



SCAN ME



STATISTIK MULTIVARIAT DALAM RISET

Tim Penulis:

Citra Savitri, Syifa Pramudita Faddila, Irmawartini, Hanif Rani Iswari, Choirul Anam,
Silvana Syah, Sri Rochani Mulyani, Pardomuan Robinson Sihombing,
Early Ridho Kismawadi, Agung Pujiyanto, Awin Mulyati, Yuhana Astuti,
Wahyu Catur Adinugroho, Rinaldi Imanuddin, Kristia, Ani Nuraini, M. Tirtana Siregar.



STATISTIK MULTIVARIAT DALAM RISET

Tim Penulis:

Citra Savitri, Syifa Pramudita Faddila, Irmawartini, Hanif Rani Iswari, Choirul Anam,
Silvana Syah, Sri Rochani Mulyani, Pardomuan Robinson Sihombing,
Early Ridho Kismawadi, Agung Pujiyanto, Awin Mulyati, Yuhana Astuti,
Wahyu Catur Adinugroho, Rinaldi Imanuddin, Kristia, Ani Nuraini, M. Tirtana Siregar.

STATISTIK MULTIVARIAT DALAM RISET

Tim Penulis:

Citra Savitri, Syifa Pramudita Faddila, Irmawartini, Hanif Rani Iswari, Choirul Anam,
Silvana Syah, Sri Rochani Mulyani, Pardomuan Robinson Sihombing,
Early Ridho Kismawadi, Agung Pujiyanto, Awin Mulyati, Yuhana Astuti,
Wahyu Catur Adinugroho, Rinaldi Imanuddin, Kristia, Ani Nuraini, M. Tirtana Siregar.

Desain Cover:

Usman Taufik

Tata Letak:

Handarini Rohana

Editor:

Dr. (c) Iskandar Ahmaddien

ISBN:

978-623-5811-15-4 (PDF)

Cetakan Pertama:

Desember, 2021

Hak Cipta 2021, Pada Penulis

Hak Cipta Dilindungi Oleh Undang-Undang

Copyright © 2021

by Penerbit Widina Bhakti Persada Bandung

All Right Reserved

Dilarang keras menerjemahkan, memfotokopi, atau memperbanyak sebagian atau seluruh isi buku ini tanpa izin tertulis dari Penerbit.

PENERBIT:

WIDINA BHAKTI PERSADA BANDUNG

(Grup CV. Widina Media Utama)

Komplek Puri Melia Asri Blok C3 No. 17 Desa Bojong Emas
Kec. Solokan Jeruk Kabupaten Bandung, Provinsi Jawa Barat

Anggota IKAPI No. 360/JBA/2020

Website: www.penerbitwidina.com

Instagram: @penerbitwidina

PRAKATA

Rasa syukur yang teramat dalam dan tiada kata lain yang patut kami ucapkan selain mengucap rasa syukur. Karena berkat rahmat dan karunia Tuhan Yang Maha Esa, buku yang berjudul “Statistik Multivariat dalam Riset” telah selesai disusun dan berhasil diterbitkan, semoga buku ini dapat memberikan sumbangsih keilmuan dan penambah wawasan bagi siapa saja yang memiliki minat terhadap pembahasan tentang Statistik Multivariat dalam Riset.

Akan tetapi pada akhirnya kami mengakui bahwa tulisan ini terdapat beberapa kekurangan dan jauh dari kata sempurna, sebagaimana pepatah menyebutkan “*tiada gading yang tidak retak*” dan sejatinya kesempurnaan hanyalah milik Tuhan semata. Maka dari itu, kami dengan senang hati secara terbuka untuk menerima berbagai kritik dan saran dari para pembaca sekalian, hal tersebut tentu sangat diperlukan sebagai bagian dari upaya kami untuk terus melakukan perbaikan dan penyempurnaan karya selanjutnya di masa yang akan datang.

Terakhir, ucapan terima kasih kami sampaikan kepada seluruh pihak yang telah mendukung dan turut andil dalam seluruh rangkaian proses penyusunan dan penerbitan buku ini, sehingga buku ini bisa hadir di hadapan sidang pembaca. Semoga buku ini bermanfaat bagi semua pihak dan dapat memberikan kontribusi bagi pembangunan ilmu pengetahuan di Indonesia.

Desember, 2021

Tim Penulis

DAFTAR ISI

PRAKATA.....	iii
DAFTAR ISI	iv
REGRESI LINEAR BERGANDA DENGAN SPSS	1
A. Pendahuluan.....	1
B. Uji Asumsi Klasik	1
C. Pengujian Hipotesis	7
ANALISA DISKRIMINAN DENGAN SPSS.....	15
A. Pengertian Analisis Diskriminan	15
B. Tujuan Analisis Diskriminan	16
C. Persyaratan Analisis Diskriminan	16
D. Fungsi dan Nilai Batas Analisis Diskriminan.....	17
E. Langkah-Langkah Analisis Diskriminan	17
F. Contoh Kasus Analisis Diskriminan	19
G. Langkah Penyelesaian Kasus dengan Analisis Diskriminan dengan Menggunakan Program SPSS	19
STRUCTURAL EQUATIONAL MODELLING DENGAN SMART PLS	33
A. Apa Itu Teknik Structural Equational Modelling?.....	33
B. Kriteria Penilaian dalam Pls-Sem	34
C. Aplikasi SEM dengan Smartpls dalam Bidang Ilmu Sosial dan Ekonomi	36
STRUCTURAL EQUATION MODELLING DENGAN AMOS.....	47
A. Pendahuluan.....	47
B. Diagram Jalur SEM.....	49
C. Langkah-Langkah Analisis dalam SEM	49
D. Kesimpulan	68
STRUCTURAL EQUATION MODELING DENGAN LISREL	69
A. Pengertian SEM	69
B. Tahapan dalam Prosedur SEM.....	71
C. Model Pengukuran (Measurement Model).....	72
D. Model Struktural.....	74
E. Uji Kecocokan Model.....	75
F. Persiapan Analisis Data dengan Lisrel	77
ANALISIS REGRESI DATA PANEL	95
VECTOR AUTOREGRESSION (VAR) DENGAN EVIEWS	113
A. Pembentukan Model VAR	113
B. Uji Kointegrasi.....	131

C. Impuls Response.....	133
D. Variance Decomposition.....	137
ANALISIS KOMPARASI (Uji Independent Sample t Test).....	145
A. Pendahuluan.....	145
B. Tahapan Penelitian.....	146
ANALISIS REGRESI SPASIAL DATA PANEL DENGAN STATA.....	159
A. Konsep Ekonometrika Spasial.....	159
B. Uji Efek Spasial.....	160
C. Matriks Pembobot Spasial.....	160
D. Evaluasi Model.....	160
ANALISA SPASIAL MENGGUNAKAN R STUDIO.....	169
A. Instalasi.....	170
B. Halaman Antarmuka Rstudio.....	171
C. Tipe Data Spasial.....	172
D. Cara Memperoleh dan Import Data Spasial.....	175
E. Aplikasi Analisa Data Spasial Data Demografi.....	180
F. Aplikasi Analisa Data Spasial Data Deforestasi.....	189
ONE-WAY ANALYSIS OF VARIANCE (ANOVA) DENGAN SPSS.....	195
REGRESI LOGISTIK DENGAN STATA.....	207
A. Regresi Logistik.....	207
B. Contoh Hasil Penelitian.....	208
C. Regresi Linier Berganda Panel.....	209
D. Regresi Panel Logit.....	210
E. Model Empiris.....	210
F. Metode Pengukuran dan Prediksi.....	211
G. Langkah-Langkah Regresi Logistik (Tahap II Penelitian).....	211
STRUCTURAL EQUATIONAL MODELLING (SEM) DENGAN AMOS.....	221
A. Pengertian CFA (Confirmatory Factor Analysis).....	221
B. Studi Kasus Second Order.....	222

REGRESI LINEAR BERGANDA DENGAN SPSS

Citra Savitri, S.E., M.M

Syifa Pramudita Faddila, S.ST., M.KM

Universitas Buana Perjuangan Kerawang

A. PENDAHULUAN

Analisis regresi linear berganda memiliki fungsi untuk mencari pengaruh dari dua atau lebih variabel independen (variabel X) terhadap variabel dependen (variabel Y). Adapun persamaan dari analisis regresi linear berganda yaitu:

$$Y = a + b_1X_1 + b_2X_2 + \dots + b_nX_n + \varepsilon$$

Keterangan:

Y = Variabel dependen

a = Nilai konstanta

b_1 = Nilai koefisien regresi variabel X_1

b_2 = Nilai koefisien regresi variabel X_2

b_n = Nilai koefisien regresi variabel X_n

X_1 = Variabel independen ke-1

X_2 = Variabel independen ke-2

X_n = Variabel independen ke-n

ε = Standar error

B. UJI ASUMSI KLASIK

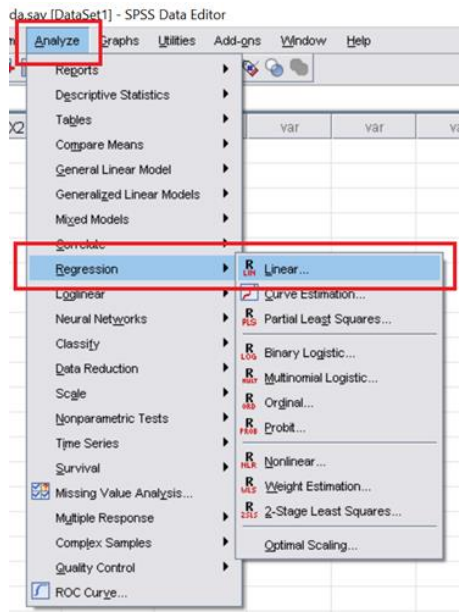
Sebelum melakukan analisis regresi linear berganda, maka beberapa uji asumsi klasik harus terpenuhi terlebih dahulu, mencakup:

1. Uji Normalitas

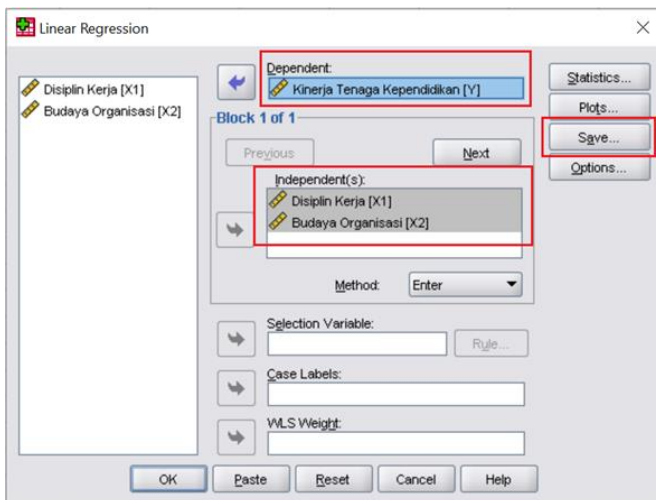
Dilakukan untuk menguji nilai residual yang dihasilkan dari model regresi terdistribusi secara normal atau tidak. Metode uji normalitas yaitu dengan melihat penyebaran data pada sumbu diagonal pada grafik P – P Plot of Regression Standardized Residual atau dengan uji One Sampel Kolmogorov Smirnov. Pengambilan keputusan menggunakan Kolmogorov Smirnov, jika nilai P-value > 0,05 maka model regresi berdistribusi normal.

Langkah Analisis pada SPSS:

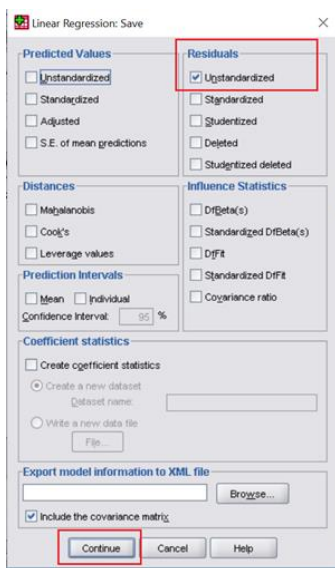
- a. Pada menu utama SPSS, pilih Analyze – Regression – Linear.



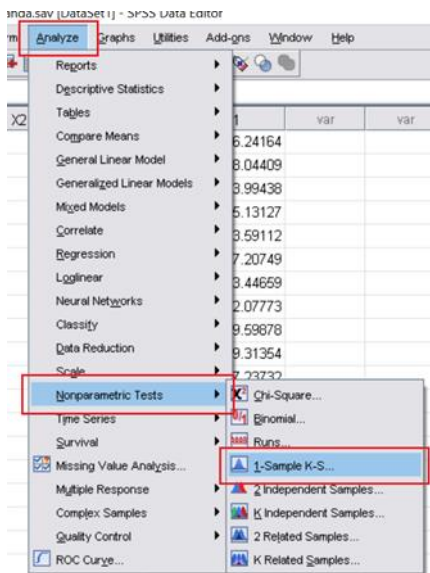
- b. Muncul jendela Linear Regression. Pada kotak Dependent masukan variabel kinerja tenaga kependidikan (variabel Y), sedangkan pada kotak Independent masukan variabel disiplin kerja (X1) dan budaya organisasi (X2).



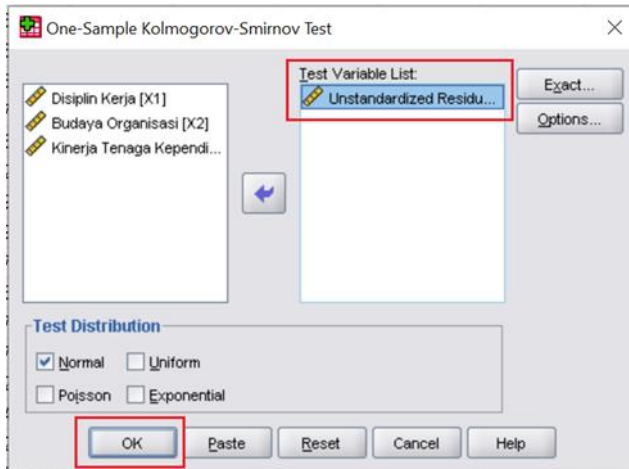
- c. Pilih menu Save, kemudian klik Unstandardized pada Residuals, lalu klik Continue dan OK. Akan muncul data RES_1 pada data view.



- d. Pada menu utama SPSS, pilih Analyze - Non Parametric Test - Legacy Dialogs - 1-Sample-KS.



- e. Masukkan Unstandardized Residual ke dalam kotak Test Variable List. Lalu pilih OK.



2. Uji Multikolinearitas

Bertujuan untuk menguji model regresi ditemukan adanya korelasi antar variabel independen (variabel bebas) atau tidak. Jika hal ini terjadi maka pengaruh variabel independen terhadap variabel dependen akan rendah walaupun nilai F model secara keseluruhan kelihatan tinggi. Hal tersebut akan berakibat H_0 pengujian koefisien akan gagal menolak H_0 walaupun peranan variabel tersebut sebetulnya penting. Oleh karena itu, dalam model regresi yang baik, adalah yang tidak terjadi multikolinearitas. Pengambilan keputusan dalam uji multikolinearitas dengan melihat nilai Tolerance dan VIF sebagai berikut:

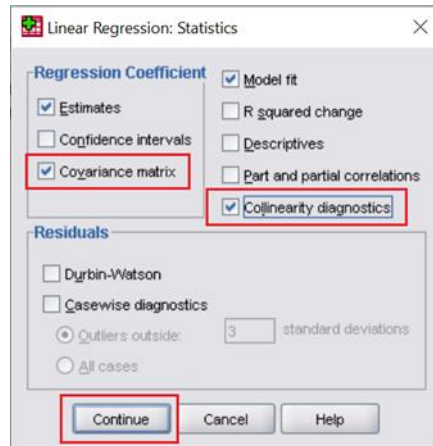
- Nilai Tolerance :
 - 1) Jika nilai Tolerance $> 0,10$ maka tidak terjadi multikolinearitas
 - 2) Jika nilai Tolerance $< 0,10$ maka terjadi multikolinearitas
- Nilai VIF (Variance Inflation Factor) :
 - 1) Jika nilai VIF $< 10,00$ maka tidak terjadi multikolinearitas
 - 2) Jika nilai VIF $> 10,00$ maka terjadi multikolinearitas

Langkah Analisis pada SPSS:

- a. Pada menu utama SPSS, pilih menu Analyze, kemudian pilih sub-menu Regression, klik Linear.
- b. Muncul jendela Linear Regression. Pada kotak Dependent masukan variabel kinerja tenaga kependidikan (variabel Y), sedangkan pada kotak

Independent masukan variabel disiplin kerja (X1) dan budaya organisasi (X2).

- c. Pada kotak Method, pilih Enter.
- d. Pilih menu Statistics, klik Covariance matrix dan Collinearity diagnostics untuk memunculkan hasil korelasi, nilai Tolerance, dan nilai VIF pada output SPSS.



- e. Klik continue, kemudian klik OK.

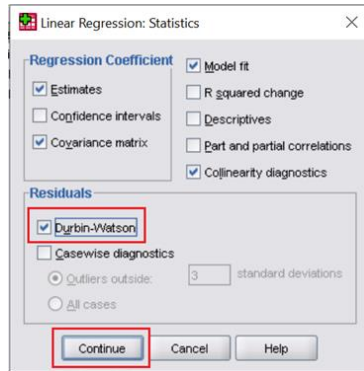
3. Uji Autokorelasi

Bertujuan untuk mengetahui ada tidaknya korelasi antar variabel pengganggu pada periode tertentu dengan variabel sebelumnya. Persamaan regresi berganda yang baik adalah yang tidak terjadi autokorelasi pada model regresinya. Jika terjadi autokorelasi, maka persamaan tersebut menjadi tidak layak untuk digunakan sebagai prediksi. Pengambilan keputusan dalam uji autokorelasi dengan melihat nilai Durbin-Watson (D-W) sebagai berikut:

- 1) Jika nilai D-W < -2 berarti ada autokorelasi positif
- 2) Jika nilai D-W terletak diantara -2 sampai +2 berarti tidak ada autokorelasi
- 3) Bila nilai D-W > +2 berarti ada autokorelasi positif

Langkah Analisis pada SPSS:

- a. Lakukan langkah analisis sesuai uji multikolinearitas.
- b. Kemudian setelah pilih menu Statistics, klik Durbin Watson untuk memunculkan nilai Durbin Watson pada output SPSS.



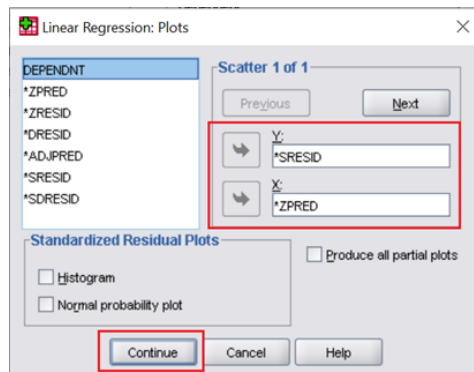
c. Klik continue, kemudian klik OK.

4. Uji Heteroskedastisitas

Heteroskedastisitas adalah keadaan dimana varians dari residual tidak sama untuk satu pengamatan ke pengamatan lainnya. Model regresi yang baik adalah yang tidak terjadi heteroskedastisitas, artinya varians dari residual data harus sama (homoskedastisitas). Salah satu cara untuk mengetahui ada tidaknya heteroskedastisitas pada suatu model regresi linier berganda, yaitu dengan melihat grafik scatterplot atau dari nilai prediksi variabel terikat yaitu SRESID dengan residual error yaitu ZPRED. Heteroskedastisitas terjadi jika titik-titik pada grafik scatterplot mempunyai pola tertentu atau tidak menyebar diatas maupun dibawah angka nol pada sumbu Y.

Langkah Analisis pada SPSS:

- Lakukan langkah analisis sesuai uji multikolinearitas.
- Kemudian pilih menu Plots. Kotak Y diisi dengan SRESID, sedangkan kotak X diisi dengan ZPRED.



c. Klik continue, kemudian klik OK.

C. PENGUJIAN HIPOTESIS

Pengujian hipotesa pada analisis regresi linear berganda, yaitu:

1. Uji Parsial (Uji t)

Pengujian pada masing-masing variabel independent terhadap variabel dependent (X_1 terhadap Y dan X_2 terhadap Y). Hubungan parsial di analisis dengan uji-T.

2. Uji Simultan (Uji F)

Pengujian pada variabel independent secara bersamaan terhadap variabel dependent (X_1 dan X_2 terhadap Y). Hubungan simultan di analisis dengan uji-F.

3. Koefisien Determinasi (R^2)

Untuk mengetahui seberapa jauh model regresi dapat menerangkan variabel dependent (variabel Y). Nilai R square (R^2) yang terdapat pada tabel Model Summary.

Contoh Kasus Analisis Regresi Linear Berganda

Judul Penelitian:

“Pengaruh disiplin kerja dan budaya organisasi terhadap kinerja tenaga kependidikan di Universitas XYZ tahun 2020”

Variabel independent (variabel X_1) = Disiplin Kerja

Variabel independent (variabel X_2) = Budaya Organisasi

Variabel dependent (variabel Y) = Kinerja Tenaga Kependidikan

Hipotesis Penelitian:

- **Pengujian parsial X_1 terhadap Y**

Ho : Tidak ada pengaruh disiplin kerja terhadap kinerja tenaga kependidikan.

Ha : Ada pengaruh disiplin kerja terhadap kinerja tenaga kependidikan.

- **Pengujian parsial X_2 terhadap Y**

Ho : Tidak ada pengaruh budaya organisasi terhadap kinerja tenaga kependidikan.

Ha : Ada pengaruh budaya organisasi terhadap kinerja tenaga kependidikan.

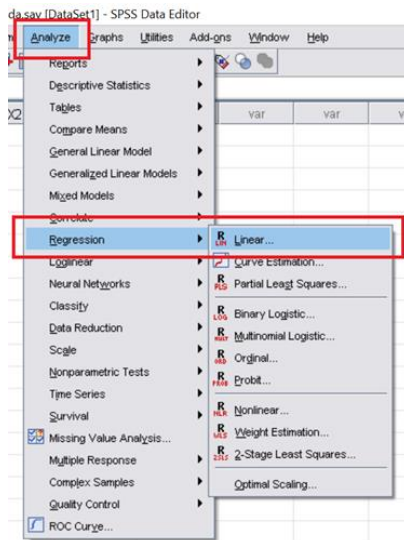
- **Pengujian simultan X_1 dan X_2 terhadap Y**

Ho: Tidak ada pengaruh disiplin kerja dan budaya organisasi terhadap kinerja tenaga kependidikan.

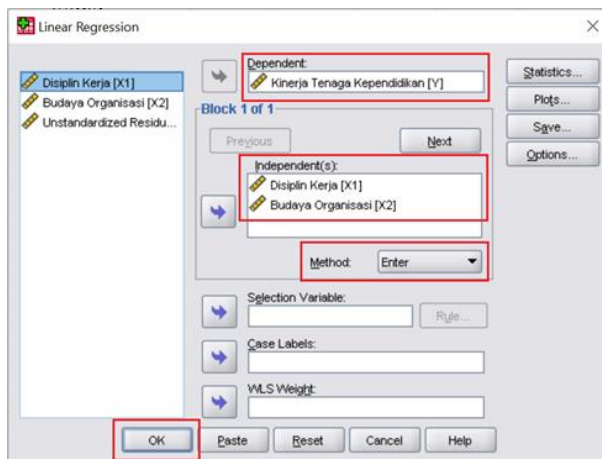
Ha: Ada pengaruh disiplin kerja dan budaya organisasi terhadap kinerja tenaga kependidikan.

Langkah Analisis SPSS

1. Pilih **Analyze – Regression – Linear**



2. Pada kolom **dependent** diisi **variabel dependent** (kinerja tenaga kependidikan) dan kolom **independent** diisi **variabel independent** (disiplin kerja dan budaya organisasi), lalu pada bagian **Method** pilih **Enter**. Kemudian **OK**.



