf{a}^2 [a + p(z, z) f(z)]^2 \mathbf{S}_x^2,$$

where $\mathbf{S}_x^2$ is the variance of x .

If $z = \mathbf{ax}_q$, then (2.11) becomes

$$(2.11b) \quad \hat{P} \sim \frac{1}{n} \sum_{i=1}^n p(x_i, \mathbf{ax}_q) I(x_i < \mathbf{ax}_q) + \mathbf{a}(x_{(r)} - \mathbf{x}_q) \{a + p(\mathbf{ax}_q, \mathbf{ax}_q) f(\mathbf{ax}_q)\} + o(n^{-\frac{1}{2}}).$$

Using the Bahadur representation (see, e.g., Bahadur (1966) and Ghosh (1971)) which states the relationship

between population quantiles and sample quantiles,

$$(2.13) \quad x_{(r)} - x_q = \frac{q - \frac{1}{n} \sum_{i=1}^n I(x_i < x_q)}{f(x_q)} + o(n^{-\frac{1}{2}}),$$

we have

$$(2.14) \quad \hat{P} \sim \frac{1}{n} \sum_{i=1}^n p(x_i, ax_q) I(x_i < ax_q) + \frac{aq \{a + p(ax_q, ax_q) f(ax_q)\}}{f(x_q)} - \frac{a \{a + p(ax_q, ax_q) f(ax_q)\}}{f(x_q)} \times \frac{1}{n} \sum_{i=1}^n I(x_i < x_q) + o(n^{-\frac{1}{2}}).$$

It is also easy to verify that $n^{\frac{1}{2}}(\hat{P} - P)$ tends to a normal distribution with mean zero and variance

$$(2.15) \quad e^2 = \left\{ \int_0^z p^2(x, z) dF(x) - P^2 \right\} - \frac{2a[a + p(z, z) f(z)]P(1 - q)}{f(x_q)} + \frac{a^2[a + p(z, z) f(z)]^2 q(1 - q)}{f^2(x_q)},$$

where $f(z)$ and $f(x_q)$ are densities at $x = z$ and $x = x_q$ respectively.

The following theorem summarizes our main results of this section on the asymptotic distributions of $\hat{P}$ with a mean poverty line or a quantile poverty line (which includes median as a special case with $q = 1/2$).

Theorem 2.1. Under the conditions that $F(x)$ is continuous and has finite first two moments and that $p_z(x, z)$ is uniformly continuous and $a = E[p_z | x < z]$ exists and is finite, then for both definitions of the poverty lines, $\hat{P}$ is a consistent estimator of $P(F; z)$ and is asymptotically normally distributed in that $n^{\frac{1}{2}}(\hat{P} - P)$ tends to a normal distribution with mean zero and variance e^2 which is given in (2.12) and (2.15).

For individual poverty measures, the variance e^2 given in (2.12) and (2.15) can be somewhat simplified. For example, for the headcount ratio, $P(F; z) = F(z)$, $p(x, z) = 1$ and $a = 0$. Thus for $z = am$,

e^2 becomes

$$(2.16) \quad e^2 = F(z) \left\{ 1 - F(z) + 2a \left[\int_0^z x dF(x) - \mathbf{m}F(z) \right] + a^2 f(z) s_x^2 \right\},$$

and for $z = \mathbf{a}x_q$,

$$(2.17) \quad e^2 = F(z)(1 - F(z)) + \frac{af(z)(1 - q)}{f(x_q)} \{af(z)q - 2F(z)\}.$$

For most other decomposable poverty measures such as the Foster et al. measure and the Watts measure,

$p(z, z) = 0$ and $a \neq 0$. Thus the variance for $z = \mathbf{a}m$ becomes

$$(2.18) \quad e^2 = \left\{ \int_0^z p^2(x, z) dF(x) - P^2 \right\} + 2a \left\{ \int_0^z x p(x, z) dF(x) - \mathbf{m}P \right\} + a^2 a^2 s_x^2,$$

and for $z = \mathbf{a}x_q$,

$$(2.19) \quad e^2 = \left\{ \int_0^z p^2(x, z) dF(x) - P^2 \right\} - \frac{2aP(1 - q)}{f(x_q)} + \frac{a^2 a^2 q(1 - q)}{f^2(x_q)}.$$

In empirical applications, researchers often consider multiple poverty lines, instead of using a single line. For example, in developing comparable poverty estimates for member countries of the European Community, O'Higgins and Jenkins (1990) specify poverty lines to be 40, 50 and 60 percent of mean equivalent disposable income of households. The asymptotic distribution of a vector of poverty estimates for multiple poverty lines can be derived from (2.11a) for $z = \mathbf{a}m$ and from (2.14) for $z = \mathbf{a}x_q$. If both types of poverty lines are employed, then one can derive the joint asymptotic distribution between the two sets of poverty estimates by considering (2.11a) and (2.14) jointly.

Finally it is worth noting the difference in the asymptotic variances between the poverty line z being absolute and being relative. If z is absolute, then the asymptotic variance is simply the first term in (2.12) and (2.15) (see, *e.g.*, Kakwani (1994)). Hence the remaining two terms in both (2.12) and (2.15) can be attributed to the nature that z is relative and $\hat{z}$ has to be estimated from the sample. Since the additional terms in e^2 are not negligible, it is important to take them into account when poverty lines are relative.

III. The Large Sample Property of the Sen Poverty Measure

For a continuous *c.d.f.* $F(x)$ and the poverty line z , the Sen poverty measure in its continuous form is

$$(3.1) \quad S(F; z) = 2 \int_0^z \left(1 - \frac{x}{z}\right) \left(1 - \frac{F(x)}{F(z)}\right) dF(x).$$

For a given *i.i.d.* random sample, $x_1, x_2, \dots, x_n$, a well known estimator of $S(F; z)$ is the poverty measure originally proposed by Sen (1976):

$$(3.2) \quad \hat{S} = \frac{2}{(n_p + 1)nz} \sum_{i=1}^{n_p} (z - x_{(i)})(n_p + 1 - i),$$

where n_p is the number of poor (people who live below the poverty line) and $x_{(i)}$ is the i th order statistic of $(x_1, x_2, \dots, x_n)$. Sen (1976) also shows that $\hat{S}$ can be expressed as

$$(3.3) \quad \hat{S} = \hat{H} \left\{ \hat{I} + (1 - \hat{I}) \hat{G}_p \frac{n_p}{n_p + 1} \right\},$$

where $\hat{H} = n_p / n$, $\hat{I} = 1 - \bar{x}_p / z$ and $\hat{G}_p = [2n_p(n_p - 1)\bar{x}_p]^{-1} \sum_{i=1}^{n_p} \sum_{j=1}^{n_p} |x_{(i)} - x_{(j)}|$ are estimators of the headcount ratio, income gap ratio and the Gini coefficient of the income among the poor with $\bar{x}_p$ being the sample mean income of the poor. The income gap ratio is defined as $I = \int_0^z \{1 - x/zF(z)\} dF(x)$.

