

TARGETING METHOD OF CONDITIONAL CASH TRANSFER PROGRAM FOR THE POOR IN INDONESIA

Rifai Afin*

Abstract

The objective of this paper is to present a methodology for combining both geographical and households targeting for the poverty reduction program in Indonesia as an alternative targeting method. The method is based on the use of data that provide information on the characteristics of the areas in which the households reside and individual household characteristic. Method of targeting for CCT program in Indonesia has to find supply side readiness and conditionality of demand side. Two methods of the targeting for CCT program in Indonesia could be used together for improving the results of the targeting. Combining the geographical and household targeting with proxy mean test (PMT) will increase accuracy for choosing benefits of the program.

Keyword: cash transfer program, geographical targeting, Proxy Mean Test (PMT)

Introduction

Public policies in developing countries are often articulated in terms of poverty reduction objectives. Resources for such purposes are invariably scarce relative to the number and magnitude of competing claims. Spending priorities must be defined, and it is often desirable to target social transfers to those beneficiaries whose needs are most urgent. Coady and Morley (2003) survey experience with such targeted transfer programs and show that errors of inclusion and exclusion are unavoidable consequences of such targeting efforts. Efforts aimed at improving targeting of public spending generally focus on reducing either one, or sometimes both, of these types of errors.

Targeting benefits to the poor first requires a precise definition of the target group. Once the target group is established, a methodology must be found for identifying individuals or households that are in that group and for excluding those who are not. For instance, if the poor are identified as a target group for a program, one must be able to make a precise judgment about the level of welfare or the means of the recipient. Targeting benefits to the poor, however simple in concept, is an inexact art in practice. Rigorous targeting requires a precise definition of the target group, which in turn may require a political consensus that is hard to solidify. This can be quite difficult technically, as well as being costly. In practice, targeting reflects the tradeoffs between the advantages of focusing program benefits on those who need them the most and the political, technical and financial difficulties to do so.

Because the precise economic circumstances of households can be difficult to ascertain it is not easy to define who should be eligible to receive a government transfer. Nor is it straightforward to design an administrative mechanism to ensure that the transfer actually reaches the intended beneficiary. In practice governments often exploit geographic variability in the design of targeting schemes: poverty is typically thought to be more concentrated in some areas of a country than others and most countries have an administrative structure that disaggregates to different levels. For example, the central government, located in the capital city, may rely on state or provincial governments to implement government policies at the state or province-level. These administrations might rely, in turn, on counties or districts, which may themselves rely on yet lower levels of administration. Resources aimed at poverty reduction can thus be directed to those localities where poverty is concentrated and administration of these transfer schemes can be carried out at the relevant local level.

A comparative study of targeting in Latin America has found that, among all targeting mechanisms, proxy means tests produce the best incidence outcomes (Grosh 1994). Proxy means tests use household or individual characteristics to proxy a means test, thus avoiding the problems involved in relying on reported income.

The first thing to notice is that interventions use a combination of targeting methods; in all cases have 253 occurrences of different targeting methods, so that the interventions in all I know sample use just over two different targeting methods on average. Just 48 interventions use a single targeting method, while 42 use two methods, 21 use three methods, and 11 use four methods.

Indonesia has experience in the transfer program for poor people, which was called Subsidi Langsung Tunai (SLT). This program is the type of Unconditional Cash Transfer (UCT). Targeting method that is applied for this program is proxy means test (PMT). Badan Pusat Statistik/ Indonesia Central Bureau of Statistic apply this method with the logit model fourteen poor variables for the poor people database and from the estimation of the model BPS got the welfare rank which is eligible or not eligible to be beneficiaries of the SLT program. Actually, Indonesia applied several programs that were provided to poor people, for instance; Raskin (Beras untuk Masyarakat Miskin), Askeskin (Asuransi Kesehatan Untuk Masyarakat Miskin), etc. All of those programs are proposed to decrease the life burden of poor people in Indonesia.

In the next generation of the cash transfer in Indonesia, central government has new plan for changing the previous cash transfer program from unconditional cash transfer becomes conditional cash transfer because some critics or problem for the previous cash transfer program. Those problems which are including targeting, socialization, payment, form of the program, as well as the impact of the program. With all of experience of those programs government hopes that in the next program we could improve and decreasing the problem that alter.

The plan of Indonesian conditional cash transfer (CCT) program is concern in two sectors those are education and health and targeting method for the program must also be provided in those two sectors. The main issue of the program is

conditionality and targeting. The conditionality in education for instance; increasing daily school attendance, increasing school enrollment, etc and conditionality for health for instance; vaccines up to date, pre and post natal care, health visits and etc.

Targeting for CCT programs in Indonesia must be fulfilled by the two components, demand side and supply side. Unlike in the developed countries, in the developing countries, public facility is not established in all of regions. It needs to be thought to solve this problem. The CCT program developed for improving demand side (improving beneficiaries' conditionality). Due to objectives of the program, supply side for supporting the increasing of demand side must be established.

