

BAB III METODE PENELITIAN

3.1 Tempat dan Waktu Penelitian

Tempat dalam penelitian ini dilaksanakan di Provinsi Nusa Tenggara Timur secara keseluruhan, Tingkat Pendidikan dan IPM melalui Badan Pusat Statistik Provinsi Nusa Tenggara Timur yang berlokasi di Jln R. Suprpto No 5 Oebobo Kota Kupang. Waktu penelitian selama Enam Bulan yaitu Bulan Januari sampai Juni 2019.

3.2 Definisi Operasional Variabel

Dalam penelitian ini terdiri dari dua variabel yaitu variabel terikat dan variabel bebas. Variabel terikat(Y) dalam penelitian ini adalah Indeks Pembangunan Manusia, variabel bebas(X1) adalah Tingkat Pendidikan, variabel bebas (X2) adalah Pertumbuhan Ekonomi.

1. IPM mengukur capaian pembangunan manusia berbasis sejumlah komponen dasar kualitas hidup. Sebagai ukuran kualitas hidup, IPM dibangun melalui pendekatan tiga dimensi dasar. Dimensi tersebut mencakup umur panjang dan sehat, pengetahuan dan kehidupan yang layak. Ketiga dimensi tersebut memiliki pengertian sangat luas karena terkait banyak faktor. Untuk mengukur dimensi kesehatan, digunakan angka umur harapan hidup. Untuk mengukur dimensi pengetahuan digunakan gabungan indikator angka melek huruf dan rata-rata lama sekolah.

2. Tingkat pendidikan adalah tahapan pendidikan yang ditetapkan berdasarkan tingkat perkembangan peserta didik, tujuan yang akan dicapai dan kemauan yang dikembangkan. Tingkat pendidikan berpengaruh terhadap perubahan sikap dan perilaku hidup sehat. Tingkat pendidikan yang lebih tinggi akan memudahkan seseorang atau masyarakat untuk menyerap informasi dan mengimplementasikannya dalam perilaku dan gaya hidup sehari-hari.

3. Pertumbuhan ekonomi adalah peningkatan output masyarakat yang disebabkan oleh semakin banyaknya jumlah faktor produksi yang digunakan dalam proses produksi, tanpa adanya perubahan “teknologi” produksi itu sendiri, misalnya kenaikan output yang disebabkan oleh pertumbuhan stok modal ataupun penambahan faktor-faktor produksi tanpa adanya perubahan pada teknologi produksi yang lama.

3.3 Jenis Dan Sumber Data

3.3.1 Jenis Data

3.3.1.1 Data Kualitatif

Data kualitatif adalah data yang bukan berbentuk angka melainkan dalam kata verbal.

3.3.1.2 Data Kuantitatif

Data Kuantitatif adalah jenis data yang dapat diukur atau dihitung secara langsung, yang berupa informasi atau penjelasan yang dinyatakan dengan bilangan atau berbentuk angka (Sugiyono, 2010). Yang menjadi data kuantitatif dari penelitian ini adalah Data Tingkat Pendidikan Provinsi NTT Tahun 2013-

2017, Data Pertumbuhan Ekonomi Provinsi NTT Tahun 2013-2017, dan Data IPM Provinsi NTT Tahun 2013-2017.

3.3.2 Sumber Data

3.3.2.1 Data Primer

Data Primer merupakan data yang diperoleh secara langsung dari obyek yang diteliti. Menurut (Sugiyono,2010;137) yang menyatakan bahwa : “sumber primer adalah sumber data langsung memberikan data kepada pengumpul data.

3.3.2.2 Data Sekunder

Menurut (Sugiyono,2010;137) data sekunder adalah “ sumber data yang tidak langsung memberikan data kepada pengumpul data misalnya lewat orang lain atau lewat dokumen.

Data sekunder antara lain disajikan dalam bentuk data data, tabel-tabel, diagram-diagram, atau mengenai topik penelitian. Data ini merupakan data yang berhubungan secara langsung dengan penelitian yang dilaksanakan dan bersumber dari dari Badan Pusat Statistik Provinsi NTT. Yang menjadi data sekunder dalam penelitian ini adalah data Tingkat Pendidikan Provinsi NTT Tahun 2013-2017, Data Pertumbuhan Ekonomi Provinsi NTT Tahun 2013-2017, dan Data IPM Provinsi NTT Tahun 2013-2017.

3.4 Populasi dan Sampel

3.4.1 Populasi

Menurut (Sugiyono,2010;215) populasi diartikan sebagai wilayah generalisasi yang terdiri atas obyek/subyek yang mempunyai kualitas dan

karakteristik tertentu yang ditetapkan oleh peneliti untuk dipelajari dan kemudian ditarik kesimpulannya.

Berdasarkan pengertian populasi diatas, maka yang akan disajikan populasi dalam penelitian ini adalah data Tingkat Pendidikan Provinsi NTT , Data Pertumbuhan Ekonomi Provinsi NTT , dan Data IPM Provinsi NTT .

3.4.2 Sampel

Menurut (Sugiyono,2010;81) sampel adalah bagian dari jumlah dan karakteristik yang dimiliki oleh populasi tersebut.

Berdasarkan pengertian sampel diatas, maka yang akan dijadikan sampel dalam penelitian ini adalah data Tingkat Pendidikan Provinsi NTT Tahun 2013-2017, Data Pertumbuhan Ekonomi Provinsi NTT Kupang Tahun 2013-2017, dan Data IPM Provinsi NTT Tahun 2013-2017.

.

3.5 Metode Pengumpulan Data

Metode pengumpulan data yang digunakan dalam penelitian ini adalah menggunakan metode Studi Kepustakaan, pengambilan data yang bersifat teori yang kemudian digunakan sebagai literatur penunjang guna mendukung penelitian yang dilakukan. Data yang diperoleh dari buku-buku sumber yang dapat dijadikan acuan yang ada kaitannya dengan masalah yang diteliti.

