

MODELING OF SOYBEAN PRODUCTION IN INDONESIA USING ROBUST REGRESSION

Susanti, Y., dan Pratiwi, H.

Department of Mathematics Sebelas Maret University Surakarta;
Jl. Ir. Sutami 36A Kentingan Surakarta 57126
E-mail: yuliana@mipa.uns.ac.id ; hpratiwi@mipa.uns.ac.id

ABSTRACT

Soybean is considered as a vital commodity for food security, but apparently the production is unable to compensate the rate of increase in community needs. Therefore, this commodity plays an important role in economic life and indirectly may affect the level of demand for other materials. So the availability of soybean plays a major role for economic stability. The aim of this paper is to construct a robust regression model for predicting the soybean production in Indonesia using M-estimation. Based on the data obtained from Susenas and BPS, we obtain a robust regression model for soybean production. Harvested area and productivity have significant influence, while production of seed does not have significant influence on soybean production. The increment of one hectare harvested area and one quintal per hectare of productivity will increase 1.34 tons and 574 tons of soybean production respectively.

Keywords: Robust regression, M-estimation, Huber function, soybean.

PEMODELAN PRODUKSI KEDELAI DI INDONESIA MENGGUNAKAN REGRESI ROBUST

ABSTRAK

Komoditi kedelai dianggap sangat vital bagi ketahanan pangan, namun ternyata produksi kedelai tidak mampu mengimbangi laju peningkatan kebutuhan masyarakat. Oleh karena itu komoditi ini memegang peranan penting di dalam kehidupan ekonomi dan secara tak langsung dapat mempengaruhi tingkat permintaan bahan-bahan lainnya. Dengan demikian besar kecilnya ketersediaan kedelai berperan besar bagi stabilitas ekonomi. Penelitian ini bertujuan untuk mengkonstruksi model regresi *robust* untuk memprediksi produksi kedelai di Indonesia dengan menggunakan estimasi M. Berdasarkan data yang diperoleh dari Susenas and BPS diperoleh model regresi *robust* untuk produksi kedelai. Luas lahan panen dan produktivitas mempunyai pengaruh yang signifikan, sedangkan produksi benih tidak mempunyai pengaruh yang signifikan terhadap produksi kedelai. Penambahan satu hektar luas lahan panen dan satu kuintal per hektar produktivitas berturut-turut akan meningkatkan 1,34 ton dan 574 ton produksi kedelai.

Kata kunci: Regresi *robust*, estimasi M, fungsi Huber, kedelai.

INTRODUCTION

Soybean is one of the important commodities in nine staples. Soybean for food processing industry in Indonesia is widely used as raw material for making tofu, tempe, soy sauce, and milk. This industry is categorized as small-medium scale industry, but in huge amount it leads to high level of soybean consumption demand. To meet demand of tofu and tempe, there are 115,000 home industries of tofu and tempe in

Indonesia based on the National Economic Census (Susenas) data 2006 by the Central Agency of Statistics (BPS).

Indonesia needs approximately 2.20 tons of soybean per year. The domestic production only meets 35–40% of the demand and the remaining 60–65% are imported from foreign countries. Therefore, through various programs, the government put strong efforts to increase soybean production toward self-sufficiency in 2010–2012 (Marwoto & Suharsono, [6]).

Food price policy is one of important instrument in creating a national food security. Because of the importance of meeting the needs of food, especially soybean, an effort is needed to determine the availability of soybeans in the future. Soybean production in Indonesia is influenced by several factors, including the harvested area, seed production and productivity.

In statistical modeling, regression analysis is a method that can be used to find and to model the relationship among variables. According to Zhang [10], the use of least squares method in regression analysis would not be appropriate in solving problems that contain outliers or extreme observations. In this case, the soybean production which exceeds other production generally can be categorized as outlier, then the use of least squares method to estimate the regression parameters are less precise.