The objective of this paper is to present a methodology for combining both geographical and households targeting. The method is based on the use of data that provide information on the characteristics of the areas in which the households reside and individual household characteristic. These data are collected from several different sources and organized as a poverty map and SUSENAS (Survey Sosial Ekonomi Nasional/ National Social Economic Survey) that identifies the target areas by their geographical coordinates and identifies the household targets. The overall goal is to evaluate the effect and cost-effectiveness of poverty alleviation programs that are targeted on small geographical areas and households. The methodology will be illustrated in the paper for evaluating the potential benefits from reducing the target areas of poverty alleviation programs from the level of the state to the level of the district. The paper is structured as follows: Section 2 provides further details on presents the methodology that is used for estimating the poverty incidence in both geographic and households and the data requirements. Section 3 presents technical notes and the econometric model simulations that can be estimated to get the best model both in probit or logit technique for geographic targeting and Proxy Mean Test (PMT) as a tool for household targeting technique.

Targeting Method for CCT Program

Briefly Review of Targeting Method for Other Social Safety Net Programs

This section discusses the targeting that has been used recently in the Indonesian Social Safety Net Program. Table 1 lists the various social safety net programs established by the Government of Indonesia to mitigate the social impact of the recent crisis. These programs were launched in early 1998, but many of them did not start until the second half of the year. These programs were intended to help protect the pre-crisis poor as well as the newly poor as a result of the crisis through a fourfold strategy: (i) ensuring the availability of food at affordable prices, (ii) supplementing purchasing power among poor households through employment creation, (iii) preserving the access of the poor to critical social services, particularly health and education, and (iv) sustaining local economic activity through regional block grant programs and the extension of small-scale credit.

In general, the targeting for these programs was based on a combination of geographic and household targeting mechanisms, except for the subsidized rice program which used only household targeting. The targeting for some programs was based on a household classification created by the National Family Planning Coordinating Agency (BKKBN). According to this classification, households are divided into four socio-economic status groups: 'pre-prosperous households' ("keluarga pra-sejahtera" or KPS), 'prosperous I households' ("keluarga sejahtera I" or KS I), KS II, and KS III. The KS I to KS III categories are often lumped together as the KS or 'prosperous' category.

A household is defined as a 'pre-prosperous' household if it fails to satisfy one of the following five conditions: (i) all household members are able to practice their religious principles, (ii) all household members are able to eat at least twice a day, (iii) all household members have different sets of clothing for home, work, school, and visits, (iv) the largest floor area of the house is not made of earth, and (v) the household is able to seek modern medical assistance for sick children and family planning services for contraceptive users. Suryahadi *et al.* (1999) find that there is a lack of correlation between this official classification and consumption-based measure of poverty. They find that while only 15 percent of the 'prosperous' households were 'poor', 75 percent of the 'pre-prosperous' households were 'non-poor'. On the other hand, 46 percent of the 'non-poor' households were 'pre-prosperous' and 38 percent of the 'poor' households were 'prosperous'.

Table 1.
Targeting Track Records for Social Safety Net Program in Indonesia

Area	Program Description and Benefits	Targeting	1998/1999	1999/2000
Food Security	OPK Program: sale of subsidized rice to targeted households Eligible Households can purchase 10-20 kg of rice at Rp 1000/kg (market price is Rp 2500-3000/kg)	Geographic	None	None
		Households	BKKBN list	BKKBN list with flexibility up dated with regional data
Community Empowerment	PDM-DKE: a community fund program that provides block grant directly to villages for either public works or revolving credit funds	Geographic	Pre-crisis data	data
		Households	Local decision making	Local decision making
Employment creation	"padat karya" a loose uncoordinated collections of several labor intensive programs in various government department	Geographic	None, various ministries	urban areas, base on employment
		Households	Weak self selection	self selection
Education	Scholarship and block grants: providing 1. scholarship of Rp 10.000/month for elementary (SD) students, Rp 20.000/month for secondary (SLTP) students, Rp 30.000/month for upper secondary (SMU) students 2. Block grants for selected schools	Geographic	old data on enrollment	poverty data updated to 1998
		Households	school committees applying criteria	school committees applying criteria
Health	JPS-BK a program providing subsidies for: 1.medical service 2.operational support for health centers 3.medicine and imported medical equipment 4.family planning service 5.nutrition (supplementary food) 6.midwife service	geographic	BKKBN pre-prosperous rates	pre-prosperous rates updated to 1999
		households	BKKBN list	BKKBN list with flexibility

There have been a number of criticisms of the use of the BKKBN lists for targeting purposes. The list does not capture transitory shocks to income as they are based on relatively fixed assets (such as the type of floor in the house, possession of changes of clothing). In addition, the lists are compiled by relatively poorly trained workers at the village level, so consistency across regions is not assured, and the composition of the list is susceptible to changes by local government officials.

The subsidized rice and the health programs explicitly used this BKKBN household classification for targeting. The selection of recipients in the scholarship program was also intended to take into account their BKKBN household status. Originally, eligible recipients for some JPS programs were only KPS card holders, but for certain programs, for example the OPK program, eligibility was extended to include KS I households as well.

The *padat karya* programs consisted of quite diverse programs and although specific programs were targeted to particular areas (such as drought areas), the lack of coordination meant that in effect there was little or no systematic geographic targeting of this set of programs. Within these labor 'intensive' programs there were a variety of disagreements about the desired characteristics of intended participants but typically the beneficiaries were not chosen according to any fixed administrative criteria. Hence, to the extent that there was targeting, it was primarily through self-selection. Only those who were willing to work should have been able to receive the benefits.

In the scholarship program, scholarship funds were at first allocated to schools so that "poorer" schools received proportionally more scholarships. In each school, the scholarships were then distributed to individual students by a school committee, which in theory consisted of the principal, a teacher representative, a student representative, the head of the parent association as the representative of community, and the village head. The selection of scholarship recipients was based on a combination of various administrative criteria, which included a number of factors, such as household data from school records, family BKKBN status, family size, and the likelihood of students dropping out of school.