Data yang diperoleh adalah data sekunder, yaitu data yang didapatkan tidak secara langsung dari objek atau subjek penelitian, data diambil dari Badan pusat Statistik Provinsi Nusa Tenggara Timur. Data yang diambil bersifat Data Kuantitatif yaitu data yang berbentuk angka pasti. Data yang dikumpulkan

merupakan data berkala/time series, yaitu data yang dikumpulkan dari waktu ke waktu untuk menggambarkan perkembangan atau kecenderungan keadaan/peristiwa/kegiatan khususnya untuk Tingkat Pendidikan, Pertumbuhan Ekonomi dan Indeks Pembangunan Manusia

3.6 Alat Analisis

3.6.1 Model Regresi Data Panel

Data panel adalah gabungan antara data runtun waktu (*time series*) dan data silang (*cross section*). Menurut Widarjono dalam Basuki dan Prawoto (2017) penggunaan data panel dalam sebuah observasi mempunyai beberapa keuntungan yang diperoleh; pertama, data panel merupakan gabungan dari dua data *time series* dan *cross section* mampu menyediakan data lebih banyak sehingga akan lebih menghasilkan *degree of freedom* yang lebih besar. Kedua, menggabungkan informasi dari data *time series* dan *cross section* dapat mengatasi masalah yang timbul ketika ada masalah penghilangan variable.

Analisis data panel adalah suatu metode mengenai gabungan dari data antar waktu (*time series*) dengan data antar individu (*cross section*). Untuk menggambarkan data panel secara singkat, misalkan pada data *cross section*, nilai dari satu variabel atau lebih dikumpulkan untuk beberapa unit sampel pada suatu waktu. Dalam data panel, unit *cross section* yang sama di survey dalam beberapa waktu (Gujarati, 2003). Regresi dengan menggunakan data panel memberikan beberapa keunggulan dibandingkan dengan pendekatan standar *cross section* dan *time series*, diantaranya:

1. Data panel dapat memberikan peneliti jumlah pengamatan yang besar, meningkatkan derajat kebebasan (*degree of freedom*), data memiliki variabilitas yang besar dan mengurangi kolinieritas antara variabel penjelas dimana dapat menghasilkan estimasi ekonometri yang efisien.
2. Data panel dapat memberikan informasi lebih banyak yang tidak dapat diberikan hanya oleh data *cross section* atau *time series* saja.
3. Data panel dapat memberikan penyelesaian yang lebih baik dalam inferensi perubahan dinamis dibandingkan data *cross section*.

Keunggulan dimiliki model data panel tersebut, ada beberapa permasalahan yang muncul dalam pemanfaatan data jenis panel yaitu permasalahan autokorelasi dan heterokedastisitas. Sementara itu ada permasalahan baru yang muncul seperti korelasi silang (*cross-correlation*) antar unit individu pada periode yang sama.

3.6.1.1 Metode Estimasi Model Regresi Panel

Penelitian ini menggunakan analisis data panel untuk mengetahui pengaruh dari variabel independen yakni Tingkat Pendidikan (X_1), Pertumbuhan Ekonomi (X_2), terhadap variabel dependen yaitu Indeks Pembangunan Manusia (Y). Data diolah menggunakan Eviews 10. Model fungsi yang digunakan untuk mengetahui Indeks Pembangunan Manusia yaitu:

$$Y = f\{X_1, X_2, \}$$

Dalam bentuk persamaan menjadi:

$$Y_{it} = \lambda_0 + \lambda_1 X_{1it} + \lambda_2 X_{2it} +$$

dimana; Y = Indeks Pembangunan Manusia

X_1 = Tingkat Pendidikan

X_2 = Pertumbuhan Ekonomi

I = *cross section*.

T = *time series*.

λ_0 = konstanta.

$\lambda_1, \dots, \lambda_4$ = koefisien.

U = *error term*.

Dalam metoda estimasi model regresi dengan menggunakan data panel dapat dilakukan melalui tiga pendekatan, antara lain:

1. *Common Effect Model*

Merupakan pendekatan model data panel yang paling sederhana karena hanya mengombinasikan data *time series* dan data *cross section*. Pada model ini tidak diperhatikan dimensi waktu maupun individu, sehingga diasumsikan bahwa perilaku data perusahaan sama dalam berbagai kurun waktu. Metode ini bisa menggunakan pendekatan *Ordinary Least Square* (OLS) atau teknik kuadrat terkecil untuk mengestimasi model data panel.

2. *Fixed Effect Model*

Model ini mengasumsikan bahwa perbedaan antar individu dapat diakomodasi dari perbedaan intersepnya. Untuk mengestimasi data panel model *Fixed Effects* menggunakan teknik *variable dummy* untuk menangkap perbedaan intersep antar perusahaan, perbedaan intersep bisa terjadi karena perbedaan budaya kerja, manajerial, dan insentif. Namun demikian, sloponya sama antar perusahaan. Model estimasi ini juga sering disebut dengan tekni *Least Dummy Variable* (LSDV).

3. *Random Effect Model*

Model ini akan mengestimasi data panel di mana variable gangguan mungkin saling berhubungan antar waktu dan antar individu. Pada model *Random Effect* perbedaan intersep diakomodasi oleh *error terms* masing-masing perusahaan. Keuntungan menggunakan model *random effect* yakni menghilangkan heteroskedastisitas. Model ini uga disebut dengan *Error Component Model* (ECM) atau teknik *Generalized Least Square* (GLS).

Untuk memilih model yang paling tepat digunakan ddalam mengelola data panel, terdapat beberapa pengujian yang dapat dilakukan, yakni:

1. Uji Chow

Chow test yakni pengujian untuk menentukan model *Fixed Effect* atau *Random effect* yang paling tepat digunakan dalam mengestimasi data panel.

