

BAB III METODE PENELITIAN

3.1 Tempat dan Waktu Penelitian

Tempat dalam penelitian ini dilaksanakan di Provinsi Nusa Tenggara Timur secara keseluruhan, Pada Badan Pusat Statistik Provinsi Nusa Tenggara Timur yang berlokasi di Jln R. Suprpto No 5 Oebobo Kota Kupang. Waktu penelitian selama 6 (enam) bulan yaitu bulan Januari sampai Juni 2019.

3.2 Defenisi 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 Kemiskinan, variabel bebas (X2) adalah Pertumbuhan Ekonomi.

1. Indeks Pembangunan Manusia 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, diukur dalam (Persen).

2. Kemiskinan (X1) adalah keadaan di mana terjadi ketidakmampuan untuk memenuhi kebutuhan dasar seperti makanan, pakaian, tempat berlindung, pendidikan, dan kesehatan. Kemiskinan dapat disebabkan oleh kelangkaan alat pemenuh kebutuhan dasar, ataupun sulitnya akses terhadap pendidikan dan pekerjaan, di ukur dalam (Persen).

3. Pertumbuhan Ekonomi (X2) 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, di ukur dalam (Persen).

3.3 Jenis Dan Sumber Data

Menurut Jenisnya data yang Digunakan Dalam Penelitian Ini adalah data Kuantitatif yaitu data berupa angka pasti. Menurut Sumbernya data yang digunakan adalah sekunder yaitu data Panel yang diperoleh langsung dari Badan Pusat Statistik Provinsi NTT dan Dinas terkait. Data Panel dalam penelitian ini adalah data Kemiskinan ,Pertumbuhan Ekonomi Dan Indeks Pembangunan Manusia Kab/Kota di Provinsi Nusa Tenggara Timur 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.

Populasi dalam penelitian ini adalah data Kemiskinan, Pertumbuhan Ekonomi Dan Indeks Pembangunan Manusia Kab/Kot di Provinsi Nusa Tenggara Timur.

2.4.2 Sampel

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

Sampel dalam penelitian ini adalah data Kemiskinan, Pertumbuhan Ekonomi Dan Indeks Pembangunan Manusia Kabupaten/Kota di Provinsi Nusa Tenggara Timur 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 Kemiskinan, Pertumbuhan Ekonomi dan Indeks Pembangunan Manusia Di Provinsi Nusa Tenggara Timur.

3.6 Teknik Analisis

3.6.1 Analisis Regresi Data Panel

Menurut Basuki (2016:276) regresi data panel merupakan teknik regresi yang menggabungkan data runtut waktu (*time series*) dengan data silang (*cross section*).

3.6.2 Metode Estimasi Model Regresi 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 ekonometrik 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 heteroskedastisitas. 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 Tenaga Kerja (X_1), Pengeluaran Pemerintah (X_2), terhadap variabel dependen yaitu Produk Domestik Regional Bruto (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 Pembnagunan Manusia

X_1 = Kemiskinan

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 enggunakan 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 serries* 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