School students in all but the lowest three grades of primary school were officially eligible. In principle, students selected to receive the scholarships were supposed to be from the poorest backgrounds. As guidance, scholarships were to be allocated at first to children from households in the two lowest BKKBN rankings. If there were more eligible students than the number of scholarships available, then additional indicators were to be used to identify the neediest students. These additional indicators included the distance from home to school, physical handicaps, and those children coming from large or single parent families. Also, a minimum of 50 percent of the scholarships, if at all possible, were to be allocated to girls.

In the health programs, meanwhile, the free medical and family planning services program was implemented by giving 'health cards' to eligible households. Eligibility was also based on BKKBN household status. A health

card given to a household could be used by all members of the household to obtain free services from designated hospitals, clinics, and health care centers for all medical and family planning purposes, including pregnancy check-ups and child-birth services.

The last social safety nets program in 2005, UCT (Unconditional Cash Transfer), as compensation of decreasing subsidy for the oil price since the oil price become high level increasing in October 2005. The beneficiaries of the program calculated base on SUSENAS database and selected by geographic and households targeting method. Eligibility Indicators of the program base on fourteen poverty indicator that has been used by BPS (Center Bureau of Statistic). Government try to develop the program from UCT become CCT. The target of beneficiaries of the program also needs to re-define because not all of the poor people will get the benefit if they do not fulfill the precondition or conditionality of the program. Base on the facts that many problems were found in previous targeting methods and reporting the field database for the beneficiaries, the better method has to develop in the CCT program.

Method of targeting for CCT program in Indonesia has to find supply side readiness and conditionality of demand side. The two method of the targeting for CCT program in Indonesia could be used together for improving the results of the targeting. The methods could be mentioned below:

1. Geographic Targeting: Geographic targeting is the first we can do to get which is region that fulfilled the supply readiness and having big amount of poor people percentage and fulfilled the conditionality. In this step we will find region that eligible to implement this program both in the supply side and demand side.
2. Household Targeting with Proxy Mean Test: Household targeting provide to select households that are eligible become beneficiaries (poor households and fulfilled conditionality).

Geographic Targeting

Geographic targeting involves allocating resources to geographic areas using information that is thought to be a good indicator of the extent of poverty in these areas. For this reason, this approach is now commonly referred to as "poverty mapping." The areas used may be political subdivisions of the country (states or counties), or they may be the catchments of specific service providers such as clinics or schools. There are a number of approaches to poverty mapping; these differ essentially according to the amount of information used and how it is combined to evaluate the extent of poverty in each area. Besides the mapping of the poor households in certain area, it is also clear need mapping for supply side readiness to catch up the improving conditionality of the poor households in certain area.

This methodology (geographic targeting steps) is based on an econometric estimation of the poverty indices in small areas by using location-specific data from a wide variety of sources (poverty map and SUSENAS). These sources include the Agricultural Survey, the Population Census, and various sources of

information on the geographical characteristics of the areas (height, distance, topography, etc.), their agro-climatic conditions, road infrastructure, public facilities, etc. The estimation methodology is based on a seven-step procedure:

1. Econometric estimation of the impact of location-specific characteristics of the areas in which the households reside on the probability that these households are poor. This estimation is based on the entire SUSENAS sample of households and on two sets of explanatory variables: (i) Household-specific variables from the SUSENAS (ii) Location-specific variables from all the other sources.
2. Estimation of the incidence of poverty in *all* the target areas (districts) in the country based on their location-specific characteristics (available from the other data sources) and on the relationships estimated in step 1.
3. Ranking areas from the poorest to the least poor according to the estimated values of the incidence of poverty in each area and grouping the areas into broad poverty groups with equal shares in the general population. The group of the *poorest* districts includes the districts that can be the target of poverty alleviation programs; the group of the *least poor* districts includes the districts that could be the target of cost recovery programs.
4. First validation of the estimations: This validation is based on a comparison of the ranking of *states* established by the econometric estimates of the values of the incidence of poverty in the states with the ranking established by the levels of poverty in the states computed directly from the SUSENAS data. High rank correlation for the states suggests that the corresponding rank correlation for the districts is also likely to be high.
5. Second validation of the predictions: This validation is based on a comparison of the predicted levels of the poverty incidence in *groups* with the actual levels of poverty in *these groups* computed directly from the SUSENAS data.
6. Clustering/standardizing the supply side readiness data from other sources for instance: education and health department (education and health facility data for each district) with demand side (predicted levels)
7. Comparing cluster/standardized the supply side with predicted level of demand side and also clustering/standardized supply side with actual demand side.

In the last results of the step we will get the regions with high density of supply side and demand side base on calculation both predicted and actual data. These regions are the most appropriate regions for the pilot program of CCT program in Indonesia because the program provide push the demand side as the objective of the program as shown in Table 2.

Table 2. Illustration Regions of Target

	Region with High Poor Households	Region with Low Poor Households
Region with High Education Health Facility	Regions of Target	
Region with Low Education and Health Facility		

In the step 1 simulation model must be tried to get the best model. Robustness of each model that is estimated becomes urgent in this case. The big questions in this simulation model step is which variables that effects the robustness of the model and what method of estimation to get the best model. From the basic model we can do some exercises to detect and select which the best model is.