2. Uji Hausman

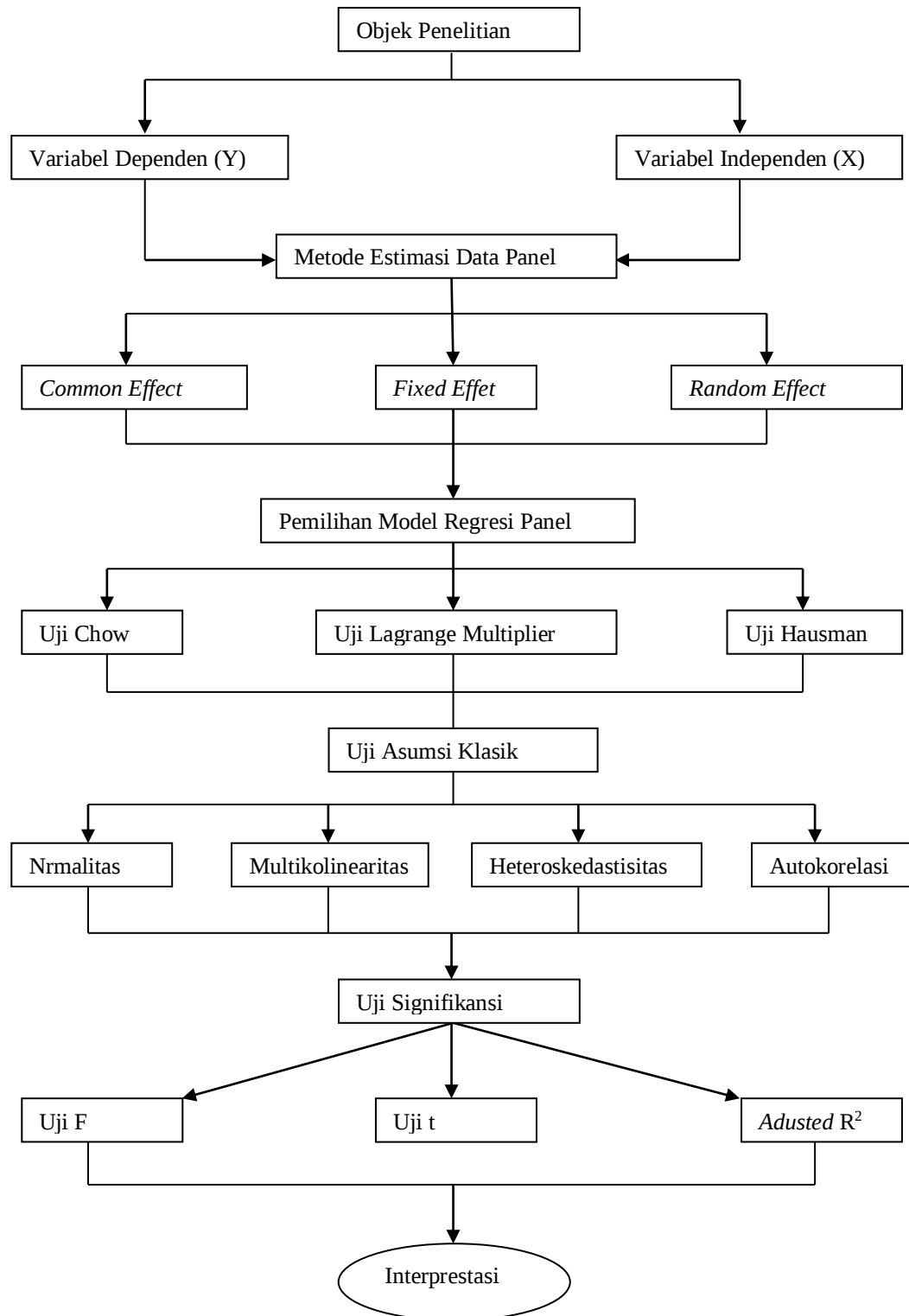
Hausman test adalah pengujian statistic untuk memilih apakah model *Fixed Effect* atau *Random Effect* yang paling teat digunakan.

3. Uji Lagrange Multiplier

Untuk mengetahui apakah model *Random Effect* lebih baik daripada metode *Common Effect* (OLS) digunakan uji Lagrange Multiplier (LM)

Model Penelitian regresi dengan data panel mengikuti pola kerangka pikir sebagaimana yang diemukakan Agus Tri Basuki dan Nano Prawoto (2017) seperti tercantum pada Gambar 3.1

Gambar 3.1
Kerangka Pemikiran Model Regresi Data Panel



1. *Common Effects Model*

Model *Common Effects* merupakan pendekatan model data panel yang paling sederhana. Model ini tidak memperhatikan dimensi waktu maupun individu, sehingga diasumsikan bahwa perilaku antar individu sama dalam berbagai kurun waktu. Model ini hanya mengombinasikan data *time series* dan *cross section* dalam bentuk *pool*, mengestimasiya menggunakan pendekatan kuadrat terkecil/*pooled least square*.

Adapun persamaan regresi dalam model *common effects* dapat ditulis sebagai berikut:

$$Y_{it} = \alpha + X_{it}\beta + \varepsilon_{it}$$

di mana:

i = Kota Kupang, Rote Ndao, ..., Malaka

t = 2013, 2014, 2015, 2016, 2017

dimana i menunjukkan *cross section* (individu) dan t menunjukkan periode waktunya. Dengan asumsi komponen *error* dalam pengolahan kuadrat terkecil biasa, proses estimasi secara terpisah untuk setiap unit *cross section* dapat dilakukan.

2. *Fixed Effects Model*

Model *Fixed Effects* mengasumsikan bahwa terdapat efek yang berbeda antar individu. Perbedaan itu dapat diakomodasi melalui perbedaan pada intersepnya. Oleh karena itu, dalam model *fixed effects*, setiap individu merupakan parameter yang tidak diketahui dan akan diestimasi dengan menggunakan teknik variable *dummy* yang dapat ditulis sebagai berikut:

$$Y_{it} = \alpha + i\alpha_{it} + X'_{it}\beta + \varepsilon_{it}$$

$$\begin{bmatrix} y_1 \\ y_2 \\ y_n \end{bmatrix} = \begin{bmatrix} \alpha \\ \alpha \\ \alpha \end{bmatrix} + \begin{bmatrix} i & 0 & 0 \\ 0 & i & 0 \\ 0 & 0 & i \end{bmatrix} \begin{bmatrix} \alpha_1 \\ \alpha_2 \\ \alpha_n \end{bmatrix} + \begin{bmatrix} x_{11} & x_{21} & x_{p1} \\ x_{12} & x_{22} & x_{p2} \\ x_{1n} & x_{2n} & x_{pn} \end{bmatrix} \begin{bmatrix} \beta_1 \\ \beta_2 \\ \beta_n \end{bmatrix} + \begin{bmatrix} \varepsilon_1 \\ \varepsilon_2 \\ \varepsilon_n \end{bmatrix}$$

Teknik seperti di atas dinamakan *Least Square Dummy Variabel* (LSDV). Selain diterapkan untuk efek tiap individu, LSDV ini juga dapat mengakomodasi efek waktu yang bersifat sistemik. Hal ini dapat dilakukan melalui penambahan variabel *dummy* waktu di dalam model.