Note that $\hat{S}$ in (3.3) can be further expressed as

$$(3.4) \quad \hat{S} = \hat{H} \left\{ \hat{I} + (1 - \hat{I}) \frac{\hat{D}_p}{2\hat{H}\bar{x}_p} \times \frac{n_p}{n_p + 1} \times \frac{n - 1}{n_p - 1} \right\},$$

where $\hat{D}_p = [n(n - 1)]^{-1} \sum_{i=1}^{n_p} \sum_{j=1}^{n_p} |x_{(i)} - x_{(j)}|$ which is an estimator of $D_p = \int_0^z \int_0^z |x_1 - x_2| dF(x_1)dF(x_2)$. By dropping the factor $n_p / (n_p + 1)$ and replacing $(n - 1) / (n_p - 1)$ with $1 / \hat{H}$ in

(3.4), we obtain:

$$(3.5) \quad \tilde{S} = \hat{H} \left\{ \hat{I} + (1 - \hat{I}) \frac{\hat{D}_p}{2\hat{H}^2\bar{x}_p} \right\}.$$

It can be easily verified that $\hat{S} \sim \tilde{S}$, hence in what follows we will derive the large sample property of $\tilde{S}$ instead of $\hat{S}$.

Using the indicator variable defined in (2.3), we may introduce the following three *U*-statistics:

$$(3.6) \quad U_1 = \frac{1}{n} \sum_{i=1}^n \hat{z} I(x_i < \hat{z}),$$

$$(3.7) \quad U_2 = \frac{1}{n} \sum_{i=1}^n \left(1 - \frac{x_i}{\hat{z}}\right) I(x_i < \hat{z}),$$

and

$$(3.8) \quad U_3 = \hat{D}_p = \frac{1}{n(n-1)} \sum_{i=1}^n \sum_{j=1}^n |x_i - x_j| I(x_i < \hat{z}) I(x_j < \hat{z}).$$

Clearly, $\tilde{S}$ in (3.5) can be expressed in terms of U_1 , U_2 and U_3 as follows:

$$(3.9) \quad \tilde{S} = U_2 + \frac{U_3}{2U_1}.$$

As a consequence, the asymptotic distribution of $\tilde{S}$ can be derived from the joint limiting distribution of U_1 , U_2 and U_3 .

By choosing $p(x, z) = z$ and $p(x, z) = 1 - x/z$ in (2.11), we obtain the following approximations

for U_1 and U_2 :

$$(3.10) \quad U_1 \sim \frac{1}{n} \sum_{i=1}^n z I(x_i < z) + (\hat{z} - z) \{F(z) + zf(z)\} + o(n^{-\frac{1}{2}})$$

and

$$(3.11) \quad U_2 \sim \frac{1}{n} \sum_{i=1}^n \left(1 - \frac{x_i}{z}\right) I(x_i < z) + (\hat{z} - z) \left\{ \frac{1}{z^2} \int_0^z x dF(x) \right\} + o(n^{-\frac{1}{2}}).$$

Denote Θ as the interval between $\hat{z}$ and z , then U_3 can be expressed as

$$(3.12) \quad \begin{aligned} U_3 &= \frac{1}{n(n-1)} \sum_{i=1}^n \sum_{j=1}^n |x_i - x_j| I(x_i < z) I(x_j < z) \\ &+ \frac{\text{sign}(\hat{z} - z)}{n(n-1)} \sum_{i=1}^n \sum_{j=1}^n |x_i - x_j| I(x_i < z) I(x_j \in \Theta) \\ &+ \frac{\text{sign}(\hat{z} - z)}{n(n-1)} \sum_{i=1}^n \sum_{j=1}^n |x_i - x_j| I(x_i \in \Theta) I(x_j < z) \\ &+ \frac{\text{sign}(\hat{z} - z)}{n(n-1)} \sum_{i=1}^n \sum_{j=1}^n |x_i - x_j| I(x_i \in \Theta) I(x_j \in \Theta) \\ &= (i) + (ii) + (iii) + (iv). \end{aligned}$$

For a large n , part (ii) in (3.12) can be approximated as

$$(3.13) \quad (ii) \sim \frac{R(\Theta)}{n(n-1)} \sum_{i=1}^n |x_i - z| I(x_i < z),$$

where $R(\Theta) = \text{sign}(\hat{z} - z) \sum_{j=1}^n I(x_j \in \Theta)$ is the (signed) number of observations between $\hat{z}$ and z . Use

Taylor's theorem, $R(\Theta)$ can be further expressed as

$$(3.14) \quad R(\Theta) = n \{F(\hat{z}) - F(z)\} + o(n^{\frac{1}{2}}) = n \left\{ f(z)(\hat{z} - z) + o(n^{-\frac{1}{2}}) \right\} + o(n^{\frac{1}{2}}) \\ = n f(z)(\hat{z} - z) + o(n^{\frac{1}{2}}).$$

Therefore, part (ii) can be approximated as

$$(3.15) \quad (ii) \sim \frac{f(z)(\hat{z} - z)}{n} \sum_{i=1}^n (z - x_i) I(x_i < z) + o(n^{-\frac{1}{2}}).$$

Similarly, part (iii) in (3.12) can be approximated as

$$(3.16) \quad (iii) \sim \frac{f(z)(\hat{z} - z)}{n} \sum_{i=1}^n (z - x_i) I(x_i < z) + o(n^{-\frac{1}{2}}).$$

For part (iv), we have

$$(3.17) \quad (iv) \leq \frac{|\hat{z} - z|}{n(n-1)} \sum_{i=1}^n \sum_{j=1}^n I(x_i \in \Theta) I(x_j \in \Theta) \\ = \frac{|\hat{z} - z|}{n(n-1)} R^2(\Theta) = f^2(z) |\hat{z} - z|^3 + o(n^{-\frac{1}{2}}),$$

where we used the fact that $|x_i - x_j| \leq |\hat{z} - z|$ if $x_i \in \Theta$ and $x_j \in \Theta$. Since $\hat{z}$ converges almost surely to z , part (iv) is negligible in the approximation of U_3 . Thus U_3 can be approximated as follows:

$$(3.19) \quad U_3 \sim \frac{1}{n(n-1)} \sum_{i=1}^n \sum_{j=1}^n |x_i - x_j| I(x_i < z) I(x_j < z) \\ + \frac{2 f(z)(\hat{z} - z)}{n} \sum_{i=1}^n (z - x_i) I(x_i < z) + o(n^{-\frac{1}{2}}) \\ \sim \frac{1}{n(n-1)} \sum_{i=1}^n \sum_{j=1}^n |x_i - x_j| I(x_i < z) I(x_j < z) \\ + 2 f(z)(\hat{z} - z) E[(z - x) I(x < z)] + o(n^{-\frac{1}{2}}),$$

where in the second approximation we utilized the fact that $n^{-1} \sum_{i=1}^n (z - x_i) I(x_i < z)$ converges almost surely to $E[(z - x) I(x < z)]$.