To overcome this problem, we need the robust estimation method where the value of the estimate is not much affected by small change in the data. According to Miguel [7] estimation by maximum likelihood estimator (MLE) method will produce an estimator that is the same as estimation by the least squares method, which means that MLE is not robust to against the influence of outliers. The well-known robust technique is M-estimation. M-estimation is an expansion of the MLE. In this method it is possible to eliminate some data, which in some cases we are not always be able to do especially if it is an important data, whose case is often encountered in agriculture (Susanti, *et al*, [9]). In this paper we apply M-estimation to construct a robust regression model of the soybean production in Indonesia.

MATERIAL AND METHOD

Linear Regression Model

Regression analysis is a statistical technique used to investigate the relationship between independent variables and the dependent variable. If we have dependent variable Y and independent variables $X_1, X_2, \dots, X_p$, then the linear regression model generally can be expressed as

$Y_i = \beta_0 + \beta_1 X_{i1} + \dots + \beta_p X_{ip} + \varepsilon_i, \quad i = 1, 2, \dots, n$
 (1) where $\beta_0, \beta_1, \dots$, and β_p are regression parameters and error ε is a normal distribution with mean zero and equal variance (Montgomery and Peck, [8]). Problem that arises in regression analysis is to determine the best estimator for $\beta_0, \beta_1, \dots$, and β_p .

Least Squares Method

Linear regression model can be estimated by least squares method which the basic idea is minimizing sum of square errors, namely

$$J = \sum_{i=1}^n \varepsilon_i^2 = \sum_{i=1}^n (y_i - \beta_0 - \beta_1 x_{i1} - \dots - \beta_p x_{ip})^2.$$

(2) Differentiating the J function with respect to the coefficients, $\beta_j, j = 0, 1, 2, \dots, p$, and setting the partial derivatives to 0, produces a system of $p+1$ estimating equations for the coefficients. So we obtain the estimation of linear regression model

$$\hat{y}_i = b_0 + b_1 x_{i1} + \dots + b_p x_{ip}. \quad (3)$$

Estimation of the error, called residual, is given by

$$e_i = \hat{y}_i - (b_0 + b_1 x_{i1} + \dots + b_p x_{ip}).$$

We sometime have data that do not follow the general pattern of the linear regression model, which is characterized by a relatively large residual. This data is called outlier. While there is no precise definition of an outlier, outliers are observations which do not follow the pattern of the other observations. This is not normally a problem if the outlier is simply an extreme observation drawn from the tail of a normal distribution, but if the outlier results from non-normal measurement error or some other violation of standard ordinary least squares assumptions, then it compromises the validity of the regression results if a non-robust regression technique is used.

In finding the best estimator, it is strongly influenced by the use of the method. For example, the use of least squares method will not be appropriate in resolving problems containing outlier or extreme observation because the assumption of normality can not be held (Zhang, [10]). In regression analysis,

a production which exceeds other production can be generally categorized as outlier, so the use of least squares method to estimate the parameters is not quite right (Barnett and Lewis, [1]).

A robust estimation method which is applicable to determine a regression model is M-estimation. Robust is defined as non sensitivity or toughness of small change from the assumption (Huber, [4]). In 2005 Miguel [7] estimated using the method of maximum likelihood estimator and found an estimator which is equal to the least squares method. It means that the maximum likelihood estimation is not robust against the influence of outlier. The M-estimation is an extension of maximum likelihood method and a robust estimation where its estimation value is not influenced by small change in data.

M-Estimation

Several researchers developed a method to overcome the impact of outlier if least squares method is used. The method is called the M-estimation (Montgomery and Peck, [8]). According to Li, *et al.* [5], the use of least squares method will not be appropriate in resolving problems containing outlier or extreme observation, because the assumption of normality can not be held. In the least squares method, the value of error which will further enlarge the square sums. The M-estimation anticipates this by defining a function of ε , $\rho(\varepsilon)$, which is called the Huber function. In the M-estimation, $b_0, b_1, \dots, b_p$ are respectively estimator of $\beta_0, \beta_1, \dots, \beta_p$ chosen such that $\sum \rho(\varepsilon)$ is minimum and Huber function is defined as.

$$\rho(\varepsilon) = \begin{cases} \varepsilon^2 & \text{if } -k \leq \varepsilon \leq k \\ 2k|\varepsilon| - k^2 & \text{if } \varepsilon < -k \text{ or } \varepsilon > k \end{cases}$$

where $k = 1.5 \hat{\sigma}$. To estimate σ we use $\hat{\sigma} = 1.483 \text{ MAD}$ which the MAD (Median of Absolute Deviation) is the median of the remaining absolute. To obtain the M-estimation value, a calculation algorithm is required. The following M-estimation algorithm is given by Birkes and Dodge [2].