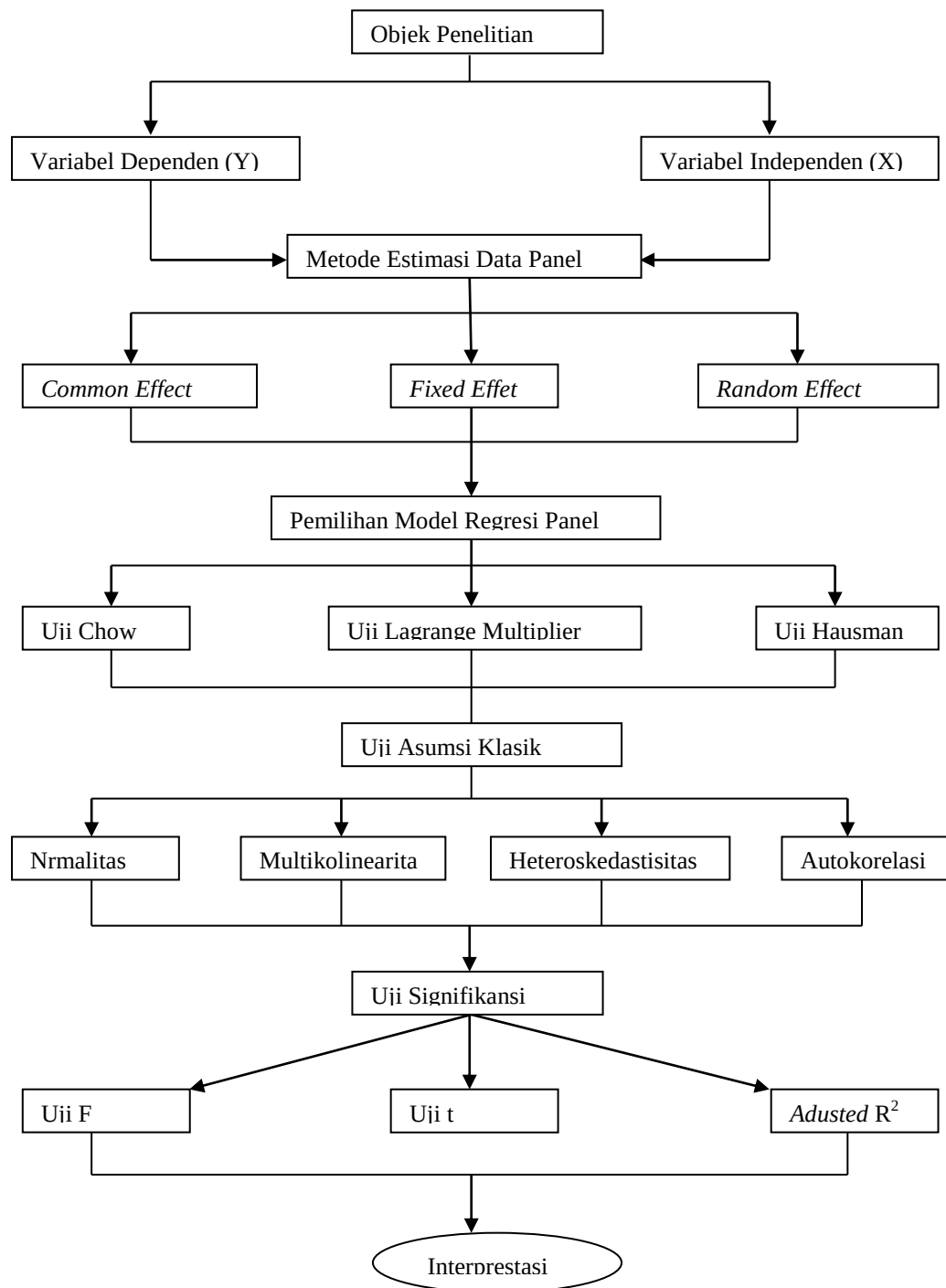
Hausman test adalah pengujian statistic untuk memilih apakah model *Fixed Effect* atau *Random Effect* yang paling teat digunakan.

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



Sumber : Agus Widardjono (2013)

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 = 2012, 2013, 2014, 2015

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 diterapka untuk efek tiap individu, LSDV ini juga data mengakomodasi efek waktu yang bersifat sistemik. Ha ini dapat dilakukan melalui penambahan variable *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 variable penjelas yang teramati, model seperti ini dinamakan *Random Effects Model* (REM). Model seperti ini sering diebut juga dengan *Error Component Model* (ECM). Dengan demikian, persamaan *model effects* dapat dituliskan sebagai berikut:

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

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

$$t = 2012, 2013, 2014, 2015$$

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_{t-1} (*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 n (jumlah data *time series*) besar dan t (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 variable spesifik individu.
2. Kemampuan mengontrol heterogenitas ini selanjutnya menjadikan data panel digunakan untuk menguji dan membangun model perilaku lebih kompleks.
3. Data panel didasarkan 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 informative, 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 *Langrange 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*.

Dari beberapa pemaparan di atas dapat disimpulkan bahwa pada model regresi data panel, uji asumsi klasik yang dipakai hanya multikolinieritas dan heteroskedastisitas saja. Tetapi berikut akan dijelaskan secara keseluruhan mengenai uji asumsi klasik .

3.6.4.1. Uji Normalitas

Widarjono (2013) uji signifikansi pangaruh 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 mendeteksi apakah residual mempunyai distribusi normal atau tidak.