Hasil Output SPSS:

1. Metode: Enter

Variables Entered/Removed^b

Model	Variables Entered	Variables Removed	Method
1	Budaya Organisasi, Disiplin Kerja ^a	.	Enter

a. All requested variables entered.

b. Dependent Variable: Kinerja Tenaga Kependidikan

Menjelaskan tentang variabel yang dianalisis, bahwa kinerja tenaga kependidikan dipengaruhi oleh disiplin kerja dan budaya organisasi. Selain itu menjelaskan juga metode yang digunakan dalam analisis regresi linear berganda, yaitu metode enter. Penggunaan metode enter dengan cara memasukkan semua variabel independen secara bersamaan dalam satu langkah, sering digunakan karena dalam pemodelan regresinya dapat melakukan pertimbangan sesuai substansi.

2. Persamaan Regresi Linear Berganda :

$$Y = a + b_1X_1 + b_2X_2 + \dots + b_nX_n$$

Coefficients^a

Model		Unstandardized Coefficients		Standardized Coefficients	t	Sig.
		B	Std. Error	Beta		
1	(Constant)	37.916	14.832		2.556	.014
	Disiplin Kerja	.278	.155	.267	1.797	.080
	Budaya Organisasi	.137	.175	.117	.784	.438

a. Dependent Variable: Kinerja Tenaga Kependidikan

Berdasarkan hasil output di atas maka model persamaan regresi linear berganda yaitu sebagai berikut:

$$Y = 37,916 + 0,278 \text{ Disiplin Kerja} + 0,137 \text{ Budaya Organisasi}$$

Untuk membuat persamaan tersebut, maka bisa dilihat pada tabel Coefficients pada kolom B di Unstrandardized Coefficients dimana hasilnya adalah:

- a = angka **Constant** pada **Unstrandardized Coefficient**.
 - ✓ Dalam penelitian ini koefisien konstanta sebesar 37,916 dengan nilai positif, ini dapat diartikan bahwa kinerja tenaga kependidikan (variabel Y) akan bernilai 37,916 apabila variabel disiplin kerja (variabel X_1) dan budaya organisasi (variabel X_2) bernilai konstan atau nol.
- b = angka koefisien regresi disiplin kerja dan budaya organisasi pada **Unstrandardized Coefficient**.
 - ✓ Dalam penelitian ini variabel disiplin kerja memiliki koefisien regresi sebesar 0,278. Nilai koefisien regresi positif menunjukkan bahwa jika setiap kenaikan satu satuan variabel disiplin kerja dengan asumsi variabel lain tetap, maka kinerja tenaga kependidikan akan menaik sebesar 0,278.
 - ✓ Dalam penelitian ini variabel budaya organisasi memiliki koefisien regresi sebesar 0,137. Nilai koefisien regresi positif menunjukkan bahwa jika setiap kenaikan satu satuan variabel budaya organisasi dengan asumsi variabel lain tetap, maka kinerja tenaga kependidikan akan menaik sebesar 0,137.

3. **Uji Hipotesis Pada Analisis Regresi Linear Berganda :**

Uji hipotesis atau uji pengaruh berfungsi untuk mengetahui apakah koefisien regresi tersebut signifikan atau tidak (variabel X berpengaruh terhadap variabel Y), maka uji hipotesis untuk analisis linear berganda dilakukan pada 2 hubungan, yaitu:

a. **Pengujian Hipotesis Hubungan Parsial**

- Pengujian parsial X_1 terhadap Y
 - Ho : Tidak ada pengaruh disiplin kerja terhadap kinerja tenaga kependidikan.
 - Ha : Ada pengaruh disiplin kerja terhadap kinerja tenaga kependidikan.
- Pengujian parsial X_2 terhadap Y
 - Ho : Tidak ada pengaruh budaya organisasi terhadap kinerja tenaga kependidikan.
 - Ha : Ada pengaruh budaya organisasi terhadap kinerja tenaga kependidikan.

Membandingkan nilai Sig. (P-value) dengan alpha (0,05)

Coefficients^a

Model		Unstandardized Coefficients		Standardized Coefficients	t	Sig.
		B	Std. Error	Beta		
1	(Constant)	37.916	14.832		2.556	.014
	Disiplin Kerja	.278	.155	.267	1.797	.080
	Budaya Organisasi	.137	.175	.117	.784	.438

a. Dependent Variable: Kinerja Tenaga Kependidikan

Kriteria Keputusan Uji Hipotesis Dari Nilai Sig. :

- Berpengaruh : Jika nilai signifikansi (Sig.) < alpha (0,05)
- Tidak berpengaruh : Jika nilai signifikansi (Sig.) > alpha (0,05)

Hasil dan Kesimpulan :

Berdasarkan hasil analisis didapatkan bahwa disiplin kerja tidak memiliki pengaruh terhadap kinerja tenaga kependidikan (nilai sig > 0,05), sama halnya dengan budaya organisasi juga tidak memiliki pengaruh terhadap kinerja tenaga kependidikan (nilai sig > 0,05).

Membandingkan nilai t-hitung (th) dengan t-tabel (ta)

Coefficients^a

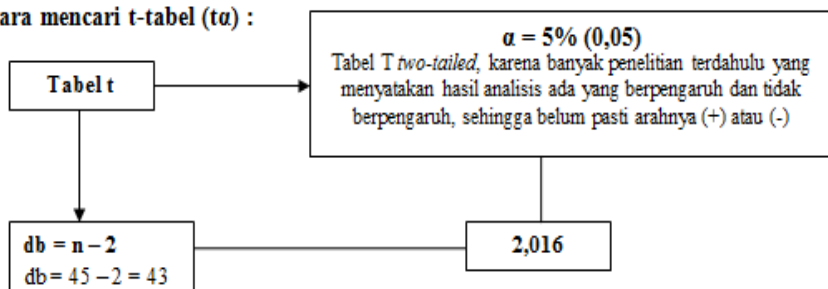
Model		Unstandardized Coefficients		Standardized Coefficients	t	Sig.
		B	Std. Error	Beta		
1	(Constant)	37.916	14.832		2.556	.014
	Disiplin Kerja	.278	.155	.267	1.797	.080
	Budaya Organisasi	.137	.175	.117	.784	.438

a. Dependent Variable: Kinerja Tenaga Kependidikan

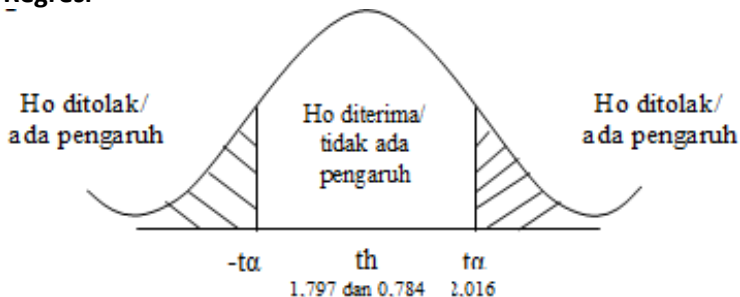
Kriteria Keputusan Uji Hipotesis / Uji t :

- Berpengaruh : Jika nilai t-hitung (th) $\geq$ t-tabel (t_{α})
- Tidak berpengaruh : Jika nilai t-hitung (th) < t-tabel (t_{α})

Cara mencari t-tabel (t_{α}) :



Kurva Regresi



Hasil dan Kesimpulan :

Berdasarkan hasil analisis didapatkan bahwa disiplin kerja tidak memiliki pengaruh terhadap kinerja tenaga kependidikan (nilai $t_h < t_{\alpha}$), begitu juga dengan budaya organisasi pun tidak memiliki pengaruh terhadap kinerja tenaga kependidikan (nilai $t_h < t_{\alpha}$).

b. Pengujian Hipotesis Hubungan Simultan

Pengujian simultan X_1 dan X_2 terhadap Y

Ho: Tidak ada pengaruh disiplin kerja dan budaya organisasi terhadap kinerja tenaga kependidikan.

Ha: Ada pengaruh disiplin kerja dan budaya organisasi terhadap kinerja tenaga kependidikan.

Membandingkan nilai Sig. (P-value) dengan alpha (0,05)

ANOVA^b

Model		Sum of Squares	df	Mean Square	F	Sig.
1	Regression	116.863	2	58.431	1.805	.177 ^a
	Residual	1359.581	42	32.371		
	Total	1476.444	44			

a. Predictors: (Constant), Budaya Organisasi, Disiplin Kerja

b. Dependent Variable: Kinerja Tenaga Kependidikan

Kriteria Keputusan Uji Hipotesis nilai Sig. :

- Berpengaruh : Jika nilai signifikansi (Sig.) $<$ alpha (0,05)
- Tidak berpengaruh : Jika nilai signifikansi (Sig.) $>$ alpha (0,05)

Hasil dan Kesimpulan :

Berdasarkan hasil analisis didapatkan bahwa disiplin kerja dan budaya organisasi secara bersama-sama tidak berpengaruh terhadap kinerja tenaga kependidikan (nilai sig > 0,05).

Membandingkan nilai F-hitung (Fh) dengan F-tabel (Fα)

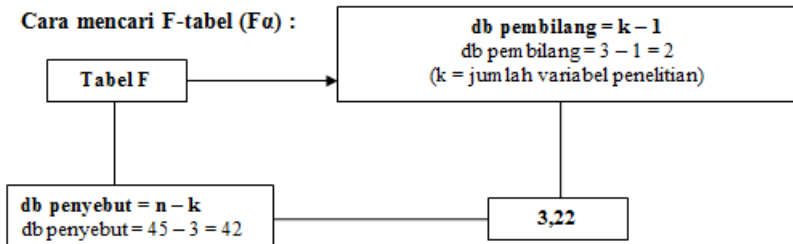
ANOVA ^b					
Model		Sum of Squares	df	Mean Square	F
1	Regression	116.863	2	58.431	1.805
	Residual	1359.581	42	32.371	
	Total	1476.444	44		

a. Predictors: (Constant), Budaya Organisasi, Disiplin Kerja

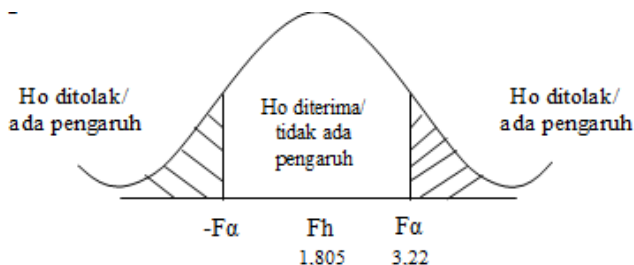
b. Dependent Variable: Kinerja Tenaga Kependidikan

Kriteria Keputusan Uji Hipotesis / Uji F :

- Berpengaruh : Jika nilai F-hitung (Fh) $\geq$ F-tabel (F α)
- Tidak berpengaruh : Jika nilai F-hitung (Fh) < F-tabel (Fα)



Kurva Regresi



Hasil dan Kesimpulan :

Berdasarkan hasil analisis didapatkan bahwa fungsi manajemen dan gaya kepemimpinan secara bersama-sama tidak berpengaruh terhadap kinerja tenaga kependidikan (nilai $F_h < F_\alpha$).

4. Analisis Koefisien Determinasi :

Model Summary				
Model	R	R Square	Adjusted R Square	Std. Error of the Estimate
1	.281 ^a	.079	.035	5.690

a. Predictors: (Constant), Budaya Organisasi, Disiplin Kerja

Dari hasil output diatas, diketahui bahwa nilai R^2 sebesar 0,079. Artinya 7,9% variabel kinerja tenaga kependidikan dapat dijelaskan oleh variasi variabel disiplin kerja dan budaya organisasi. Sedangkan sisanya ($100\% - 7,9\% = 92,1\%$) dijelaskan oleh variabel lain yang tidak diteliti.

DAFTAR PUSTAKA

- Arikunto, Suharsimi. (2010). *Prosedur Penelitian Suatu Pendekatan Praktek*. Jakarta: Rineka Cipta.
- Ghozali, Imam. (2006). *Aplikasi Analisis Multivariate dengan Program SPSS*. Semarang: Badan Penerbit UNDIP.
- Narbuko, Cholid. Abu Achmadi. (2005). *Metodologi Penelitian*. Jakarta: PT. Bumi Aksara.
- Prasetyo, Bambang. Lina Miftahul Jannah. (2006). *Metode Penelitian Kuantitatif: Teori dan Aplikasi*. Jakarta: PT raja Grafindo Persada.
- Priyatno, Duwi. (2009). *SPSS untuk Analisis Korelasi, Regresi, dan Multivariate*. Yogyakarta: Gaya Media.

ANALISA DISKRIMINAN DENGAN SPSS

Dr. Irmawartini, S.Pd., M.K.M
Poltekkes Kemenkes

A. PENGERTIAN ANALISIS DISKRIMINAN

Pemilihan uji statistik dalam menganalisis data penelitian ditentukan oleh banyak faktor, antara lain adalah tujuan penelitian dan skala data dari variabel yang diteliti. Salah satu uji statistik multivariat adalah analisis diskriminan. Analisis diskriminan merupakan metode analisis multivariat dependensi (variabel yang diteliti sudah dikelompokkan menjadi variabel independen dan dependen). Analisis ini akan menghubungkan beberapa variabel independen (lebih dari satu) dengan minimal satu variabel dependen.

Analisis diskriminan akan mengelompokkan data sesuai dengan karakteristik data. Analisis diskriminan menghasilkan kombinasi linier dari variabel independen dan menggunakannya sebagai dasar untuk mengklasifikasikan kasus ke dalam salah satu kelompok yang saling eksklusif/berbeda secara tegas. Analisis diskriminan berusaha membedakan antara beberapa kelompok secara eksklusif. Kombinasi linier yang terbentuk disebut dengan fungsi diskriminan. Perlu dicatat bahwa hanya satu fungsi diskriminan yang diperlukan untuk membedakan antara dua kelompok, dua fungsi diskriminan untuk membedakan antara tiga kelompok (Aljandali, 2017).

Analisis diskriminan digunakan jika tujuan analisis adalah pengelompokan data, dimana analisis diskriminan akan mengklasifikasikan pengamatan-pengamatan ke dalam kelompok-kelompok yang sudah ditentukan. Pengelompokan ini didasarkan kepada skor dari data-data pengamatan tersebut. Data-data pengamatan tersebut haruslah berskala numerik dengan distribusi normal. Tujuan data pengamatan harus berdistribusi normal adalah agar analisis diskriminan menjadi optimal, dimana kesalahan klasifikasi dapat diminimalkan. Jika data variabel independen tidak berdistribusi normal, maka pilihan analisis adalah regresi logistik.

Analisis Diskriminan adalah alat statistik multivariat yang digunakan untuk mengklasifikasikan sekumpulan objek atau individu ke dalam salah satu kelompok yang telah ditentukan. Individu-individu tersebut dikelompokkan berdasarkan lebih dari satu variabel independen. Analisis diskriminan dipakai

untuk menjawab pertanyaan bagaimana individu dapat dimasukkan ke dalam kelompok berdasarkan beberapa variabel. Persamaan fungsi diskriminan yang dihasilkan untuk memberikan prediksi yang paling tepat untuk mengklasifikasi individu ke dalam kelompok pada kategori variabel dependen berdasarkan skor-skor pada variabel independen. Analisis diskriminan digambarkan dengan banyaknya kategori yang terdapat pada variabel dependen.

B. TUJUAN ANALISIS DISKRIMINAN

Tujuan utama penggunaan analisis diskriminan adalah untuk memahami perbedaan kelompok dan untuk memprediksi individu atau objek berada pada kelompok tertentu berdasarkan beberapa variabel independen yang berdata numerik (Joseph F. Hair Jr, William C. Black, Barry J. Babin, 2010). Secara luas tujuan penggunaan analisis diskriminan adalah untuk mendapatkan fungsi diskriminan yang merupakan kombinasi linier dari variabel independen, untuk pengelompokan individu sesuai kelompok pada variabel dependen. Selain itu analisis diskriminan juga bertujuan membuat pengelompokan baru berdasarkan nilai koefisien dari model determinan yang terbentuk. Tujuan selanjutnya adalah menentukan variabel independen yang paling berkontribusi terhadap terbentuknya pengelompokan(diskriminan) pada variabel dependen dan menguji signifikansi pengelompokan variabel dependen berdasarkan variabel independen.

C. PERSYARATAN ANALISIS DISKRIMINAN

Sama dengan penggunaan uji statistik lainnya, analisis diskriminan juga mempunyai persyaratan-persyaratan yang harus dipenuhi agar hasil analisis menjadi tepat. Persyaratan yang harus dipenuhi untuk analisis diskriminan adalah :

- a. Terdiri dari lebih dari satu variabel independen dengan data numerik dan minimal satu variabel dependen berdata kategorik.
- b. Multivariat Normality. Semua data variabel independen harus berdistribusi normal. Jika terdapat data variabel tidak berdistribusi normal, maka akan terjadi ketidaktepatan fungsi diskriminan. Regresi logistik dapat dipilih sebagai alternatif jika data variabel independen tidak berdistribusi normal.
- c. Matrik covarians semua variabel independen harus sama (equal).
- d. Tidak terdapat data yang nilainya sangat ekstrim (outlier). Data ekstrim akan menyebabkan ketidaktepatan fungsi diskriminan dalam pengklasifikasian.

- e. Tidak terdapat multikolinearitas sesama variabel independen. Multikolinearitas disebabkan karena adanya hubungan yang sangat erat antar variabel independen. Multikolinearitas dideteksi dengan melihat nilai korelasi (nilai r). Terdapat multikolinearitas jika nilai r antar variabel independen $>0,8$.

D. FUNGSI DAN NILAI BATAS ANALISIS DISKRIMINAN

Fungsi diskriminan merupakan model kombinasi linier dengan rumus sebagai berikut:

$$Z_{jk} = \alpha + W_1X_{1k} + W_2X_{2k} + \dots + W_nX_{nk}$$

Keterangan:

Z_{jk} = Nilai (skor) diskriminan dari fungsi diskriminan j untuk objek k .

α = Intercept

W_n = Skor diskriminan untuk variabel independen

X_{nk} = Variabel independen n untuk objek

Nilai batas (Cut Of Point) ditentukan dengan rumus sebagai berikut:

$$Z_{CU} = \frac{N_A Z_B + N_B Z_A}{N_A + N_B}$$

Keterangan :

Z_{CU} = Angka kritis, yang menjadi skor cut off

N_A = Jumlah sampel di kelompok A

N_B = Jumlah sampel di kelompok B

Z_A = Angka centroid pada subjek kelompok A

Z_B = Angka centroid pada subjek kelompok B

E. LANGKAH-LANGKAH ANALISIS DISKRIMINAN

Langkah-langkah analisis diskriminan adalah sebagai berikut (Sarma and Vardhan, 2019) & (Idrus, 2013)&(Garson, 2016), (Härdle and Hlavka, 2015):

a. Menentukan variabel independen dan dependen.

Analisis diskriminan termasuk dalam analisis dependensi yaitu analisis statistik yang sudah mengelompokkan variabel menjadi variabel independen dan dependen. Variabel independen berdata numerik dan variabel dependen berdata kategorik.

Berdasarkan jumlah kelompok variabel dependen analisis diskriminan dibedakan menjadi dua, yaitu:

- 1) Analisis diskriminan dua kelompok/kategori (two-group discriminant analysis); dimana variabel dependen bersifat dikotom (dua kelompok) dan akan menghasilkan satu fungsi diskriminan.
- 2) Analisis diskriminan berganda (multiple discriminant analysis); dimana variabel dependen dikelompokkan lebih dari 2 kelompok. Fungsi diskriminan yang terbentuk sebanyak $k-1$, dimana k adalah banyak kelompok/kategori.

b. Menguji kenormalan data variabel indenpenden.

Kenormalan data variabel independen dapat diuji dengan beberapa cara antara lain:

- 1) Membandingkan nilai skewness dengan standar error. Jika nilainya ≤ 2 , maka data berdistribusi normal.
- 2) Menggunakan grafik histogram, normal Q-Q plot, boxplot, steam and leaf dan sejenisnya. Data dikatakan berdistribusi normal jika sebaran data cenderung berada di tengah-tengah baik pada grafik histogram, normal Q-Q plot, boxplot, steam and leaf.
- 3) Menggunakan uji statistik

Ada 2 uji statistik yang bisa digunakan untuk kenormalan data yaitu Uji Kolmogorov-Smirnov atau Shapiro-Wilk. Uji Kolmogorov-Smirnov digunakan untuk ukuran sampel besar (≥ 30), sedangkan uji Shapiro-Wilk digunakan untuk ukuran sampel kecil (< 30).

Data dikatakan berdistribusi normal nilai sig (nilai p) pada uji statistik $> 0,05$.

c. Deteksi multikolinearitas antar variabel independen

Deteksi multikolinearitas variabel independen dilakukan dengan uji korelasi pearson product moment. Nilai $r > 0,8$ mengindikasikan adanya multikolinieritas. Jika terdapat variabel dengan nilai $r > 0,8$, maka variabel tetap dapat dimasukkan pada tahapan analisis selanjutnya, karena secara otomatis analisis diskriminan dapat mendeteksi yaitu adanya test tolerance dan secara otomatis akan mengeluarkan variabel tersebut.

d. Melakukan uji Equity

Asumsi Equity semua variabel independen dapat dilihat dari nilai signifikansi (nilai p) dari Wilk's Lambda . Jika nilai p dari Wilk's Lambda $> 0,05$ berarti asumsi equity terpenuhi. Selain itu Equity dapat dilihat dari nilai signifikansi univariat anova. Jika nilai p $< 0,05$ maka variabel tersebut signifikan dan dapat dilanjutkan sebagai kandidat dalam membuat fungsi diskriminant.

e. Membuat fungsi diskriminan

Fungsi diskriminan dapat dibuat dengan dua metode yaitu Direct method (semua variabel dimasukkan ke dalam model secara bersama-sama untuk membentuk fungsi diskriminan dan dengan metode Stepwise Discriminant Analysis (variabel dimasukkan satu per satu ke dalam model diskriminan).

Fungsi diskriminan dinyatakan signifikan jika nilai p pada uji F $\leq 0,05$. Ketepatan klasifikasi fungsi diskriminan dilakukan uji dengan Casewise Diagnostics. Jika fungsi diskriminan mempunyai ketepatan mengklasifikasi kasus $\geq 50\%$, ketepatan model dianggap tinggi.

F. CONTOH KASUS ANALISIS DISKRIMINAN

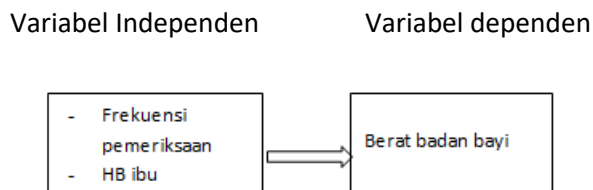
Contoh kasus analisis diskriminan :

Penelitian mengelompokkan berat badan bayi berdasarkan Hb Ibu dan frekuensi kunjungan pemeriksaan kehamilan. Berat badan bayi dikelompokkan menjadi BBLR dan tidak BBLR.