1. Let b^0 is an initial estimates selected from the least squares method and calculate

$$\hat{y}_i = b_0^0 + b_1^0 x_{i1} + \dots + b_p^0 x_{ip}, \quad e_i^0 = y_i - \hat{y}_i^0$$

and then calculate $\hat{\sigma} = 1.483 \text{ MAD}$.

2. Cut in e_i^0 order to get the value e_i^* where $e_i^* = 1.5 \hat{\sigma}$ if $e_i^0 > 1.5 \hat{\sigma}$ and $e_i^* = -1.5 \hat{\sigma}$ if $e_i^0 < -1.5 \hat{\sigma}$.
3. Calculate $y_i^* = \hat{y}_i^0 + e_i^*$ and then find b^0 value by using least squares method. Iteration process continues until the obtained value b^0 equal to the previous iteration. The b^0 value is obtained from regression with M-estimation as in equation (3).

Significant Test for Linear Regression Model

Now we determine whether independent variables have a significant impact on dependent variables using hypothesis test. The hypothesis are

$$H_0 : \beta_j = 0 \quad \forall j, j = q+1, q+2, \dots, p$$

$$H_1 : \exists j \ni \beta_j \neq 0, j = q+1, q+2, \dots, p$$

where q is the number of independent variables included in the reduced model

$$Y_i = \beta_0 + \beta_1 X_{i1} + \dots + \beta_q X_{iq} + \varepsilon_i. \quad (4)$$

The statistic is

$$F_M = \frac{STR_{reduced} - STR_{full}}{(p-q)\hat{\lambda}} \quad \text{with} \quad \hat{\lambda} = \frac{(n/m) \sum (e_i^*)^2}{n-p-1} \quad (5)$$

where STR (*Sum of Transformed Residuals*) is the amount remaining transformed. STR_{full} and $STR_{reduced}$ are obtained from the full model and reduced model. STR_{full} algorithm can be described as follows:

1. Calculate $e_p \hat{\sigma} = 1.483 \text{ MAD}$.
2. Cut e_i values to get the values of $\rho(e_i)$, i.e. $\rho(e_i) = e_i^2$ for $1.5 \hat{\sigma} \leq e_i \leq 1.5 \hat{\sigma}$, and

$\rho(e_i) = 2k|e_i| - k^2$ for $e_i < -1,5\hat{\sigma}$ or $e_i > 1,5\hat{\sigma}$. We obtain $STR_{full}^{reduced}$ using the same calculation as STR_{full} derived from the full model. The value is obtained by cutting e_i at $-1,5\hat{\sigma}$ if $e_i < -1,5\hat{\sigma}$ and $1,5\hat{\sigma}$, if $e_i > 1,5\hat{\sigma}$, while M is the number of e_i which is not cut. The null hypothesis is rejected if the $F_M > F$ table with degree of freedom ($p-q$, $n-p-1$) and the significance level α (Birkes and Dodge, [2]).

RESULTS AND DISCUSSION

In this research, data of soybean production (Y), harvested area (X_1), productivity (X_2), and seed production (X_3) in Indonesia on 2009 are obtained from the Susenas and BPS (Heriawan, [3]). In Figure 1 we can see that the soybean production could be represented as linear function of harvested area, productivity and seed production. Based on the least squares method we know that with the p -value < 0.01 the assumption of normality does not be held and there are two outliers in data (Figure 2). So we will apply the M-estimation to find