Since $\hat{z}$ converges almost surely to z , we know from (3.10), (3.11) and (3.19) that U_1 converges almost surely to $zF(z)$, U_2 converges almost surely to $\int_0^z (1 - x/z) dF(x)$ and U_3 converges almost surely to $\int_0^z \int_0^z |x_1 - x_2| dF(x_1) dF(x_2)$. Therefore, by Slutsky's theorem, $\tilde{S} = U_2 + U_3 / 2U_1$ converges in probability to $S(F; z)$ for either definition of relative poverty line. This establishes $\tilde{S}$ as a consistent esti-

mator of $S(F; z)$.

Define the following four simple U -statistics $\hat{f}_1 = n^{-1} \sum_{i=1}^n I(x_i < z)$, $\hat{f}_2 = \hat{z}$, $\hat{f}_3 = n^{-1} \sum_{i=1}^n (1 - x_i/z) I(x_i < z)$, and $\hat{f}_4 = [n(n-1)]^{-1} \sum_{i=1}^n \sum_{j=1}^n |x_i - x_j| I(x_i < z) I(x_j < z)$, which clearly are consistent estimators of $f_1 = F(z)$, $f_2 = z$, $f_3 = \int_0^z (1 - x/z) dF(x)$, and $f_4 = D_p$ respectively.

Thus U_1 , U_2 and U_3 can be approximated as functions of these statistics and so can $\tilde{S}$:

$$(3.20) \quad \tilde{S} \sim \hat{f}_3 + \left\{ \frac{1}{z^2} \int_0^z x dF(x) (\hat{f}_2 - z) \right\} + \frac{\hat{f}_4 + 2zf(z)f_3(\hat{f}_2 - z)}{2\{z\hat{f}_1 + [f_1 + zf(z)](\hat{f}_2 - z)\}} + o(n^{-\frac{1}{2}}).$$

Therefore, we may obtain the large sample property of $\tilde{S}$ by deriving the joint asymptotic distribution of $\hat{\Phi} = (\hat{f}_1, \hat{f}_2, \hat{f}_3, \hat{f}_4)'$ instead of $(U_1, U_2, U_3)'$.

The following lemma provides the asymptotic distribution of $\hat{\Phi}$ and can be directly verified by applying Hoeffding's theorem (Theorem 7.1 of Hoeffding (1948)) which concerns the joint asymptotic distribution of several U -statistics.

Lemma 3.1. Under the assumption that $F(x)$ is continuous and has finite first two moments,

$\hat{\Phi} = (\hat{f}_1, \hat{f}_2, \hat{f}_3, \hat{f}_4)'$ is asymptotically normal in that $n^{\frac{1}{2}}(\hat{\Phi} - \Phi)$ tends to a normal distribution with mean zero and variance-covariance matrix $\Pi_{4 \times 4}$:

if $z = am$,

$$(3.21) \quad \Pi_{4 \times 4} = \begin{bmatrix} f_1(1-f_1) & a \left\{ \int_0^z x dF(x) - mf_1 \right\} & f_3(1-f_1) & 2f_4(1-f_1) \\ a \left\{ \int_0^z x dF(x) - mf_1 \right\} & a^2 s_x^2 & h_{23} & 2h_{24} \\ f_3(1-f_1) & h_{23} & h_{33} & 2h_{34} \\ 2f_4(1-f_1) & 2h_{24} & 2h_{34} & 4h_{44} \end{bmatrix};$$

if $z = ax_q$,

$$(3.22) \quad \Pi_{4 \times 4} = \begin{bmatrix} f_1(1-f_1) & \frac{af_1(1-q)}{f(x_q)} & f_3(1-f_1) & 2f_4(1-f_1) \\ \frac{af_1(1-q)}{f(x_q)} & \frac{a^2q(1-q)}{f^2(x_q)} & \frac{af_3(1-q)}{f(x_q)} & \frac{2af_4(1-q)}{f(x_q)} \\ f_3(1-f_1) & \frac{af_3(1-q)}{f(x_q)} & h_{33} & 2h_{34} \\ 2f_4(1-f_1) & \frac{2af_4(1-q)}{f(x_q)} & 2h_{34} & 4h_{44} \end{bmatrix},$$

where

$$(3.23) \quad h_{23} = a \left\{ \int_0^z x \left(1 - \frac{x}{z}\right) dF(x) - mf_3 \right\},$$

$$(3.24) \quad h_{24} = a \left\{ \int_0^z \int_0^z x_1 |x_1 - x_2| dF(x_1) dF(x_2) - mf_4 \right\},$$

$$(3.25) \quad h_{33} = \int_0^z \left(1 - \frac{x}{z}\right)^2 dF(x) - f_3^2,$$

$$(3.26) \quad h_{34} = \int_0^z \int_0^z \left(1 - \frac{x_1}{z}\right) |x_1 - x_2| dF(x_1) dF(x_2) - f_3 f_4, \text{ and}$$

$$(3.27) \quad h_{44} = \int_0^z \left\{ \int_0^z |x_1 - x_2| dF(x_2) \right\}^2 dF(x_1) - f_4^2.$$

Note that in deriving the variance-covariance structure, we have used $\hat{f}_2 = a\bar{x}$ for $z = am$ and $\hat{f}_2 = a \left\{ n^{-1} \sum_{i=1}^n I(x_i < z) - q \right\} / f(x_q) + ax_q + o(n^{-\frac{1}{2}})$ for $z = ax_q$, which is obtained from the

Bahadur representation (2.13).