G. LANGKAH PENYELESAIAN KASUS DENGAN ANALISIS DISKRIMINAN DENGAN MENGGUNAKAN PROGRAM SPSS

a. Penentuan variabel independen dan dependen

Variabel independen dan dependen digambarkan dalam kerangka konsep sebagai berikut:

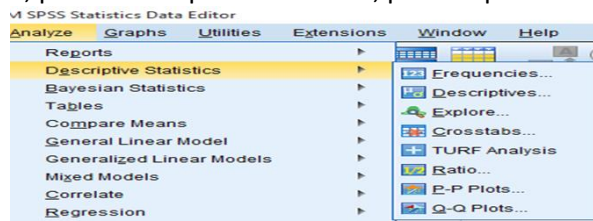


b. Menguji kenormalan data variabel independen

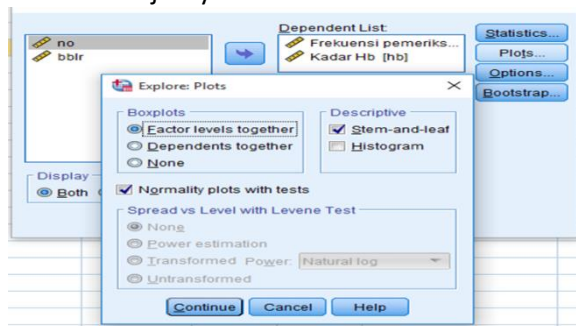
Uji kenormalan data variabel independen

Langkah dengan menggunakan SPSS adalah :

- 1) Klik “Analyze”, pilih “Descriptive Statistics”, pilih “Explore”



- 2) Masukkan variabel Frekuensi pemeriksaan dan Hb ibu ke kotak “dependent list”. Klik “Plots”. Selanjutnya klik “ Normality plots with tests”. Klik “Continue” dan selanjutnya klik “OK”



- 3) Output yang dibaca adalah hasil test Kolmogorov Smirnov pada tabel “Test of Normality”

	Kolmogorov-Smirnov ^a			Shapiro-Wilk		
	Statistic	df	Sig.	Statistic	df	Sig.
Frekuensi pemeriksaan	.134	40	.068	.932	40	.019
Kadar Hb	.131	40	.081	.933	40	.021

a. Lilliefors Significance Correction

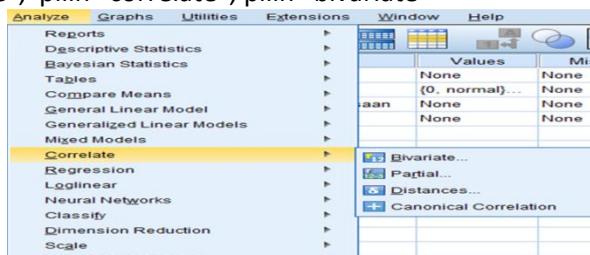
Nilai p hasil uji Kolmogorove Smirnov untuk variabel frekuensi pemeriksaan = 0,068 dan untuk kadar Hb 0,081, berarti $>0,05$. Dengan demikian data kedua variabel berdistribusi normal. Dengan demikian asumsi normality terpenuhi.

c. Deteksi multikolinearitas antar variabel independen

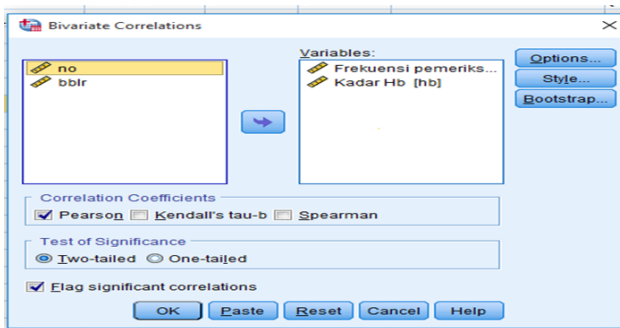
Dilakukan dengan mengkorelasikan sesama variabel independen.

Langkah SPSS nya adalah :

- 1) Klik “analyze”, pilih “correlate”, pilih “bivariate”



- 2) Masukkan semua variabel independen yaitu frekuensi pemeriksaan dan kadar Hb ke dalam kotak “variables”. Pada “correlation coefficients” pilih “Pearson”, klik “ok”



- 3) Outputnya adalah sebagai berikut:

Correlations

		Frekuensi pemeriksaan	Kadar Hb
Frekuensi pemeriksaan	Pearson Correlation	1	.364*
	Sig. (2-tailed)		.021
	N	40	40
Kadar Hb	Pearson Correlation	.364*	1
	Sig. (2-tailed)	.021	
	N	40	40

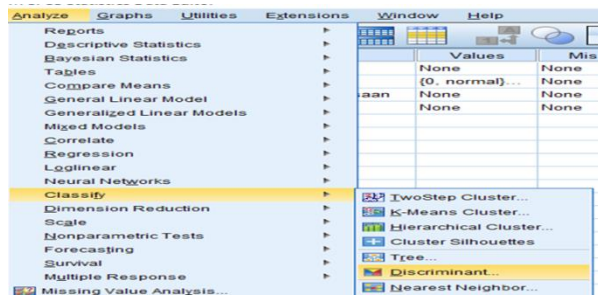
*. Correlation is significant at the 0.05 level (2-tailed).

Terlihat bahwa nilai r antara frekuensi pemeriksaan dengan kaadr Hb adalah $= 0,364$ ($<0,0$), tidak terdapat multikolinieritas antar variabel independen. Dengan demikian asumsi multikolineritas terpenuhi.

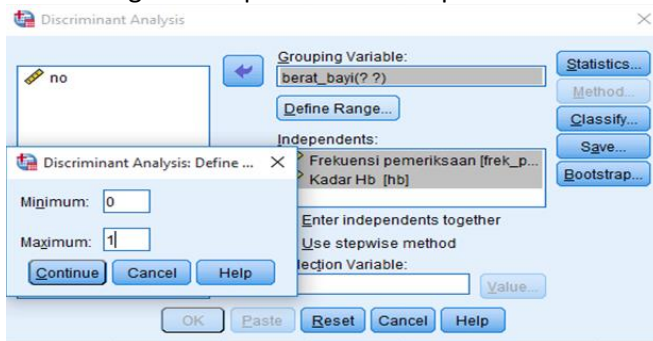
d. Menguji Equity

Pengujian Equity dilihat dari nilai signifikansi Wilk's Lambda dengan langkah-langkah SPSS sebagai berikut:

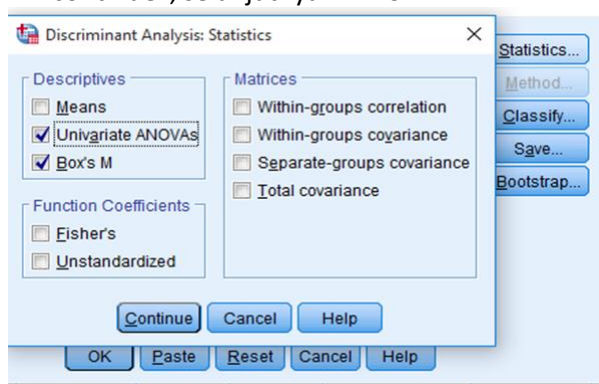
- 1) Klik “Analyze” pilih “classify” dan pilih “discriminat”



- 2) Pada kotak “grouping Variables” masukkan variabel dependen yaitu Berat bayi. Pada kotak “Independents” masukkan variabel independen, yang dalam hal ini adalah frekuensi pemeriksaan dan kadar Hb. Klik “Define Range” Masukkan angka 0 ke “minimum” dan angka 1 ke “maximum”. Angka ini sesuai dengan kode pada variabel dependen.



- 3) Klik “statistics”, pada “descriptives” pilih “Box’s M” dan “Univariate ANOVAs”, klik “continue”, selanjutnya klik “ok”



4) Pada output nilai signifikansi Box's M di baca pada "tabel Test Result"

Test Results		
Box's M		1.148
F	Approx.	.360
	df1	3
	df2	115051.185
	Sig.	.782

Tests null hypothesis of equal
population covariance matrices.

Terlihat nilai p pada Box's M adalah 0,782 ($>0,05$). Dengan demikian asumsi equity terpenuhi. Selain itu juga dilihat dari nilai Wik's Lambda dan nilai signifikansi uji Anov (Uji F)

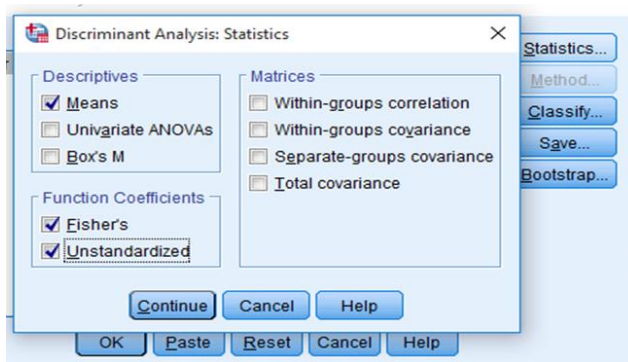
Tests of Equality of Group Means					
	Wilks' Lambda	F	df1	df2	Sig.
Frekuensi pemeriksaan	.678	18.088	1	38	.000
Kadar Hb	.637	21.619	1	38	.000

Terlihat bahwa semua variabel mempunyai nilai Wik's Lambda diatas 0,05 dan nilai p pada uji F $< 0,05$. Dengan demikian asumsi equity sudah terpenuhi

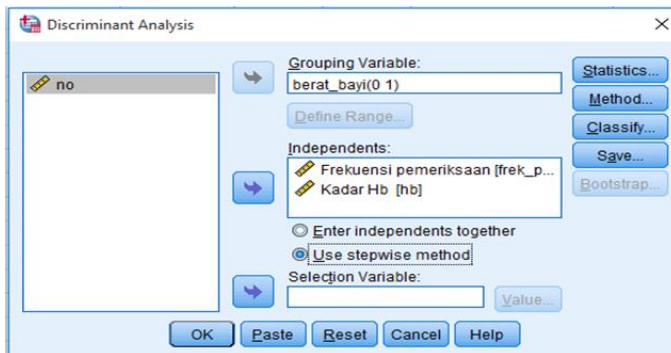
e. Membuat Fungsi diskriminan

Langkah-langkah SPSS nya adalah sebagai berikut:

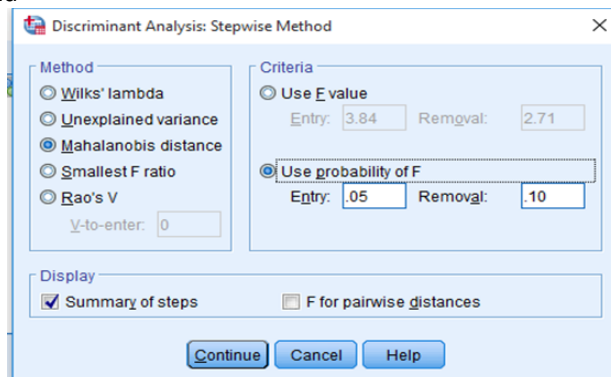
- 1) Pilih "Analyze", pilih "Classify", lalu pilih "Discriminant"
- 2) Pada kotak " Grouping variable" masukkan variabel dependen dan pada kotak "independents, masukkan semua variabel independen
- 3) Klik "Define Range" kemudian masukkan kode 0 dan 1, lalu klik "continue" maka akan kembali ke menu utama
- 4) Pilih "Statistics", Pada "descriptives" pilih "means". Pada bagian "Function Coefficients" pilih "Fisher's dan "Unstandardized". Klik "Continue"



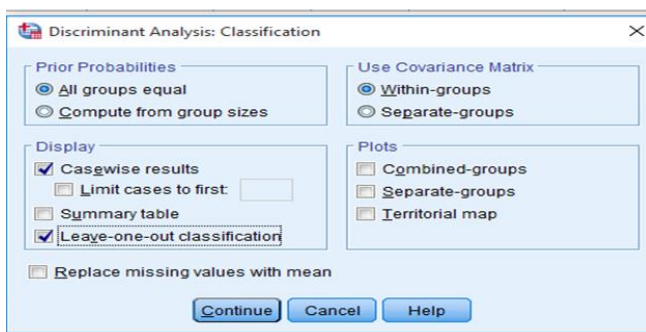
- 5) Pada menu utama klik “use stepwise method” sehingga menu “method” menjadi aktif



- 6) Klik “Method” pilih “Mahalanobis distance” dan pada bagian “criteria” pilih “Use probability of F” dan selanjutnya klik “Continue” sehingga kembali ke menu utama



- 7) Klik "Classify". Pada bagian "Display" pilih "Casewise results" dan "Leave-one-out classification". Selanjutnya klik "Continue" dan "ok"



- 8) Output-output yang diperhatikan adalah :

- a) Tabel Group Statistik untuk emlihat deskripsi dari variabel

Group Statistics

berat_bayi		Mean	Std. Deviation	Valid N (listwise)	
				Unweighted	Weighted
normal	Frekuensi pemeriksaan	4.8696	1.65980	23	23.000
	Kadar Hb	12.0652	.89727	23	23.000
bblr	Frekuensi pemeriksaan	2.7059	1.49016	17	17.000
	Kadar Hb	10.8353	.71932	17	17.000
Total	Frekuensi pemeriksaan	3.9500	1.90748	40	40.000
	Kadar Hb	11.5425	1.02254	40	40.000

Interpretasi :

1. Terlihat terdapat perbedaan rata-rata antara dua kelompok yaitu antara berat badan bayi normal dengan BBLR.
2. Jumlah sampel sebanyak 40 yang terdiri dari 23 bayi dengan berat badan normal dan 17 bayi dengan BBLR.
3. Tidak terdapat missing data (data yang hilang/tidak lengkap) sehingga untuk proses selanjutnya jumlah sampel adalah 40 .

- b) Tabel Variables Entered/Removed untuk melihat variabel-variabel yang masuk ke dalam fungsi diskriminan.

Variables Entered/Removed ^{a,b,c,d}							
Step	Entered	Statistic	Between Groups	Min. D Squared			
				Statistic	df1	df2	Sig.
1	Kadar Hb	2.212	normal and bblr	21.619	1	38.000	0.039
2	Frekuensi pemeriksaan	3.931	normal and bblr	18.708	2	37.000	0.02

Data pada Tabel “Variables Entered/Removed” menunjukkan variabel-variabel yang masuk ke dalam fungsi diskriminan yaitu variabel dengan nilai $p \leq 0,05$. Dengan demikian semua variabel independen masuk ke dalam fungsi diskriminan.

- c) Tabel Variables in the Analysis untuk melihat variabel-variabel yang masuk ke dalam analisis diskriminan dalam setiap tahapan

Variables in the Analysis				
Step	Tolerance	Sig. of F to Remove	Min. D Squared	Between Groups
1 Kadar Hb	1.000	.000		
2 Kadar Hb	.999	.001	1.850	normal and bblr
Frekuensi pemeriksaan	.999	.003	2.212	normal and bblr

Pada tahap 1: variabel Kadar Hb masuk analisis diskriminan karena mempunyai nilai $p < 0,05$. Pada tahap 2 variabel kadar Hb dan frekuensi pemeriksaan juga masuk ke dalam analisis diskriminan karena mempunyai nilai $p < 0,05$.

- d) Tabel Eigenvalues untuk melihat keeratan hubungan skor diskriminan dengan group

Eigenvalues				
Function	Eigenvalue	% of Variance	Cumulative %	Canonical Correlation
1	1.011 ^a	100.0	100.0	.709

a. First 1 canonical discriminant functions were used in the analysis.

Dilihat dari nilai Canonical Correlation = 0,709. Keeratan hubungan berkisar dari 0 sampai 1. Nilai 0,709 menunjukkan hubungan yang sangat erat. Kuadrat dari nilai Canonical Correlation, menunjukkan seberapa besar variabel mampu menjelaskan variasi pembeda antara berat badan bayi normal dengan BBLR. $0,709^2 = 0,5027 = 50,27\%$. Dengan demikian fungsi diskriminan ini mampu menjelaskan variasi pembeda antara berat badan bayi normal dengan BBLR sebesar 50,27%.

- e) Tabel Wilks' Lambda untuk mengetahui signifikansi dari fungsi diskriminan

Wilks' Lambda				
Test of Function(s)	Wilks' Lambda	Chi-square	df	Sig.
1	.497	25.854	2	.000

Nilai Wilks' Lambda menunjukkan signifikansi dari fungsi diskriminan. Diperoleh nilai $p < 0,05$, dengan demikian fungsi diskriminan yang terbentuk signifikan sebagai pembeda.

- f) Tabel Canonical Discriminant Coefficients untuk mengetahui koefisien determinan

**Canonical Discriminant Function
Coefficients**

	Function 1
frekuensi pemeriksaan kehamilan	.416
Kadar Hb	.880
(Constant)	-11.803

Unstandardized coefficients

Tabel ini menunjukkan fungsi diskriminan yang terbentuk. Sesuai dengan rumus fungsi diskriminan, maka fungsi diskriminan yang terbentuk adalah

$$D = -11,803 + (0,416 \times \text{frekuensi pemeriksaan kehamilan}) + (0,880 \times \text{Kadar HB})$$

Guna dari fungsi ini adalah untuk mengetahui individu yang dalam data ini adalah ibu hamil, akan masuk kelompok yang mana apakah kelompok berat bayi normal atau BBLR.

- g) Tabel Prior Probabilitas for Groups untuk mengetahui komposisi responden.

Prior Probabilities for Groups

berat_bayi	Prior	Cases Used in Analysis	
		Unweighted	Weighted
normal	.500	23	23.000
bblr	.500	17	17.000
Total	1.000	40	40.000

Tabel ini menunjukkan komposisi dari responden yang berjumlah 40. Sebanyak 23 responden berada pada kelompok berat bayi normal dan sebanyak 17 responden berada pada kelompok BBLR.

- h) Tabel Function of Group Centroid untuk mengetahui nilai centroid dalam menghitung cut off

**Functions at Group
Centroids**

	Function
berat_bayi	1
normal	.843
bblr	-1.140

Unstandardized
canonical discriminant
functions evaluated at
group means

Tabel ini menunjukkan centroid dari setiap group dimana berat bayi normal centroidnya 0,843 dan BBLR mempunyai centroid -1,140. Nilai centroid ini berguna untuk membuat cut off dari masing-masing kelompok.

Berdasarkan rumus cut off maka nilai cut off dari data ini adalah:

$$Z_{CU} = \frac{N_A Z_B + N_B Z_A}{N_A + N_B}$$

$$Z_{CU} = \frac{23 \times 0,843 + 17 \times -1,140}{23 + 17}$$

$$= 19,389 - 19,38/40 = 0$$

Berarti pada data ini cut off nya adalah 0

Penggunaannya adalah sebagai berikut

1. Jika skor dari fungsi diskriminan $> Z_{CU}$, maka individu tersebut akan masuk kelompok normal (kode 0)
2. Jika skor dari fungsi diskriminan $< Z_{CU}$, maka individu tersebut akan masuk kelompok BBLR (kode 1)

Contoh pemakaian fungsi diskriminan dan cut off

Seorang ibu hamil dengan kadar Hb 12 dan melakukan pemeriksaan kehamilan sebanyak 5 kali. Data tersebut dimasukkan ke dalam fungsi diskriminan sehingga di peroleh nilai fungsi diskriminan sebagai berikut:

$$\begin{aligned} D &= -11,803 + 0,416 \times 5 + 0,880 \times 12 \\ &= -11,803 + 2,08 + 10,56 \\ &= 0,837 \end{aligned}$$

Perhitungan sebelumnya diperoleh cut off (Z_{CU}) = 0, maka skor fungsi diskriminan lebih besar dari Z_{CU} . Sehingga kecenderungan ibu hamil tersebut akan melahirkan bayi dengan berat badan lahir normal (kode 1).

i) Tabel Classification Results untuk menguji ketepatan diskriminan

Classification Results^{a,c}

		Predicted Group Membership			
		berat_bayi	normal	bblr	Total
Original	Count	normal	19	4	23
		bblr	4	13	17
	%	normal	82.6	17.4	100.0
		bblr	23.5	76.5	100.0
Cross-validated ^b	Count	normal	19	4	23
		bblr	4	13	17
	%	normal	82.6	17.4	100.0
		bblr	23.5	76.5	100.0

a. 80.0% of original grouped cases correctly classified.

b. Cross validation is done only for those cases in the analysis. In cross validation, each case is classified by the functions derived from all cases other than that case.

c. 80.0% of cross-validated grouped cases correctly classified.

Dari output diatas, bahwa fungsi diskriminan yang dihasilkan dapat memprediksi berat bayi lahir sebesar 80,0%. Untuk memperhitungkan kemungkinan berbagai bias dilakukan uji kekuatan prediksi dengan metode Leave-one-out-cross validation, dan diperoleh hasil 80,0%. Terlihat bahwa

ketepatan prediksi dari fungsi yang dihasilkan tinggi, sehingga fungsi ini dapat digunakan untuk memprediksi apakah ibu hamil masuk kelompok melahirkan bayi dengan berat normal atau BBLR.

DAFTAR PUSTAKA

- Aljandali, A. (2017) *Multivariate Methods and Forecasting with IBM® SPSS® Statistics*. Edited by D. Ruppert et al. London: Springer International Publishing. doi: 10.1007/978-3-319-56481-4.
- Garson, G. D. (2016) *Discriminat Function Analysis*. USA: Statistical Publishing Association.
- Härdle, W. K. and Hlávka, Z. (2015) *Multivariate statistics: Exercises and solutions, second edition*. Edited by S. Edition. Germany: Springer Heidelberg. doi: 10.1007/978-3-642-36005-3.
- Idrus, M. S. (2013) *Multivariate Data Analisis Dan Non Parametrik Statistik Untuk Penelitian Bidang Manajemen*. Sidoarjo: Zifatama.
- Joseph F. Hair Jr, William C. Black, Barry J.Babin, R. E. A. (2010) *Multivariate Data Analysis.pdf*. Sevent h e. Pearson Prentice Hall.
- Sarma, K. V. S. and Vardhan, R. V. (2019) *Multivariate Statistics Made Simple*. London: CRC Press. doi: 10.1201/9780429465185.

STRUCTURAL EQUATIONAL MODELLING DENGAN SMART PLS

Hanif Rani Iswari, S.E., M.M., CAPF., CAPM., CIBA., CERA., CNPHRP

Choirul Anam, S.E., M.M

Universitas Widyagama Malang

A. APA ITU TEKNIK STRUCTURAL EQUATIONAL MODELLING?

Teknik Structural Equational Modelling adalah model yang menjelaskan hubungan antara variabel terukur dan variabel laten, serta hubungan antar variabel laten. Teknik ini tergolong dalam teknik analisis statistik multivariat yang digunakan untuk menganalisis hubungan struktural; yakni kombinasi dari analisis faktor dan analisis regresi berganda.

Ada beberapa pendekatan berbeda untuk SEM:

- Pendekatan pertama adalah SEM berbasis komponen yang dikenal sebagai Generalized Structured Component Analysis (GSCA); itu diimplementasikan melalui VisualGSCA atau aplikasi berbasis web yang disebut GeSCA.
- Pendekatan kedua untuk melakukan SEM lainnya adalah dengan Nonlinear Universal Structural Relational Modeling (NEUSREL), menggunakan perangkat lunak Causal Analytics NEUSREL.
- Pendekatan ketiga adalah SEM berbasis Covariance yang diterapkan secara luas (CB-SEM)6, menggunakan paket perangkat lunak seperti AMOS, EQS, LISREL dan MPLus.
- Cara lainnya adalah Partial Least Squares (PLS), yang berfokus pada analisis varians. Dengan demikian PLS mampu menganalisis variabel laten, variabel indikator, kesalahan pengukuran secara langsung. Oleh karenanya PLS dapat menyelesaikan regresi berganda apabila terjadi permasalahan spesifik terhadap data, seperti adanya data yang hilang (missing values), ukuran sampel penelitian kecil dan multikolinearitas (Abdillah dan Jogiyanto, 2015). Analisis ini bisa diterapkan untuk seluruh skala data, sehingga tidak perlu banyak asumsi dan untuk ukuran sampelnya tidak harus besar. PLS juga dapat digunakan untuk pemodelan struktural dengan indikator yang bersifat reflektif (indikator sebagai refleksi variasi variabel laten) maupun formatif (variabel laten sebagai kombinasi indikator-indikatornya) (Ghozali dan Latan, 2015). Software yang digunakan dalam

PLS bermacam-macam yakni PLS-Graph, VisualPLS, SmartPLS, dan WarpPLS. Ini juga dapat digunakan menggunakan modul PLS dalam paket perangkat lunak statistik "r".

- Dihadapkan dengan berbagai pendekatan untuk pemodelan jalur, yang harus dipertimbangkan adalah kelebihan dan kekurangan untuk memilih pendekatan yang sesuai. Untuk kesempatan kali ini akan dibahas lebih mendalam SEM dengan pendekatan Partial Least Squares (PLS) menggunakan SmartPLS.