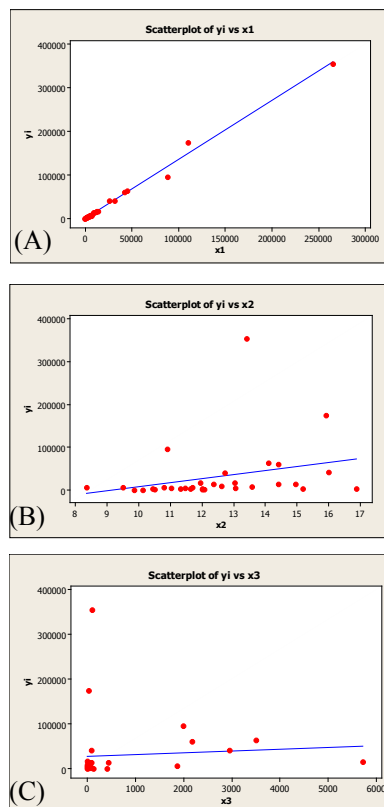


Figure 1. Plot of Y (a) Y versus X_1 (b) Y versus X_2 (c) Y versus X_3

a robust regression model. The process of calculating the M-estimation iteratively begins by determining an initial estimate of regression coefficients obtained from the least squares method, $b^0 = (-19209 ; 1.35 ; 1593 ; -0.886)$.

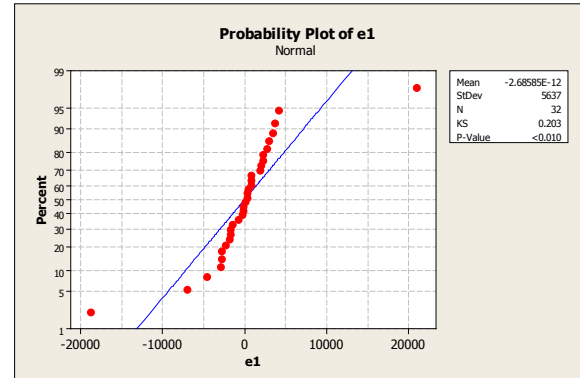


Figure 2. Probability plot of residual

Based on the M-estimation algorithm, we calculate the values $\hat{y}_i^0$ and $e_i^0 = y_i - \hat{y}_i^0$. Median of absolute deviation $|e_i^0|$ is $MAD = 1931.67$, then $1.5\hat{\sigma} = 4297.01$. There is one residual value greater than $1.5\hat{\sigma}$, then this residual is cut into 4297.01 . Meanwhile, three residuals values are smaller than $-1.5\hat{\sigma}$, then these residuals are cut into -4297.01 and the values $y_i^* = \hat{y}_i^0 + e_i^*$ are obtained. These y_i^* values are used to calculate b^0 on the next iteration.

Iteration continues until the score is same as b^0 in previous iteration. The result in each iteration is presented in Table 1. The process stops at the 23th iteration because the new value b^0 is equal to the previous: $b^0 = (-7267 ; 1.34 ; 574 ; -0.114)$. So the linear regression model is

$$\hat{y} = -7267 + 1.34x_1 + 574x_2 - 0.114x_3 \quad (6)$$

Regression model (6) shows that the increment of one hectare of harvested area and one quintal per hectare of productivity will increase 1.34 tons and 574 tons of soybean production respectively. The increment of one ton of seed production will decrease 0.114 tons of soybean production.