Estimation basic model can be done by several methods. Firstly, we can do Probit or Logit Model Estimation. This is the most usual method to get the best model in geographic targeting. Lot of technical paper of geographic targeting uses this method. In the first simulation we can use the dependent variable of the model as probability of per capita consumption expenditure for each household and then we collect and rank the household in each region and do the next steps of geographic targeting. The probability that the level of per capita consumption of an individual household with the characteristics specified by the explanatory variables falls below the poverty line is measured by equation (II.2.1) below:

$$Prob(Y^H \leq Z) = Prob\{(X^H)' \beta^H + (X^A)' \beta^A \leq Z\} = F\{Z - [(X^H)' \beta^H + (X^A)' \beta^A]\}$$

Where Y^H denotes the household's per capita consumption expenditure, X^H is the vector of explanatory variables that describe the household's information, and X^A is the vector of explanatory variables that describe the characteristics of the "area" or area information- the district (or the region) in which the households resides, and Z is the poverty line. F is a cumulative distribution function, which is standard normal in the case of probit and logistic in the case of logit regression.

The regression analysis is conducted over the entire data set of the SUSENAS after incorporating the vector X^H of individual households' characteristics from the SUSENAS and the vector X^A of the area characteristics from the Population Census and all the other sources. For a given poverty line Z and a given set of observations on X^H and X^A , the estimates of β^H and β^A can be obtained by maximizing the corresponding likelihood function. Two equations were estimated in the empirical analysis, one where the explanatory variables are both X^H and X^A , and the other where only X^A are the explanatory variables. The former equation estimates the marginal impact of the location-specific variables, whereas the latter equation estimates their overall impact.

To identify the group of districts that should be the target of the poverty alleviation program while minimizing the prediction errors, the districts are ranked in step 3 according to the values of the poverty incidence estimates from the *poorest* district, in which the value of the estimated incidence of poverty is the highest, to the *least poor* district, in which the value of that estimate is the lowest. The districts were then divided into several target groups that have approximately equal share in the *general* population. The districts in the first group, in which the estimates are the highest, are also the districts that should have the highest priority in the implementation of poverty alleviation programs; the districts in the last group, in which the estimates are the lowest, could be the target of cost recovery programs. For simplicity, we divide the districts into four groups and refer to them as "*poorest*", "*highly poor*", "*moderately poor*" and "*least poor*." The number of households in the SUSENAS in each of these groups

is sufficient to provide statistically significant estimates of the *actual* incidence of poverty in each group on the basis of the SUSENAS data. We divided the districts into four groups in order to minimize as much as possible the probability that districts that were classified due to estimation error in the group of “*poorest*” districts – and thus would be entitled to the benefits of poverty alleviation programs – should, in fact, have been classified in the group of the “*least poor*” districts.

The first validation test in step 4 draws conclusions on the reliability of the district ranking by comparing the ranking of *states* established by the econometric estimates with the ranking established by the SUSENAS data. The higher the coefficient of rank correlation for the states, the higher the likelihood that the ranking of districts established by these estimates will also be highly correlated with the ranking established by the SUSENAS data. The second validation test in step 5 compares the actual values of the incidence of poverty in the four groups that were calculated directly from the SUSENAS survey with the values calculated from the estimates for the individual districts calculated in the econometric analysis. Since the sample of households in each of these groups of districts is sufficiently large, we can obtain statistically significant estimates of the incidence of poverty from the SUSENAS survey data. If the difference between the SUSENAS estimates of the poverty incidence and those based on econometric estimation is not very large, we can conclude that, despite the possibly high prediction error at the *individual* district level, the predictions for groups of districts are sufficiently reliable. After we do the entire steps that were examined above, then we can do the next steps to get the regions of the target.

Secondly, we use the Probit or Logit model estimation to get probability of the region with eligibility on demand side directly. This way is done by running the model in regional aggregation not at household level. Household information variables are generated in the regional form and dependent variable is changed by the average consumption expenditure in each region. The probability of the dependent variable is 1 if the average consumption expenditure on one region is below poverty line and 0 if above the poverty line. From this estimation technique we get the predicted value of the regional consumption expenditure and then we rank the region and decide where the cut off point is and the last we can do the next step.

Thirdly, we can do the econometric exercises by the Ordinary Least Square (OLS) to run with different model from the previous model in the Probit and Logit model estimation. In this way we change the model, especially variables included in the model. Ratio of the poor incidence as dependent variable, and explanatory variables is accorded to the dependent variables. Basically, explanatory variable in this model is similar with the second model estimation but the difference is the dependent variable. In the second model dependent variable is the probability of the mean of per capita consumption expenditure in the certain region or district but in the third model dependent variable is the ratio of the poor incidence in the certain area or region with all population in that certain region. To get the best model we have to try the entire model above and then compare all of the goodness of fit of the model with several indicators such as Adjusted R-Squared, Akaike Information Criteria (AIC), Schwartz

Criteria, and Final Prediction Error (FPE) as well as Hannan-Quinn indicator. Each of the goodness of fit criteria usually has little difference with other criteria but sometime high different level. We could estimate and simulation of the model which is showed that all or almost all of the goodness of fit model indicator is best.

In step 6 and 7 we try to standardize the supply readiness database with the predicted value of the econometric exercises and also the actual value of the poor incidences in each region/district. The region with high supply readiness and high poor incidence both predicted and actual could be the pilot area of the CCT program in Indonesia.