B. KRITERIA PENILAIAN DALAM PLS-SEM

Dalam PLS, model hubungan dapat diasumsikan bahwa variabel laten dan indikator atau manifes variable di skala zero means dan unit variance (nilai standardized) sehingga parameter lokasi (konstanta) dapat dihilangkan dalam model tanpa mempengaruhi nilai generalisasi. Teknik parametrik untuk menguji signifikansi parameter tidak diperlukan karena PLS tidak menghasilkan adanya distribusi tertentu untuk estimasi parameter (Chin, et al, 2010; Ghazali dan Latan, 2015). Kriteria penilaian model dalam PLS-SEM dapat dilihat pada Tabel 1

Tabel 1 Parameter Model Pengukuran PLS – Outer Model

Pengujian	Parameter	Rule of Thumb
Evaluasi Model Pengukuran Refleksif		
<i>Convergent Validity</i>	Loading Factor	• >0.70 untuk Confrimatory Research • >0.60 untuk Exploratory Research
	Average Variance Extracted (AVE)	• >0.50 untuk Confrimatory maupun Exploratory
	Communalitiy	• >0.50 untuk Confrimatory maupun Exploratory Research
<i>Discriminant Validity</i>	Cross Loading	• >0.70 untuk setiap variabel
	$\sqrt{\text{AVE}}$ dan korelasi antar Konstruk laten	• $\sqrt{\text{AVE}} >$ korelasi antar konstruk variabel
Reliabilitas	<i>Composite Reability</i>	• > 0.70 untuk Confrimatory Research • > 0.60 masih dapat diterima untuk Exploratory Research

	<i>Cronbach's Alpha</i>	• >0.70 untuk Confrimatory Research • 0.60-0.70 masih diterima untuk Exploratory Research
--	-------------------------	---

Evaluasi Model Pengukuran Formatif		
	Signifikansi Weight	• >1,65 (significance level = 10%) • >1.96 (significance level = 5%) • >2.58 (significance level = 1%)
	Multicollinearity	• VIF < 10 atau < 5 • Tolerance > 0,10 atau > 0,20
	<i>R-Square</i> (R^2)	0,67 = model kuat, 0,33 = model moderate, 0,19 = model lemah (Chin, 1998) 0,75 = model kuat, 0,50 = model moderate, 0,25 = model lemah (Hair, et al., 2011)
	<i>Effect Size</i> (f^2)	0,02 = kecil, 0,15 = menengah, 0,35 = besar
	<i>Q-Square</i> (Q^2) <i>predictive relevance</i>	• Q^2 > model mempunyai <i>predictive relevance</i> • Q^2 < model kurang memiliki <i>predictive relevance</i>
	q^2 <i>predictive relevance</i>	• 0,02 = lemah • 0,15 = moderate • 0,35 = kuat
	Signifikansi (two-tailed)	t-value • 1,65 (significance level = 10%) • 1.96 (significance level = 5%) • 2,58 (significance level = 1%)

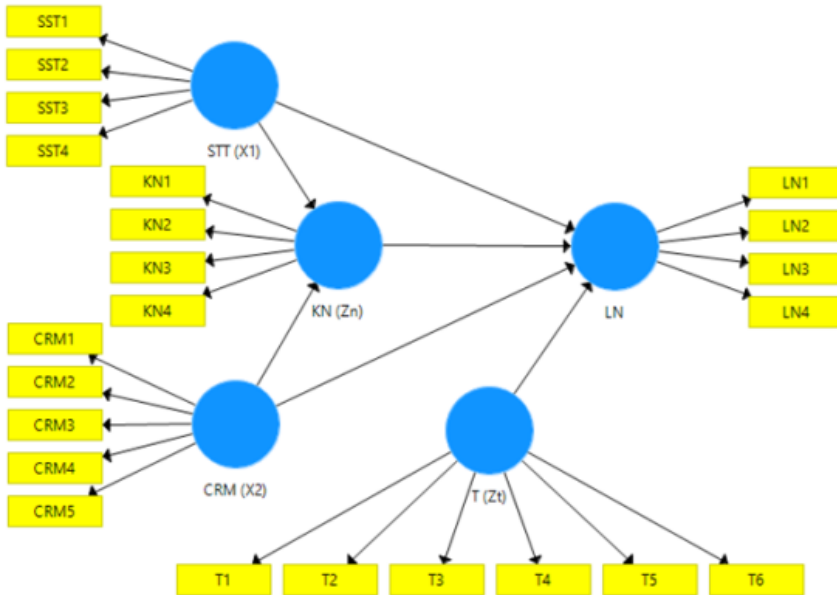
Sumber: Ghazali dan Latan (2015)

C. APLIKASI SEM DENGAN SMARTPLS DALAM BIDANG ILMU SOSIAL DAN EKONOMI

Permasalahan ilmu sosial, ekonomi semakin kompleks permasalahannya karena meningkatnya berbagai faktor yang perlu diambil dalam pengambilan keputusan. Oleh karenanya pendekatan yang cocok digunakan adalah Teknik SEM dengan menggunakan SmartPLS. Adapun aplikasi yang dilakukan dalam sebuah kasus terdiri atas kasus di bidang ilmu ekonomi dengan data primer dan data sekunder.

1. Kasus di bidang ilmu ekonomi dengan data primer

Menggunakan sebuah contoh struktur penelitian pemasaran dengan data primer dengan judul: **Pengaruh Self Service Technology dan Customer Relationship Marketing Terhadap Kepuasan Nasabah dan Loyalitas Nasabah dengan Kepercayaan Nasabah Sebagai Moderasi Pada Nasabah Bank X.** Adapun kerangka konsep atas judul tersebut dapat dilihat pada gambar 1:



Gambar 1 Kerangka Konsep Penelitian

Dari kerangka konsep tersebut berikut langkah-langkah pengolahan menggunakan software SmartPLS:

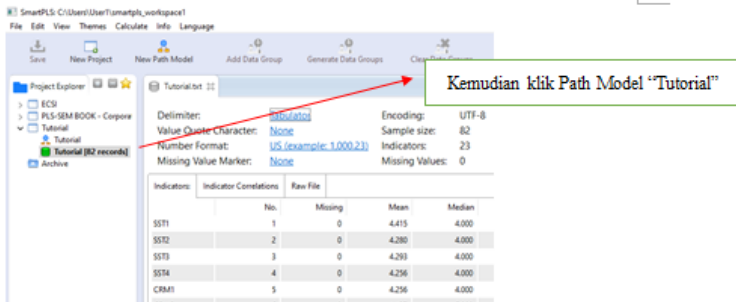
2. Proses Analisis Data

- 1) Menyiapkan data dengan cara merubah format dari *xls copy – paste menjadi *txt. Seperti nampak pada gambar dibawah ini.

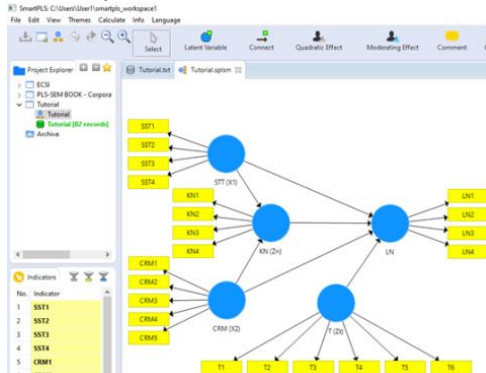
Gambar 2 File data_mediasi dan moderasi.xls dengan responden 82 orang

Gambar 3 File data_mediasi dan moderasi.xls dengan responden 82 orang

4) Selanjutnya, muncul seperti gambar dibawah ini :



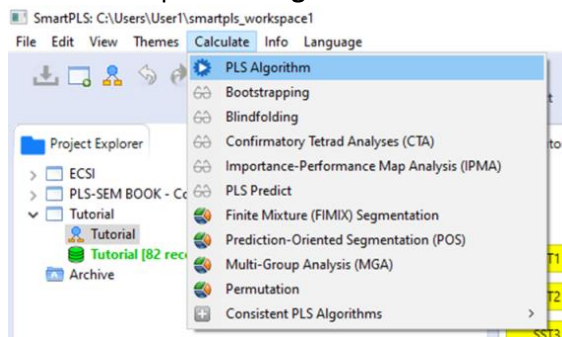
5) Maka, muncul gambar seperti ini :



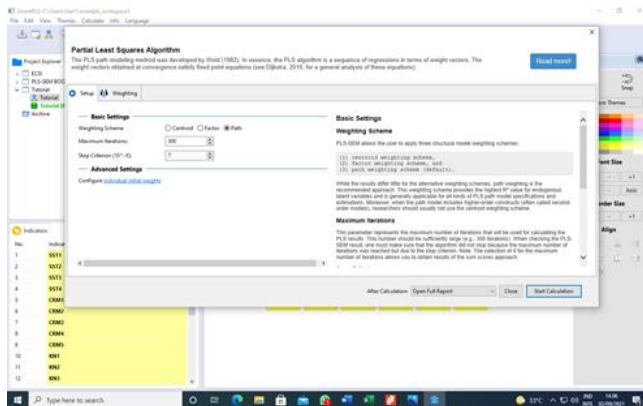
Masing-masing variabel memiliki item, jadi indikator (item) yang terdapat disamping, kemudian drag ke variabel sesuai dengan jumlah notasi.

OUTER MODEL

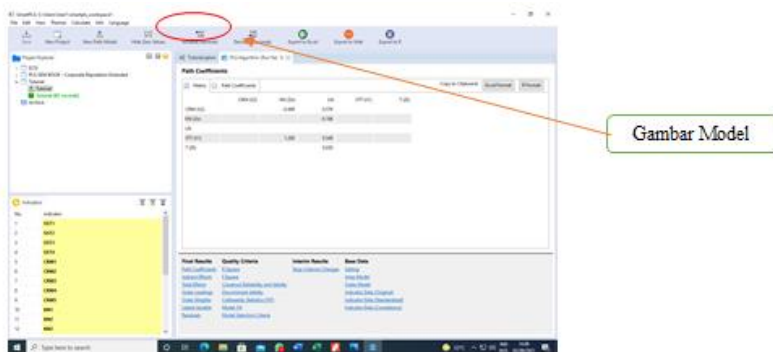
- Langkah 1: klik **calculate** > pilih **PLS Algorithm**



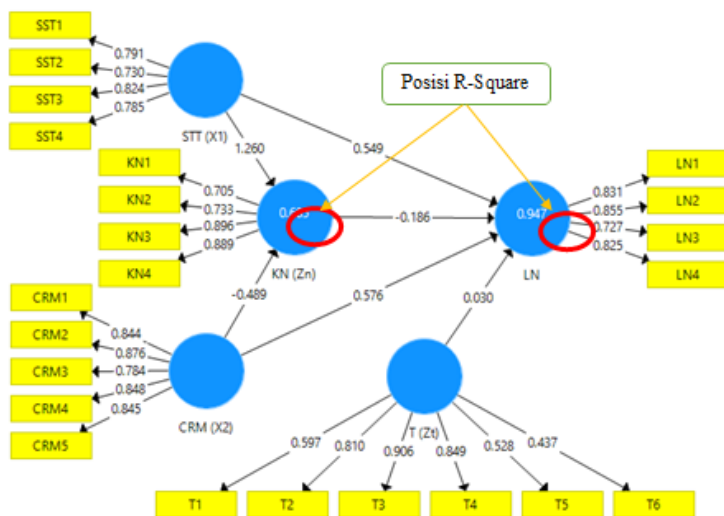
- Langkah 2: Selanjutnya, akan muncul seperti nampak digambar, dan pilih **Start Calculation**



- **Langkah 3:** Maka akan muncul hasil analisis dari **PLS Algorithm**, seperti gambar berikut:



Untuk tampilan model gambar PLS Algorithm dapat dipilih tab sebelahny
Bagian outer model yang dilihat adalah convergent validity, discriminant validity, composite validity, dan cronbach alpha.



Gambar Hasil Outer Model

Dari gambar tersebut untuk mengetahui nilai dari convergent validity, discriminant validity, composite validity, dan cronbach alpha dapat dilihat pada Final Result dan Quality Criteria melalui tampilan berikut :

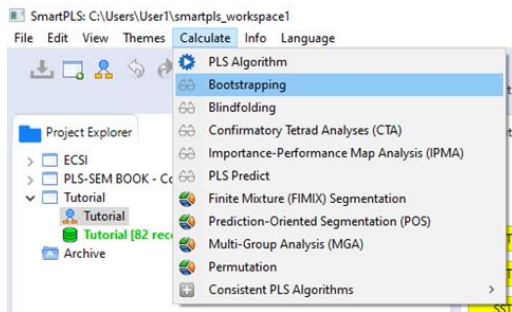
Final Results	Quality Criteria	Interpretation
Path Coefficients	R Square	Menunjukkan hasil <i>composite validity</i> , dan <i>cronbach alpha</i>
Indirect Effects	f Square	
Total Effects	Construct Reliability and Validity	Menunjukkan hasil <i>discriminant validity</i>
Outer Loadings	Discriminant Validity	
Outer Weights	Collinearity Statistics (VIF)	Menunjukkan hasil <i>convergent validity</i>
Latent Variable	Model Fit	
Residuals	Model Selection Criteria	

Catatan:

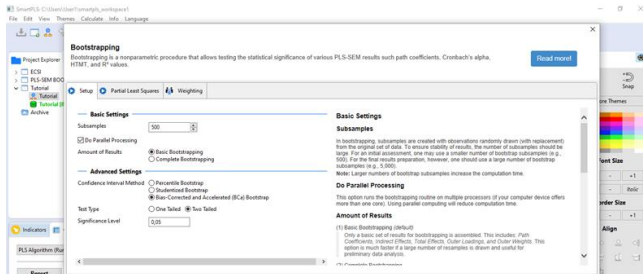
- ❖ Outer Weights diperuntukkan jika variabel yang diteliti memiliki indikator yang bersifat formatif.
- ❖ R-Square: menunjukkan kuat atau lemahnya pengaruh yang ditimbulkan oleh variabel dependen terhadap variabel independen.
- ❖ Gambar model dapat disimpan dengan cara: pilih tab gambar model > Klik File > Export as Image to File.

INNER MODEL

- **Langkah 1:** klik **calculate** > pilih **Bootstrapping**



- **Langkah 2:** Selanjutnya, akan muncul seperti nampak digambar, dan pilih **Start Calculation**

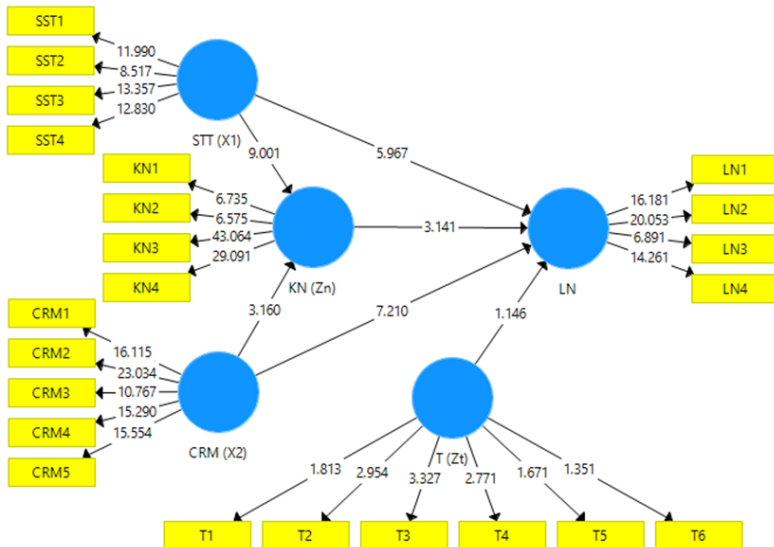


- **Langkah 3:** Maka akan muncul hasil analisis dari **Bootstrapping**, seperti gambar berikut:



Untuk tampilan model gambar **Bootstrapping** dapat dipilih tab sebelahnya.

Bagian outer model yang dilihat adalah direct effect, indirect effect, dan moderating effect.



Diketahui bahwa lingkaran merah menunjukkan hasil kausalitas (sebab-akibat) secara langsung (direct) memiliki tingkat signifikansi diatas 1,96, maka dapat dikatakan signifikan. Sedangkan pada indirect effect dapat ditunjukkan pada lingkaran merah pada gambar berikut:

Final Results	Histograms	Base Data
Path Coefficients	Path Coefficients Histogram	Setting
Total Indirect Effects	Indirect Effects Histogram	Inner Model
Specific Indirect Effects	Total Effects Histogram	Outer Model
Total Effects		Indicator Data (Original)
Outer Loadings		Indicator Data (Standardized)
Outer Weights		

Tutorial.splsm

PLS Algorithm (Run No. 1)

Bootstrapping (Run No. 1)

Specific Indirect Effects

Mean, STDEV, T-Values, P-Values

Confidence Intervals

Confidence Intervals Bias Corrected

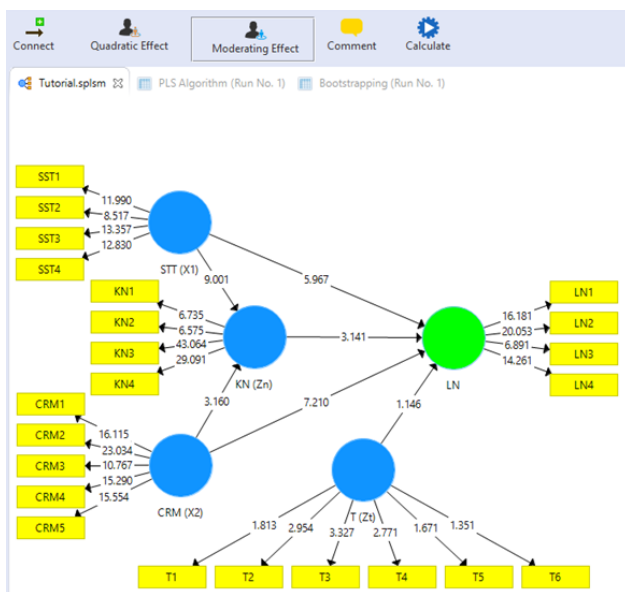
Samples

	Original Sampl...	Sample Mean (...)	Standard Devia...	T Statistics (O/...	P Values
CRM (X2) -> KN (Zn) -> LN	0.091	0.106	0.046	1.973	0.049
STT (X1) -> KN (Zn) -> LN	-0.234	-0.263	0.074	3.155	0.002

Pada hasil secara tidak langsung (indirect effect) menunjukkan hasil kausalitas (sebab-akibat) memiliki tingkat signifikansi diatas 1,96, maka dapat dikatakan signifikan.

Moderating Effect

- Langkah 1: klik **Moderating Effect**, pilih variabel yang dihitung nilai moderating nya, seperti pada gambar berikut:



Selanjutnya, keluar menu seperti nampak digambar :

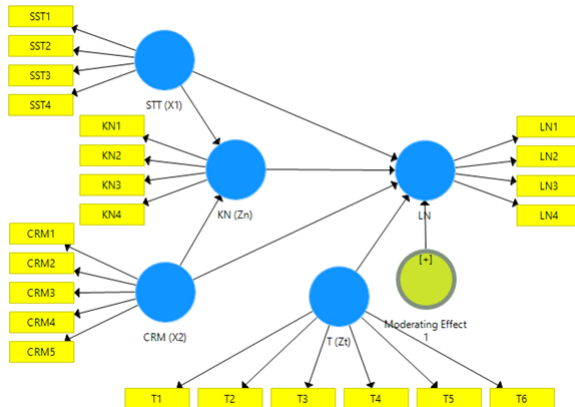
The 'Moderating Effect' dialog box is shown with the following settings:

- Basic Settings:**
 - Dependent Variable: LN
 - Moderator Variable: KN (Zn)
 - Independent Variable: (circled in red)
 - Calculation Method: Two Stage (selected)
- Advanced Settings:**
 - Product Term Generation: Mean Centered (selected)
 - Weighing Mode: Automatic (selected)

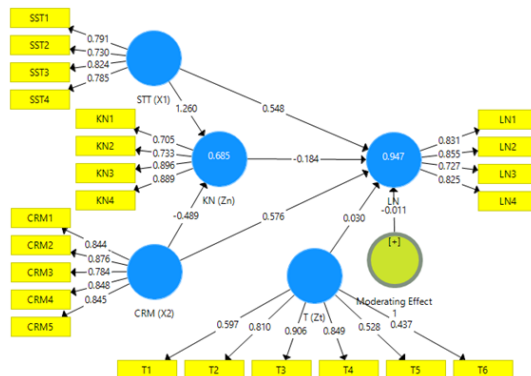
Annotations in the image:

- Pilih variabel moderasi dengan notasi Zt (points to the Independent Variable field)
- Pilih variabel independen dengan notasi Zn (points to the Calculation Method field)

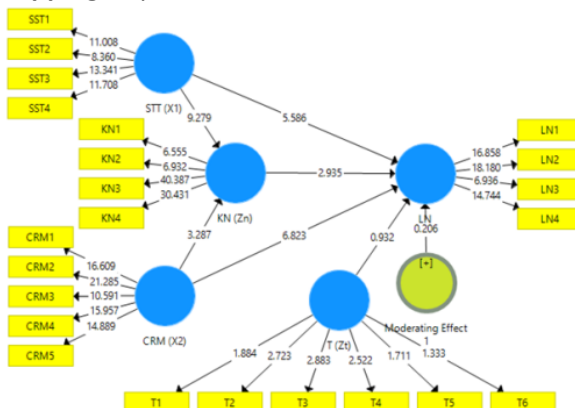
Setelah selesai, klik OK, maka hasil gambar seperti dibawah:



Proses analisis hampir sama dengan klik PLS Algorithm dan Bootstrapping, namun disini yang dicari adalah moderating effectnya, maka di **PLS Algorithm**, seperti berikut:



Untuk **Bootstrapping**, seperti berikut:



Diketahui bahwa moderating effect yang ditunjukkan pada variabel KN Terhadap LN dimoderasi oleh T tidak signifikan, hal ini karena efek moderasi yang ditimbulkan lebih kecil dari 1,96.

Secara singkat hasil penafsiran, pengujian secara langsung maupun tidak langsung serta pengujian moderasi dapat di asumsikan sebagai berikut:

Pada uji inner model, mulai dari convergent validity dan discriminant validity memiliki nilai diatas 0,7 sehingga dapat dikatakan valid, hanya tiga item dari variabel T yang memiliki nilai dibawah 0,7. Lebih lanjut, composite validity, dan cronbach alpha memiliki nilai diatas 0,7 sehingga dapat dikatakan reliabel. R-Square yang diperoleh variabel KN adalah 0,685 (68,5%) artinya bahwa variabel SST dan CRM memiliki pengaruh pada variabel KN sebesar 68,5% sisanya dipengaruhi variabel lain, begitu juga variabel LN sebesar 0,947 (94,7%) artinya bahwa variabel SST, CRM, KN memiliki pengaruh pada variabel LN sebesar 94,7%, sisanya dipengaruhi variabel lain.

Pengujian hipotesis, secara langsung diketahui bahwa variabel SST, CRM, dan KN memiliki pengaruh signifikan terhadap LN, hal ini dibuktikan dengan t-statistik >1,96. Lantas, secara tidak langsung diketahui bahwa KN memiliki peran mediasi yang bersifat partial mediation pada pengaruh SST terhadap LN melalui KN dan pengaruh CRM terhadap LN melalui KN, hal ini dibuktikan dengan t-statistik >1,96. Selanjutnya, pada uji moderasi diketahui bahwa variabel T belum mampu menjadi penguat dari pengaruh KN terhadap LN, hal ini dibuktikan dengan t-statistik <1,96.

3. Kasus di bidang ilmu ekonomi dengan data sekunder (diupdate lebih lanjut)

DAFTAR PUSTAKA

- Ghozali, Imam. (2006). *Structural Equation Modeling Metode Alternatif dengan Partial Least Square (PLS)*. Semarang: Badan Penerbit Universitas Diponegoro.
- Hussein, A.S. (2015). *Penelitian Bisnis dan Manajemen Menggunakan Partial Least Square (PLS) dengan smartPLS 3.0*. Fakultas Ekonomi dan Bisnis Universitas Brawijaya.
- Indriantoro, Nur. Bambang, Supomo. (2002). *Metodologi Penelitian Bisnis, Cetakan Kedua*. Yogyakarta;: Penerbit BFEE UGM.
- Iqbaria, M. Zinatelli. (1997). *Personal Computing Acceptance Factors in Small Firm: A Structural Equation Modelling*. *MIS Quarterly*. 21 (3).
- Ringle, C. M. Wende, S., dan Becker, J.-M. (2015). "SmartPLS 3." Boenningstedt: *SmartPLS GmbH*. <http://www.smartpls.com>.
- Willy, Abdilah. dan Hartono, Jogiyanto. (2015). *Partial Least Square (PLS) – Alternatif Structural Equation Modeling (SEM) dalam penelitian bisnis*. Yogyakarta: Penerbit Andi.

