

BAB III

METODOLOGI PENELITIAN

A. Tujuan Penelitian

Berdasarkan masalah-masalah yang telah peneliti rumuskan, maka tujuan penelitian ini adalah untuk:

1. Mengetahui besarnya pengaruh jumlah penduduk terhadap pengangguran terbuka di provinsi Jambi.
2. Mengetahui besarnya pengaruh investasi terhadap pengangguran terbuka di provinsi Jambi.

B. Objek Penelitian

Objek dari penelitian ini adalah pengangguran di Provinsi Jambi serta pengaruh dari jumlah penduduk dan investasi terhadap pengangguran terbuka. Ruang lingkup penelitian ini adalah seluruh kabupaten/kota yang ada di Provinsi Jambi dengan rentang waktu antara tahun 2007 - 2012 untuk jumlah penduduk, investasi dan pengangguran terbuka. Wilayah dipilih karena terjangkau dan ketersediaan data-data yang relevan dengan penelitian.

C. Metode Penelitian

Peneliti dalam melakukan penelitian ini menggunakan metode penelitian *expos facto*, yang merupakan suatu penelitian yang dilakukan untuk meneliti

peristiwa yang telah terjadi untuk mengetahui faktor-faktor yang dapat menimbulkan kejadian tersebut. Sesuai dengan tujuan penelitian yaitu untuk mengetahui apakah terdapat pengaruh jumlah penduduk dan investasi terhadap pengangguran terbuka provinsi Jambi. Sedangkan untuk pendekatan korelasional yang digunakan adalah dengan menggunakan korelasi ganda, yang dipilih karena dapat menunjukkan arah pengaruh faktor-faktor penentu (jumlah penduduk dan investasi) terhadap pengangguran terbuka yang ada dalam penelitian ini.

D. Jenis dan Sumber Data

Jenis data yang digunakan dalam penelitian ini adalah data sekunder yang bersifat kuantitatif yaitu data yang telah tersedia dalam bentuk angka. Bentuk data yang digunakan adalah data panel. Data panel adalah data yang berstruktur urut waktu sekaligus *cross section*³⁶. Penggunaan data tahunan pada kabupaten/kota di Provinsi Jambi dipilih untuk melihat jumlah pengangguran dan jumlah investasi. Data sekunder tersebut diperoleh dari Badan Pusat Statistik.

E. Operasionalisasi Variabel Penelitian

1. Pengangguran Terbuka

a. Definisi Konseptual

Pengangguran adalah orang yang tidak bekerja, atau orang yang sedang mencari atau menunggu panggilan pekerjaan.

³⁶ Moch. Doddy Ariefianto, *Esensi dan Aplikasi dengan Menggunakan Eviews* (Jakarta: Erlangga, 2010), hal. 148

b. Definisi Operasional

Pengangguran merupakan orang dalam usia kerja yang tidak memiliki pekerjaan. Jumlah pengangguran menunjukkan jumlah angkatan kerja yang tidak memiliki pekerjaan. Data yang akan digunakan adalah data di Provinsi Jambi tahun 2007-2012.

2. Jumlah Penduduk

a. Definisi Konseptual

Jumlah penduduk adalah sekumpulan orang yang berada di suatu wilayah dan bertempat tinggal atau menetap yang dipengaruhi oleh kelahiran, kematian dan migrasi.

b. Definisi Operasional

Jumlah penduduk adalah total penduduk pada suatu wilayah dari umur 0 – 65 tahun pada kurun waktu 2007-2012 yang diambil dari data jumlah penduduk berdasarkan tiap-tiap daerah kabupaten/ kota di provinsi Jambi.