3. *Random Effects Model*

Berbeda dengan *fixed effects model*, efek spesifik dari masing-masing individu diperlakukan sebagai bagian dari komponen *error* yang bersifat acak dan tidak berkorelasi dengan variabel penjelas yang teramati, model seperti ini dinamakan *Random Effects Model* (REM). Model seperti ini sering disebut juga dengan *Error Component Model* (ECM). Dengan demikian, persamaan *model effects* dapat dituliskan sebagai berikut:

$$Y_{it} = \alpha + Y'_{it}\beta + w_{it}$$

$$i = \text{Kota Kupang, Rote Ndao, ..., Malaka}$$

$$t = 2013, 2014, 2015, 2016, 2017$$

di mana:

$$w_{it} = \varepsilon_{it} + u_i; E(w_{it}) = 0; E(w_{it}^2) = \sigma^2 + \sigma_u^2;$$

$$E(w_{it} \cdot w_{jt-1}) = 0; i \neq j; E(u_i \cdot \varepsilon_{it}) = 0;$$

$$E(\varepsilon_i \cdot \varepsilon_{is}) = E(\varepsilon_{it} \cdot \varepsilon_{jt}) = E(\varepsilon_{it} \cdot \varepsilon_{js}) = 0$$

Meskipun komponen error w_t bersifat homoskedastik, nyatanya terdapat korelasi antara w_t dan w_{it-s} (*equicorrelation*), yakni:

$$\text{Corr}(w_{it}, w_{i(t-1)}) = \alpha_u^2 / (\alpha^2 + \alpha_u^2)$$

Karena itu, metode OLS tidak bisa digunakan untuk mendapatkan estimator yang efisien bagi model *random effects*. Metode yang tepat untuk mengestimasi model *random effects* adalah *Generalized Least Squares* (GLS) dengan asumsi homoskedatik dan tidak ada *cross-sectional correlation*.

Judge (1980) dalam Fadly (2011), menyatakan ada perbedaan mendasar untuk menentukan pilihan antara FEM (*Fixed Effect Model*) dan ECM (antara lain sebagai berikut (Gujarati, 2004):

1. Jika T (jumlah data *time series*) besar dan N (jumlah unit *cross section*) kecil, perbedaan antara FEM dan ECM adalah sangat tipis. Oleh karena itu, dapat dilakukan perhitungan secara konvensional. Pada keadaan ini, FEM mungkin lebih disukai.
2. Ketika N besar dan T kecil, estimasi diperoleh dengan dua metode dapat berbeda secara signifikan. Pada ECM, di mana adalah komponen random *cross-section* dan pada FEM, ditetapkan dan tidak acak. Jika sangat yakin dan percaya bahwa individu, atau unit *cross-section* sampel adalah tidak acak, maka FEM lebih cocok digunakan. Jika unit *cross-section* sampel adalah random/acak, maka ECM lebih cocok digunakan.
3. Komponen error individu dan satu atau lebih regresor berkorelasi, estimator yang berasal dari ECM adalah bias, sedangkan yang berasal dari FEM adalah *unbiased*.
4. Jika N besar dan T kecil, serta jika asumsi untuk ECM terpenuhi, maka estimator ECM lebih efisien dibanding estimator FEM.

Keunggulan regresi data anel menurut Wibisono (2005) antara lain:

1. Panel data mampu memperhitungkan heterogenitas individu secara eksplisit dengan mengizinkan variabel spesifik individu.
2. Kemampuan mengontrol heterogenitas ini selanjutnya menjadikan data panel data digunakan untuk menguji dan membangun model perilaku lebih kompleks.
3. Data panel didasarkan diri pada observasi *cross-section* yang berulang-ulang (*time series*), sehingga metode data panel cocok digunakan sebagai *study of dynamic adjustment*.
4. Tingginya jumlah observasi memiliki implikasi pada data yang lebih informatif, lebih variatif, dan kolinearitas (multikoliner) antara data semakin berkurang, dan derajat kebebasan (*degree of freedom/df*) lebih tinggi sehingga dapat diperoleh hasil estimasi yang lebih efisien.
5. Data panel dapat digunakan untuk mempelajari model-model perilaku yang kompleks.
6. Data panel dapat digunakan untuk meminimalkan bias yang mungkin ditimbulkan oleh agregasi data individu.

Secara formal, ada tiga prosedur pengujian yang akan digunakan, yaitu uji statistik F yang digunakan untuk memilih antara:

1. Model *common effects* atau *fixed effects*;
2. Uji *Lagrange Multiplier* (LM) yang digunakan untuk memilih antara model *common effects* atau model *random effects*;

3. Uji Hausman yang digunakan untuk memilih antara model *fixed effects* atau model *random effects*.

3.6.2 Statistik Inferensial

3.6.2.1 Uji Asumsi Klasik

Dalam penelitian ini, peneliti akan melakukan uji statistik regresi dalam mempelajari hubungan yang ada diantara variabel-variabel tidak bebas jika variabel bebasnya diketahui atau sebaliknya. Dalam prakteknya ada empat uji asumsi klasik yang paling sering digunakan yaitu:

3.6.2.1.1 Normalitas

(Widarjono,2013) Uji signifikan pengaruh variabel independen terhadap variabel dependen melalui uji t hanya akan valid jika residual yang kita dapatkan mempunyai distribusi normal. Ada beberapa metode yang bisa digunakan untuk apakah residual mempunyai distribusi normal atau tidak.

1. Histogram Residual. Histogram residual adalah merupakan metode yang paling sederhana digunakan untuk mengetahui apakah bentuk dari Probability Distribution Function dari variabel random berbentuk distribusi normal atau tidak. Jika histogram residual menyeruapi grafik distribusi normal maka bisa dikatakan bahwa residual mempunyai distribusi normal. Bentuk grafik distribusi normal ini menyerupai lonceng seperti distribusi t sebelumnya dimana jika grafik distribusi normal tersebut dibagi dua akan mempunyai bagian yang sama.
2. Uji- Jarque-Bere. Uji normalitas residual metode OLS secara formal dapat dideteksi dengan metode yang dikembangkan oleh Jarque- Bere. Metode ini

didasarkan pada sampel besar yang diasumsikan bersifat asympfotic. Uji statistic ini menggunakan perhitugn skewness dari kurtosis sebagai berikut:

$$JB = n\left(\frac{s}{6} + \frac{(k-3)}{24}\right)$$

Dimana S= Koefisien Skewness dan K= Koefisien Kurtosis.