Table 1 Values of b^0 , MAD, cut residual, and y_i^* in each iteration

Iteration	B^0	MAD	residual cut
1	(-12969 ; 1.34 ; 1072 ; -0.318)	1931.67	3, 12, 17, 28
2	(-10816 ; 1.34 ; 886 ; -0.180)	1370.90	3, 12, 17
3	(-9628; 1.34 ; 782 ; -0.123)	983.458	3, 11, 12, 17, 25, 28
4	(-8849; 1.34 ; 714 ; -0.098)	799.607	1,3,9,11,12,13,17,25,28
5	(-8342 ; 1.34 ; 669 ; -0.094)	680.038	1,3,8,9,10,11,12,13,17,25,28
6	(-8002 ; 1.34 ; 639 ; -0.097)	579.361	1,3,8,9,10,11,12,13,17,25,28
7	(-7766 ; 1.34 ; 618 ; -0.102)	562.143	1,3,8,9,10,11,12,13,17,25,28
8	(-7607 ; 1.34 ; 604 ; -0.104)	565.323	1,3,8,9,10,11,12,13,17,25
9	(-7502 ; 1.34 ; 594 ; -0.107)	541.353	1,3,8,9,10,11,12,13,17,25
10	(-7431; 1.34 ; 588 ; -0.109)	524.654	1,3,8,9,10,11,12,13,17,25
11	(-7381 ; 1.34 ; 584 ; -0.110)	512.982	1,3,8,9,10,11,12,13,17,25
12	(-7346 ; 1.34 ; 581 ; -0.111)	504.816	1,3,8,9,10,11,12,13,17,25
13	(-7322 ; 1.34 ; 578 ; -0.112)	499.095	1,3,8,9,10,11,12,13,17,25
14	(-7305 ; 1.34 ; 577 ; -0.113)	495.082	1,3,8,9,10,11,12,13,17,25
15	(-7293 ; 1.34 ; 576 ; -0.113)	492.264	1,3,8,9,10,11,12,13,17,25
16	(-7285 ; 1.34 ; 575 ; -0.114)	490.285	1,3,8,9,10,11,12,13,17,25
17	(-7279 ; 1.34 ; 575 ; -0.114)	488.894	1,3,8,9,10,11,12,13,17,25
18	(-7275 ; 1.34 ; 574 ; -0.114)	487.917	1,3,8,9,10,11,12,13,17,25
19	(-7272 ; 1.34 ; 574 ; -0.114)	487.229	1,3,8,9,10,11,12,13,17,25
20	(-7270 ; 1.34 ; 574 ; -0.114)	486.746	1,3,8,9,10,11,12,13,17,25
21	(-7268 ; 1.34 ; 574 ; -0.114)	846.406	1,3,8,9,10,11,12,13,17,25
22	(-7267 ; 1.34 ; 574 ; -0.114)	486.167	1,3,8,9,10,11,12,13,17,25
23	(-7267 ; 1.34 ; 574 ; -0.114)	485.999	1,3,8,9,10,11,12,13,17,25

From the linear regression model (6), hypothesis testing is conducted to determine whether the harvested area, productivity and production of seed have influences on soybean production in Indonesia in 2009. There are seven possible reduced models on a linear regression model with three independent variables, i.e.

$$Y_i = \beta_0 + \varepsilon_i \quad (7)$$

$$Y_i = \beta_0 + \beta_1 X_{i1} + \varepsilon_i \quad (8)$$

$$Y_i = \beta_0 + \beta_2 X_{i2} + \varepsilon_i \quad (9)$$

$$Y_i = \beta_0 + \beta_3 X_{i3} + \varepsilon_i \quad (10)$$

$$Y_i = \beta_0 + \beta_1 X_{i1} + \beta_2 X_{i2} + \varepsilon_i$$

$$(11)$$

$$Y_i = \beta_0 + \beta_1 X_{i1} + \beta_3 X_{i3} + \varepsilon_i$$

$$(12)$$

$$Y_i = \beta_0 + \beta_2 X_{i2} + \beta_3 X_{i3} + \varepsilon_i$$

$$(13)$$

Before testing hypothesis for each reduced model, STR_{full} value is calculated from the regression model (6). Based on the calculation algorithm of STR_{full} , the first step is to find the errors of e_i from the regression model (6). We obtain $MAD = 485.999$ and $k = 1.5 \hat{\sigma} = 1081.11$. There are 22 residuals in interval $-1.5 \hat{\sigma} \leq e_i \leq 1.5 \hat{\sigma}$ and 10 others are not in interval $-1.5 \hat{\sigma} \leq e_i \leq 1.5 \hat{\sigma}$, so we obtain $r(e_i) = e_i^2 + 2k|e_i| - k^2 2k|e_i| - k^2$. Sum of $\rho(e_i)$ is $STR_{full} = 5,318,967 + 130,461,409 =$

135,780,376. To obtain value of $\hat{\lambda}$, we count $\sum (e_i^*)^2 = 17,014,957$ and we have

$$\hat{\lambda} = \frac{(n/m)\sum(e^*)^2}{n-p-1} = \frac{(32/22)17014957}{32-3-1} = 883,894$$

Furthermore we count $STR_{reduced}$ from the reduced models (7) – (13). This can be calculated if we estimate all reduced models, same as the full model but the value of $\hat{\sigma}$ is fixed ($1.5 \hat{\sigma} = 1081.11$). The $STR_{reduced}$ of models (7) – (13) are presented in the Table 2.