3. Investasi

a. Definisi Konseptual

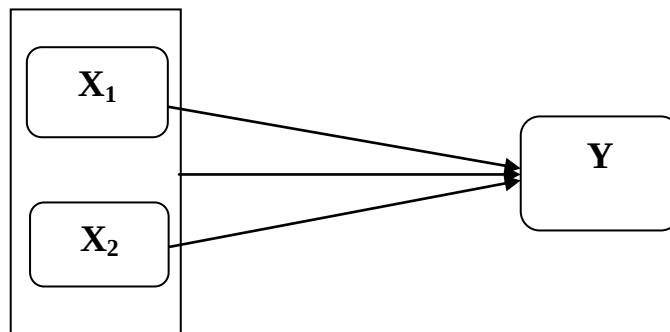
Investasi juga lazim disebut dengan istilah penanaman modal atau pembentukan modal. Penanaman modal merupakan komponen yang menentukan tingkat pengeluaran agregat. Apabila seorang pengusaha menggunakan uang untuk membeli barang-barang modal maka pengeluaran tersebut dinamakan investasi.

b. Definisi Operasional

Penanaman modal merupakan data sekunder yang diambil dari Badan Koordinasi Promosi dan Penanaman Modal (BKPM) yang diterbitkan secara berkala. Data yang akan digunakan adalah data penjumlahan realisasi penanaman modal asing (PMA) dan penanaman Modal dalam negeri (PMDN) yang dinyatakan dalam satuan Rupiah di Provinsi Jambi tahun 2007 hingga tahun 2012.

F. Konstelasi Pengaruh Antar Variabel

Konstelasi pengaruh antar variabel didalam penelitian ini bertujuan untuk memberikan gambaran dari penelitian ini, yang dapat digambarkan sebagai berikut:



Keterangan:

X_1 = Jumlah Penduduk (variabel bebas)

X_2 = Investasi (variabel bebas)

Y = Pengangguran Terbuka (variabel terikat)

→ = arah pengaruh

G. Teknik Analisis Data

1. Analisis Data Panel

Data yang digunakan dalam analisis ekonometrika dapat berupa data *time series*, data *cross section*, atau data panel. Data panel (*panel pooled data*) merupakan gabungan data *time series* dan data *cross section*. Dengan kata lain, data panel merupakan unit-unit individu yang sama yang diamati dalam kurun waktu tertentu. Jika kita memiliki T periode waktu ($t = 1, 2, \dots, T$) dan N jumlah individu ($i = 1, 2, \dots, N$), maka dengan data panel kita akan memiliki total unit observasi sebanyak NT . Jika jumlah unit waktu sama untuk setiap individu, maka data disebut *balanced panel*. Jika sebaliknya, yakni jumlah unit waktu berbeda untuk setiap individu, maka disebut *unbalanced panel*. Penggunaan data panel pada dasarnya merupakan solusi akan ketidaktersediaan data *time series* yang cukup panjang untuk kepentingan analisis ekonometrika.

Gujarati (2001) berdasarkan uraian dari Baltagi, keunggulan penggunaan data panel dibanding data runtun waktu dan data lintas sektor adalah:

- 1) Estimasi data panel dapat menunjukkan adanya heterogenitas dalam tiap unit.
- 2) Dengan data panel, data lebih informatif, mengurangi kolinearitas antara variabel, meningkatkan derajat kebebasan dan lebih efisien.
- 3) Data panel cocok digunakan untuk menggambarkan adanya dinamika perubahan.
- 4) Data panel dapat lebih mampu mendeteksi dan mengukur dampak.

- 5) Data panel bisa digunakan untuk studi dengan model yang lebih lengkap.
- 6) Data panel dapat meminimumkan bias yang mungkin dihasilkan regresi.

Terdapat tiga metode pada estimasi model menggunakan data panel, yaitu model *common effect*, *fixed effect*, dan *random effect*.

a. **Model Common Effect**

Model *common effect* merupakan model regresi data panel yang paling sederhana. Model ini pada dasarnya mengabaikan struktur panel dari data, sehingga diasumsikan bahwa perilaku antar individu sama dalam berbagai kurun waktu atau dengan kata lain pengaruh spesifik dari masing-masing individu diabaikan atau dianggap tidak ada. Dengan demikian, akan dihasilkan sebuah persamaan regresi yang sama untuk setiap unit *cross section*. Persamaan regresi untuk model *common effect* dapat dituliskan sebagai berikut:

$$Y_{it} = \alpha + \beta X_{it} + \varepsilon_{it} \quad i = 1, 2, \dots, N \quad t = 1, 2, \dots, T$$