Jika suatu variabel didistribusikan secara normal maka nilai koefisien S=0 dan K=3. Oleh karena itu, jika residual terdistribusikan secara normal maka diharapkan nilai statistic JB akan sama dengan nol. Nilai statistik Jb ini didasarkan pada distribusi Chi Squares dengan derajat kebebasan(df)=2. Jika nilai probabilitas p dari statistik JB besar atau dengan kata lain jika nilai statistik dari JB ini tidak signifikan maka kita gagal menolak hipotesis residual mempunyai distribusi normal karena nilai statistik Jb mendekati nol. Sebaliknya jika nilai probabilitas p dari statistik JB kecil atau signifikan maka kita menolak hipotesis bahwa residual mempunyai distribusi normal karena nilai statistik JB tidak sama dengan Nol.

3.6.2.1.2 Multikolinearitas

(Widarjono, 2013) Uji Multikolinearitas bertujuan untuk mengetahui ada hubungan linier antara variabel independen dalam suatu regresi.

Deteksi Multikolinearitas atau korelasi yang tinggi antar variabel independen:

1. Nilai R^2 Tinggi tetapi hanya sedikit variabel independen yang signifikan.
Salah satu cirri adanya multikolinieritas adalah model mempunyai koefisien determinasi yang sangat tinggi(R^2) katakanlah diatas 0,8 tetapi hanya sedikit variabel independen yang signifikan mempengaruhi variabel

dependen melalui uji t^2 . Namun berdasarkan uji F secara statistic signifikan yang berarti semua variabel independen secara bersama sama mempengaruhi variabel dependen. Dalam hal ini terjadi suatu kontradiktif dimana berdasarkan uji t secara individual variabel independen tidak berpengaruh terhadap variabel dependen, namun secara bersama sama variabel independen mempengaruhi variabel dependen.

2. Korelasi Parsial antarvariabel independen. Sebagaimana yang sudah dijelaskan bahwa multikolinieritas adalah hubungan linier antara variabel independen di dalam regresi. Oleh karena itu kenapa kita tidak mendeteksi multikolinieritas engan menguji koefisien korelasi (r) antarvariabel independen. Sebagian aturan main yang kasar (rule of thumb), jika koefisien korelasi cukup tinggi katakanlah diatas 0,85 maka akan diduga ada multikolinieritas dalam model. Sebaliknya jika koefisien korelasi relatif rendah maka kita duga model tidak mengandung unsur multikolinieritas. Namun deteksi dengan menggunakan metode ini diperlukan hati-hati. Masalah ini timbul terutama pada data time series dimana korelasi antarvariabel independen cukup tinggi. Korelasi yang tinggi ini terjadi karena kedua data mengandung unsur tren yang sama yaitu data naik dan turun secara bersamaan.
3. Regresi Auxiliary. Pada uji korelasi kita menguji multikolinieritas hanya dengan melihat hubungan secara individual antara suatu variabel independen dengan satu variabel independen yang lain. Tetapi multikolinieritas bisa juga muncul karena satu atau lebih variabel

independen merupakan kombinasi linier dengan variabel independen lainnya. Untuk mengetahui apakah variabel independen X_1 berhubungan dengan variabel independen X_2 adalah dengan regresi yang kita gunakan yang disebut regresi auxiliary. Setiap koefisien determinasi (R^2) dari regresi auxiliary ini kita gunakan untuk menghitung distribusi F dan kemudian digunakan untuk mengevaluasi apakah model mengandung multikolinearitas atau tidak. Adapun formula menghitung nilai F yaitu:

$$F_i = \frac{R^2 X_1 X_2 X_3 \dots X_K}{(K-1)}$$

Didalam persamaan ini menunjukkan jumlah variabel independen termasuk konstanta dan $R^2 X_1 X_2 X_3 \dots X_K$ adalah koefisien determinasi setiap variabel independen X_i dengan sisa variabel X yang lain sedangkan nilai kritis dari distribusi F didasarkan pada derajat kebebasan $k-1$ dan $n-k$. Keputusan ada tidaknya unsur multikolinieritas dalam model ini sebagaimana biasanya adalah dengan membandingkan nilai F hitung dengan F Kritis. Jika nilai F Hitung lebih besar dari nilai hitung F kritis dengan tingkat signifikan α dan derajat kebebasan tertentu maka dapat disimpulkan model mengandung unsur multikolinieritas yakni terdapat hubungan linier antara satu variabel dan variabel lainnya. sebaliknya jika nilai F hitung lebih kecil dari nilai F kritis maka tidak ada hubungan linier antara suatu variabel X dan variabel X lainnya.

4. Metode Deteksi Klien. Dengan mendapatkan determinasinya $R^2 X_1 X_2 X_3 \dots X_K$ Klien menyarankan untuk mendeteksi masalah multikolinieritas dengan hanya membandingkan koefisien determinasi

auxiliary dengan koefisien determinasi (R^2) model regresi aslinya yaitu Y dengan variabel independen X^4 . Jika $R^2 X_1 X_2 X_3 \dots X_K$ lebih besar dari R^2 maka model mengandung unsure multikolinieritas antara variabel independen begitupun sebaliknya.

5. Variance Inflation Factor dan Tolerance. Jika kita mempunyai sejumlah k variabel independen tidak termasuk konstanta didalam sebuah model maka varian dari koefisien regresi determinasi regresi parsial .

$$Var(\beta_j) = \left(\frac{\alpha^2}{\sum x_j^2} \right) \left(\frac{1}{1 - r^2} \right)$$

Atau dapat ditulis

$$Var(\beta_j) = \left(\frac{\alpha^2}{\sum x_j^2} \right) VIF_j$$

Dimana R_j^2 merupakan R^2 yang diperoleh dari regresi Auxiliary antara variabel independen dengan variabel independen lainnya. Sedangkan VIF adalah Variance Inflation Factor. Ketika R_j^2 mendekati satu atau dengan kata lain ada kolineritas antara variabel independen maka VIF akan naik mendekati tak terhingga jika nilai $R_j^2 = 1$. Dengan demikian kita bisa menggunakan VIF untuk mendeteksi masalah Multikolinieritas di dalam sebuah regresi linier berganda. Jika nilai VIF semakin membesar maka diduga ada multikolinieritas.