Table 2 $STR_{reduced}$

Model	$STR_{reduced}$
(7)	1,894,025,520
(8)	155,507,503
(9)	1,811,996,138
(10)	1,764,778,707
(11)	136,212,934
(12)	155,427,069
(13)	1,687,062,464

Now we will determine the independent variable which influences to the dependent variable. The hypothesis for the reduced model (7) is $H_0: \beta_i = 0 \forall i, i = 1, 2, 3$ (harvested area, productivity and production of seed do not influence the production of soybean significantly) $H_1: \exists i \ni \beta_i \neq 0, i = 1, 2, 3$ (at least one of harvested area, productivity or production of seed influences the production of soybean significantly) From (5) we have the test statistic

$$F_M = \frac{STR_{reduced} - STR_{full}}{(p-q)\hat{\lambda}} = \frac{1894025520-135780376}{3(883894)} = 663.068$$

The $F_{0.05;3,28} = 2.95$ and $F_M > 663.068$, then H_0 is rejected. So, at least one of harvested area, productivity or production of seed influences the production of soybean significantly. The hypothesis for the reduced

model (11) is $H_0: \beta_3 = 0$ (production of seed does not influence the production of soybean significantly) $H_1: \beta_3 \neq 0$ (production of seed influences the production of soybean significantly) The test statistic is

$$F_M = \frac{STR_{reduced} - STR_{full}}{(p-q)\hat{\lambda}} = \frac{136212934-135780376}{1(883894)} = 0.489$$

The $F_{0.05;2,28} = 4.20$ and $F_M < 0.489$, then H_0 is not rejected. So, the production of seed does not influence the production of soybean significantly. The significance test result of reduced models is shown in Table 3.

From the all hypothesis tests of the reduced models we conclude that the harvested area and the productivity influence soybean production significantly, but the production of seed does not influence soybean production significantly. So regression model (6) shows that the increment of one hectare of harvested area and one quintal per hectare of productivity will increase 1.34 tons and 574 tons of soybean production respectively. This model has $R^2 = 100\%$, means all the total variation could be explained by the independent variables.

CONCLUSIONS

As the result we conclude that the robust regression model for production of soybean in Indonesia using M-estimation is $\hat{y} = -7,267 + 1.34x_1 + 574x_2 - 0.114x_3$. The increment of one hectare of harvested area and one quintal per hectare of productivity will increase 1.34 tons and 574 tons of soybean productivity respectively.

The model has $R^2 = 100\%$, nevertheless the convergence in estimating the parameter needs a long time. By using estimation-M the convergence was reached in 23th iteration. Moreover this model could not reduce all outliers, so the use of other robust estimation method, for example MM-estimation or Least Trimmed Squares (LTS) estimation, could be considered in order to get the better model.