Keterangan:

- Y : variabel dependen
- α : koefisien regresi
- X : variabel independen
- β : estimasi parameter (koefisien)
- ε : *error term*
- N : jumlah (individu)
- T : jumlah periode waktu

Berdasarkan asumsi struktur matriks varians-kovarians residual, maka pada model *common effect* terdapat empat model estimasi yang dapat digunakan, yaitu:

- 1) *Ordinary Least Square (OLS)*, jika struktur matrik varians-kovarians residualnya diasumsikan bersifat homoskedastik dan tidak ada *cross sectional correlation*.
- 2) *Generalized Least Square (GLS)/Weighted Least Square (WLS): Cross Sectional Weight*, jika struktur matrik varians-kovarians residualnya diasumsikan bersifat heteroskedastik dan tidak ada *cross sectional correlation*.
- 3) *Feasible Generalized Least Square (FGLS)/ Seemingly Uncorrelated Regression (SUR)* atau *Maximum Likelihood Estimator (MLE)*, jika struktur matriks varians-kovarians residualnya diasumsikan bersifat heteroskedastik dan ada *cross sectional correlation*.
- 4) *Feasible Generalized Least Square (FGLS)* dengan proses *Auto Regressive (AR)* pada *error term*-nya, jika struktur matriks varians-kovarians residualnya diasumsikan bersifat heteroskedastik dan ada korelasi antar waktu pada residualnya.

b. Model Fixed Effect

Jika model *common effect* cenderung mengabaikan struktur panel dari data dan pengaruh spesifik masing-masing individu, maka model *fixed effect* adalah sebaliknya. Pada model ini, terdapat efek spesifik individu α_i dan diasumsikan berkorelasi dengan variabel penjelas yang teramati X_{it} . Ekananda (2005) menyatakan bahwa berdasarkan asumsi struktur matriks varians-kovarians residual, maka pada model *fixed effect* terdapat tiga metode estimasi yang dapat digunakan, yaitu:

- 1) *Ordinary Least Square* (OLS/LSDV), jika struktur matriks varians-kovarians residualnya diasumsikan bersifat homoskedastik dan tidak ada *cross sectional correlation*.
- 2) *Weight Least Square* (WLS), jika struktur matriks varians-kovarians residualnya diasumsikan bersifat heteroskedastik dan tidak ada *cross sectional correlational*.
- 3) *Seemingly Uncorrelated Regression* (SUR), jika struktur matriks varians-kovarians residualnya diasumsikan bersifat heteroskedastik dan ada *cross sectional correlation*.

Secara umum, pendekatan *fixed effect* dapat dituliskan dalam persamaan sebagai berikut:

$$Y_{it} = \alpha + \beta X_{it} + \gamma_i W_{it} + \delta Z_{it} + \varepsilon_{it}$$

Keterangan:

Y_{it} : variabel terikat untuk individu ke-i dan waktu ke-t

X_{it} : variabel bebas untuk individu ke-i dan waktu ke-t

W_{it} dan Z_{it} : variabel *dummy* yang didefinisikan sebagai berikut:

W_{it} : 1; untuk individu i; $i = 1, 2, \dots, N = 0$; lainnya

Z_{it} : 1; untuk individu t; $t = 1, 2, \dots, N = 0$; lainnya

ε_{it} : *error term* untuk individu ke-i dan waktu ke-t

c. **Model Random Effect**

Pada model *random effect*, efek spesifik dari masing-masing individu α_i diperlakukan sebagai bagian dari komponen *error* yang bersifat acak dan tidak

berkorelasi dengan variabel penjelas yang teramati X_{it} . Dengan demikian, persamaan model *random effect* dapat dituliskan sebagai berikut:

$$Y_{it} = \alpha + \beta X_{it} + \varepsilon_{it}; \varepsilon_{it} = u_i + v_i + w_{it}$$

Keterangan:

- u_i : komponen *error cross section*
- v_i : komponen *error time series*
- w_{it} : komponen *error gabungan*

Asumsi dasar model ini adalah perbedaan nilai intersep antar unit *cross section* dimasukkan ke dalam *error*. Karena hal ini, model *random effect* sering disebut dengan *Error Component Model* (ECM). Model ini diestimasi dengan metode *Generalized Least Square* (GLS). Intersep model ini bervariasi terhadap individu dan waktu, namun slopenya konstan terhadap individu dan waktu. Penggunaan pendekatan *random effect* tidak mengurangi derajat kebebasan sebagaimana terjadi pada model *fixed effect* yang akan berakibat pada parameter hasil estimasi akan menjadi lebih efisien.

d. Uji Kesesuaian Model

Untuk mengetahui kesesuaian atau kebaikan model dari ketiga metode pada teknik estimasi model dengan data panel digunakan *Chow test*, *LM test*, dan *Hausman test* yang digunakan untuk menguji kesesuaian model antara model yang diperoleh dari *common effect* dengan model yang diperoleh dari metode *fixed effect*. *LM test* digunakan untuk menguji kesesuaian model yang diperoleh dari *common effect* dengan model yang terbaik yang diperoleh dari metode

random effect. Sedangkan *Hausman test* digunakan terhadap model yang terbaik yang diperoleh dari metode *fixed effect* dengan model yang diperoleh dari metode *random effect*.³⁷

Tabel III.1
Pengujian Signifikansi Model Panel Terbaik

No	Pengujian Signifikansi Model	Hipotesis Pengujian	Rumus Uji	Ket
1	CE atau FE	H_0 : CE lebih baik dari FE	Uji Chow	Tolak H_0 $F_{hitung} > F_{tabel}$
		H_1 : FE lebih baik dari CE		
2	CE atau RE	H_0 : CE lebih baik dari RE	Uji LM	Tolak H_0 $LM > \chi^2_{tabel}$
		H_1 : RE lebih baik dari CE		
3	FE atau RE	H_0 : RE lebih baik dari FE	Uji Hausman	Tolak H_0 $\chi^2_{hitung} > \chi^2_{tabel}$
		H_1 : FE lebih baik dari RE		

Sumber: Wing W. Winarno, *Analisis Ekonometrika dan Statistika*, 2011.

Keterangan:

CE: *Common Effect*

FE: *Fixed Effect*

RE: *Random Effect*

³⁷ Winarno, *Analisis Ekonometrika dan Statistika dengan Eviews*, (Yogyakarta: UPP STIM YKPM, 2007), hal. 21.

1) *Chow test*

Chow test dapat dilakukan dengan uji statistik F untuk mengetahui apakah model yang digunakan *common effect* atau *fixed effect*. Sebagaimana yang diketahui bahwa terkadang asumsi yang menyatakan bahwa setiap unit *cross section* memiliki perilaku yang sama cenderung tidak realistis mengingat dimungkinkan setiap unit *cross section* memiliki perilaku berbeda. Dalam pengujian ini dilakukan dengan hipotesis sebagai berikut:

H_0 : Model *common effect*

H_1 : Model *fixed effect*

Dasar penolakan terhadap hipotesa nol (H_0) adalah dengan menggunakan F-statistik seperti yang dirumuskan sebagai berikut:

$$F = \frac{(RSS_1 - RSS_2)/(n - 1)}{RSS_2/(nT - n - k)}$$

Keterangan:

RSS_1 : *Residual Sum Square* hasil pendugaan model *fixed effect*

RSS_2 : *Residual Sum Square* hasil pendugaan model *common effect*

n : jumlah data *cross section*

T : jumlah data *time series*

k : jumlah variabel penjelas

Statistik *Chow test* mengikuti distribusi F-statistik dengan derajat bebas $(n - 1, nT - n - k)$. Jika nilai *chow statistics* (F statistik) hasil pengujian lebih besar dari F_{tabel} , maka cukup bukti untuk melakukan penolakan terhadap

hipotesa nol, sehingga model yang digunakan adalah model *fixed effect*, dan begitu juga sebaliknya.