3.6.3.1.3 Heteroskedasitas

(Widarjono, 2013) Sifat heteroskedasitas yaitu jika variabel gangguan tidak mempunyai rata-rata nol maka tidak mempengaruhi slope, hanya akan mempengaruhi slope, hanya akan mempengaruhi intersep. Hal ini tidak membawa

konsekuensi serius karena perhatian dalam aplikasi ekonometrika bukan pada intersep tetapi pada slope. Deteksi Heteroskedastisitas:

1. Metode informal. Cara yang paling tepat dan dapat digunakan untuk menguji masalah Heteroskedastisitas adalah dengan mendeteksi pola residual melalui sebuah grafik. Jika residual mempunyai varian yang sama maka kita tidak mempunyai pola yang pasti dari residual. Sebaliknya jika residual mempunyai sifat Heteroskedastisitas residual ini akan menunjukkan pola yang tertentu.
2. Metode Park. Setelah membahas deteksi heteroskedastisitas secara informal dengan metode grafis, maka selanjutnya kita akan membahas uji deteksi Heteroskedastisitas dimulai dari metode yang dikembangkan oleh Park. Menurut Park, variabel –variabel gangguan yang tidak konstan atau masalah Heteroskedastisitas muncul karena residual ini tergantung dari variabel independen yang ada dalam model. Menurutnya fungsi variabel gangguan adalah sebagai berikut:

$$\sigma_i^2 = \sigma^2 x_i^\beta e^{u_i}$$

Dalam transformasi logaritma dapat ditulis sebagai berikut

$$\ln \sigma^2 + \beta \ln X_i + v_i$$

Dimana $\ln$ = Logaritma Natural dan v_i = variabel gabungan.

Karena varian variabel gangguan populasi tidak diketahui maka Park menyarankan menggunakan residual dari hasil regresi sebagai proksi dari

residual. Dengan demikian menggunakan residual dari hasil regresi dengan menggunakan persamaan sebagai berikut:

$$\ln e_i^2 = \ln \sigma^2 + \beta \ln X_i + v_i$$

Keputusan ada tidaknya masalah Heteroskedastisitas berdasarkan uji statistik estimator. Jika tidak signifikan melalui uji t maka dapat disimpulkan tidak ada heteroskedastisitas karena varian residualnya tidak tergantung dari variabel independen. Sebaliknya jika signifikan secara statistik maka model mengandung unsur Heteroskedastisitas karena besar kecilnya varian residual ditentukan oleh variabel independen. Maka prosedur dari uji Park adalah :Melakukan regresi terhadap model yang ada dengan metode OLS dan kemudian mendapatkan residualnya. Selanjutnya adalah melakukan regresi terhadap residual kuadrat sebagaimana pada persamaan, Jika nilai statistik t hitung lebih kecil dari nilai kritis tabel t maka tidak ada masalah Heteroskedastisitas dan jika sebaliknya maka mengandung masalah Heteroskedastisitas.

3. Metode Glejser

Sejalan dengan Park, ahli ekonometrika yang lain yakni Glejser mengatakan bahwa varian variabel gangguan nilainya tergantung dari variabel independen yang ada di dalam model. Berbeda dengan Park, agar kita bisa mengetahui apakah pola variabel gangguan mengandung heteroskedastisitas atau tidak maka Glejser menyarankan untuk melakukan regresi nilai absolut residual independennya. Glejser menyarankan untuk melakukan regresi fungsi-fungsi residual sebagai berikut:

$$|e_i| = \beta_o + \beta_1 X_i + v_i$$

$$|e_i| = \beta_o + \beta_1 \sqrt{X_i} + v_i$$

$$|e_i| = \beta_o + \beta_1 \frac{1}{X_i} + v_i$$

$$|e_i| = \beta_o + \beta_1 \frac{1}{\sqrt{X_i}} + v_i$$

$$|e_i| = \sqrt{\beta_o + \beta_1 X_i} + v_i$$

$$|e_i| = \sqrt{\beta_o + \beta_1 X_i^2} + v_i$$

Sebagaimana Park jika b_1 , tidak signifikan melalui uji t maka dapat disimpulkan tidak ada Heteroskedastisitas dan sebaliknya jika b_1 signifikan secara statistik maka model mengandung masalah Heteroskedastisitas. Glejser dalam penelitiannya menemukan bahwa untuk sampel besar, model fungsi residual. Metode Park dan Glejser merupakan metode sederhana dan mudah dilakukan. Namun kedua model mengandung kelemahan yakni berkaitan dengan masalah residual di persamaan.

4. Metode Korelasi Spearmen

Metode berikutnya untuk mendeteksi masalah Heteroskedastisitas adalah metode yang dikembangkan oleh Spearmen. Sebelum membahas metode korelasi dari Spearmen, kita definisikan terlebih dahulu korelasi yang dikembangkan oleh Spearmen. Formulasi korelasi dari spearmen adalah

$$r_s = 1 - 6 \frac{\sum d_i^2}{n(n^2 - 1)}$$

Dimana d adalah perbedaan rank antara residual dengan variabel independen x dan n adalah jumlah observasi. Metode deteksi Heteroskedastisitas dengan korelasi spearman ini dapat dijelaskan dengan

$$Y^i = \beta_0 + \beta_1 X_i + e_i$$

Langkah yang harus dilakukan untuk menguji ada tidaknya masalah Heteroskedastisitas dalam hasil regresi dengan menggunakan korelasi Spearman.

5. Metode Goldfeld-Quandt

Adanya kelemahan metode Park dan Glejser menginspirasi ahli ekonometrika lain untuk mengembangkan metode deteksi Heteroskedastisitas. GoldFeld-Quant kemudian mengembangkan metode deteksi Heteroskedastisitas lebih lanjut. Ide itu dapat dijelaskan sebagai berikut:

$$Y_i = \beta_0 + \beta_1 X_i + e_i$$

Dapat ditulis juga

$$\sigma_i^2 = \sigma^2 X_i^2$$

Hal ini berarti bahwa semakin besar X kuadrat semakin besar juga σ^2 .

Dengan demikian jika variabel gangguan terus meningkat secara substansial maka diduga masalah heteroskedastisitas ada didalam model. Namun jika varian variabel gangguan hampir sama maka diduga tidak ada masalah Heteroskedastisitas di dalam model. Namun jika varian variabel gangguan mempunyai karakteristik homoskedastisitas. Namun jika varian variabel

gangguan menunjukkan tren yang meningkat maka model mengandung heteroskedastisitas.