To establish the asymptotic distribution of $\tilde{S}$, we apply Theorem 7.5 of Hoeffding (1948) on limiting distributions of differentiable functions of U -statistics. Denoting

$$g(y) = y_3 + g_1(y_2 - z) + \{y_4 + g_2(y_2 - z)\} / 2 \{zy_1 + g_3(y_2 - z)\} \text{ with } g_1 = z^{-2} \int_0^z x dF(x),$$

$g_2 = 2zf(z)f_3$, and $g_3 = f_1 + zf(z)$, we can present our second main result of the paper in the follow-

ing theorem.⁶

Theorem 3.1. Under the assumption of Lemma 3.1 and for both definitions of the poverty line, $\hat{S}$ is a consistent estimator of the Sen poverty measure $S(F; z)$ and is asymptotically normally distributed in that $n^{\frac{1}{2}}(\hat{S} - S)$ tends to a normal distribution with mean zero and variance $U^2 = T \Pi T'$ where Π is the 4×4 variance-covariance matrix given in (3.21) and (3.22) and

$$(3.28) \quad T = \left[\frac{\partial g(y)}{\partial y_i} \Big|_{y=\Phi} \right] = \left[-\frac{f_4}{2z f_1^2}, g_1 + \frac{g_2 z f_1 - g_3 f_4}{2z^2 f_1^2}, 1, \frac{1}{2z f_1} \right].$$

Note that in the theorem we have substituted $\tilde{S}$ with $\hat{S}$ since they have the same limiting distribution.

IV. Distribution-Free Statistical Inferences of Poverty Measures

The variances (e^2 and u^2) derived in the previous two sections depend upon the underlying distribution, hence, the estimates of poverty measures ($\hat{P}$ and $\hat{S}$) are not nonparametric distribution-free. However, if the variances can be consistently estimated then asymptotic nonparametric distribution-free inference procedures can be established. For example, if $\hat{e}^2$ is a consistent estimator of e^2 , then by Slutsky's theorem, statistic

$$(4.1) \quad w = \frac{n^{\frac{1}{2}}(\hat{P} - P)}{\hat{e}}$$

has a limiting standard normal distribution and is asymptotically distribution-free over the class of all continuous distributions with finite variances. Therefore, if consistent estimators of e^2 and u^2 are found, distribution-free inference procedures can be established in a straightforward manner. In what follows we show that both e^2 and u^2 for either poverty line can be consistently estimated.

First note that e^2 and u^2 contain density $f(am)$ for a mean poverty line and contain $f(ax_q)$ and

⁶ The asymptotic distribution for the measure proposed by Thon (1979) can be similarly derived. The Thon measure simply replaces $n_p + 1$ with n in the Sen measure (3.2).

$f(X_q)$ for a quantile poverty line. In the literature, there exist several nonparametric approaches to density estimation. Silverman (1986) provides a comprehensive survey on various methods ranging from the oldest method of histogram to some quite sophisticated ones. In this paper we adopt the kernel method which was introduced by Rosenblatt (1956) and is discussed thoroughly in Silverman (1986). We choose this method because it is very popular and, more importantly, because the consistency of kernel estimation has been well established in the literature.

For an *i.i.d.* random sample $x_1, x_2, \dots, x_n$, the kernel estimator of population density $f(z)$ is given as

$$(4.2) \quad \hat{f}(z) = \frac{1}{nh} \sum_{i=1}^n K\left(\frac{z - x_i}{h}\right),$$

where K is a kernel function and h is a "window width" which depends on the sample size. Under the conditions that

$$(4.3) \quad \int |K(x)| dx < \infty, \int K(x) dx = 1, \text{ and } |xK(x)| \rightarrow 0 \text{ as } |x| \rightarrow \infty,$$

and

$$(4.4) \quad h \rightarrow 0 \text{ and } nh \rightarrow \infty \text{ as } n \rightarrow \infty,$$

Parzen (1962) shows that $\hat{f}(z)$ converges in probability to $f(z)$, provided f is continuous at z . Stronger consistency of $\hat{f}(z)$ has also been established in the literature (see, for example, Silverman (1978)).

In computing $\hat{f}(z)$, one needs to choose a specific kernel function K and a window width h . Silverman (1986) documents several kernel functions and window width functions. In our empirical application of the next section, we will use the well known kernel function introduced by Epanechnikov (1969),

$$(4.5) \quad K(x) = \begin{cases} \frac{3}{4\sqrt{5}} \left(1 - \frac{1}{5}x^2\right) & \text{if } -\sqrt{5} \leq x \leq \sqrt{5} \\ 0 & \text{otherwise} \end{cases},$$

along with the window width function recommended by Silverman (1986, p. 48),

$$(4.6) \quad h = 0.9A n^{-\frac{1}{5}},$$

where $A = \min(\text{standard deviation}, \text{interquartile range}/1.34)$. One can easily verify that the above kernel function and window width function satisfy both conditions (4.3) and (4.4). Thus $\hat{f}(am)$, $\hat{f}(ax_q)$ and

$\hat{f}(X_q)$ using (4.5) and (4.6) are consistent estimators of $f(am)$, $f(ax_q)$ and $f(x_q)$ respectively, provided f is continuous at am , ax_q and x_q .

It is easy to see that all other elements besides densities in $\mathbf{e}^2$ can also be consistently estimated, hence, by Slutsky's theorem, a consistent estimator, $\hat{\mathbf{e}}^2$, can be obtained by substituting each element in $\mathbf{e}^2$ with its consistent estimator. For example, one may substitute P with $\hat{P}$ of (2.2) and a with

$$(4.7) \quad \hat{a} = \frac{1}{n} \sum_{i=1}^n p_z(x_i, \hat{z}) I(x_i < \hat{z}).$$

To show that $\mathbf{u}^2$ can be consistently estimated, it is sufficient to show that each element of matrix $\Pi_{4 \times 4}$ and vector T can be consistently estimated. Slutsky's theorem again ensures that a consistent estimator of $\mathbf{u}^2$ can be obtained by substituting each element of $\mathbf{u}^2$ with its consistent estimator.