2) *Lagrange Multiplier Test (LM Test)*

Pengujian signifikansi *common effect* atau *random effect* melalui pengujian *Lagrange Multiplier* untuk mengetahui signifikansi dari *random effect* berdasarkan residual dari OLS (*common effect*). LM mengikuti sebaran *chi-square* dengan derajat bebas satu. Secara matematis, statistik uji untuk *LM test (Lagrange Multiplier)* dapat dituliskan sebagai berikut:

$$LM = \frac{nT}{2(T-1)} \left[\frac{\sum_{t=1}^n (\sum_{i=1}^n e_{it})^2}{\sum_{t=1}^n \sum_{i=1}^T e_{it}^2} - 1 \right]^2$$

Hipotesis pengujian ini adalah:

H_0 : Model *common effect*

H_1 : Model *random effect*

Jika nilai *LM test* lebih besar dari χ^2_{tabel} , maka hipotesis nol ditolak. Sehingga model yang akan diterima dan digunakan adalah model *random effect* dan sebaliknya.

3) *Hausman test*

Setelah menguji signifikansi antara *common effect* atau *fixed effect* serta *common effect* atau *random effect*, maka selanjutnya jika terbukti *fixed effect* dan *random effect* sama-sama lebih baik dari *common effect* adalah melakukan uji Hausman untuk pengujian signifikansi mana yang lebih baik

fixed effect atau *random effect*. Uji ini dilakukan dengan membandingkan nilai statistik Hausman dengan nilai kritis statistik *chi-square*. Secara sistematis dengan menggunakan notasi matriks, statistik uji Hausman (H) dapat dituliskan sebagai berikut:

$$H = (\hat{\beta}_{FE} - \hat{\beta}_{RE})[\text{var}(\hat{\beta}_{FE}) - \text{var}(\hat{\beta}_{RE})]^{-1}(\hat{\beta}_{FE} - \hat{\beta}_{RE})$$

Hipotesis nul pada *Hausman test* adalah pendugaan parameter dengan menggunakan *random effect* adalah konsisten dan efisien, sedangkan pendugaan dengan *fixed effect* meskipun tetap konsisten tetapi tidak lagi efisien. Hipotesis alternatif estimasi dengan *random effect* menjadi tidak konsisten, sebaliknya estimasi dengan *fixed effect* tetap konsisten.

Hipotesis pengujian ini adalah:

H_0 : Model *random effect*

H_1 : Model *fixed effect*

Jika nilai *Hausman test* hasil pengujian lebih besar dari $\text{chi-square}_{\text{tabel}}$, maka hipotesis nol ditolak, yang berarti estimasi yang tepat untuk regresi data panel adalah model *fixed effect* dan sebaliknya.

Sementara itu, Judge *et.al.* dalam Gujarati (2003) memberikan sejumlah pertimbangan terkait pilihan, apakah menggunakan model *fixed effect* (FE) atau model *random effect* (RE). Pertimbangan-pertimbangan itu adalah sebagai berikut:

- a) Jika jumlah data *time series* (T) besar dan jumlah data *cross section* (N) kecil, ada kemungkinan perbedaan nilai parameter yang diestimasi dengan FE dan RE cukup kecil. Karena itu, pilihan ditentukan berdasarkan kemudahan perhitungan. Dalam hal ini adalah model FE.
- b) Ketika N besar dan T kecil, estimasi kedua metode dapat berbeda secara signifikan. Pada kondisi seperti ini, pilihan ditentukan berdasarkan keyakinan apakah individu yang diobservasi merupakan sampel acak yang diambil dari populasi tertentu atau tidak. Jika observasi bukan merupakan sampel acak, maka digunakan model FE. Jika sebaliknya, maka digunakan model RE.
- c) Jika efek individu tidak teramati α_i berkorelasi dengan satu atau lebih variabel bebas, maka estimasi dengan RE bias, sedangkan estimasi dengan FE tidak bias.
- d) Jika N besar T kecil, serta semua asumsi yang disyaratkan oleh model RE terpenuhi, maka estimasi dengan menggunakan RE lebih efisien dibanding estimasi dengan FE.