Household Targeting (Proxy Means Test/PMT)

Proxy means tests use a relatively small number of household characteristics to calculate a score that indicates the household's economic welfare. This score is used to determine eligibility for receipt of program benefits and possibly also the level of benefits. The Proxy Mean Test of in this step is done from the previous geographical targeting. The steps for the proxy Means Test can be explained below:

1. Data Selected for the Analysis

The data used for the CCT exercise is SUSENAS, conducted by the BPS. This is a multi topic household survey in the style of a PSE.RT, with modules on consumption, income, employment, health, nutrition, fertility, education, and living conditions. It also includes information on benefits received from existing welfare programs, and was designed to be representative at the national, provincial, and district levels.

2. Selecting an Indicator for Actual Household Welfare

The second step in designing a proxy means test is to select a few variables that are well correlated with poverty and have three characteristics:

- Variables should be few enough that it is feasible to apply the proxy means test to the significant share of the population that may apply for the program, possibly as much as one third.
- Variables selected must be easy to measure or observe.
- Variables should be relatively difficult for the household to manipulate just to get into a program. These variables are typically drawn from the data sets of detailed household surveys, for example, a household budget survey or a multi topic survey that include detailed information on consumption, employment, education, health, housing, and family structure.

There are also issues of conceptual notions of poverty. Economists traditionally have focused on income or consumption as a measure of welfare, the latter typically being interpreted as a better proxy for "permanent" or lifetime income. In contrast, much of the history of poverty mapping has used a "basic needs" approach with poverty defined in terms of access to basic services. The indicators used are often interpreted using one of these approaches.

In most cases the variables selected include indicators of the location of the family's home, the quality of its dwelling, its ownership of durable goods, the demographic structure of the household, labor force status, occupation or sector of work for the adults, and sometimes partial measures of income as well as for the welfare indicator it is better to use consumption expenditure than income.

In development literature, consumption expenditure is generally considered a more accurate measure of welfare than income for several reasons. First, because consumption expenditures tend to be less variable than income over seasons, it is more likely to indicate the household's "true" economic status, as a result of households with sporadic incomes smoothing their consumption patterns over time. Second, in practice, consumption is generally measured with far greater accuracy than income in a household survey, primarily because households' sources of income may include home-based production, own farms and businesses. Calculating the flow of *net* incomes from these sources turn out to be a big problem since the flow of costs and returns from these activities are often inaccurately reported by households.

3. Predicting Welfare: The Choice of Ordinary Least Squares (OLS)

To predict welfare, the consumption variable is regressed, using OLS method, on different sets of explanatory variables. The case for using OLS as the model for predicting welfare is driven primarily by convenience and ease of interpretation. The first problem with using an OLS model is that many of the explanatory variables are likely to be endogenous to (and thus not independent explanators of) household welfare. This problem is however is of less concern to us, since our objective is solely to *identify* the poor and not to explain the *reasons* for their poverty. Second, Grosh and Baker (1995) points out that strictly speaking, OLS is inappropriate for predicting poverty since the technique minimizes the squared errors between the "true" and the predicted levels of welfare, which is a different theoretical problem from that of minimization of poverty. That said, OLS is considered convenient and useful by these authors when a large numbers of predictor variables, including continuous variables, are available. Moreover, using OLS has the advantage of being able to intuitively interpret the coefficients of the predictors on welfare – a feature that is likely to appeal to a policymaker and more amenable to achieving political consensus in the country.

4. Predicting Welfare: The Choice of Variables

Selection of variables to predict welfare as measured by per capita consumption should take into account two separate criteria: *correlation* between the welfare measure and the predictor, which will determine accuracy of the prediction, and *verifiability* of the predictor, which will determine the accuracy of information used to impute welfare. The types of predictors used for this exercise, discussed below, were arrived at after judging all possible predictors on the basis of these two criteria.

- *Location* variables are obviously the most easily verifiable, and the same is true for *characteristics of the community*, when it is defined in simple terms like the presence of a bank or administrative offices. *Housing quality* may also be easily verified by a social worker visiting

the home. *Household characteristics*, such as the number of members and dependents, and age, education and occupation of the household head, are less easy to verify. However, it is generally felt that these information, firstly, are not overly difficult to verify, and secondly, that households are less likely to misrepresent such information. Using program officers, who live in the same community as the applicant households to collect the information, also makes it more likely that such information will be reported correctly.

- *Ownership of durable goods or farm equipment* is verifiable by inspection – however they can be misrepresented by the household removing the goods from the home during an expected visit by the social worker, which is easier to do with small or mobile items than for items such as stoves or refrigerators. The general presumption in the literature is also that people are more willing to lie about ownership of such items than they are about household characteristics. However, these variables tend to have high predictive power for welfare, and therefore including them can reduce mis-targeting substantially.
- *Ownership of productive assets* is again not easy to verify. The presence of livestock is verifiable to some extent. As for land ownership, while it may not be measured perfectly, one can reasonably expect that program officers who belong to the community will have local knowledge about whether a household owns a large amount of land or not, which will deter misrepresentation. The fact that these variables are likely to have high correlations with poverty in rural areas makes a strong case for including them as predictors of welfare.