By inspecting matrix $\Pi_{4 \times 4}$ and vector T , we can easily find consistent estimators for all elements except for h_{24} , h_{34} and h_{44} . Using results from the theory of U -statistics, we can also construct consistent estimators for h_{24} , h_{34} and h_{44} as follows:

$$(4.8) \quad \hat{h}_{24} = \frac{a}{n(n-1)} \sum_{i < j} |x_i^2 - x_j^2| I(x_i < \hat{z}) I(x_j < \hat{z}) - a \bar{x} \hat{f}_4,$$

$$(4.9) \quad \hat{h}_{34} = \hat{f}_4 - \frac{\hat{h}_{24}}{a \hat{z}} - \bar{x} \hat{f}_4 - \hat{f}_3 \hat{f}_4,$$

and

$$(4.10) \quad \hat{h}_{44} = \frac{2}{n(n-1)(n-2)} \sum_{i < j < k} g(x_i, x_j, x_k) - \hat{f}_4^2$$

with

$$(4.11) \quad g(x_i, x_j, x_k) = \left\{ |x_i - x_j| |x_i - x_k| + |x_j - x_i| |x_j - x_k| + |x_k - x_i| |x_k - x_j| \right\} \\ \times I(x_i < \hat{z}) I(x_j < \hat{z}) I(x_k < \hat{z}).$$

Using the same method as in proving the consistency of U_3 and the fact that $h(x_1, x_2)$

$= (1/2) |x_1^2 - x_2^2| I(x_1 < z) I(x_2 < z)$ is a symmetric kernel of degree 2 for parameter

$\int_0^z \int_0^z x_1 |x_1 - x_2| dF(x_1) dF(x_2)$, we can prove that $[n(n-1)]^{-1} \sum_{i < j} |x_i^2 - x_j^2| I(x_i < \hat{z}) I(x_j < \hat{z})$ is a

consistent estimator of $\int_0^z \int_0^z x_1 |x_1 - x_2| dF(x_1) dF(x_2)$. Similarly we can show that

$[n(n-1)(n-2)]^{-1} \sum_{i < j < k} g(x_i, x_j, x_k)$ is a consistent estimator of $\int_0^z \left\{ \int_0^z |x_1 - x_2| dF(x_2) \right\}^2 dF(x_1)$. It

follows from Slutsky's theorem that $\hat{h}_{24}$, $\hat{h}_{34}$ and $\hat{h}_{44}$ are consistent estimators of h_{24} , h_{34} and h_{44} respectively.

The following theorem summarizes the main results of this section:

Theorem 4.1. Under the conditions that $F(x)$ is continuous with finite first two moments and that $f(x)$ is continuous, variances Θ^2 and U^2 can be consistently estimated.

Therefore asymptotic distribution-free inferences procedures on poverty measures can be established to test for poverty changes over time and cross section. For example, if one wants to compare decomposable poverty indices between two countries in a given year or between two years of a given country, then under the assumption of independence the (asymptotic) standard normal test statistic is

$$(4.12) \quad s = \frac{\hat{P}_2 - \hat{P}_1}{\left(\hat{\Theta}_2^2 / n_2 + \hat{\Theta}_1^2 / n_1 \right)^{\frac{1}{2}}}$$

where $\hat{P}_1$ and $\hat{P}_2$ are estimates of poverty indices, $\hat{\Theta}_1^2$ and $\hat{\Theta}_2^2$ are estimates of variances and n_1 and n_2 are sample sizes (subscripts 1 and 2 may denote two countries or two periods of a given country). The test statistic for the Sen poverty measure can also be similarly computed.

V. An Application: International Comparisons of Poverty Across Ten Countries

To illustrate the application of the statistical inferences developed in this paper, we compare poverty across ten countries and over two periods. In applying the inference procedures, we treat the income samples drawn from each country as if they were independently and identically distributed. Given this assumption and the other problems associated with income survey data, it is necessary to regard the investigation presented

below as an illustration, rather than a comprehensive empirical study⁷.

The data we use are from the Luxembourg Income Study (LIS) database which currently contains income data from more than twenty countries⁸. The countries we selected are Canada, Denmark, Israel, the Netherlands, Norway, Poland, Sweden, Taiwan, the United Kingdom and the United States. The periods we consider are 1986/1987 and 1991/1992, that is, Period 1 is either 1986 or 1987 and Period 2 is either 1991 or 1992. To be specific, Period 1 is 1986 for Israel, Norway, Poland, Taiwan, the United Kingdom and the United States, and is 1987 for other countries; Period 2 is 1991 for Canada, the Netherlands, Norway, Taiwan, the United Kingdom and the United States, and is 1992 for other countries. The selection of the ten countries and the choice of the two time periods are largely determined by the availability of data. In what follows we like to compare poverty across the ten countries in each period and to examine changes in poverty over time for each country.

The income concept we use in calculating poverty indices is disposable income (after-tax income). The poverty measures we employ are the headcount ratio, the Watts measure and the Sen measure; the last two are distribution-sensitive measures. The poverty lines we use are one-half of mean income and one-half of median income. To account for differences in family composition, we employ the equivalence scale recommended by the Panel on Poverty and Family Assistance (1995),

$$(5.1) \quad (D + BK)^G,$$

where D is the number of adults in a family, K is the number of children under age 18, B is the proportion that each child should be treated as an adult, and G is a scale factor. The panel further recommends that B and G be close to 0.70. In our illustration, we let $B = G = 0.70$. The income after being adjusted by equivalence scale is the so-called "adult equivalent income" and the poverty lines are estimated from the distribution of "adult equivalent income."

⁷ The statistical inferences developed in this paper can be generalized to the case of dependent samples, albeit somewhat complicated. The inferences can also be extended to other types of samples such as stratified samples and cluster samples. However, we pursue these extensions elsewhere.

⁸ For a list of the countries and a description of the data, one may visit the LIS web site at <dpls.dacc.wisc.edu/apdu/lis>. The site also provides information for accessing the database.

Table 1 reports poverty estimates of three poverty measures and two poverty lines for the ten countries over two periods. In the table, symbol 'H' stands for the headcount ratio, 'W' stands for the Watts poverty measure and 'S' stands for the Sen poverty measure. The standard errors are given in parentheses underneath poverty estimates. By inspection, it is clear that poverty exists in every country in each period and at each poverty line (the w-value calculated by use of (4.1) is very large for each poverty estimate).

Table 2 presents cross-time statistical poverty comparisons of the ten countries with two poverty lines. The first number in each cell is the difference $(\hat{P}_2 - \hat{P}_1)$ between the poverty estimate of Period 1 and that of Period 2; the second number in [] is the t-value calculated from (4.12). Since for each country, we compare six changes in poverty indices (three for each poverty line), it is necessary to test them simultaneously. In this paper we adopt a relatively conservative joint test--the Student Maximum Modulus (SMM) Test.⁹ The critical value for six comparisons is 2.378 at the 10% significance level. Hence if the t-value is greater than 2.378 or less than -2.378, then we reject the null hypothesis that there is no change in poverty index. For example, for Canada, while each poverty measure at each poverty line indicates that poverty has decreased from Period 1 to Period 2, none of these changes is significant at the 10% level (all t-values are negative and greater than -2.378). In this manner, we can perform statistical inference for each country by each poverty measure and at each poverty line.