2. Analisis Regresi

a. Persamaan Regresi Linear Berganda

Penelitian ini menggunakan teknik analisis regresi linear berganda.

Persamaan regresi yang digunakan adalah sebagai berikut:

$$Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \varepsilon$$

Rumus menentukan β_0 , β_1 , dan β_2 :

$$\beta_0 = \bar{Y} - \beta_1 \bar{X}_1 - \beta_2 \bar{X}_2$$

$$\beta_1 = \frac{(\sum x_2^2)(\sum x_1 y) - (\sum x_1 x_2)(\sum x_2 y)}{(\sum x_1^2)(\sum x_2^2) - (\sum x_1 x_2)^2}$$

$$\beta_2 = \frac{(\sum x_1^2)(\sum x_2 y) - (\sum x_1 x_2)(\sum x_1 y)}{(\sum x_1^2)(\sum x_2^2) - (\sum x_1 x_2)^2}$$

Keterangan:

- Y : pengangguran terbuka (variabel terikat)
 β_0 : koefisien titik potong intersep
 β_1 : koefisien regresi jumlah penduduk
 β_2 : koefisien regresi investasi
 X_1 : jumlah penduduk (variabel bebas)
 X_2 : investasi (variabel bebas)
 ε : *error/disturbance* (variabel pengganggu)

b. Uji Hipotesis

1) Uji t

Uji t digunakan untuk mengetahui apakah variabel bebas secara parsial berpengaruh signifikan terhadap variabel terikat. Dengan demikian, bagi setiap nilai koefisien regresi dapat dihitung nilai t-nya. Untuk mencari t_{hitung} dapat dicari dengan menggunakan rumus:

$$t_{hitung} = \frac{r\sqrt{n-2}}{\sqrt{1-r^2}}$$

Hasil yang diperoleh kemudian dibandingkan dengan nilai t_{tabel} dengan taraf signifikansi (α) adalah 0,05 dan derajat kebebasan ($n-K$).

Hipotesis pengujian 1:

$$H_0: \beta_1 = 0$$

$$H_i: \beta_1 \neq 0$$

Kriteria pengujian 1:

- H_0 ditolak jika $t_{hitung} > t_{tabel}$, maka koefisien regresi dikatakan signifikan.
Artinya jumlah penduduk mempunyai pengaruh yang signifikan terhadap pengangguran terbuka.
- H_i diterima jika $t_{hitung} > t_{tabel}$, maka koefisien regresi dikatakan tidak signifikan, artinya jumlah penduduk mempunyai pengaruh yang tidak signifikan terhadap jumlah pengangguran terbuka.

Hipotesis pengujian 2:

$$H_0: \beta_2 = 0$$

$$H_i: \beta_2 \neq 0$$

Kriteria pengujian 2:

- H_0 ditolak jika $t_{hitung} > t_{tabel}$, maka koefisien regresi dikatakan signifikan.
Artinya investasi mempunyai pengaruh yang signifikan terhadap pengangguran terbuka
- H_i diterima jika $t_{hitung} > t_{tabel}$, maka koefisien regresi dikatakan tidak signifikan, artinya investasi mempunyai pengaruh yang tidak signifikan terhadap pengangguran.

2) Uji F

Uji F digunakan untuk menguji apakah variabel independen secara bersama-sama berpengaruh secara signifikan terhadap variabel dependen. Metode yang digunakan dalam uji ini adalah dengan cara membandingkan antara F_{hitung} dengan F_{tabel} atau $F_{(\alpha; n+k-1; nT-n-k)}$ pada tingkat kesalahan 5% dengan hipotesis:

$$H_0: \beta_1 + \beta_2 = 0$$

$$H_1: \beta_1 + \beta_2 \neq 0$$

Hipotesis nol ditolak jika $F_{hitung} > F_{tabel}$, maka seluruh variabel independen berpengaruh signifikan terhadap variabel dependen secara simultan dan sebaliknya. Untuk menguji kedua hipotesis ini digunakan nilai statistik F yang dihitung dengan rumus sebagai berikut:

$$F = \frac{R^2/(k-1)}{(1-R^2)/(n-k)}$$

Keterangan:

R^2 : koefisien determinasi

k : jumlah variabel bebas

n : jumlah data

3) Koefisien Korelasi Berganda

Korelasi berganda (*multiple correlation*) merupakan angka yang menunjukkan arah dan kuatnya hubungan antara dua variabel independen secara bersama-sama atau lebih dengan satu variabel dependen. Rumus korelasi berganda adalah sebagai berikut:

$$r_{y \cdot X_1 X_2} = \sqrt{\frac{r_{yX_1}^2 + r_{yX_2}^2 - 2r_{yX_1}r_{yX_2}r_{X_1X_2}}{1 - r_{X_1X_2}^2}}$$

Keterangan:

$r_{y \cdot X_1 X_2}$: korelasi antara variabel X_1 dengan X_2 secara bersama-sama dengan variabel Y

r_{yX_1} : korelasi *Product Moment* antara X_1 dengan Y

r_{yX_2} : korelasi *Product Moment* antara X_2 dengan Y

$r_{X_1X_2}$: korelasi *Product Moment* antara X_1 dengan X_2

Jadi untuk dapat menghitung korelasi ganda, maka harus dihitung terlebih dahulu korelasi sederhananya melalui korelasi *Product Moment* dengan rumus sebagai berikut:

$$r_{xy} = \frac{N \sum XY - (\sum X)(\sum Y)}{\sqrt{\{N \sum X^2 - (\sum X)^2\}\{N \sum Y^2 - (\sum Y)^2\}}}$$

Keterangan:

r_{xy} : koefisien korelasi *Product Moment*

N : jumlah responden

$\sum X$: jumlah skor dalam sebaran X

$\sum Y$: jumlah skor dalam sebaran

$\sum XY$: jumlah hasil kali dari X dan Y

4) Koefisien Determinasi (R^2)

Koefisien determinasi (R^2) digunakan untuk mengukur seberapa besar kemampuan model dalam menerangkan variasi variabel dependen. Atau dengan kata lain koefisien determinasi mengukur seberapa baik model yang dibuat

mendekati fenomena variabel dependen yang sebenarnya. Besarnya R^2 dihitung dengan rumus:

$$KD = r^2 \times 100\%$$

Keterangan:

KD : nilai koefisien determinasi

100% : pengali yang menyatakan dalam persentase

Penggunaan R^2 sebagai ukuran kelayakan suatu model adalah bahwa R^2 tidak pernah menurun dengan penambahan regresor, sebaliknya justru meningkat. Untuk mengatasi permasalahan ini, suatu instrumen mengukur nilai kebaikan model telah dikembangkan. Ukuran tersebut merupakan modifikasi dari R^2 ini memberikan *penalty* bagi penambahan variabel penjelas yang tidak menurunkan residual secara signifikan. Ukuran tersebut adalah *adjusted* R^2 yang dapat dihitung dengan rumus:

$$R^2 = 1 - (1 - R^2) \frac{nT - 1}{nT - n - k}$$

Keterangan:

R^2 : koefisien determinasi

n : jumlah observasi

T : jumlah waktu

k : banyaknya variabel tanpa intersep

3. Pengujian Asumsi Klasik

Untuk membangun persamaan regresi panel yang terbaik dari kriteria ekonometrika, perlu dilakukan penyelidikan dan penanganan adanya masalah-masalah yang berkaitan dengan pelanggaran asumsi dasar. Berikut ini adalah asumsi-asumsi yang diperlukan dalam analisis regresi:

a. Uji Normalitas

Uji normalitas digunakan untuk mengetahui apakah data yang digunakan berdistribusi normal ataukah tidak. Untuk melakukan pengujian normalitas dapat dihitung dengan melakukan uji Jarque-Bera. Uji Jarque-Bera, yang merupakan uji yang mengukur normalitas data dengan menghitung perbedaan *skewness* dan *kurtosis* data. Adapun formula uji statistik J-B adalah sebagai berikut:

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

dimana S = koefisien *skewness* dan K = koefisien *kurtosis*

Hipotesis

- H_0 : Error berdistribusi normal
- H_1 : Error tidak berdistribusi normal

Statistik pengujian : Jarque-Bera

Alfa pengujian : 5%

Jika hasil perhitungan menunjukkan $p\text{-value Jarque-Bera} > 0,05$ maka H_0 diterima, artinya eror mengikuti fungsi distribusi normal.

b. Uji Multikolinearitas

Multikolinearitas adalah keadaan dimana antara dua variabel independen atau lebih pada model regresi terjadi hubungan linear yang sempurna atau mendekati sempurna. Model regresi yang baik mensyaratkan tidak adanya masalah multikolinearitas.

Untuk mengetahui ada atau tidaknya multikolinearitas dengan melihat nilai *Tolerance and Variance Factor* (VIF). Semakin kecil nilai *Tolerance* dan semakin besar nilai VIF maka akan semakin terjadinya masalah multikolinearitas. Nilai yang dipakai jika nilai *Tolerance* lebih dari 0,1 dari VIF kurang dari 10 maka tidak terjadi multikolinearitas.

- Kriteria pengujian $VIF > 10$, maka terjadi multikolinearitas.
- Kriteria pengujian $VIF < 10$, maka artinya tidak terjadi multikolinearitas.

Sedangkan kriteria pengujian statistic dengan melihat nilai *Tolerance*, yaitu:

- Jika nilai *Tolerance* $< 0,1$ maka artinya terjadi multikolinearitas.
- Jika nilai *Tolerance* $> 0,1$ maka artinya tidak terjadi multikolinearitas.

c. Heteroskedastisitas

Uji Heterokedastisitas bertujuan menguji apakah di dalam model regresi terjadi ketidaksamaan varians dari residual satu pengamatan kepengamatan lain. Jika varians residual dari setiap pengamatan berbeda, berarti telah terjadi heterokedastisitas. Regresi yang baik adalah model regresi yang homoskedastisitas atau tidak terjadi heterokedastisitas. Heteroskedastisitas

muncul karena terjadinya variansi data yang digunakan tidak konstan. Masalah heteroskedastisitas biasanya muncul pada data *cross section* yang sangat heterogen. Uji heteroskedastisitas bertujuan menguji apakah dalam model regresi tidak terjadi ketidaksamaan varians dari residual satu pengamatan ke pengamatan yang lain.

Hipotesis

- H_0 : Varians error bersifat homoskedastisitas
- H_1 : Varians error bersifat heteroskedastisitas

Statistik pengujian : Uji White

Alfa pengujian : 5%

Jika hasil p -value Prob. Chi Square $> 0,05$ maka H_0 diterima, artinya varians error bersifat homoskedastisitas.

untuk asumsi yang dipakai di dalam uji White adalah sebagai berikut:

- Bila nilai Obs*R – Squared (probabilitas) $< \alpha$ (5%) maka ada gejala heteroskedastisitas.
- Bila nilai Obs*R – Squared (probabilitas) $> \alpha$ (5%) maka tidak ada gejala heteroskedastisitas.

d. Autokorelasi

Pengujian autokorelasi digunakan untuk mengetahui apakah dalam suatu model terdapat korelasi antara residual satu observasi dengan residual observasi lainnya. Salah satu cara untuk melihat ada atau tidaknya masalah autokorelasi

dapat dilakukan dengan uji Breusch-Godfrey, dengan alfa pengujian 5%, untuk asumsi yang dipakai di dalam uji Breusch-Godfrey adalah sebagai berikut:³⁸

- Bila nilai $\text{Obs} \cdot R^2$ (probabilitas) $< \alpha$ (5%) maka ada gejala autokorelasi.
- Bila nilai $\text{Obs} \cdot R^2$ (probabilitas) $> \alpha$ (5%) maka tidak ada gejala autokorelasi.

³⁸Wing Wahyu Winarno, *ibid.*, hal .5.31.