Very briefly, the steps in the procedure for arriving at the PMTF run as follows. The original set of variables belonging to the six broad categories is identified based on the two criteria mentioned above. Dichotomous variables are then created for some of the continuous variables in order to identify those characteristics that discriminate between poor and rich households. The set of selected predictors are then introduced in a weighted OLS regression of (log of) per capita monthly consumption expenditure. Different subsets of variables are checked for possible multicollinearity, and a few variables are adjusted or dropped as necessary to reduce such problems. A stepwise regression is then used with the remaining set of variables because it is designed to eliminate from the regression variables that are not statistically significant and do not increase the model's overall explanatory power. From this process, different models evolve based on the subset of variables entering into the regression.

5. Determining Eligibility

Each model predicts a certain level of welfare, as measured by (log of) per capita monthly consumption expenditure. These predicted welfare levels are used to assign individuals to eligible or ineligible groups, based on an eligibility cutoff point.

The selection of the cutoff point is essentially a policy, and not a technical decision. By simulating a wide range of scenarios corresponding to different cutoff points for each model, we seek to achieve two objectives. Firstly, the exercise will show the sensitivity of the model and its attendant

errors in targeting to changes in cutoff points. Second, the simulations will help the government make a policy decision on what the cutoff point should be, taking into account the trade offs inherent in choosing a relatively higher cutoff vis-à-vis a low one.

6. Evaluating the Targeting Formulae

As with all regression analyses, different specifications of the model and different samples of the population yield different results and it is not always easy to say which specification is superior. However, a variety of tests can be conducted, which, taken together, can be used to select one model over another. We use two types of criteria to evaluate alternate options for the PMTF. The first criterion is the regression's R^2 , which is the proportion of the variation in consumption that is explained by the regression model. Higher the R^2 , the better are a particular set of variables in predicting welfare.

The second criterion involves looking at measures that indicate the ability of various models to identify the poor properly. No matter what model is used, given that it can predict welfare only with some imperfection, it is likely that some truly eligible people will be left out, while others who are not eligible will benefit. Following Grosh and Baker (1995) and related literature for other countries, we evaluate targeting accuracy of alternate models using Type I and II errors, from which rates of under coverage and leakage are derived, and incidence of benefits across income/consumption groups. Individuals are categorized in four groups according to whether their *true* and *predicted* (by the regression model) welfare levels fall above or below the defined eligibility cutoff point. Those whose true welfare falls below the eligibility threshold constitute the "target" group, while those with predicted welfare below the eligibility threshold constitute the "eligible" group. Individuals whose true and predicted welfare measures put them on the same side of the cutoff line are targeting "successes".

Table 3. Illustration of Type I and II error

	Target Group	Non-Target Group	Total
Eligible: predicted by PMT	Targeting Success (s1)	Type II error (e2)	m1
Ineligible: predicted by PMT	Type I error (e1)	Targeting Success (s2)	m2
Total	n1	n2	n

While it would be preferable to have low levels of leakage *and* under coverage, in reality one may face tradeoffs between these two objectives. In general, the higher the priority assigned to raising the welfare of the poor, the more important it is to eliminate under coverage. Conversely, if saving program costs is a higher priority, it is more important to minimize leakage. Lowering leakage, besides being cost-efficient, can also be welfare increasing in the presence of a budget constraint – lower the leakage of benefits to ineligible individuals; higher would be the amount available for transfers to those who are eligible.

The last criterion to evaluate targeting efficiency is by looking at how a specific PMTF allocates potential beneficiaries across the expenditure distribution. It is preferred that a model has good *incidence*, i.e. most of the identified beneficiaries belong to the bottom of the consumption (income) distribution, and relatively few, if any, from the top of the distribution.

Technical Notes and Econometric Model Simulations

Technical Notes

As mentioned in section II.3.3 it should be noted that the estimated coefficients may not be consistent if the disturbances are heterocedastic. Further, the fact that some of the explanatory variables, notably ownership of durable goods, are endogenous—that is, determined by the income level of household, and hence implicitly by the poverty status—may add problem of inconsistency and bias of the estimated coefficients. The latter problem is commonly incurred in studies in which regression analysis is used to combine poverty indicators. As pointed out by Minot (2000) however, in the present context the methodology may be at least partially justified by the fact that the overarching objective is to use regression analysis to develop a descriptive tool which will enable us to identify the poor, rather than study the determinants of poverty or the magnitude of the coefficients. These issues will be found especially in the Proxy Mean Test but sometimes we find in the geographic targeting also with the Probit or Logit model.

The second question with our model that we will be estimated above in the geographic targeting that is which technique that we will be estimated for results the best predicted value for geographic targeting step is. Many of literature say that Linear Probability Model (LPM), Logit, and Probit give qualitatively similar results; we will confine our attention to Logit and Probit models because of the problems with the LPM. We know that LPM plagued by several problems, such as (1) non normality of u_i (error term), (2) heteroscedasticity, (3) possibility of predicted value lying outside the 0-1 range, and generally low R^2 values. Between Logit and Probit, which model is preferable? In most applications the models are quite similar, the main difference being that the logistic distribution has slightly fatter tails, which can be seen in Figure 3.1

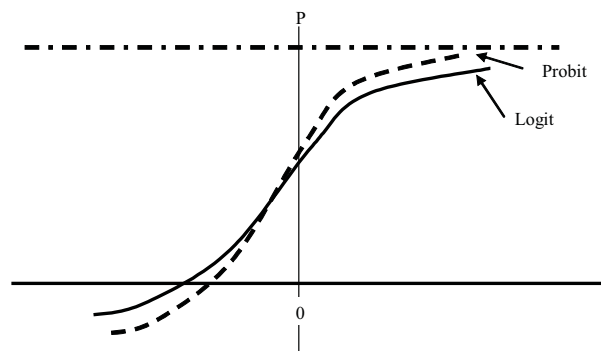


Figure 1. Probit and Logit Cumulative Distribution

That is to say, the conditional probability P_1 approaches zero or one at a slower rate in Logit than in Probit. Therefore, there is no compelling reason to choose one over the other. In practice many researchers choose the Logit model because of its comparative mathematical simplicity.