From Table 2, we can see that from Period 1 to Period 2, Canada, Israel, Sweden and the United States did not experience significant changes in poverty by all three poverty measures and at two poverty lines. For other countries, changes in poverty are somewhat dependent upon poverty measures used and/or poverty lines specified: (1) For Denmark and Norway, the conclusions on poverty comparison are measure-dependent. Both H and the S indicate a significant increase in Denmark's poverty but W indicates no significant changes at either poverty line. While H indicates a significant decrease in Norway's poverty, both W and S indicate opposite changes which are not statistically significant. (2) For Taiwan, the poverty comparison is line-dependent. When the poverty line is one-half of mean income, none of the three measures indicates a significant

⁹ For a detailed description of the test, see Savin (1984) or Miller (1981).

cant change, but at one-half of median income all three measures indicate a significant increase in poverty.

(3) For the Netherlands, Poland and United Kingdom, poverty comparisons are both measure-dependent and line-dependent.

Poverty estimates and standard errors in Table 1 also enable us to compare poverty for any two countries at each poverty line.¹⁰ Table 3 reports poverty rankings among the ten countries in two periods at two alternative poverty lines. In each column of Table 3 countries are ranked from the highest poverty level to the lowest. The value given in [] is the s-value from (4.12) for the comparison between two adjacent countries (here $\hat{P}_1$ and $\hat{P}_2$ are poverty estimates of two adjacent countries). For example, the first column ranks countries by H with the poverty line being one-half of mean income in Period 1. The value 4.536 beneath Israel (IS) is the s-value of the comparison between the U.S. and Israel, it indicates that the U.S. has a significant higher poverty level than Israel.¹¹ On the other hand, the value 0.122 underneath Sweden (SW) reveals that there is no significant difference between Norway (NW) and Sweden.

By inspecting Table 3, we can see that the United States is ranked the highest in poverty in both periods for all three poverty measures and two poverty lines. The United Kingdom, Canada and Israel are ranked next to the U.S. for most of comparisons. On the other hand, Norway is virtually the country having the lowest poverty level (the only exception is H at one-half of median income in Period 1). The Netherlands and Sweden are ranked closely above Norway for most of comparisons. The rankings for most countries are also sensitive to the poverty measure and/or the poverty line used. For example, in Period 1 and at the mean poverty line, W shows that the United Kingdom is not significant different from Israel but S indicates that Israel's poverty level is significantly higher than the United Kingdom. Also in Period 1, W shows Taiwan has a significant higher poverty level than Poland at the one-half of mean income, but the direction of ranking is reversed at one-half of median income.

VI. Summary and Conclusion

¹⁰ Results on pairwise comparisons are available from the author upon request.

¹¹ The critical SMM value for poverty rankings at the 10% significance level is 2.512.

In recent poverty studies, relative poverty lines such as one-half of mean income and one-half of median income have been widely used. A relative poverty line, as argued by the Panel on Poverty and Family Assistance, provides a way to keep the poverty line up to date with overall economic changes in a society. It is also easy to understand, easy to calculate and easy to update. Besides, an absolute approach such as "expert budgets" also contains large elements of relativity.

This paper contributes to the literature by developing statistical inferences for two popular classes of poverty measures with relative poverty lines. The poverty measures we considered are the popular decomposable poverty measures and the famous Sen poverty measure. The poverty lines we considered are percentages of mean income and percentages of quantiles which include median income as a special case. Under certain regularities and assumptions, we show that the poverty indices can be consistently estimated and the poverty estimates are asymptotically normally distributed; the variances can also be consistently estimated. Therefore distribution-free statistical inferences can be established in a straightforward manner.

The paper also applies the developed inference procedures to compare poverty across ten countries and over two time periods. The data we used are from the Luxembourg Income Study database. The poverty measures we use are the headcount ratio, the Watts measure and the Sen measure; the poverty lines we set are one-half of mean income and one-half of median income. Between the two periods, we find that changes in poverty indices of Canada, Israel, Sweden and the United States are not statistically significant; for other countries, conclusions are somewhat dependent upon measures used and the line specified. In both periods, we find that the U.S. is unambiguously the country with the highest poverty level for all three measures and two poverty lines; Norway, on the other hand, is virtually the country having the lowest poverty level in both periods.

References

- Bahadur, R. (1966), A Note on Quantiles in Large Samples, *Annals of Mathematical Statistics*, 37, 577-580.
- Bishop, J. K. V. Chow and B. Zheng (1995): Statistical Inference and Decomposable Poverty Measures, *Bulletin of Economic Research* 47, 329-340.
- Bishop, J. J. Formby and B. Zheng (1997): Statistical Inference and the Sen Index of Poverty, *International Economic Review* 38, 381-387.
- Blackburn, M. (1994): International Comparisons of Poverty, *American Economic Review* 84 (Papers and Proceedings), 371-374.
- (1990): Trends in Poverty in the United States, 1967-84, *Review of Income and Wealth* 36, 53-66.
- Chakravarty, S. R. (1990): *Ethical Social Index Numbers*, New York: Springer-Verlag.
- (1983a): A New Index of Poverty, *Mathematical Social Science* 6, 307-313.
- Clark, S., R. Hemming and D. Ulph (1981): On Indices for the Measurement of Poverty, *Economic Journal* 91, 515-526.
- Donaldson, D., and J. A. Weymark (1986): Properties of Fixed-Population Poverty Indices, *International Economic Review* 27, 667-688.
- Epanechnikov, V. A. (1969): Nonparametric Estimation of a Multidimensional Probability Density, *Theor. Probab. Appl.* 14, 153-158.
- Expert Committee on Family Budget Revisions (1980): *New American Family Budget Standards*, Madison: Institute for Research on Poverty, University of Wisconsin.
- Foster, J. (1984): On Economic Poverty: A Survey of Aggregate Measures, In: R. L. Basmann and G. F. Rhodes (eds.), *Advances in Econometrics* 3, Connecticut: JAI Press.
- , J. Greer and E. Thorbecke (1984): A Class of Decomposable Poverty Measures, *Econometrica* 52, 761-766.
- Fuchs, V. R. (1967): Redefining Poverty and Redistributing Income, *Public Interest* 5, 88-95.
- Ghosh, J. K. (1971): A New Proof of the Bahadur Representation of Quantiles and an Application, *The Annals of Mathematical Statistics* 42, 1957-1961.
- Hoeffding, W. (1948): A Class of statistics with asymptotically normal distribution, *Annals of Mathematical Statistics* 19, 293-325.
- Jäntti, M. (1992): *Poverty Dominance and Statistical Inference*, Research Report, Stockholm University.
- Johnson, P. and S. Webb (1992): Official Statistics on Poverty in the United Kingdom, In: *Poverty Measurement for Economies in Transition in Eastern European Countries*, Warsaw: Polish Statistical Association and Polish Central Statistical Office.
- Kakwani, N. (1994): Poverty Measurement and Hypothesis Testing, In: J. Creedy (ed.), *Taxation, Poverty and Income Distribution*, Aldershot: Edward Elgar.
- Miller, R. (1981): *Simultaneous Statistical Inference*, New York: John Wiley & Sons.