6. Metode Breusch-Pagan. Metode ini bisa dijelaskan dengan model regresi sederhana sebagai berikut:

$$Y_i = \beta_0 + \beta_1 X_i + e_i$$

Diasumsikan bahwa varian dari variabel gangguan mempunyai fungsi sebagai berikut:

$$\sigma_i^2 = f(\alpha_0 + \alpha_1 z_i)$$

7. Metode White

Tidak seperti metode Breusch-Pagan yang sangat tergantung pada asumsi tentang normalitas pada variabel gangguan. Hal White mengembangkan sebuah metode yang tidak memerlukan asumsi tentang adanya normalitas pada variabel gangguan. Model dari metode White:

$$Y_i = \beta_0 + \beta_1 X_{1i} + \beta_2 X_{2i} + e_i$$

Langkah uji White sebagai berikut:

1. Estimasi persamaan dan dapatkan residualnya(e)
2. Lakukan regresi pada persamaan berikut yang disebut regresi auxiliary

Regresi auxiliary tanpa perkalian antarvariabel independen

$e_i^2 = \alpha_0 + \alpha_1 X_{1i} + \alpha_2 X_{2i} + \alpha_3 X_{3i} + v_i$ merupakan residual kuadrat yang kita peroleh dari persamaan. Jika kita mempunyai lebih dari dua variabel independen dalam persamaan akan lebih.

3.6.3.1.4 Autokorelasi

(Widarjono, 2013) Sifat dan konsekuensi dari autokorelasi berarti adanya korelasi antara anggota observasi satu dengan observasi lain yang berlainan waktu. Dengan kaitannya dengan asumsi OLS autokorelasi merupakan korelasi antara satu variabel gangguan dengan variabel gangguan yang lainnya. Deteksi masalah Autokorelasi:

1. Metode Durbin Waston (DW). Salah satu uji yang populer yang biasa digunakan didalam ekonometrika adalah metode yang dikemukakan oleh Durbin Wiston. Prodesur uji yang dikembangkan oleh Durbin Wiston dapat dijelaskan dengan model regresi: $y_t = \beta_0 + \beta_1 X_1 + e_t$. Hubungan antarvariabel gangguan hanya tergantung dari variabel gangguan sebelumnya disebut model AR:

$$e_t = p e_{t-1} + v_t \quad -1 < p < 1$$

2. Metode Breusch-Godfrey. Walaupun uji autokorelasi dari Durbin Wiston mudah dilakukan karena informasi nilai statistik hitung d selalu diinformasikan setiap program computer, namun uji ini mengandung beberapa kelemahan. Pertama, uji ini hanya berlaku jika variabel independen bersifat random atau stokastik. Jika uji ini memasukkan variabel independen yang bersifat nonstokastik seperti memasukkan variabel keambanan dari variabel dependen sebagai variabel depende sebagai variabel independen yang disebut dengan model autoregresif maka uji Durbon Wiston tidak bisa digunakan. Kedua, uji Durbin Wiston hanya berlaku jika hubungan

autokorelasi antara residual dalam order pertama autoregresif yang lebih tinggi seperti AR(2) AR(3) dan seterusnya. Ketiga, model ini juga tidak bisa digunakan dalam kasus rata rata bergerak dari residual yang lebih tinggi. Contoh dalam model regresi:

$$y_t = \beta_0 + \beta_1 X_1 + e_t$$

Maka uji autokorelasi dengan AR. Berdasarkan kelemahan-kelemahan diatas maka Breus dan Godfrey mengembangkan uji autokorelasi yang lebih umum dan dikenal dengan Uji langrange Multiplier. Sebagai catatan kta bisa memasukkan lebih dari satu variabel independen namun untuk memudahkan kita menggunakan regresi sederhana.

3.6.3.2 Analisis Regresi

(Imam Ghozali, 2009) Analisis regresi adalah studi mengenai ketergantungan variabel dependen (variabel terikat) dengan satu atau lebih variabel independen (variabel bebas), dengan tujuan untuk mengestimasi dan atau memprediksi rata-rata populasi atau nilai rata-rata variabel dependen berdasarkan nilai variabel independen yang diketahui. Hasil regresi adalah berupa koefisien untuk masing-masing variabel independen. (Sugiono, 2007) hasil analisis regresi bermanfaat untuk membuat keputusan apakah naik dan menurunnya variabel dependen dapat dilakukan melalui peningkatan variabel independen atau tidak.

3.6.3.2.1 Analisis Regresi Linear Berganda

(Widarjono, 2013) Dalam kenyataannya model regresi sederhana tidak mencerminkan kondisi perilaku variabel ekonomi yang sebenarnya. Model regresi

linear yang terdiri dari lebih dari satu variabel independen dikenal dengan model regresi berganda. Bentuk umum regresi linear berganda.

$$Y_i = \beta_0 + \beta_1 X_{1i} + \beta_2 X_{2i} + \dots + \beta_k X_{ki} + e_i$$

Dimana Y adalah variabel dependen X1 dan X2 adalah variabel independen e_i adalah variabel gangguan. Subkrip menunjukkan observasi ke- i untuk data cross section dan jika kita gunakan dalam data time series biasanya kita beri subkrip t yang menunjukkan waktu. Didalam persamaan sebagaimana pada regresi sederhana. Ada beberapa asumsi OLS yang digunakan dalam regresi berganda seperti dalam regresi sederhana. Karena ada lebih dari satu variabel independen maka pada asumsi ditambah tidak ada hubungan linier antara variabel independen atau tidak ada multikolinieritas. Dalam kasus regresi berganda berarti tidak ada multikolinieritas antara X1 dan X2.

$$E(Y_i | X_{1i}, X_{2i}) = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + e_i$$

Arti persamaan tersebut adalah nilai harapan (Expected Value) atau rata-rata dari Y pada nilai tertentu pada variabel independen X1 dan X2.

3.6.3.3 Uji Hipotesis

Di gunakan untuk menentukan apakah ada pengaruh keterkaitan antara (X_1 dengan Y, X_2 dengan Y,) yang dapat dilihat dari besarnya t hitung terhadap t tabel dengan uji 2 sisi menurut Sujarweni (2015: 158-164).