STRUCTURAL EQUATION MODELLING DENGAN AMOS

Silvana Syah, S.Si, M.Si

Institut Bisnis Dan Informatika Kosgoro 1957

Salah satu alasan mengapa kita perlu mempelajari dan memahami metode analisis data, salah satunya SEM ini adalah agar hasil penelitian kita atau informasi yang kita peroleh adalah valid. Salah satu aspek dari validitas penelitian khususnya validitas internal penelitian adalah terkait dengan pemilihan metode analisis data yang tepat dan aplikasinya secara benar. Kalau kita memilihnya tidak tepat, itu otomatis hasil penelitiannya atau informasi yang kita peroleh adalah tidak valid. Jika memilih metodenya benar atau tepat tapi aplikasinya tidak benar maka akan menghasilkan informasi yang bias.

Statistika itu terkait dua hal yaitu metode pengumpulan data dan metode analisis data. Sasarannya untuk mendapatkan informasi catatan informasi ini valid dan akurat. Kemudian juga kita cermati dimana peran statistika? Tentunya kalau kita menelisik tentang jenis-jenis pendekatan penelitian maka statistika itu adalah merupakan salah satu metode khusus pada bidang kuantitatif. Pada saat penelitian itu kualitatif maka manfaat atau peranan statistika sangat kecil bahkan kadangkadang tidak diperlukan sama sekali. Itu pun pada saat kita melakukan penelitian empiris. Tapi jika penelitian kita mix method maka statistika akan berperan. Tantangan bagi seorang statistikawan adalah bagaimana statistika juga bisa berperan pada bidang normatif, khususnya SEM ini. Ini adalah tantangan yang menarik.

Faktor-Faktor yang Mempengaruhi Dan Dampaknya pada Kinerja Selling-In
(Studi pada Outlet Binaan PT. Indosat Semarang)

A. PENDAHULUAN

Dalam era globalisasi tantangan yang harus dihadapi oleh PT Indosat akan semakin berat, tidak hanya bertujuan untuk dapat survive melainkan harus mampu memiliki keunggulan bersaing dibandingkan dengan perusahaan lain. Michman dalam Wahyudi (2002) berpendapat bahwa perusahaan mempunyai keterbatasan-keterbatasan dalam menjual produknya sehingga diperlukan

perantara sebagai saluran distribusi untuk menjangkau konsumen akhir. Selling-in merupakan kegiatan distribusi yang diarahkan pada upaya untuk melakukan penjualan pada semua pedagang perantara sehingga mempermudah pencapaian suatu tingkat market coverage yang optimal, yaitu menggunakan perantara outlet untuk menjangkau konsumen akhir (Ferdinand, 2000).

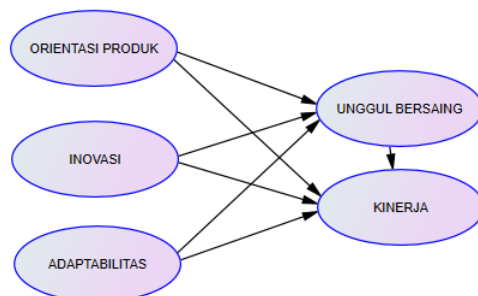
Dari penelitian-penelitian sebelumnya, terdapat research gap yaitu menurut Mustafa (2005) bahwa kinerja selling-in mempengaruhi keunggulan bersaing, tetapi menurut Rahmad Rialdi (2010) bahwa keunggulan bersaing yang mempengaruhi kinerja suatu perusahaan. Penelitian Rahmat Rialdi juga didukung oleh Fengki Octora Kurniawan (2005) yang menyatakan bahwa keunggulan bersaing mempengaruhi kinerja penjualan. Sedangkan menurut Asa, Ismeth dan Latief (2008) menyatakan bahwa diferensiasi tidak mempengaruhi keunggulan bersaing bila produknya merupakan produk standar. Pendapat ini berbeda dengan Fengki Octora Kurniawan (2005) yang menyatakan bahwa adanya hubungan positif antara diferensiasi produk terhadap keunggulan bersaing produk.

Penelitian ini bertujuan menguji pengaruh orientasi produk, inovasi dan adaptabilitas terhadap keunggulan bersaing. Penelitian ini juga menguji pengaruh keunggulan bersaing terhadap kinerja selling-in di outlet-outlet binaan PT. Indosat Semarang.

PERTANYAAN PENELITIAN

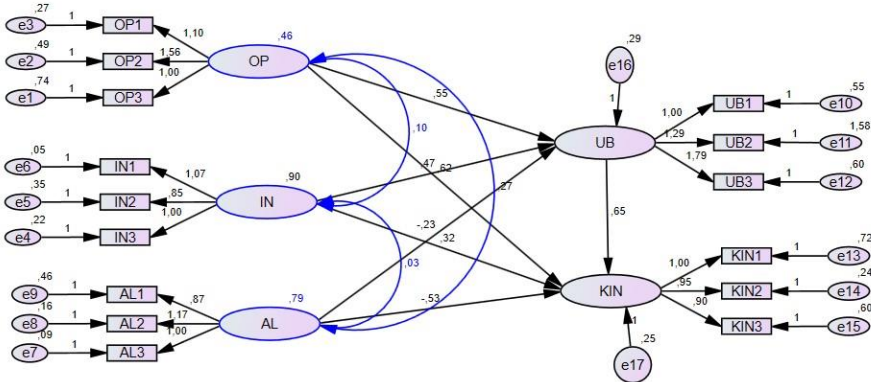
1. : Apakah ada pengaruh Orientasi Produk terhadap Keunggulan Bersaing
2. : Apakah ada pengaruh Inovasi terhadap Keunggulan Bersaing
3. : Apakah ada Pengaruh Adaptabilitas terhadap Keunggulan Bersaing
4. : Apakah ada pengaruh Keunggulan Bersaing terhadap Kinerja
5. : Apakah ada pengaruh Orientasi Produk terhadap Kinerja
6. : Apakah ada pengaruh Inovasi terhadap Kinerja
7. : Apakah ada Pengaruh Adaptabilitas terhadap Kinerja

Model Faktor-faktor yang mempengaruhi Keunggulan Bersaing Serta Dampaknya Terhadap Kinerja



B. DIAGRAM JALUR SEM

Diagram jalur SEM berfungsi untuk menunjukkan pola hubungan antar variabel yang kita teliti. Dalam SEM pola hubungan antar variabel akan diisi dengan variabel yang diobservasi, variabel laten dan indikator. Di bawah ini Faktor-faktor yang mempengaruhi Keunggulan Bersaing Serta Dampaknya Terhadap Kinerja



Gambar 1 Model Persamaan Struktural Faktor-faktor yang mempengaruhi Keunggulan Bersaing Serta Dampaknya Terhadap Kinerja

C. LANGKAH-LANGKAH ANALISIS DALAM SEM

Pada gambar 1 terdapat 3 variabel eksogen nya dan ada 1 intervening dan variabel endogen nya 1. Masing-masing memiliki 3 konstruk atau 3 pengukuran jika ditotalkan semuanya ada 15 konstruk. Model ini judulnya adalah “Faktor-faktor yang mempengaruhi keunggulan bersaing serta dampaknya terhadap kinerja” nanti juga akan kita tampilkan goodness-of-fit nya di dalam gambar kita.

Untuk melakukan analisis SEM diperlukan langkah-langkah sebagai berikut:

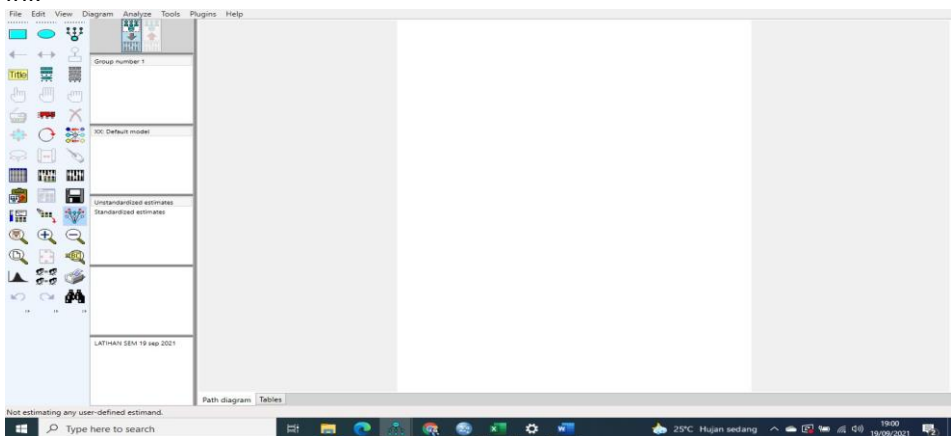
- **Langkah Pertama,**

Untuk mengolah data menggunakan Amos ini ada dua cara yaitu : datanya bisa disajikan di Excel atau juga data bisa disajikan di dalam SPSS.

	OP1	OP2	OP3	IN1	IN2	IN3	AL1	AL2	AL3	UB1	UB2	UB3	KN1	KN2	KN3
330	4.00	3.00	3.00	3.00	3.00	3.00	3.00	3.00	3.00	3.00	3.00	4.00	3.00	3.00	3.00
331	5.00	3.00	4.00	3.00	3.00	3.00	3.00	4.00	3.00	3.00	3.00	3.00	3.00	3.00	4.00
332	5.00	4.00	3.00	3.00	3.00	3.00	3.00	3.00	3.00	2.00	3.00	3.00	3.00	2.00	3.00
333	3.00	3.00	3.00	3.00	3.00	3.00	3.00	3.00	3.00	3.00	4.00	4.00	3.00	3.00	3.00
334	3.00	3.00	3.00	4.00	4.00	3.00	3.00	3.00	3.00	3.00	3.00	4.00	3.00	3.00	3.00
335	3.00	3.00	3.00	3.00	3.00	3.00	3.00	3.00	3.00	3.00	3.00	4.00	3.00	3.00	4.00
336	3.00	3.00	3.00	3.00	3.00	3.00	4.00	3.00	4.00	4.00	4.00	4.00	4.00	4.00	3.00
337	3.00	3.00	3.00	3.00	3.00	3.00	4.00	3.00	3.00	4.00	4.00	4.00	4.00	4.00	4.00
338	3.00	4.00	3.00	3.00	4.00	3.00	3.00	3.00	3.00	3.00	3.00	4.00	3.00	3.00	4.00
339	3.00	3.00	3.00	3.00	3.00	3.00	3.00	3.00	3.00	3.00	4.00	4.00	3.00	3.00	4.00
340	5.00	3.00	4.00	3.00	3.00	3.00	4.00	4.00	4.00	5.00	5.00	3.00	4.00	4.00	4.00
341	3.00	4.00	3.00	4.00	4.00	3.00	4.00	3.00	3.00	3.00	3.00	3.00	4.00	3.00	3.00
342	3.00	3.00	3.00	3.00	3.00	3.00	3.00	3.00	3.00	3.00	4.00	3.00	3.00	3.00	4.00
343	3.00	4.00	3.00	4.00	4.00	3.00	4.00	3.00	3.00	4.00	4.00	3.00	4.00	3.00	3.00
344	3.00	4.00	3.00	4.00	4.00	3.00	4.00	3.00	3.00	3.00	3.00	3.00	4.00	3.00	3.00
345	2.00	3.00	3.00	3.00	3.00	3.00	3.00	3.00	3.00	3.00	3.00	2.00	3.00	3.00	3.00
346	3.00	4.00	3.00	4.00	4.00	3.00	4.00	3.00	3.00	4.00	4.00	3.00	4.00	4.00	3.00
347	3.00	3.00	3.00	3.00	3.00	3.00	3.00	3.00	3.00	3.00	4.00	3.00	3.00	3.00	3.00
348	3.00	4.00	4.00	4.00	3.00	4.00	3.00	3.00	3.00	3.00	3.00	3.00	3.00	3.00	3.00
349	3.00	3.00	3.00	4.00	4.00	3.00	3.00	3.00	3.00	3.00	3.00	3.00	3.00	3.00	3.00
350	5.00	4.00	3.00	3.00	3.00	3.00	4.00	3.00	3.00	4.00	4.00	3.00	4.00	4.00	4.00
351															
352															

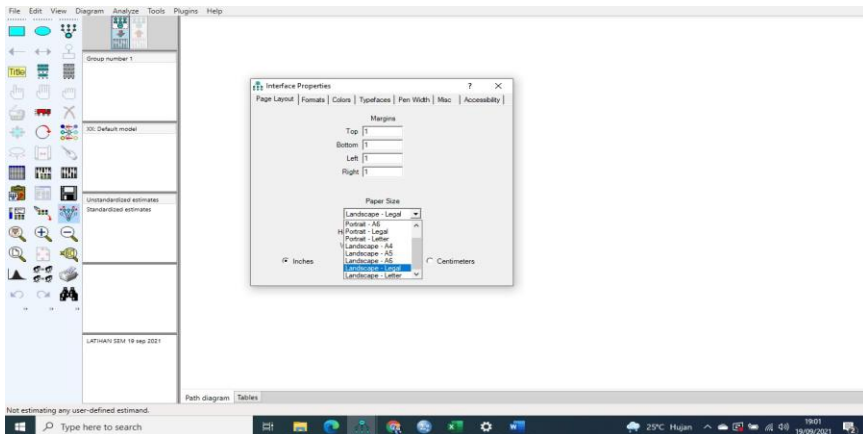
Gambar 2 Data SEM yang disimpan dalam SPSS

- **Langkah Kedua,**
Masuk ke software Amos nya standar. Software Amos tampilannya seperti ini.



Gambar 3 Tampilan standar Amos

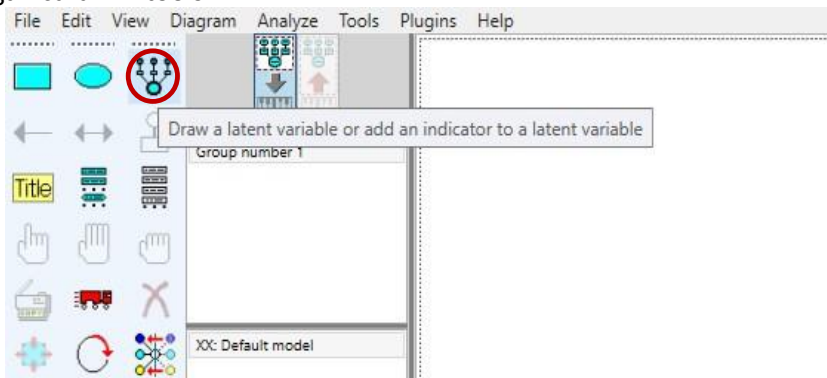
Namun banyak yang suka bekerja bidang gambarnya landscape. Maka ganti tampilannya dengan cara klik view → interface → paper size → pilih landscape legal → silang/close → tampilan akan berubah menjadi gambar 4



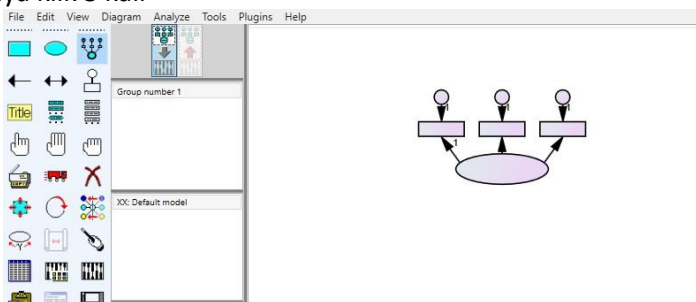
Gambar 4 Tampilan Landscape Legal

- **Langkah Ketiga,**

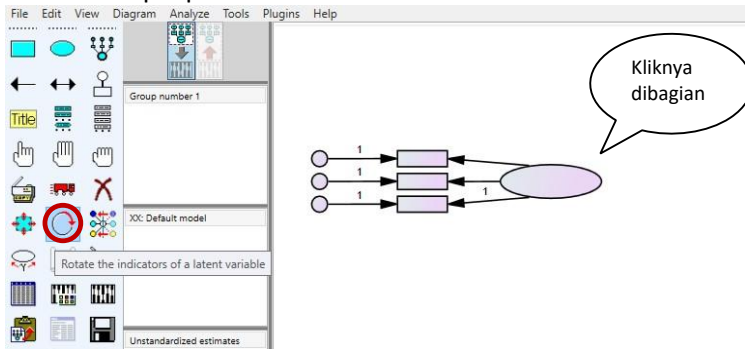
Menggambar diagram dengan menarik variable dengan indikatornya dengan cara klik tools ini



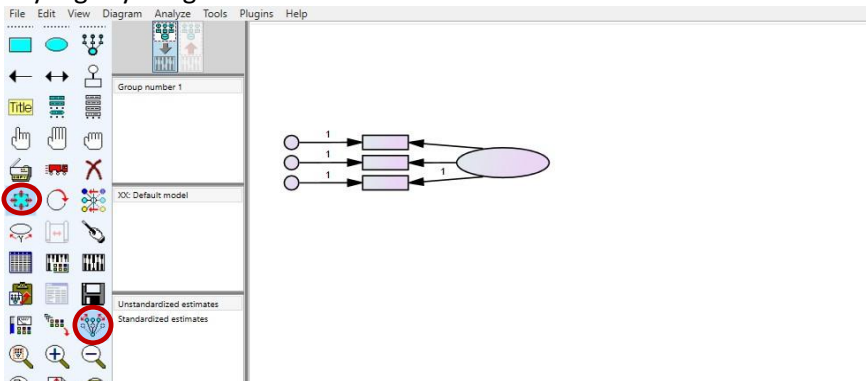
lalu ini manifesnya akan muncul di bidang gambar. Karena ada 3 indikator maka saya klik 3 kali



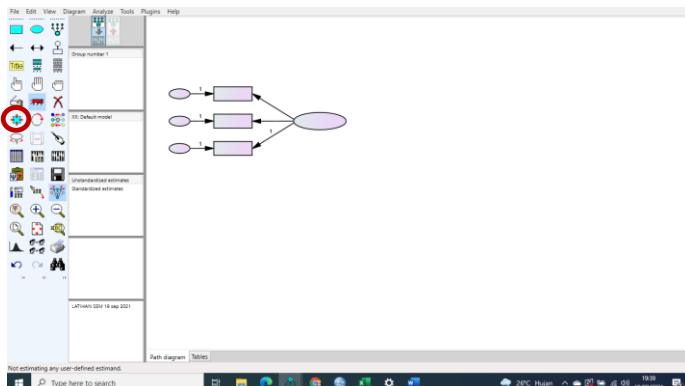
Kemudian lakukan rotasi agar posisi kakinya kekiri lalu klik dibagian lingkaran besar sampai posisi kaki ke kiri



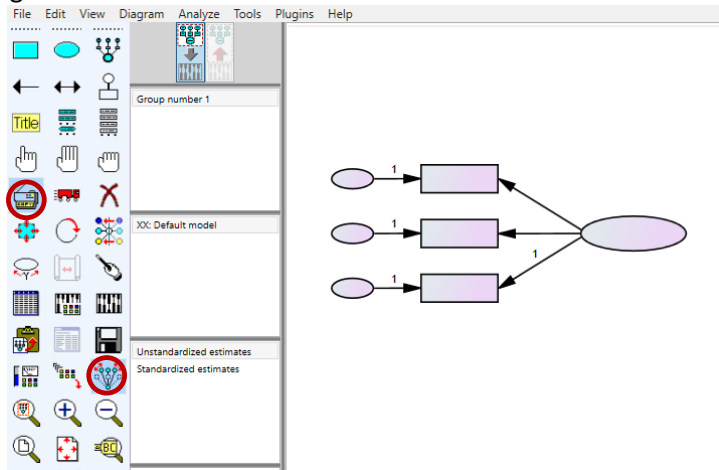
kemudian besarnya variable bisa dibesarkan bentuknya dengan cara klik tools yang saya lingkari dibawah ini.



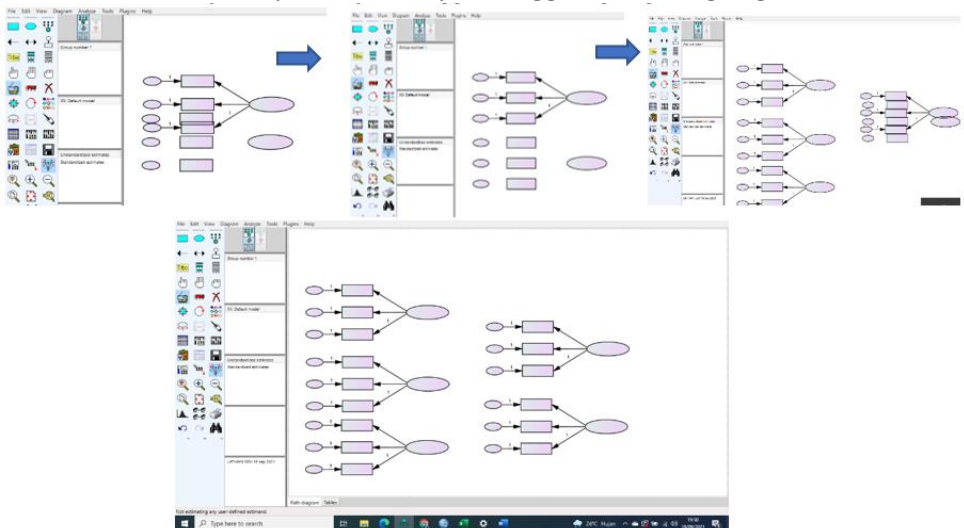
Setelah di klik maka merubah satu yang lainnya akan berubah. Jangan lupa untuk memberi jarak/memindahkan bisa menggunakan tools yang dilingkari ini



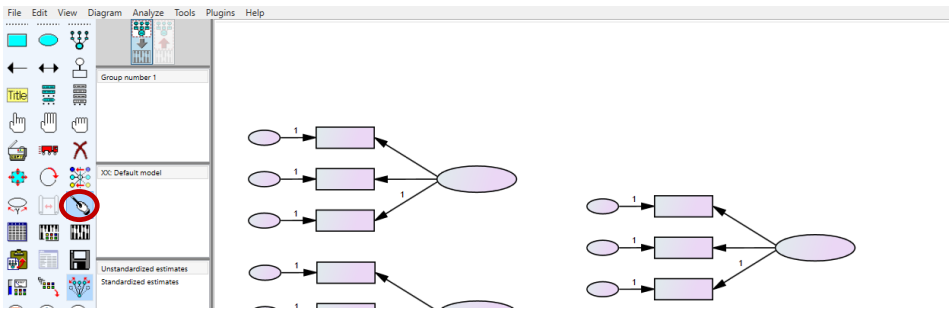
Kemudian untuk mengkopi menjadi 4 kali maka perlu menggunakan tools yang dilingkari merah



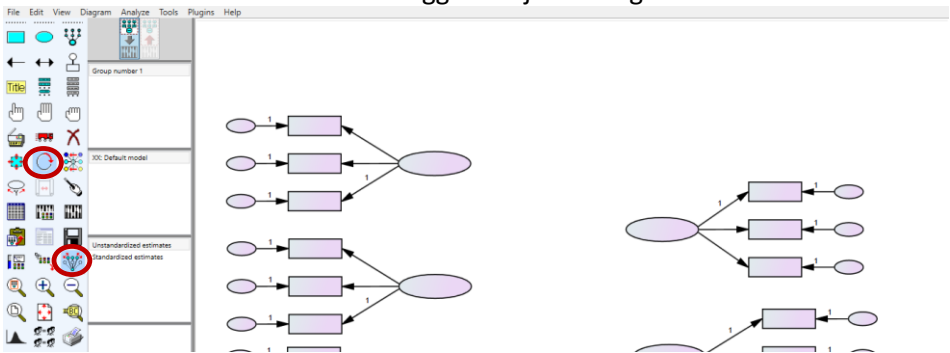
Kemudian lakukan proses fotokopi sehingga sesuai dengan gambar 1