- OECD (1982): *The OECD List of Social Indicators*, Paris.
- O'Higgins, M. and S. Jenkins (1990): Poverty in the EC: Estimates for 1975, 1980, and 1985, In: R. Teekins and B. M. S. van Praag (eds.), *Analysing Poverty in the European Community: Policy Issues, Research Options, and Data Sources*, Luxembourg: Office of Official Publications of the European Communities.
- Panel on Poverty and Family Assistance (1995): *Measuring Poverty: a New Approach*, Washington, D. C.: National Academy Press.
- Parzen, E. (1962): On Estimation of a Probability Density Function and Mode, *Annals of Mathematical Statistics* 33, 1065-1076.
- Rosenblatt, M. (1956): Remarks on Some Nonparametric Estimates of a Density Function, *Annals of Mathematical Statistics* 27, 832-837.
- Savin, N. (1984): Multiple Hypothesis Testing, In: Z. Griliches and M. Intriligator (eds.), *Handbook of Econometrics* 2, Elsevier Science Publishers.
- Sen, A. K. (1976): Poverty: an Ordinal Approach to Measurement, *Econometrica* 44, 219-231.
- Serfling, R. J. (1980): *Approximation Theorems of Mathematical Statistics*, New York: John Wiley & Sons.
- Silverman, B. W. (1986): *Density Estimation for Statistics and Data Analysis*, London: Chapman and Hall.
- (1978): Weak and Strong Uniform Consistency of the Kernel Estimate of a density Function and Its Derivative, *Annals of Statistics* 6, 177-184.
- Smeeding, T. M. (1991): Cross-National Comparisons of Inequality and Poverty Position, in L. Osberg (ed.), *Economic Inequality and Poverty: International Perspectives*, Armonk: M. E. Sharpe, Inc.
- Thon, D. (1979): On Measuring Poverty, *Review of Income and Wealth* 25, 429-440.
- Watts, H. (1968): An Economic Definition of Poverty, In D. P. Moynihan (ed.), *On Understanding Poverty*, New York: Basic Books.
- Wolfson, M. and J. M. Evans (1989): *Statistics Canada's Low Income Cut-Offs: Methodological Concerns and Possibilities -- A Discussion Paper*, Research Paper Series, Ottawa: Statistical Canada.
- Zheng, B. (1997): Aggregate Poverty Measures, *Journal of Economic Surveys* 11, 123-163.

Table 1. Poverty Estimates of Ten Countries

	Poverty Line = 1/2 Mean Income						Poverty Line = 1/2 Median Income					
	Period 1			Period 2			Period 1			Period 2		
	<i>H</i>	<i>W</i>	<i>S</i>	<i>H</i>	<i>W</i>	<i>S</i>	<i>H</i>	<i>W</i>	<i>S</i>	<i>H</i>	<i>W</i>	<i>S</i>
Canada	0.1405 (0.0028)	0.1199 (0.0030)	0.1056 (0.0019)	0.1353 (0.0021)	0.0535 (0.0013)	0.0536 (0.0011)	0.1090 (0.0027)	0.0448 (0.0018)	0.0442 (0.0016)	0.1046 (0.0020)	0.0410 (0.0012)	0.0411 (0.0011)
Denmark	0.1015 (0.0025)	0.0477 (0.0031)	0.0383 (0.0014)	0.0641 (0.0021)	0.0445 (0.0033)	0.0306 (0.0013)	0.0822 (0.0025)	0.0429 (0.0031)	0.0331 (0.0015)	0.0545 (0.0020)	0.0418 (0.0033)	0.0275 (0.0013)
Israel	0.1887 (0.0054)	0.0558 (0.0025)	0.0616 (0.0024)	0.1719 (0.0051)	0.0494 (0.0022)	0.0547 (0.0022)	0.1184 (0.0048)	0.0316 (0.0020)	0.0359 (0.0025)	0.1039 (0.0045)	0.0264 (0.0017)	0.030 (0.0022)
Nether-lands	0.0688 (0.0038)	0.0286 (0.0037)	0.0254 (0.0021)	0.0731 (0.0042)	0.0438 (0.0043)	0.0352 (0.0025)	0.0377 (0.0030)	0.0216 (0.0035)	0.0171 (0.0020)	0.0501 (0.0033)	0.0362 (0.0041)	0.0274 (0.0023)
Norway	0.0679 (0.0033)	0.0209 (0.0020)	0.0221 (0.0017)	0.0578 (0.0026)	0.0230 (0.0020)	0.0215 (0.0014)	0.0492 (0.0031)	0.0173 (0.0019)	0.0178 (0.0017)	0.0393 (0.0022)	0.0198 (0.0020)	0.0174 (0.0014)
Poland	0.1136 (0.0029)	0.0328 (0.0012)	0.0369 (0.0012)	0.1351 (0.0040)	0.0498 (0.0023)	0.0515 (0.0020)	0.0839 (0.0027)	0.0242 (0.0011)	0.0273 (0.0013)	0.0938 (0.0035)	0.0343 (0.0022)	0.0356 (0.0020)
Sweden	0.0674 (0.0024)	0.0358 (0.0020)	0.0323 (0.0014)	0.0641 (0.0020)	0.0367 (0.0019)	0.0320 (0.0013)	0.0611 (0.0024)	0.0338 (0.0020)	0.0301 (0.0015)	0.0565 (0.0020)	0.0332 (0.0019)	0.0287 (0.0013)
Taiwan	0.1544 (0.0035)	0.0401 (0.0013)	0.0451 (0.0013)	0.1568 (0.0033)	0.0428 (0.0013)	0.0482 (0.0013)	0.0784 (0.0023)	0.0191 (0.0009)	0.0216 (0.0010)	0.0942 (0.0024)	0.0237 (0.0009)	0.0270 (0.0011)
United Kingdom	0.1297 (0.0040)	0.0568 (0.0036)	0.0496 (0.0021)	0.2077 (0.0056)	0.0731 (0.0037)	0.0726 (0.0026)	0.0755 (0.0034)	0.0434 (0.0034)	0.0348 (0.0023)	0.1297 (0.0043)	0.0443 (0.0031)	0.0432 (0.0026)
United States	0.2168 (0.0031)	0.1199 (0.0030)	0.1056 (0.0019)	0.2199 (0.0027)	0.1233 (0.0030)	0.1059 (0.0017)	0.1789 (0.0031)	0.0947 (0.0029)	0.0842 (0.0022)	0.1766 (0.0028)	0.0979 (0.0029)	0.0836 (0.0021)

Data Source: LIS database. Period 1 is 1986 for Israel, Norway, Poland, Taiwan, the United Kingdom and the United States, and is 1987 for other countries; Period 2 is 1991 for Canada, the Netherlands, Norway, Taiwan, the United Kindom and the United States, and is 1992 for other countries. Standard errors are in parentheses. *H* is the headcount ratio, *W* is the Watts measure and *S* is the Sen measure.