1. Histogram Residual

Histogram residual merupakan metode grafis yang paling sederhana digunakan untuk mengetahui apakah bentuk dari *Probability Distribution Function* (PDF) dari variabel random berbentuk distribusi normal atau tidak. Jika histogram residual menyerupai 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-Bera

Uji normalitas residual metode OLS secara formal dapat dideteksi dari metode yang dikembangkan oleh Jarque-Bera (J-B). Metode JB ini didasarkan pada sampel besar yang diasumsikan bersifat *asymptotic*. Uji statistik dari J-B ini menggunakan perhitungan *skewness* dan kurtosis. Adapun formula uji statistic J-B adalah sebagai berikut:

$$JB = n \left[\frac{s^2}{6} + \left(\frac{K - 3}{24} \right)^2 \right]$$

Dimana S = koefisien skweness dan K = koefisien kurtosis

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

3.6.4.2 Uji Multikolinearitas

Uji untuk melihat ada atau tidaknya korelasi yang tinggi antara variabel-variabel bebas dalam suatu model regresi linear berganda. Jika ada korelasi yang

tinggi diantara variabel-variabel bebasnya, maka hubungan antar variabel bebas terhadap variabel terikatnya menjadi terganggu.

Alat statistik yang digunakan dalam penelitian ini untuk menguji ada tidaknya multikolinearitas adalah dengan menggunakan *Auxiliary Regression*, dengan membuat 5 model regresi dengan Variabel dependen yang berbeda, secara berurut variabel dependen regresi 1-5 yakni; NAV Reksadana Campura Syariah, Inflasi, Kurs Rupiah, BI Rate dan *Gross Domestic Product* (GDP). Data penelitian dinyatakan terbebas dari multikolinearitas apabila nilai *R-square* model 1 lebih besar dari nilai *R-square* yang lainnya. (Winarno, 2011:5.3)

3.6.4.3 Uji Heteroskedastisitas

Uji untuk melihat apakah terdapat ketidaksamaan varian dari residual satu pengamatan ke pengamatan yang lain. Model regresi yang memenuhi persyaratan adalah dimana terdapat kesamaan varians dari residual satu pengamatan ke pengamatan yang lain tetap atau disebut homoskedastisitas. Metode yang digunakan untuk uji heteroskedastisitas adalah Uji White, Glejser, Breusch-Pagan-Godfrey, Harvey, dan ARCH. Model memenuhi persyaratan apabila nilai probabilitas *chi-square* nyanya melebihi nilai alpha 0,5. (Winarno, 2011:5.14) .

3.6.4.4 Uji 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: $y_e = \beta_o + \beta_1 X_1 + et$. 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 nonstokatik 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:

$$y_t = \beta_\sigma + \beta_1 X_1 + et$$

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 kita bisa memasukkan lebih dari satu variabel independen namun untuk memudahkan kita menggunakan regresi sederhana.

3.6.5 Uji Hipotesis

3.6.5.1 Uji Koefisien Determinasi (*Adjusted*)

Menurut Ghazali (2013:97), Koefisien determinasi mengukur seberapa jauh kemampuan model dalam menerangkan variasi variabel dependen. Nilai yang kecil berarti kemampuan variabel independen dalam menjelaskan variasi variabel dependen amat terbatas. Nilai yang mendekati satu berarti variabel independen memberikan hampir semua informasi yang di butuhkan untuk memprediksi variasi variabel dependen. Dalam penelitian ini pengukuran menggunakan *Adjusted* karena lebih akurat untuk mengevaluasi model regresi tersebut.

3.6.5.2 Uji Simultan (Uji *F*)

Menurut Ghazali (2013:98), uji *F* pada dasarnya bertujuan untuk menunjukkan apakah semua variabel bebas atau independen yang di masukkan dalam model mempunyai pengaruh secara bersama-sama terhadap variabel terikat atau dependen. Uji *F* ini dilakukan dengan menggunakan nilai signifikansi. Rumusan hipotesis sebagai berikut:

Ho : variabel independen secara simultan tidak berpengaruh terhadap variabel dependen.

Ha : variabel independen secara simultan berpengaruh terhadap variabel dependen.

Adapun kinerja pengujiannya sebagai berikut:

Ho diterima jika tingkat signifikansi $> 0,05$

Ha diterima jika tingkat signifikansi $< 0,05$

3.6.5.3 Uji Parsial (Uji t)

Menurut Ghazali (2013:98), uji t pada dasarnya bertujuan untuk menunjukkan seberapa jauh pengaruh satu variabel penjelas atau independen secara individual dalam menerangkan variabel dependen. Rumusan hipotesis yang digunakan sebagai berikut:

Ho : variabel independen tidak berpengaruh signifikansi terhadap variabel dependen.

Ha : variabel independen berpengaruh secara signifikan terhadap variabel dependen.

Adapun kriteria pengujiannya sebagai berikut:

Ho diterima jika tingkat signifikansi $> 0,05$.

Ha diterima jika tingkat signifikansi $< 0,05$.