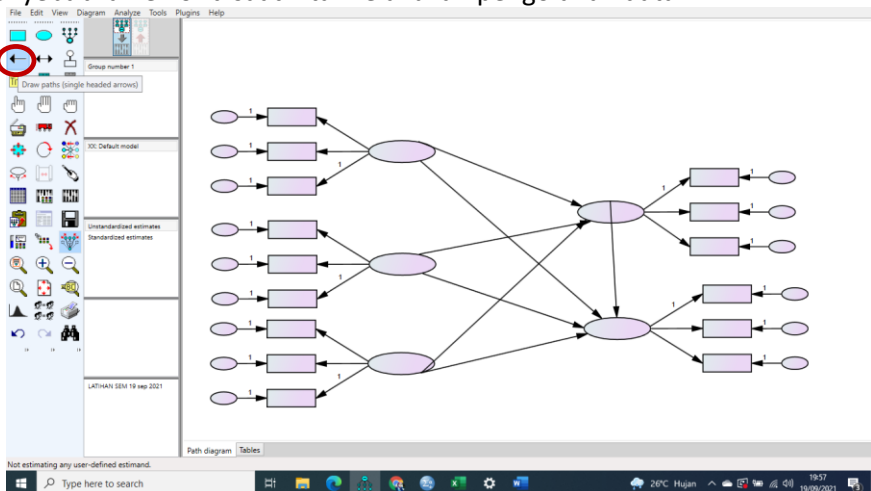
Kemudian lakukan touch up agar lebih rapi dengan cara mengklik tools yang dilingkari merah ini



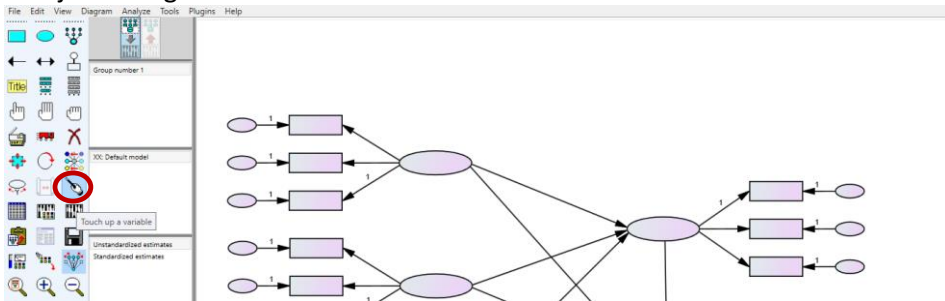
Kemudian lakukan rotate sehingga menjadi sebagai berikut



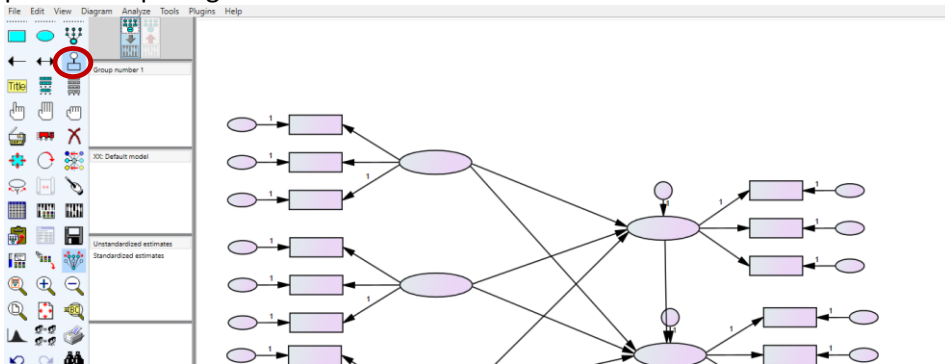
Setelah itu menarik pengaruhnya. Jangan sampai salah menarik dan jangan terlepas di bidang dan variabelnya jangan tercecer karena nanti bisa menyebabkan error disaat kita melakukan pengolahan data.



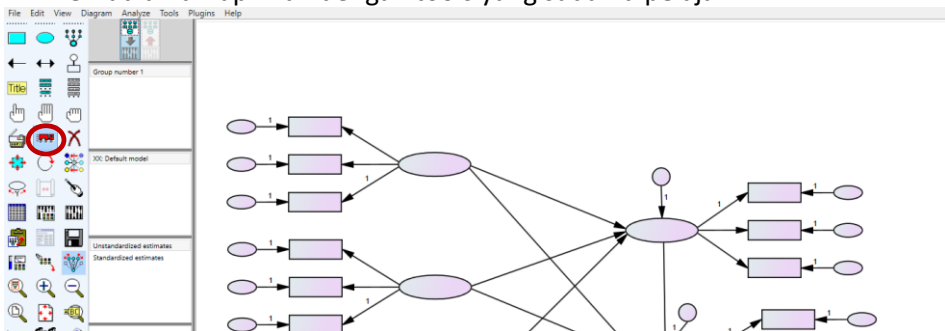
Biar saja berantakan, nanti bisa di touch up. Setelah di touch up, hasilnya menjadi sebagai berikut:



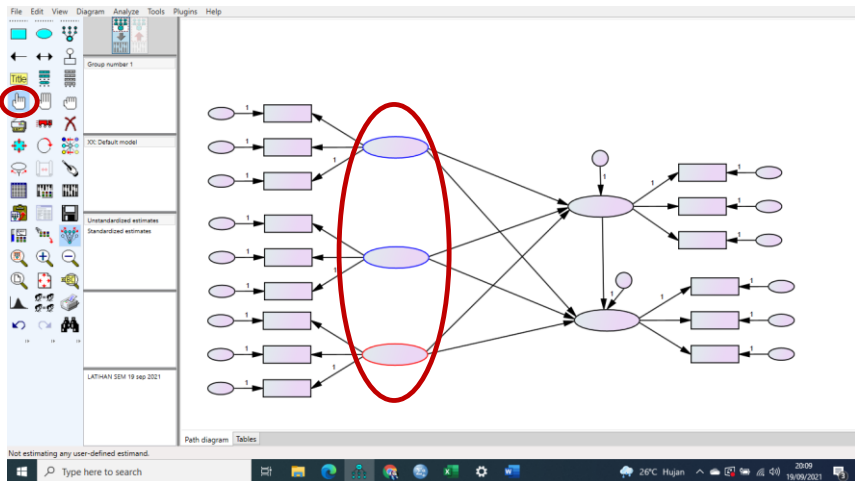
Kemudian yang paling penting adalah kita harus membuat kesalahan residual, akan ada 2 kesalahan residual. Klik tools yang dilingkari merah dan perhatikan pada gambar berikut



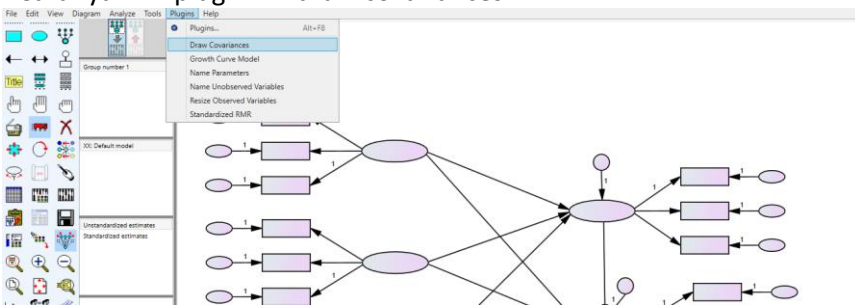
Kemudian di rapihkan dengan tools yang sudah dipelajari.



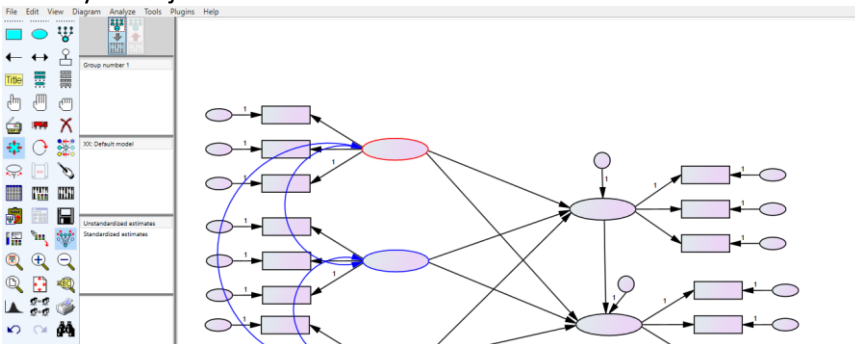
Kemudian select 3 variabel dengan tools yang dilingkari merah dan variable yg di select akan berubah menjadi biru.



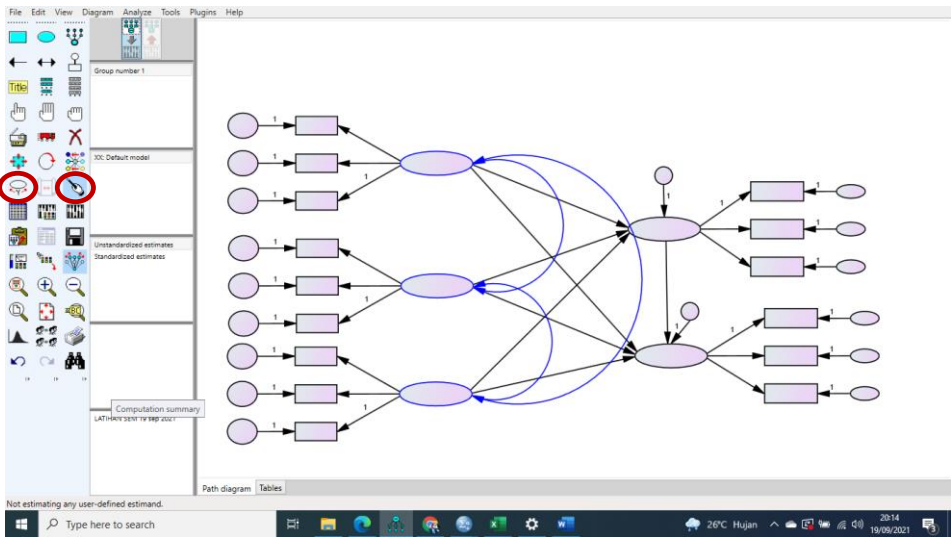
Kemudian dalam persamaan yang seperti ini harus dibuat dari covariance nya. Caranya klik plug in → draw covariances.



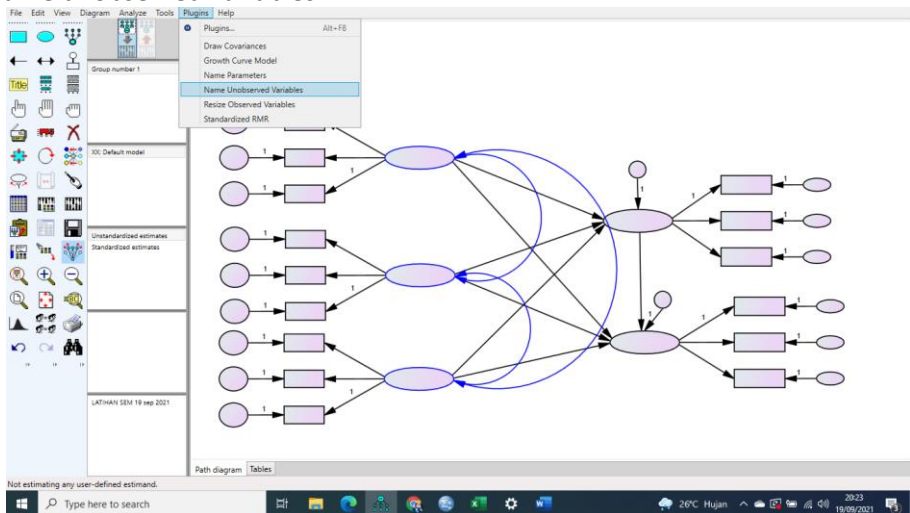
Hasilnya menjadi



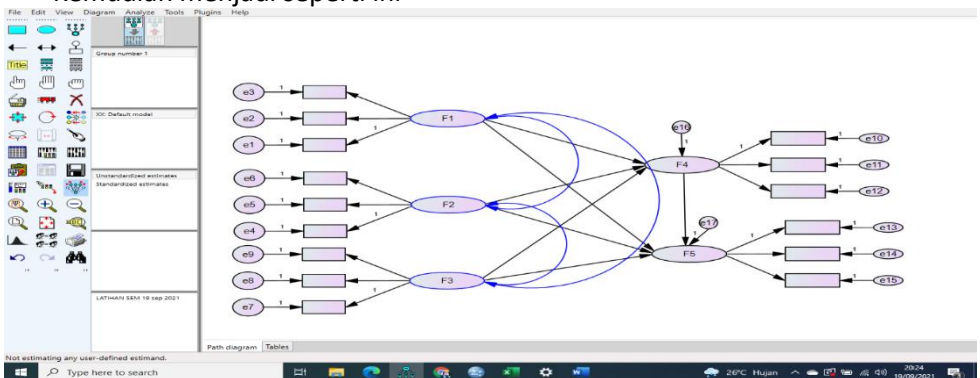
boleh model seperti ini ini. Tapi saya lebih sering menggunakan sebelah kanan maka saya tariknya seperti ini.



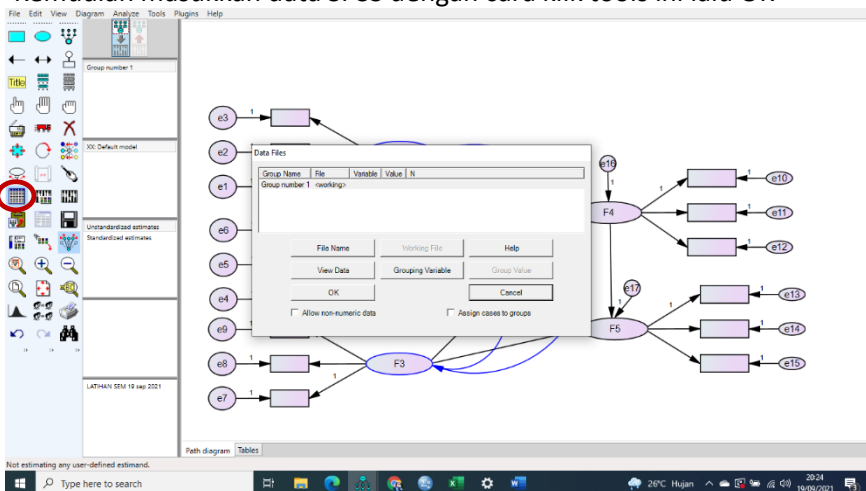
Kemudian namakan semua variable unobserved nya dengan cara plugin → name unobserved variables



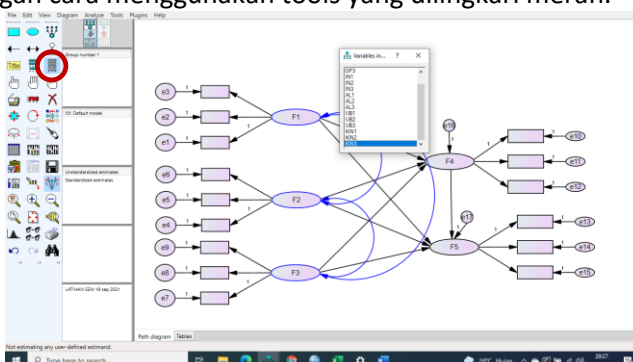
Kemudian menjadi seperti ini



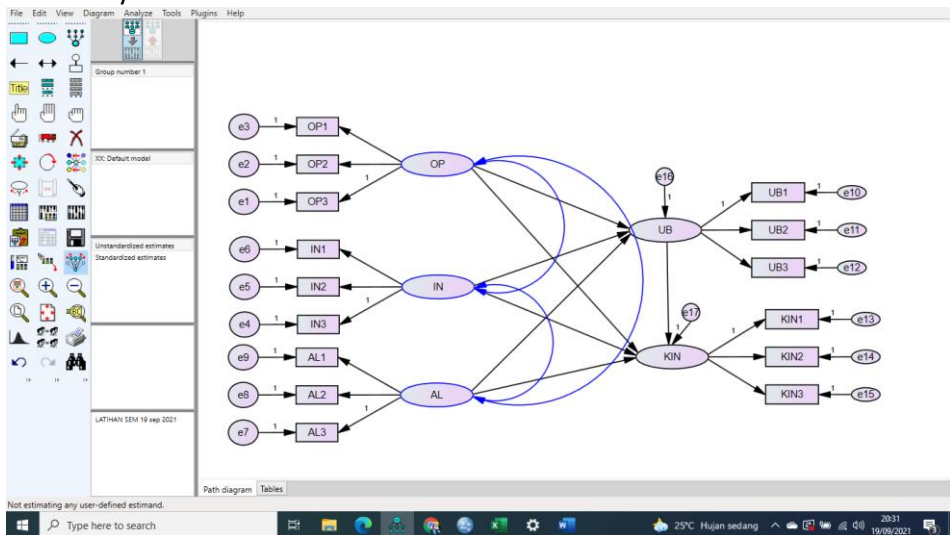
Kemudian masukkan data SPSS dengan cara klik tools ini lalu OK



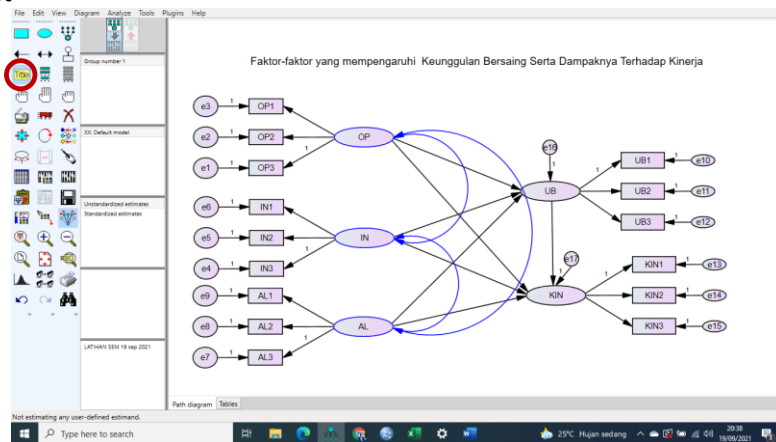
Setelah data SPSS masuk ke Amos maka selanjutnya adalah memasukkan variable dengan cara menggunakan tools yang dilingkari merah.



Setelah itu masukkan list variable satu per satu sesuai grand theory penelitian kita. Beri nama variable latennya dengan cara mengklik ke variabelnya.



Selanjutnya bisa memberi Judul Diagram dengan menggunakan tools berikut



Bisa juga membuat output Goodness of fit di bawah diagram dengan cara sebagai berikut:

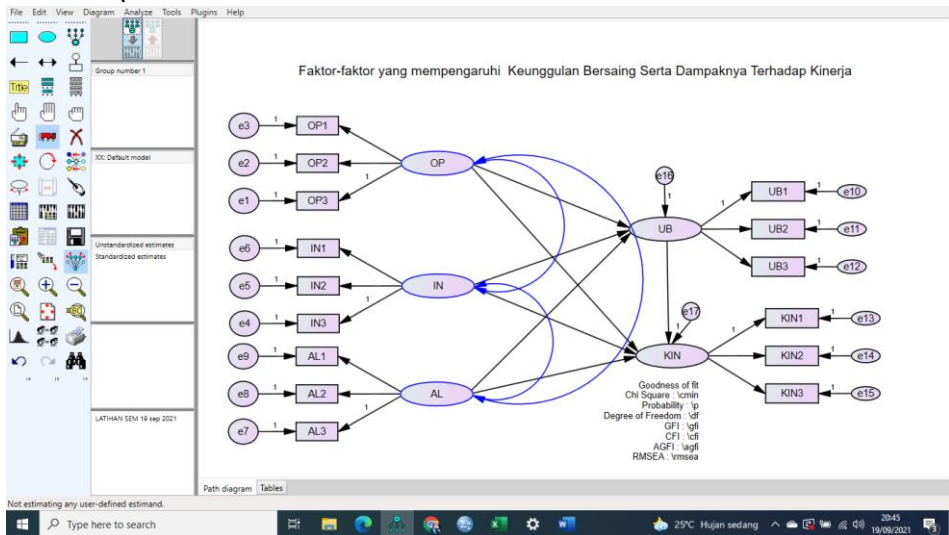
Goodness of fit

Chi Square : \cmin

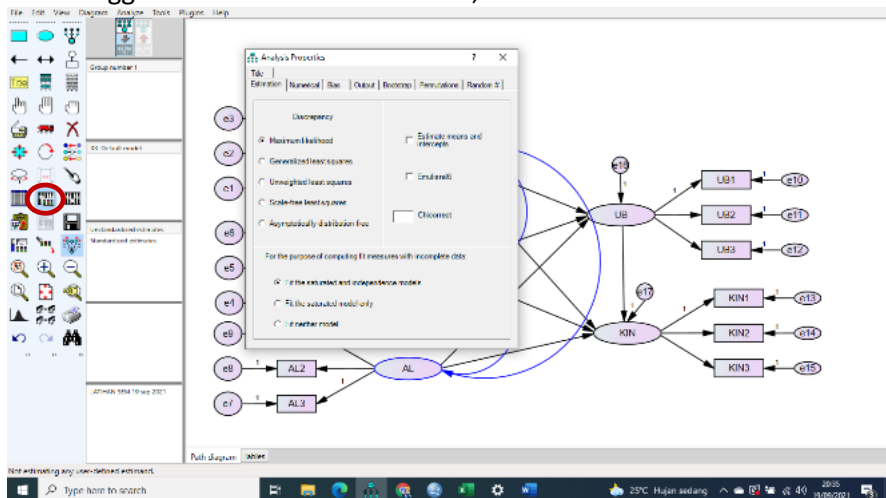
Probability : \p

Degree of Freedom : \df

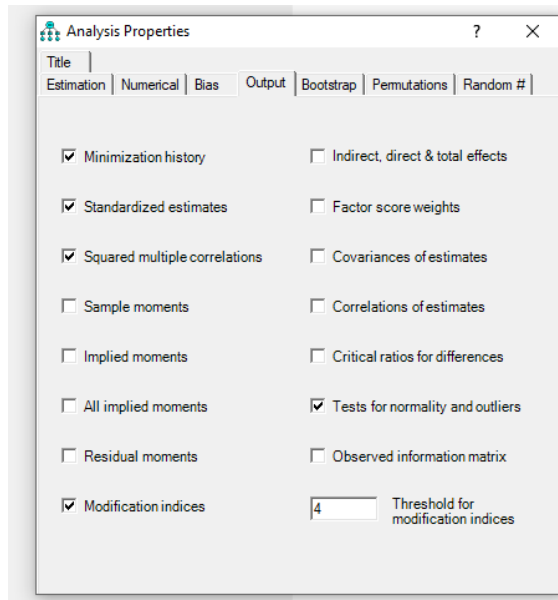
GFI : \gfi
 CFI : \cfi
 AGFI : \agfi
 RMSEA : \rmsea



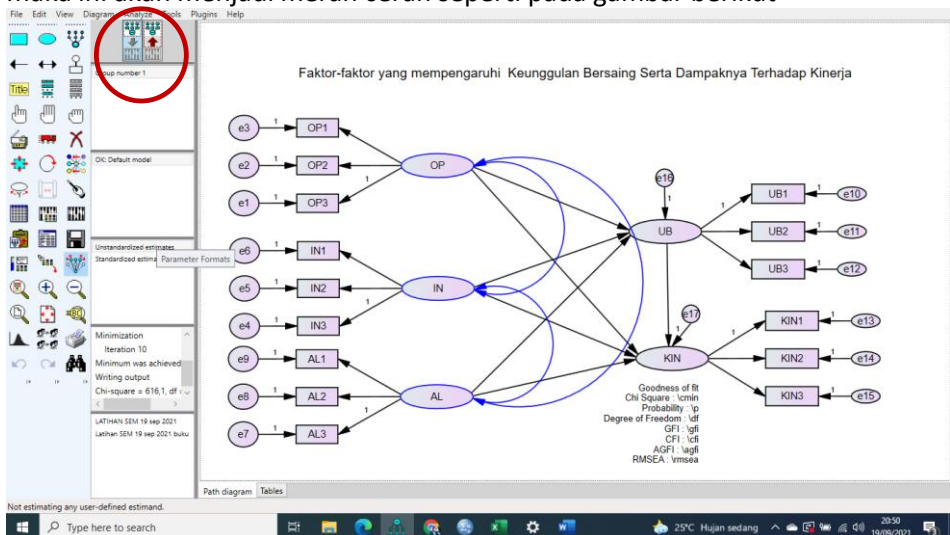
- **Langkah Keempat,**
 Memilih analysis properties kemudian pilih estimation yang sesuai dalam hal ini menggunakan maximum likelihood,



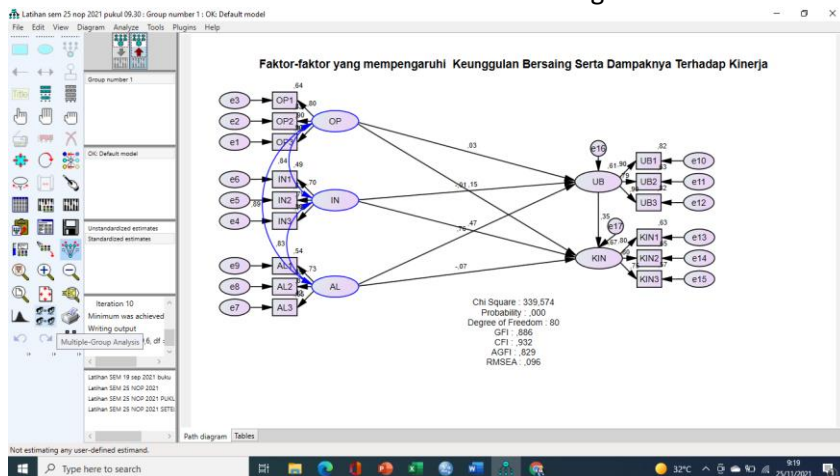
Kemudian ke output serta melakukan calculate lalu klik yang ada di gambar ini.