Table 3 The significance test result of reduced models

Reduced Model	Hypothesis	F_M and F table	Conclusion
$Y_i = \beta_0 + \varepsilon_i$	$H_0: \beta_i = 0 \forall i \ i = 1,2,3$ $H_1: \exists i \ni \beta_i \neq 0 \ i = 1,2,3$	$F_M = 663.068$ $F \text{ table} = 2.95$	H_0 is rejected $\exists i \ni \beta_i \neq 0 \ i = 1,2,3$
$Y_i = \beta_0 + \beta_1 X_{i1} + \varepsilon_i$	$H_0: \beta_i = 0 \forall i \ i = 2,3$ $H_1: \exists i \ni \beta_i \neq 0 \ i = 2,3$	$F_M = 11.1592$ $F \text{ table} = 3.34$	H_0 is rejected $\exists i \ni \beta_i \neq 0 \ i = 2,3$
$Y_i = \beta_0 + \beta_2 X_{i2} + \varepsilon_i$	$H_0: \beta_i = 0 \forall i \ i = 1,3$ $H_1: \exists i \ni \beta_i \neq 0 \ i = 1,3$	$F_M = 948.20$ $F \text{ table} = 3.34$	H_0 is rejected $\exists i \ni \beta_i \neq 0 \ i = 1,3$
$Y_i = \beta_0 + \beta_3 X_{i3} + \varepsilon_i$	$H_0: \beta_i = 0 \forall i \ i = 1,2$ $H_1: \exists i \ni \beta_i \neq 0 \ i = 1,2$	$F_M = 921.5$ $F \text{ table} = 3.34$	H_0 is rejected $\exists i \ni \beta_i \neq 0 \ i = 1,2$
$Y_i = \beta_0 + \beta_1 X_{i1} + \beta_2 X_{i2} + \varepsilon_i$	$H_0: \beta_3 = 0$ $H_1: \beta_3 \neq 0$	$F_M = 0.489$ $F \text{ table} = 4.20$	H_0 is not rejected $\beta_3 = 0$
$Y_i = \beta_0 + \beta_1 X_{i1} + \beta_3 X_{i3} + \varepsilon_i$	$H_0: \beta_2 = 0$ $H_1: \beta_2 \neq 0$	$F_M = 22.23$ $F \text{ table} = 4.20$	H_0 is rejected $\beta_2 \neq 0$
$Y_i = \beta_0 + \beta_2 X_{i2} + \beta_3 X_{i3} + \varepsilon_i$	$H_0: \beta_1 = 0$ $H_1: \beta_1 \neq 0$	$F_M = 1,755.05$ $F \text{ table} = 4.20$	H_0 is rejected $\beta_1 \neq 0$

ACKNOWLEDGEMENT

The authors would like to thank to Sebelas Maret University and The Indonesian Ministry of National Education for giving the financial support to finish our work under Small Business Innovation Research No. 05/UN279/PI/2011.

REFERENCES

- Barnett, J. & Lewis.1978. *Outliers in Statistical Data*, John Wiley & Sons Inc., New York.
- Birkes, D. & Dodge, Y., 1993. *Alternative Methods of Regression*, John Wiley & Sons Inc., New York.
- Heriawan, R., 2010. *Statistik Indonesia*, Badan Pusat Statistik Indonesia, Katalog BPS: 1101001, ISSN: 0126-2912.
- Huber, P.J., 1981. *Robust Statistic*, John Wiley & Sons. Inc., New York.
- Li, S.Z., Wang, H., & Soh, W.Y.C., 1998. Robust Estimation of Rotation Angle from Image Sequences Using the Annelling M Estimator, *Jurnal of Mathematical Imaging and Vision*, Vol 8. No 2, pp. 181-192.
- Marwoto & Suharsono. 2008. Strategi dan Kom- ponen Teknologi Pengendalian Ulat Grayak (Spodoptera Litura Fabricius) pada Tanaman Kedelai, *Jurnal Litbang Pertanian*, 27(4).
- Miguel, A.A., 2005. Convergence of the Optimal M-estimator over a Parametric Family of M-estimators, *An Official Journal of the Spanish Society of Statistics and Operations Research*, vol. 14, No 1, pp. 281-315.
- Montgomery, D.C. & Peck, 2006. E.A., *Intro- duction to Linear Regression Analysis*, John Wiley & Sons Inc., New York.

- Susanti, Y., Pratiwi, H. & Liana, T., 2009. *Aplication of M-estimation to Predict Paddy Production in Indonesia*, Presented at IndoMS International Conference on Mathematics and Its Applications IICMA, Yogyakarta
- Zhang, Z., 1996. *Robust Estimations*, <http://www.sop.Inria.fr/robust/personel/zhang/Publis/tutorial-Estinu/modezo.html>.