Table 2. Poverty Changes between Two Periods

Countries	Difference in Poverty Indices					
	Poverty Line = 1/2 Mean Income			Poverty Line =1/2 Median Income		
	H	W	S	H	W	S
Canada	-0.0052 [-1.487]	-0.0042 [-1.841]	-0.0036 [-1.963]	-0.0044 [-1.306]	-0.0038 [-1.772]	-0.0031 [-1.589]
Denmark	-0.0374 [-11.34]*	-0.0032 [-0.701]	-0.0077 [-3.983]*	-0.0277 [-8.811]*	-0.0011 [-0.241]	-0.0056 [-2.780]*
Israel	-0.0168 [-2.264]	-0.0064 [-1.942]	-0.0069 [-2.161]	-0.0145 [-2.225]	-0.0052 [-1.986]	-0.0059 [-1.749]
Nether- lands	0.0043 [0.757]	0.0152 [2.685]*	0.0098 [3.022]*	0.0124 [2.782]*	0.0146 [2.708]*	0.0103 [3.361]*
Norway	-0.0101 [-2.400]*	0.0021 [0.739]	-0.0006 [-0.278]	-0.0099 [-2.602]*	0.0025 [0.906]	-0.0004 [-0.185]
Poland	0.0215 [4.385]*	0.0170 [6.598]*	0.0146 [6.391]*	0.0099 [2.222]	0.0101 [4.129]*	0.0083 [3.504]*
Sweden	-0.0033 [-1.057]	0.0009 [0.322]	-0.0003 [-0.157]	-0.0046 [-1.477]	-0.0006 [-0.217]	-0.0014 [-0.729]
Taiwan	0.0024 [0.497]	0.0027 [1.474]	0.0031 [1.712]	0.0158 [4.825]*	0.0046 [3.695]*	0.0054 [3.670]*
United Kingdom	0.0780 [11.38]*	0.0163 [3.188]*	0.0230 [6.859]*	0.0542 [9.967]*	0.0009 [0.194]	0.0084 [2.451]*
United States	0.0031 [0.753]	0.0034 [0.799]	0.0003 [0.120]	-0.0023 [-0.557]	0.0032 [0.775]	-0.0006 [-0.200]

Data Source: LIS data base. Symbol * indicates the difference is significant at the 10% level; the s-value is given in [] and the SMM critical value for 6 comparisons at the 10% significance level is 2.378.

Table 3. Poverty Rankings among Ten Countries

Period 1						Period 2					
Poverty Line = 1/2 Mean Income			Poverty Line = 1/2 Median Income			Poverty Line = 1/2 Mean Income			Poverty Line = 1/2 Median Income		
<i>H</i>	<i>W</i>	<i>S</i>	<i>H</i>	<i>W</i>	<i>S</i>	<i>H</i>	<i>W</i>	<i>S</i>	<i>H</i>	<i>W</i>	<i>S</i>
US	US	US	US	US	US	US	US	US	US	US	US
IS	CN	IS	IS	CN	CN	UK	UK	UK	UK	UK	UK
[4.536]*	[17.62]*	[14.67]*	[10.68]*	[14.73]*	[14.78]*	[1.963]	[10.58]*	[10.79]*	[9.240]*	[12.48]*	[12.35]*
TW	UK	CN	CN	UK	IS	IS	CN	IS	CN	DK	CN
[5.326]*	[0.224]	[1.576]	[1.714]	[0.364]	[2.769]*	[4.731]*	[5.037]*	[5.318]*	[5.338]*	[0.549]	[0.753]
CN	IS	UK	PL	DK	UK	TW	PL	CN	IS	CN	PL
[3.075]*	[0.230]	[2.924]*	[6.625]*	[0.108]	[0.321]	[2.487]	[1.404]	[0.458]	[0.144]	[0.226]	[2.413]
UK	DK	TW	DK	SW	DK	CN	IS	PL	TW	NL	IS
[2.216]	[2.028]	[1.807]	[0.464]	[2.436]	[0.617]	[5.536]*	[0.127]	[0.944]	[1.926]	[1.118]	[1.883]
PL	TW	DK	TW	IS	SW	PL	DK	TW	PL	PL	SW
[3.288]*	[2.235]	[3.517]*	[1.138]	[0.770]	[1.423]	[0.044]	[1.235]	[1.418]	[0.094]	[0.408]	[0.509]
DK	SW	PL	UK	PL	PL	NL	NL	NL	SW	SW	DK
[3.181]*	[1.176]	[0.749]	[0.712]	[3.249]*	[1.438]	[10.67]*	[0.128]	[4.666]*	[9.182]*	[0.382]	[0.660]
NL	PL	SW	SW	NL	TW	DK	TW	SW	DK	IS	NL
[7.159]*	[1.267]	[2.472]	[3.483]*	[0.711]	[3.642]*	[1.907]	[0.222]	[1.150]	[0.711]	[2.698]*	[0.037]
NW	NL	NL	NW	TW	NW	SW	SW	DK	NL	TW	TW
[0.178]	[1.090]	[2.725]*	[3.051]*	[0.694]	[1.947]	[0.000]	[2.666]*	[0.776]	[1.153]	[1.424]	[0.156]
SW	NW	NW	NL	NW	NL	NW	NW	NW	NW	NW	NW
[0.122]	[1.848]	[1.239]	[2.673]*	[0.845]	[0.269]	[1.935]	[4.915]*	[4.799]*	[2.733]*	[1.814]	[5.497]*

1. CN = Canada, DK = Denmark, IS = Israel, NL = the Netherlands, NW = Norway, PL = Poland, SW = Sweden, TW = Taiwan, UK = the United Kingdom, and US = the United States.

2. Period 1 is 1986 for Israel, Norway, Poland, Taiwan, the United Kingdom and the United States, and is 1987 for other countries; Period 2 is 1991 for Canada, the Netherlands, Norway, Taiwan, the United Kingdom and the United States, and is 1992 for other countries.

3. Symbol * indicates that the difference is significant at the significance 10% level. The s-value given in [] represents the poverty comparison of the country with the one above it.