Rapikan kembali setelah itu dilakukan calculate. Ketika model sudah sesuai dan tidak terkendala maka nanti akan ada save dan beri nama. Nah ini nanti adalah outputnya ini lambangnya seperti hijau agak kabur. Nanti kalau berhasil maka ini akan menjadi merah cerah seperti pada gambar berikut



Lalu kita klik output yang merah sehingga tampilannya menjadi seperti ini
 Hasil untuk unstandardized variables adalah sebagai berikut



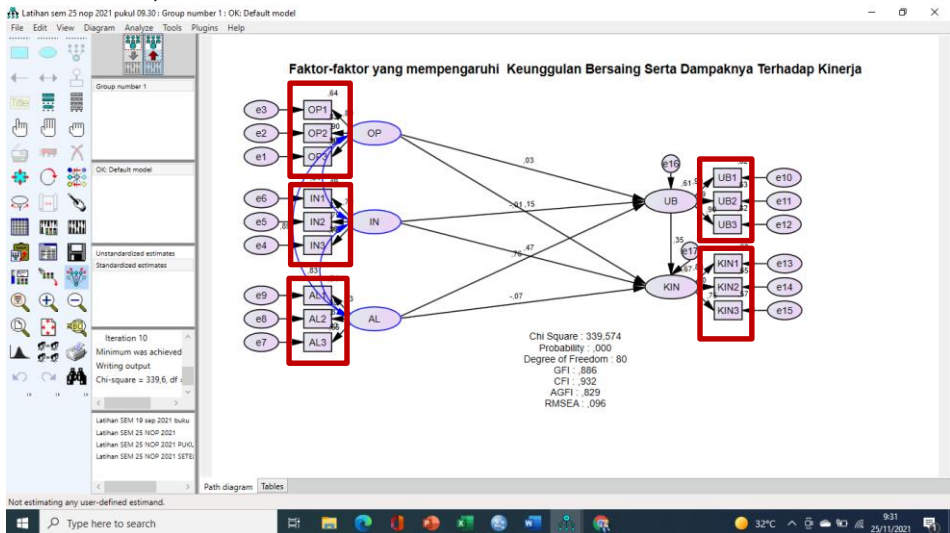
PEMBAHASAN

1. Uji validitas dengan uji CFA atau uji validitas konstruk (indicator) yaitu mengukur apakah konstruk indicator mampu atau tidak merefleksikan variable latennya. Hasilnya memenuhi kriteria yaitu nilai Critical Ratio (CR) > 1,96 dengan probability (P) < 0,05. Tanda *** adalah signifikan <0,001

Regression Weights: (Group number 1 - Default model)

			Estimate	S.E.	C.R.	P	Label
OP3	<---	OP	1,000				
OP2	<---	OP	,945	,032	29,431	***	
OP1	<---	OP	,893	,041	21,776	***	
IN3	<---	IN	1,000				
IN2	<---	IN	,999	,069	14,501	***	
IN1	<---	IN	,857	,068	12,660	***	
AL3	<---	AL	1,000				
AL2	<---	AL	1,305	,102	12,810	***	
AL1	<---	AL	1,279	,110	11,642	***	
UB1	<---	UB	1,000				
UB2	<---	UB	,883	,046	19,323	***	
UB3	<---	UB	,924	,038	24,322	***	
KIN1	<---	KIN	1,000				
KIN2	<---	KIN	1,066	,070	15,237	***	
KIN3	<---	KIN	,931	,065	14,264	***	

Selain itu uji validitas bisa juga dilihat dari nilai loading faktornya (ada pada gambar SEM nya) dimana nilai yang di loading factor atau standardized loading estimate $> 0,5$



2. Uji reliabilitas dengan uji Construct Reliability, yaitu menguji keandalan dan konsistensi data. Memenuhi kriteria apabila Construct Reliability $> 0,7$. Nilai Construct Reliability diantara $0,6 - 0,7$ masih dapat diterima dengan syarat validitas konstruk (indicator) dalam model adalah baik. Hasilnya semua diatas $0,7$

Standardized Regression Weights: (Group number 1 - Default model)

	Estimate
UB <--- OP	,025
UB <--- IN	-,008
UB <--- AL	,763
KIN <--- OP	,155
KIN <--- IN	,472
KIN <--- AL	-,074
KIN <--- UB	,351
OP3 <--- OP	,950
OP2 <--- OP	,903
OP1 <--- OP	,803
IN3 <--- IN	,762

	Estimate
IN2 <--- IN	,794
IN1 <--- IN	,697
AL3 <--- AL	,648
AL2 <--- AL	,831
AL1 <--- AL	,732
UB1 <--- UB	,905
UB2 <--- UB	,795
UB3 <--- UB	,904
KIN1 <--- KIN	,796
KIN2 <--- KIN	,803
KIN3 <--- KIN	,754

$$AVE = \frac{\sum \text{Standardized Loading}^2}{\sum \text{Standardized Loading}^2 + \sum e_j}$$

$$CR = \frac{(\sum \text{Standardized Loading})^2}{(\sum \text{Standardized Loading})^2 + \sum e_j}$$

VARIABEL	Variabel manifest	λ	λ^2	e	CR	VE
OP	OP1	0.803	0.645	0.355	0.975	0.928
	OP2	0.903	0.815	0.185		
	OP3	0.950	0.903	0.098		
IN	IN1	0.697	0.486	0.514	0.932	0.821
	IN2	0.794	0.630	0.370		
	IN3	0.762	0.581	0.419		
AL	AL1	0.732	0.536	0.464	0.940	0.842
	AL2	0.831	0.691	0.309		
	AL3	0.648	0.420	0.580		
UB	UB1	0.905	0.819	0.181	0.949	0.860
	UB2	0.795	0.632	0.368		
	UB3	0.904	0.817	0.183		
KIN	KIN1	0.796	0.634	0.366	0.940	0.839
	KIN2	0.803	0.645	0.355		
	KIN3	0.754	0.569	0.431		

Semua nilai Construct Reliability $CR \geq 0,7$ dan $VE \geq 0,5$ mengindikasikan bahwa validitas konstruk (indicator) dalam model adalah baik.

3. Uji Normalitas

Dari table assessment of normality di dapatkan nilai sebesar multivariate normality 17,752. Nilai ini lebih besar dari cut point multivariate normality yaitu sebesar 2,56. Hal ini berarti data tidak berdistribusi normal.

Karena data tidak berdistribusi normal maka perlu dibuat outlier data. Caranya adalah dengan melihat nilai table observations farthest from the centroid (Mahalanobis distance)

Assessment of normality (Group number 1)

Variable	min	max	skew	c.r.	kurtosis	c.r.
KIN3	2,000	5,000	,073	,557	-,685	-2,614
KIN2	2,000	5,000	,067	,512	-,864	-3,301
KIN1	2,000	5,000	-,198	-1,516	-,639	-2,440
UB3	2,000	5,000	,079	,606	-,535	-2,045
UB2	2,000	5,000	,036	,275	-,746	-2,851
UB1	2,000	5,000	,193	1,477	-,778	-2,971
AL1	2,000	5,000	,285	2,178	-,743	-2,837
AL2	2,000	5,000	,212	1,620	-,781	-2,983
AL3	3,000	5,000	,182	1,389	-1,090	-4,163
IN1	3,000	5,000	,026	,198	-,755	-2,882
IN2	3,000	5,000	,098	,747	-,857	-3,272
IN3	3,000	5,000	,291	2,222	-1,008	-3,849
OP1	2,000	5,000	,026	,196	-1,232	-4,703
OP2	1,000	5,000	-,191	-1,455	-,221	-,844
OP3	2,000	5,000	,073	,556	-1,016	-3,878
Multivariate					42,858	17,752

Observations farthest from the centroid (Mahalanobis distance) (Group number 1)

(Tidak semua output ditampilkan)

Observation number	Mahalanobis d-squared	p1	p2
214	77,479	,000	,000
20	49,014	,000	,000
33	43,311	,000	,000
12	41,881	,000	,000
340	39,946	,000	,000
2	38,743	,001	,000
180	37,980	,001	,000
80	37,581	,001	,000
102	35,064	,002	,000
9	34,485	,003	,000
350	34,403	,003	,000
51	33,816	,004	,000
57	32,942	,005	,000
149	32,461	,006	,000
164	31,944	,007	,000
27	31,574	,007	,000
8	31,422	,008	,000
135	31,378	,008	,000
98	31,092	,009	,000
28	29,167	,015	,000
17	28,978	,016	,000

Agar supaya data berdistribusi normal maka Amos memberikan rekomendasi untuk membuang semua Data observation number yang ada pada table observations farthest from the centroid (Mahalanobis distance). Kemudian setelah dikeluarkan maka dilakukan input data ulang dari awal dan melakukan kembali analisis SEM dengan AMOS.

Untuk pembahasan mengenai model terbaik dengan mengeluarkan outlier akan dibahas di pembahasannya selanjutnya. Untuk selanjutnya kami akan menggunakan data awal saja (tetap memasukkan data yang outlier).

4. Menilai Goodness of Fit Model

Goodness of fit index	Cut of Value	Hasil Analisis	Evaluasi Model
Chi Square	$\leq 199,24$, dimana Chi Square untuk df 168; Taraf Sig 5%	339,574 P= 0,0000	Marginal
GFI	>0,90	0,886	Baik
AGFI	>0,90	0,829	Baik
IFI	>0,90	0,932	Marginal
CFI	>0,90	0,932	Marginal
NFI	>0,90	0,931	Marginal
RMSEA	<0,08	0,096	Marginal

Perlu dilakukan treatment

5. Pengujian Hipotesis

Regression Weights: (Group number 1 - Default model)

	Estimate	S.E.	C.R.	P	Label
UB <--- OP	,027	,162	,167	,867	Tidak Sig
UB <--- IN	-,011	,179	-,061	,952	Tidak Sig
UB <--- AL	1,252	,302	4,152	***	Sig
KIN <--- OP	,146	,132	1,110	,267	Tidak Sig
KIN <--- IN	,594	,162	3,661	***	Sig
KIN <--- AL	-,107	,272	-,394	,694	Tidak Sig
KIN <--- UB	,311	,076	4,098	***	Sig

D. KESIMPULAN

1. Tidak terdapat pengaruh Orientasi Produk terhadap Keunggulan Bersaing
2. Tidak terdapat pengaruh Inovasi terhadap Keunggulan Bersaing
3. Terdapat Pengaruh positif Adaptabilitas terhadap Keunggulan Bersaing
4. Terdapat pengaruh Keunggulan Bersaing terhadap Kinerja
5. Tidak terdapat pengaruh Orientasi Produk terhadap Kinerja
6. Terdapat Pengaruh positif Inovasi terhadap Kinerja
7. Tidak terdapat Pengaruh Adaptabilitas terhadap Kinerja

DAFTAR PUSTAKA

- Arikunto, Suhasimi. (2010). *Prosedur Penelitian Suatu Pendekatan Praktik*. Jakarta: Rineka Cipta. Ghozali,
- Imam. (2005). *Aplikasi Multivariate dengan Program SPSS*. Semarang: Badan Penerbit Universitas Diponegoro. Ghozali,
- Imam. (2011). *Model Persamaan Struktural Konsep dan aplikasi dengan Program AMOS 21.0*. Semarang: Badan Penerbit Universitas Diponegoro.
- Nasution, S. (2012). *Metode Research (Penelitian Ilmiah)*. Jakarta: PT Bumi Aksara.
- Singgih, Santoso. (2011). *Structural Equation Modeling (SEM) Konsep dan Aplikasi dengan AMOS 18*. Jakarta.: PT Elex Media Komputindo.

STRUCTURAL EQUATION MODELING DENGAN LISREL

Dr. Sri Rochani Mulyani, S.E., M.Si
Universitas Sangga Buana YPKP

A. PENGERTIAN SEM

Menurut Bollen (2011) sebagaimana dikutip oleh Latan (2013: 5), “SEM are sets of equations that encapsulate the relationships among the latent variables, observed variables, and error variables”. SEM dapat digunakan untuk menjawab berbagai masalah riset (research question) dalam suatu set analisis secara sistematis dan komprehensif. Menurut Ramadiani (2010), SEM adalah singkatan structural equation model yang merupakan model persamaan struktural generasi kedua teknik analisis multivariat yang memungkinkan peneliti untuk menguji hubungan antara variabel.

Menurut Ghazali & Fuad (2008: 3), model persamaan struktural (Structural Equation Modeling) adalah generasi kedua teknik analisis multivariat (Bagozzi dan Fornell, 1982) yang memungkinkan peneliti untuk menguji hubungan antara variabel yang kompleks baik recursive maupun nonrecursive untuk memperoleh gambaran menyeluruh mengenai keseluruhan model. Dengan demikian SEM adalah salah satu teknik analisis multivariat yang digunakan untuk menganalisis hubungan antar variabel yang lebih kompleks dibandingkan dengan analisis regresi berganda dan analisis faktor.

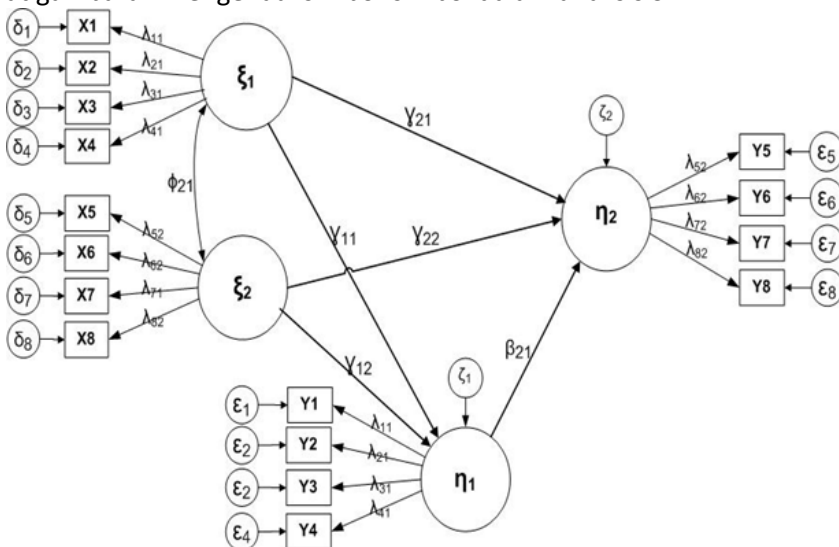
Model persamaan struktural (Structural equation modeling) adalah generasi kedua teknik analisis multivariate (Bagozzi dan Fornell 1982) yang memungkinkan peneliti menguji hubungan antara variabel yang kompleks. SEM dapat menguji secara bersama-sama:

- 1) Model struktural
- 2) Model pengukuran
- 3) Menguji kesalahan pengukuran

Umumnya LISREL digunakan untuk memecahkan masalah model persamaan struktural (SEM). Hair, Anderson, Tatham dan Black mengemukakan 7 langkah dalam mengaplikasikan SEM sebagai berikut :

- 1) Merumuskan model berbasis teori : teridentifikasi variabel laten endogen dan eksogen, argumentasi hubungan kausal antar variabel laten, serta mengidentifikasi indikator atau variabel manifest eksogen dan endogen.
- 2) Mengkonstruksi diagram jalur : tergambar dengan jelas setting atau adegan hubungan antar variabel laten dan teridentifikasi jumlah parameter yang akan diestimasi
- 3) Merumuskan persamaan : model pengukuran dan model struktural
- 4) Menentukan data input dan metode estimasi : menentukan matrik korelasi atau matriks kovarians dan metode estimasi maximum likelihood, dll
- 5) Identifikasi model : model dapat menghasilkan estimasi yang bersifat unik (tunggal) atau tidak. Syarat suatu model menghasilkan estimasi unik adalah model tersebut bersifat just-identified atau over identified. Model dikatakan just-identified apabila derajat bebas model tersebut nol (0) dan dikatakan over identified apabila derajat bebas model lebih dari nol (> 0)
- 6) Uji kesesuaian model : pendekatan dua tahap, secara individual digunakan uji t dan secara keseluruhan digunakan kriteria Goodness of Fit (GOF).
- 7) Interpretasi dan modifikasi model : menjawab masalah penelitian dan memodifikasi model berdasarkan justifikasi teoritis tertentu.

Dalam analisis SEM, terdapat tiga bagian yang akan di uji yaitu, Uji Model Pengukuran, Uji Model Struktural dan Uji Kecocokan Model (Goodness of Fit Model). Sehingga dalam bagian ini akan dibahas mengenai ketiga cara penghitungan dan ketentuannya. Sebelum memasuki perhitungan, maka dibuat gambaran mengenai simbol-simbol dalam analisis SEM.



Gambar 8 Contoh Paradigma SEM beserta Simbol

Tabel 2 Simbol-Simbol dalam SEM

Notasi	Keterangan
ξ (ksi)	Variabel laten eksogen (variabel independen), digambarkan sebagai lingkaran pada model struktural SEM
η (eta)	Variabel laten endogen (variabel dependen, dan juga dapat menjadi variabel independen pada persamaan lain, juga digambarkan sebagai lingkaran)
γ (gamma)	Hubungan langsung variabel eksogen terhadap variabel endogen
β (beta)	Hubungan langsung variabel endogen terhadap variabel endogen
Y	Indikator variabel endogen
X	Indikator variabel eksogen
λ (lambda)	Hubungan antara variabel laten eksogen maupun endogen terhadap indikator-indikatornya
δ (delta)	Kesalahan pengukuran dari indikator variabel eksogen
ε (epsilon)	Kesalahan pengukuran dari indikator variabel endogen
ζ (Zeta)	Kesalahan dalam persamaan yaitu antara variabel eksogen dan atau endogen terhadap variabel endogen

B. TAHAPAN DALAM PROSEDUR SEM

Menurut Bollen dan Long dalam (Wijanto, 2008) prosedur SEM mengandung tahap-tahap sebagai berikut.

1) Spesifikasi model (model specification)

Tahap ini berkaitan dengan pembentukan model awal persamaan struktural, sebelum dilakukan estimasi. Model awal ini diformulasikan berdasarkan suatu teori atau penelitian sebelumnya. Definisikan model menjadi variabel-variabel laten yang ada dalam penelitian. Definisikan variabel-variabel teramati dan definisikan hubungan antara setiap variabel laten dengan variabel – variabel teramati yang terkait. Kemudian definisikan hubungan kausal diantara variabel-variabel laten tersebut.

2) Identifikasi (Identification)

Tahap ini berkaitan dengan pengkajian tentang kemungkinan diperolehnya nilai yang unik untuk setiap parameter yang ada di dalam model dan kemungkinan persamaan simultan tidak ada solusinya. Definisi identifikasi di dalam SEM sebagai berikut:

- Under-Identified model adalah model dengan jumlah parameter yang diestimasi lebih besar dari jumlah data yang diketahui (data tersebut merupakan variance dan covariance dari variabel-variabel teramati).
- Just-Identified model adalah model dengan jumlah parameter yang diestimasi sama dengan data yang diketahui.
- Over-Identified model adalah model dengan jumlah parameter yang diestimasi lebih kecil dari jumlah data diketahui.

3) Estimasi (Estimation)

Tahapan ini berkaitan dengan estimasi terhadap model untuk menghasilkan nilai-nilai parameter dengan menggunakan salah satu metode estimasi yang tersedia. pemilihan metode estimasi yang digunakan seringkali ditentukan berdasarkan karakteristik dari variabel-variabel yang dianalisis.

4) Uji Kecocokan (Testing fit)

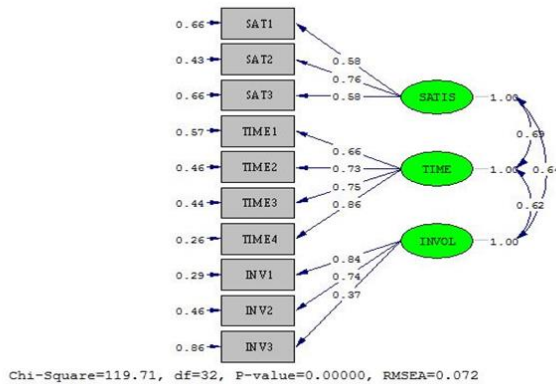
Tahap ini berkaitan dengan pengujian kecocokan antara model dengan data. Beberapa kriteria ukuran kecocokan atau Goodness of fit (GOF) dapat digunakan untuk melaksanakan langkah ini.

4) Respesifikasi (Respesification)

Tahap ini berkaitan dengan respesifikasi model berdasarkan atas hasil uji kecocokan tahap sebelumnya untuk memperoleh nilai kecocokan model yang lebih baik.

C. MODEL PENGUKURAN (MEASUREMENT MODEL)

Model pengukuran Tahap ini dilakukan uji unidimensional dari konstruk-konstruk eksogen dan endogen dengan teknik Confirmatory Factor Analysis (CFA). CFA dilakukan dengan membuat hubungan korelasi antara masing-masing konstruk dengan indikatornya. Secara visual dapat dilihat contoh untuk perhitungan model pengukuran adalah sebagai berikut.



Gambar 9 Contoh Model Pengukuran

Menurut (Wijanto, 2008), evaluasi ini akan dilakukan terhadap setiap konstruk atau model pengukuran (hubungan antara sebuah variabel laten dengan beberapa variabel teramati/indikator) secara terpisah melalui:

- Evaluasi terhadap validitas (validity) dari model pengukuran
- Evaluasi terhadap reliabilitas (reliability) dari model pengukuran

Menurut Igbaria et.al (1997) dalam Wijanto (2008) menyatakan bahwa muatan faktor (loading factor) standar adalah ≥ 0.50 adalah very significant. Kemudian setelah diketahui nilai muatan faktornya, maka dilanjutkan dengan uji reliabilitas. Untuk mengukur reliabilitas dalam SEM akan digunakan composite reliability measure (ukuran reliabilitas komposit) dan variance extracted measure (ukuran ekstrak varian). Reliabilitas komposit suatu konstruk dihitung sebagai:

$$CR = \frac{(\sum std. Loading)^2}{(\sum std. Loading)^2 + \sum e_j}$$

Dimana std. loading (standardized loadings) yang diperoleh secara langsung dari keluaran program LISREL dan e_j adalah measurement error untuk setiap indikator atau variabel teramati (Fornel dan Larker, 1981).