Model Simulation

In this section we try to investigate an alternative model for both geographic and household targeting estimation formula. In the geographic targeting there are several models that can be estimated to result in the best predicted value. For all models, stepwise regressions can be used to eliminate insignificant variables, and retain only those whose statistical significance but ordinary regression could be used also for comparisons. The alternative model for geographic targeting can be seen below:

Table 4. Alternative Explanatory Variables included in Geographic Model

Variables	Model 1	Model 2	Model 3
Explanatory Variables:			
Household Information			
1. Household Characteristics			
Number of Children (0-5)	V	V	
Number of Children (6-15)	V	V	
Having Micro credit	V	V	
Households Size	V	V	
2. Housing Characteristic			
Per capita Floor	V	V	
Type of Floor is not Land	V		
Toilet Facility is Private	V	V	
Clean Water Source	V	V	
Electricity is PLN	V		
3. Household Head Characteristics			
Households Head Junior High School	V	V	
Households Head Senior High School or Above	V	V	
Sex of Households Head	V	V	
Age of Households Head (Productive or not)	V	V	
4. Households Ownership			
Telephone	V	V	
Gas Stove	V		
Computer	V		
Refrigerator	V		
Television	V		
Radio	V		
Video	V		
Area Information:			
Life Expectancy	V	V	V
Adult Literacy	V	V	V
Female Literacy	V	V	V
Medical Infrastructure	V	V	V
Medical Worker	V	V	V
Education Infrastructure	V	V	V
Education Worker	V	V	V

Table 4 gives us the alternatives or choices to get the best model estimation for predicting poverty. Model 1 contains the full set of predictors. These include selected variables household characteristics (number of children (0-5), number of children (6-15), having micro credit, household size), housing characteristic (per capita floor, type of floor is not land, toilet facility is not private, clean water sources, and electricity is PLN), household head characteristic (junior or senior high school graduates, type of sex, age), household ownership (all of the household durable goods), and area characteristic where the household resides (life expectancy, adult literacy, female literacy, health infrastructure, health worker, education infrastructure, and education worker).

In model 2 we try to dropping several problematic variables such as type floor is not land, electricity is PLN, and several; durable goods. Type floor is not land is problematic because in Indonesia several regions have tradition that floor of their house from wood and sands although they are not poor. Electricity is PLN does not mention whether those households have the electric meter by their own or not. In several places most of poor people does not electric meter but they join or share with their neighborhood. Durable goods does not included in the model 2 because several cases show that people bring their some durable goods to their neighborhood when the survey staff comes to them. In model 3 as we mentioned in the section 2, this estimation are made on the basis of the relationship between 'area' characteristics and the probability that households residing in these areas are poor. In other words, in this step the probability that households in a given district are poor is estimated on the basis of the district characteristics alone, and this estimate is given by:

$$Prob(Y^H = Z) = Prob\{(X^A)' \beta^A = Z\} = F\{Z - (X^A)' \beta^A\}$$

β^A denote the coefficients from equation (II.2.1) that were estimated in step I. The prediction error in these estimates depends, on the one hand, on how detailed and how accurate the available location-specific information is and, on the other hand, on the explanatory power of the location-specific variables with respect to the level of households' consumption. When the information available is not very detailed or when the explanatory power of location-specific variables is relatively small, the prediction error can be quite large. In the subsequent steps, we design the analysis so as to take this possibility into account. Notice also that the Probit regression is used to predict the probability of the household, rather than the individual, being poor. Using the information on household size, this estimated probability could be extended to estimate also the probability of an individual being poor.

In the Table 5 we try to describe the explanatory variable in the household targeting model (Proxy Mean Test).

Table 5. Alternative Model of Household Targeting

Variables	Model 1	Model 2
Explanatory Variables:		
Household Information		
1. Household Characteristics		
Number of Children (0-5)	V	V
Number of Children (6-15)	V	V
Having Micro credit	V	V
Households Size	V	V
2. Housing Characteristic		
Per capita Floor	V	V
Type of Floor is not Land	V	
Toilet Facility is Private	V	V
Clean Water Source	V	V
Electricity is PLN	V	
3. Household Head Characteristics		
Households Head Junior High School	V	V
Households Head Senior High School or Above	V	V
Sex of Households Head	V	V
Age of Households Head (Productive or not)	V	V
4. Households Ownership		
Telephone	V	V
Gas Stove	V	
Computer	V	
Refrigerator	V	
Television	V	
Radio	V	
Video	V	

From the Table 5 we can try two models for getting the best predicted value in each region or eligible region for receiving CCT programs. This is the last section of the targeting exercises. The results from this step are eligible households who reside in the eligible regions or districts.

In econometric modeling both geographic and household model, not only the best predicted value that has to be thought but also including econometric rules. The econometric rules include endogenous variable, efficiency, specification, and goodness of fit of the model. For these problems, targeting is not easy to do; so many exercises must be done to get the best result.