3.6.3.3.1 Uji signifikansi Simultan (Uji Statistik F)

(Widarjono,2013) Uji statistik F digunakan untuk uji signifikan model. Uji F ini bisa dijelaskan menggunakan analisis varian (analysis of Variance=ANOVA). Dengan rumus sebagai berikut

$$Y_i = \beta_0 + \beta_1 X_{1i} + \beta_2 X_{2i} + e_i$$

Koefisien determinasinya = $TSS = ESS + SSR$. TSS mempunyai $df = n - 1$, ESS mempunyai df sebesar $k - 1$ sedangkan SSR mempunyai $df = n - k$. Analisis varian ini bisa ditampilkan dalam tabel. Dengan hipotesis bahwa semua variabel independen tidak berpengaruh terhadap variabel dependen yakni:

$\beta_0 = \beta_1 = \dots = \beta_k = 0$ maka uji F dapat diformulakan sebagai berikut:

$$F = \frac{ESS(k - 1)}{SSR(n - k)}$$

Dimana n jumlah observasi dan k = jumlah parameter estimasi termasuk intersep atau konstanta. Formula uji statistik ini bisa dinyatakan dalam bentuk formula yang lain dengan cara memanipulasikan persamaan:

$$F = \frac{ESS/(k - 1)}{(TSS - ESS)/(n - k)}$$

Dari persamaan diatas hipotesis nol terbukti, maka kita harapkan nilai dari ESS dan R^2 akan sama dengan nol sehingga F akan sama dengan nol. Dengan demikian, kita akan gagal menolak hipotesis nol karena variabel independen hanya sedikit menjelaskan varian variasi dependen di sekitar rata ratanya.

Walaupun uji F menunjukkan adanya penolakan hipotesis nol yang menunjukkan bahwa secara bersama-sama semua variabel independen mempengaruhi variabel dependen namun hal ini bukan berarti secara individual variabel independen mempengaruhi variabel dependen melalui uji t. Keadaan ini terjadi karena kemungkinan ada korelasi yang tinggi antara variabel independen. Kondisi ini menyebabkan standar error sangat tinggi dan rendahnya nilai t hitung meskipun model secara umum mampu menjelaskan data secara baik.

$$Y_i = \beta_0 + \beta_1 X_{1i} + \beta_2 X_{2i} + e_i$$

Untuk menguji apakah koefisien regresi secara bersama sama atau secara menyeluruh berpengaruh terhadap variabel dependen, uji F dapat dijelaskan sebagai berikut:

1. Membuat hipotesis nol(H_0) dan hipotesis alternatif H_a
2. Mencari nilai F hitung dengan formula seperti pada persamaan diatas dan nilai F kritis dari tabel distribusi frekuensi F. Nilai kritis berdasarkan besarnya a dan df dimana besarnya ditentukan oleh numerator(k-1) dan df untuk denominator(n-k)
3. Keputusan menolak atau gagal menolak H_0 sebagai berikut:

Jika F hitung > F kritis, maka kita menolak H_0 dan sebaliknya jika F hitung < F kritis maka gagal menolak H_0 .

3.6.3.3.2 Uji Signifikan Parameter Individual (Uji Statistik t)

(Widarjono, 2013) Prodesur uji t pada koefisien regresi parsial pada regresi berganda sama dengan prosedur uji koefisien regresi sederhana. Lagkah langkah uji t sebagai berikut:

1. Membuat hipotesis melalui uji satu atau dua sisi

Uji hipotesis positif satu sisi.

$$H_0: \beta_1 = 0$$

$$H_a: \beta_1 > 0$$

Uji hipotesis negatif satu sisi

$$H_0: \beta_1 = 0$$

$$H_a: \beta_1 < 0$$

Atau uji dua sisi

$$H_0: \beta_1 = 0$$

$$H_a: \beta_1 \neq 0$$

2. Kita ulangi langkah pertama.
3. Menghitung nilai t hitung untuk β_1 dan β_2 dan mencari nilai t kritis dari table distribusi t. Nilai t hitung dicari dengan formula sebagai berikut:

$$t = \frac{\beta_1 - \beta_i}{se(\beta_i)}$$

Dimana β_1 merupakan nilai pada hipotesis nol.

4. Bandingkan nilai t hitung untuk masing masing estimator dengan t kritisnya dari tabel. Keputusan menolak atau gagal menolak H_0 sebagai berikut:

Jika nilai t hitung $>$ nilai kritis maka H_0 ditolak atau menerima H_a

Jika nilai t hitung $<$ nilai t kritis maka H_0 gagal ditolak.

3.6.3.3 Analisis Koefisien Determinasi (R^2)

(Widarjono, 2013) koefisien determinasi untuk menjelaskan proporsi variasi variabel dependen dijelaskan oleh variabel independen. Di dalam regresi berganda kita juga akan menggunakan koefisien determinasi untuk mengukur seberapa baik garis regresi yang kita punyai. Dalam hal ini kita mengukur seberapa besar proporsi variasi variabel dependen dijelaskan oleh semua variabel independen. Formula untuk menghitung koefisien determinasi regresi berganda sama dengan regresi sederhana yaitu sebagai berikut:

$$R^2 = \frac{ESS}{TSS} = \frac{TSS - SSR}{TSS} = 1 - \frac{SSR}{TSS}$$

$$R^2 = 1 - \frac{Ee_i^2}{E y_i^2}$$

$$R^2 = 1 - \frac{Ee_i^2}{E(Y_i - \bar{Y})^2}$$

Koefisien determinasi tidak pernah menurun terhadap jumlah variabel independen. Artinya koefisien determinasi akan semakin besar jika terus menambah variabel independen di dalam model. Mengingat bahwa nilai koefisien determinasi tidak pernah menurun maka kita harus berhati-hati membandingkan dua regresi yang mempunyai variabel independen Y sama tetapi berbeda dalam jumlah variabel independen X . Kehatian-hatian ini perlu karna tujuan regresi

metode OLS adalah mendapatkan nilai koefisiensi determinasi yang tinggi. Salah satu persoalan besar penggunaan koefisien determinasi R dengan demikian adalah nilai R selalu naik ketika kita menambah variabel independen X. dalam model walaupun penambahan variabel independen X belum tentu mempunyai Justifikasi atau pembenaran dari teori ekonomi ataupun logika ekonomi. Para ekonometrika telah mengembangkan alternatif lain agar nilai R^2 tidak merupakan fungsi dari variabel independen. Sebagai alternatif digunakan R^2 yang disesuaikan (Adjusted R^2 dengan rumus sebagai berikut:

$$R^2 = 1 - \frac{Ee^2(n-k)}{E(Y_i - \bar{Y})(n-1)}$$