Conclusion

Targeting benefits to the poor first requires a precise definition of the target group. Once the target group is established, a methodology must be found for identifying individuals or households that are in that group and for excluding those who are not.. Targeting benefits to the poor, however simple in concept, is an inexact art in practice. Rigorous targeting requires a precise definition of the target group, which in turn may require a political consensus that is hard to solidify.

The plan of Indonesian conditional cash transfer (CCT) program is concern in two sectors those are education and health and targeting method for the program must also be provided in those two sectors. The main issue of the program is conditionality and targeting. The conditionality in education for instance; increasing daily school attendance, increasing school enrollment, etc and conditionality for health for instance; vaccines up to date, pre and post natal care, health visits and etc.

Targeting for CCT programs in Indonesia must be fulfilled by the two components, demand side and supply side. Unlike in the developed countries, in the developing countries, public facility is not established in all of regions. It needs to be thought to solve this problem. The CCT program developed for improving demand side (improving beneficiaries' conditionality). Due to objectives of the program, supply side for supporting the increasing of demand side must be established.

Method of targeting for CCT program in Indonesia has to find supply side readiness and conditionality of demand side. The two method of the targeting for CCT program in Indonesia could be used together for improving the results of the targeting. Combining the geographical and household targeting with proxy mean test (PMT) will increase accuracy for choosing benefits of the program. Geographic targeting involves allocating resources to geographic areas using information that is thought to be a good indicator of the extent of poverty in these areas. For this reason, this approach is now commonly referred to as "poverty mapping." Proxy means tests use a relatively small number of household characteristics to calculate a score that indicates the household's economic welfare. This score is used to determine eligibility for receipt of program benefits and possibly also the level of benefits.

References

- Baker, J., and Grosh, M. 1994. Poverty Reduction Through Geographic Targeting: How Well Does It Work? *World Development*, 22(7): 983-995.
- Coady, D. and Morley, S. 2003. *From Social Assistance to Social Development: Targeted Education Subsidies in Developing Countries* (Center for Global Development and International Food Policy Research Institute).
- Demombynes, G., Elbers, C., Lanjouw, J.O., Lanjouw, P., Mistiaen, J. and Özler, B. 2002 .Producing an Improved Geographic Profile of Poverty: Methodology and Evidence from Three Developing Countries. In van der Hoeven, R. and Shorrocks, A. (eds) *Growth, Inequality and Poverty: Prospects for Pro-Poor Economic Development* (Oxford: Oxford University Press).

- Duclos, J., Makdissi, P., and Wodon, Q. 2003. Poverty-Efficient Transfer Programs: the Role of Targeting and allocation Rules. *CIRPÉE Working Paper 03-05*. Elbers, C., Lanjouw, J.O. and Lanjouw, P. (2002). Micro-Level Estimation of Welfare. Policy Research Working Paper 2911, Development Research Group, the World Bank, Washington D.C.
- Elbers, C., Lanjouw, J.O. and Lanjouw, P. 2003. Micro-Level Estimation of Poverty and Inequality. *Econometrica*, 71(1): 355-64.
- Elbers, C., Lanjouw, P., Mistiaen, J., Özler, B., and Simler, K. 2004. Unequal Inequality of Poor Communities. *World Bank Economic Review* (forthcoming).
- Fujii, T. 2003. Commune-level Estimation of Poverty Measures and its Application in Cambodia. *Unpublished manuscript, Dept. of Agriculture and Resource Economics*, UC Berkeley.
- Gelbach, J. and Pritchett, L. 2002. Is More for the Poor Less for the Poor? The Politics of Means-Tested Targeting. *Topics in Economic Analysis and Policy*, Vol 2(1).
- Glewwe, P. 1992. Targeting Assistance to the Poor. *Journal of Development Economics*, 38: 297-321.
- Hentschel, J. and Lanjouw, P. 1996. Constructing an Indicator of Consumption for the Analysis of Poverty: Principles and Illustrations with Reference to Ecuador. *LSM Working Paper No.124*, DECRG-World Bank: Washington DC.

Appendix

We introduced the Logit and Probit model in previous section as a tool for the geographic targeting tool. As we know the Logit and Probit model have dependent variable takes value only between 0 and 1 (or between 0 and 100, if it is in percentage form). We first describe the Logit model and then the Probit model. The Logit or Logistic model has the following functional form:

$$\ln \frac{P}{1-P} = X' \beta + u \quad (1)$$

P denotes the value of the dependent variable between 0 and 1. The rationale for this form can be seen by solving the equation for P (by exponentiation of both sides). We then obtain the probability that the dependent variable takes the value P, as follows:

$$P = \frac{1}{1 + e^{-(X' \beta + u)}} \quad (2)$$

It is easy to see that if $X = 1$, P is 1 and when $X = 0$, P takes the value 0. Thus, P never is outside the range $[0, 1]$.

The estimation procedure depends on whether the observed P is between 0 and 1, or whether it is binary and takes the value 0 or the value 1. In the case in which P is strictly between 0 and 1 the method is simply to transform P and obtain $Y = \ln[P/(1-P)]$. Then regress Y against a constant, and X (more explanatory variable are easily added). If, however, P is binary, the logarithm of $P/(1-P)$ is undefined when P is either 0 or 1.

***Rifai Afin** is Lecturer at Faculty Economic and Business of Trunojoyo University and Associate Researcher at Regional Economic Development Institute (REDI).