Ekstrak varian mencerminkan jumlah varian keseluruhan dalam indikator-indikator (variabel-variabel teramati) yang dijelaskan oleh variabel laten. Ukuran varian dapat dihitung sebagai berikut:

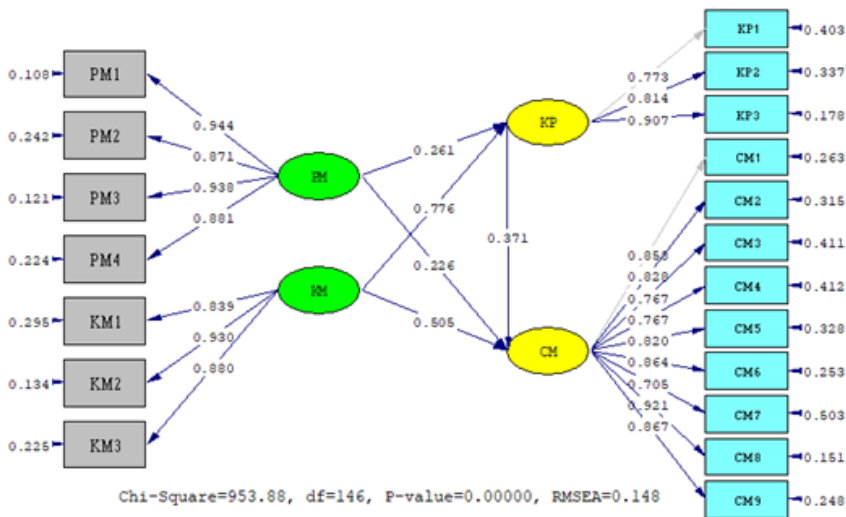
$$VE = \frac{\sum std. Loading^2}{\sum std. Loading^2 + \sum e_j}$$

Sebuah konstruk mempunyai reliabilitas yang baik jika:

- Nilai Construct Reliability (CR) ≥ 0.70 dan
- Nilai Variance Extracted (VE) ≥ 0.50 .

D. MODEL STRUKTURAL

Langkah ini dilakukan untuk melihat kesesuaian model dan hubungan kausalitas yang dibangun dalam model yang diuji. Model ini dilakukan dengan mengganti hubungan korelasi (dua anak panah) dengan satu anak panah pada masing-masing variabel laten. Secara visual, contoh pengujian model struktural adalah sbegaia berikut.



Gambar 10 Contoh Model Struktural

Dalam model struktural, yang akan diuji adalah pengaruh dari satu variabel terhadap variabel lainnya. Menurut wijanto (2008), evaluasi terhadap model struktural mencakup pemeriksaan terhadap signifikansi koefisien-koefisien yang diestimasi. Metode SEM dan LISREL tidak saja menyediakan nilai-nilai koefisien-koefisien yang diestimasi tetapi juga nilai t-hitung untuk setiap koefisien. Dengan menspesifikasikan tingkat signifikan (lazimnya 0.05), maka setiap koefisien yang mewakili hubungan kausal yang dihipotesiskan dapat diuji signifikansi secara statistik. Koefisien-koefisien tersebut serupa dengan koefisien beta pada regresi berganda, yaitu nilai koefisien yang mendekati nol menandakan pengaruh yang semakin kecil. Sebagai ukuran menyeluruh pada persamaan struktural, overall coefficient of determination (R^2) dihitung seperti regresi berganda.

E. UJI KECOCOKAN MODEL

Tujuannya adalah untuk mendeteksi ada tidaknya masalah identifikasi berdasarkan evaluasi terhadap hasil estimasi yang dilakukan. Masalah yang diidentifikasi : program komputer tidak menghasilkan matrik informasi yang harus disajikan, standard error yang besar untuk satu atau lebih, munculnya angka yang aneh seperti adanya varians eror yang negatif. Sehingga dibutuhkan **evaluasi kriteria Goodness Of the Fit (GOF)**. Tujuan adalah untuk mengevaluasi pemenuhan asumsi yang disyaratkan SEM, dan kesesuaian model berdasarkan kriteria Goodness Of Fit (GOF) tertentu. Pembahasan tentang uji kecocokan serta batas-batas nilai yang menunjukkan tingkat kecocokan yang baik (good-fit) untuk setiap GOF, maka dibuat ringkasan berikut.

Tabel 3 Perbandingan Ukuran-Ukuran GOF

Ukuran GOF	Tingkat Kecocokan yang Bisa Diterima
Absolute -Fit Measure	
Statistic Chi Square (χ^2)	Mengikuti uji statistik yang berkaitan dengan persyaratan signifikan. Semakin kecil semakin baik
Non Centrally Parameter (NCP)	Dinyatakan dalam bentuk spesifikasi ulang dari chi square. Penilaian didasarkan atas perbandingan dengan model lain. Semakin kecil semakin baik.
Scaled NCP (SNCP)	NCP dinyatakan dalam bentuk rata-rata perbedaan setiap observasi dalam rangka perbandingan antar model. Semakin kecil semakin baik.
Goodness of fit Index (GFI)	Nilai berkisar antara 0-1 dengan nilai lebih tinggi adalah lebih baik. $GFI \geq 0,90$ adalah good fit, sedang $0.80 \leq GFI \leq 0.90$ adalah marginal fit.
Root Mean Square Residual (RMR)	Residual rata-rata antara matrik (korelasi atau kovarian) teramati dan hasil estimasi. Standardized RMR ≤ 0.05 adalah good fit.
Root Mean Square Error of Approximation (RMSEA)	Rata-rata perbedaan per degree of freedom yang diharapkan terjadi dalam populasi dan bukan dalam sampel. $RMSEA \leq 0.08$ adalah good fit, sedang $RMSEA < 0.05$ adalah close fit.
Expected Cross-Validation Index (ECVI)	Digunakan untuk perbandingan antar model. Semakin kecil semakin baik

Ukuran GOF	Tingkat Kecocokan yang Bisa Diterima
Incremental Fit Measures	
Tucker Lewis Index atau Non Normed Fit Index (TLI/NNFI)	Nilai berkisar antara 0-1 dengan nilai lebih tinggi adalah lebih baik. $TLI \geq 0,90$ adalah good fit, sedang $0.80 \leq TLI < 0.90$ adalah marginal fit.
Normed Fit Index (NFI)	Nilai berkisar antara 0-1 dengan nilai lebih tinggi adalah lebih baik. $NFI \geq 0,90$ adalah good fit, sedang $0.80 \leq NFI < 0.90$ adalah marginal fit.
Adjusted Goodness of fit Index (AGFI)	Nilai berkisar antara 0-1 dengan nilai lebih tinggi adalah lebih baik. $AGFI \geq 0,90$ adalah good fit, sedang $0.80 \leq AGFI < 0.90$ adalah marginal fit.
Relative fit Index (RFI)	Nilai berkisar antara 0-1 dengan nilai lebih tinggi adalah lebih baik. $RFI \geq 0,90$ adalah good fit, sedang $0.80 \leq RFI < 0.90$ adalah marginal fit.
Incremental Fit Index (IFI)	Nilai berkisar antara 0-1 dengan nilai lebih tinggi adalah lebih baik. $IFI \geq 0,90$ adalah good fit, sedang $0.80 \leq IFI < 0.90$ adalah marginal fit.
Comparative Fit Index (CFI)	Nilai berkisar antara 0-1 dengan nilai lebih tinggi adalah lebih baik. $CFI \geq 0,90$ adalah good fit, sedang $0.80 \leq CFI < 0.90$ adalah marginal fit.
Parsimonous Fit Measures	
Parsimonous Goodness of Fit Index (PGFI)	Spesifikasi ulang dari GFI, dimana nilai lebih tinggi menunjukkan parsimoni yang lebih besar. Ukuran ini digunakan untuk perbandingan diantara model-model.
Normed chi square	Rasio antara chi square dibagi degree of freedom. Nilai yang disarankan: batas bawah 1.0, batas tengah: 2.0 atau 0.3 dan yang lebih longgar 5.0.
Parsomonous Normed Fot Indeks (PNFI)	Nilai tinggi menunjukkan kecocokan lebih baik, hanya digunakan untuk perbandingan antar model alternatif.
Akaike Information Criterion (AIC)	Nilai positif lebih kecil menunjukkan parsimoni lebih baik, digunakan untuk perbandingan antar model. pada model tunggal, nilai AIC dari model yang mendekati nilai saturated AIC menunjukkan good fit.
Consistent Akaike Information Criterion (CAIC)	Nilai positif lebih kecil menunjukkan parsimoni lebih baik, digunakan untuk perbandingan antar model. pada model tunggal, nilai CAIC dari model yang mendekati nilai saturated CAIC menunjukkan good fit.

Sumber: Wijanto, 2008

F. PERSIAPAN ANALISIS DATA DENGAN LISREL

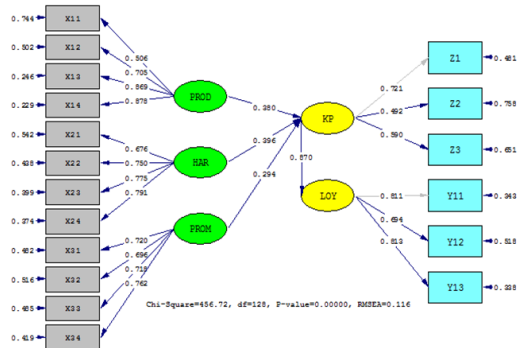
1) Data Excel (Variabel dan Indikator)

Ini merupakan data pentah hasil dari lapangan

2) SPSS (Data Input)

Contoh Model Penelitian

Judul : Pengaruh Produk, Harga Dan Promosi Terhadap Keputusan Pembelian Dan Implikasinya Pada Loyalitas Pelanggan



Tahapan Lisrel 8.72

1) Masukkan data excel ke SPSS

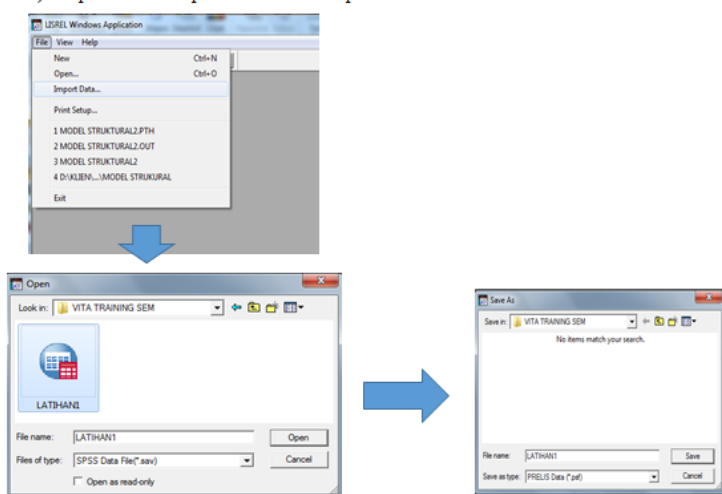
Data Excel

	A	B	C	D	E	F	G	H	I	J	K	L	M	N	O	P	Q	R	S
1	No	PRODUK				HARGA				PROMOSI				KEPUTUSAN PEMBELIAN			LOYALITAS		
2		X11	X12	X13	X14	X21	X22	X23	X24	X31	X32	X33	X34	Z1	Z2	Z3	Y11	Y12	Y13
3	1	4	4	3	3	2	3	2	2	3	3	3	3	3	3	3	3	3	3
4	2	3	3	3	3	2	2	2	2	2	3	2	2	3	3	3	3	4	4
178	176	4	4	4	4	3	3	2	3	4	3	3	3	3	2	3	4	3	4
179	177	4	3	3	3	3	3	3	2	4	3	3	4	4	3	3	3	3	3
180	178	3	3	3	3	2	3	2	3	3	4	3	4	3	3	4	3	2	3

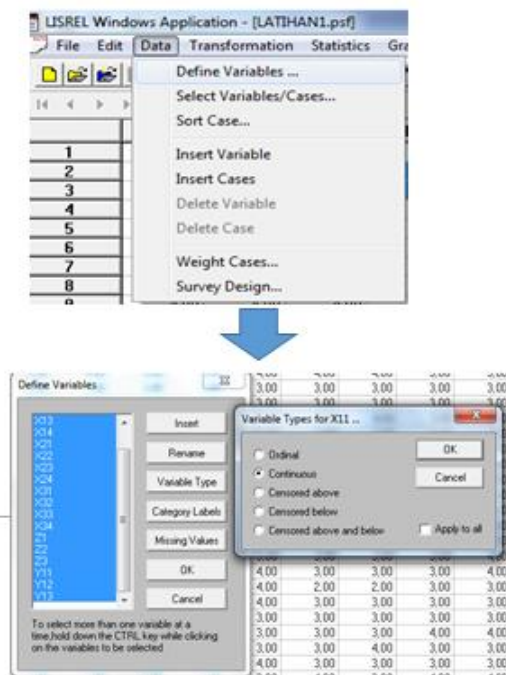
Data SPSS

		X11	X12	X13	X14	X21	X22	X23
1		4.00	4.00	3.00	3.00	2.00	3.00	
2		3.00	3.00	3.00	3.00	2.00	2.00	
3		3.00	3.00	3.00	3.00	2.00	3.00	
4		3.00	3.00	3.00	3.00	2.00	3.00	
5		4.00	3.00	4.00	4.00	3.00	3.00	
6		3.00	3.00	3.00	3.00	2.00	3.00	
7		4.00	3.00	3.00	3.00	3.00	3.00	
8		4.00	4.00	3.00	3.00	3.00	3.00	
9		4.00	4.00	4.00	4.00	3.00	3.00	
10		3.00	3.00	3.00	3.00	2.00	3.00	
11		3.00	3.00	3.00	3.00	2.00	2.00	
12		3.00	3.00	3.00	3.00	2.00	3.00	
176		4.00	4.00	4.00	4.00	3.00	3.00	
177		4.00	3.00	3.00	3.00	3.00	3.00	
178		3.00	3.00	3.00	3.00	2.00	3.00	

2) Import data dari spss ke lisrel dan simpan

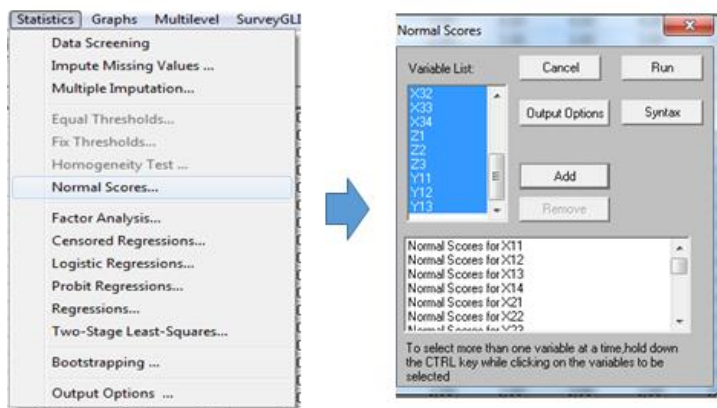


3) Ubah data menjadi Continuous

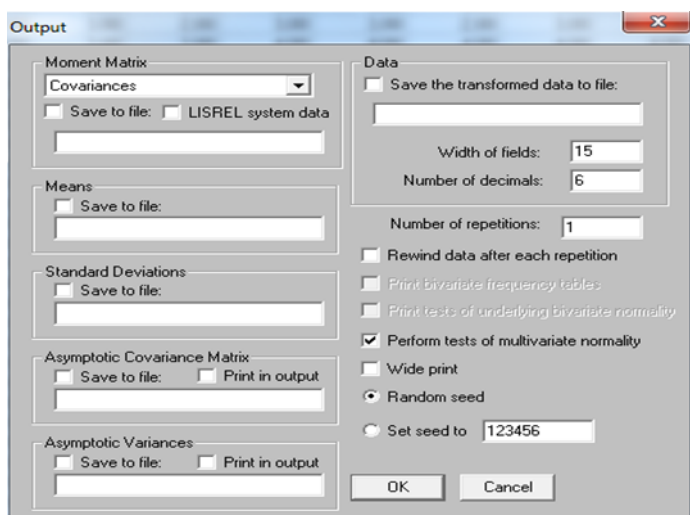


Prosedur Uji Normalitas Data

1. Pilih statistics, lalu klik normal score



2. Kemudian muncul tampilan seperti ini



- 4) Kemudian muncul output hasil uji normalitas (normal jika P-Value > 0,05)

Test of Univariate Normality for Continuous Variables

Variable	Skewness		Kurtosis		Skewness and Kurtosis	
	Z-Score	P-Value	Z-Score	P-Value	Chi-Square	P-Value
X11	-4.498	0.000	-10.615	0.000	132.921	0.000
X12	4.609	0.000	-9.237	0.000	106.563	0.000
X13	6.065	0.000	-0.862	0.388	37.530	0.000
X14	3.828	0.000	-97.606	0.000	9541.544	0.000
X21	0.625	0.532	-1.054	0.292	1.500	0.472
X22	2.620	0.009	3.689	0.000	20.473	0.000
X23	0.087	0.930	-3.083	0.002	9.515	0.009
X24	0.770	0.441	-2.183	0.029	5.359	0.069
X31	0.170	0.865	-0.601	0.548	0.389	0.823
X32	1.512	0.130	33.434	0.000	1120.135	0.000
X33	-0.015	0.988	1.231	0.218	1.516	0.469
X34	-1.195	0.232	-1.782	0.075	4.603	0.100
Z1	0.185	0.853	2.000	0.045	4.036	0.133
Z2	0.122	0.903	0.324	0.746	0.120	0.942
Z3	0.955	0.340	2.875	0.004	9.175	0.010
Y11	2.457	0.014	35.677	0.000	1278.892	0.000
Y12	0.240	0.810	2.127	0.033	4.582	0.101
Y13	5.951	0.000	-1.244	0.213	36.966	0.000

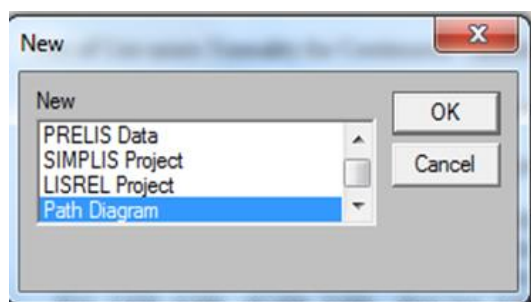
Relative Multivariate Kurtosis = 1.045

Test of Multivariate Normality for Continuous Variables

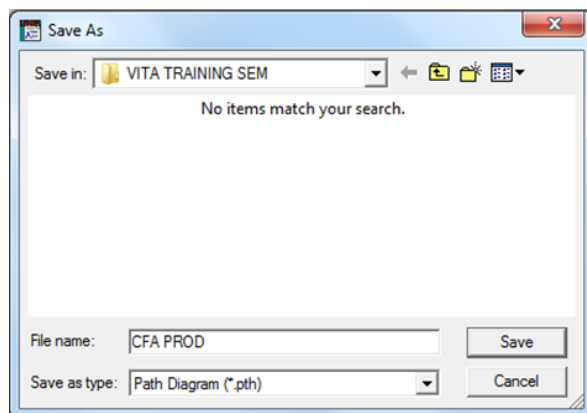
Skewness			Kurtosis			Skewness and Kurtosis	
Value	Z-Score	P-Value	Value	Z-Score	P-Value	Chi-Square	P-Value
110.830	33.130	0.000	376.375	4.279	0.000	1115.890	0.000

Cara Membuat Pengujian CFA (Model Pengukuran)

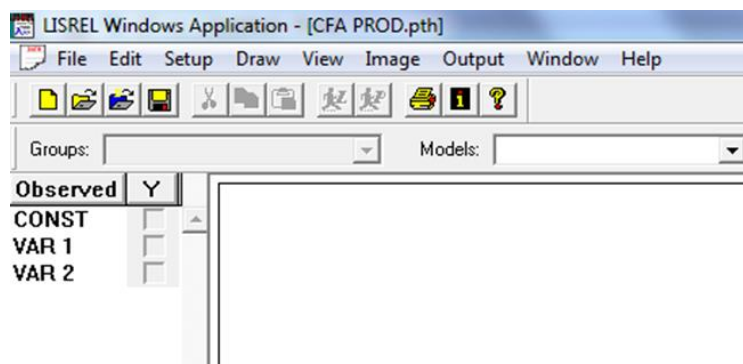
1. Klik menu New kemudian pilih path diagram → OK



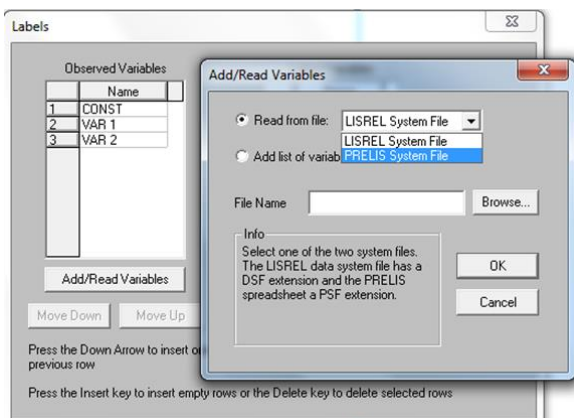
2. Simpan file pada menu Save as



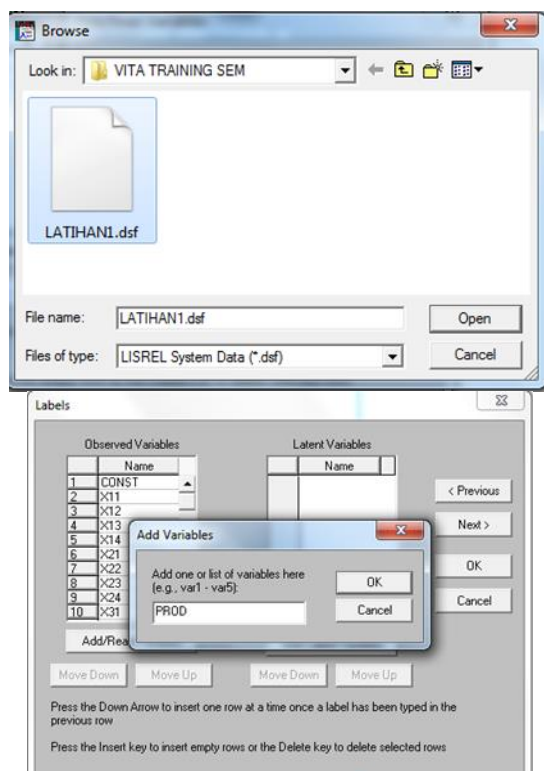
3. Akan muncul jendela seperti ini



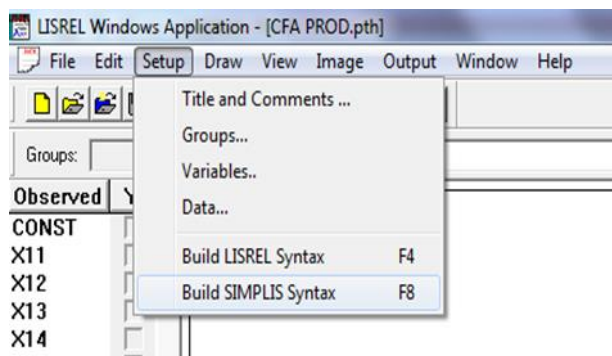
4. Tahapan Pemodelan CFA: Klik Setup-Variable-akan muncul tampilan berikut



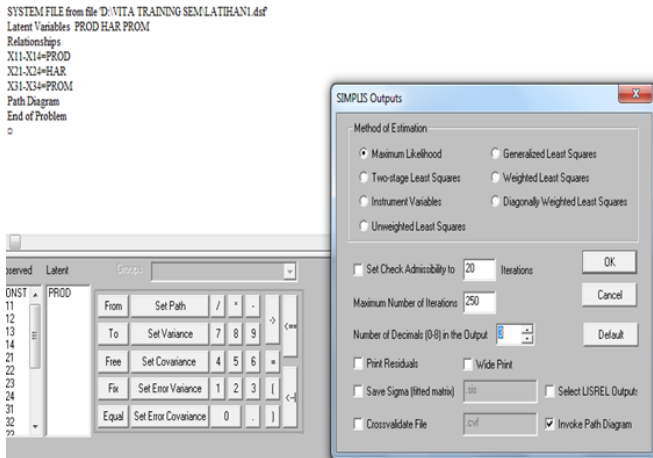
5. Pilih Browse untuk mengambil data yang sudah diimport—OPEN



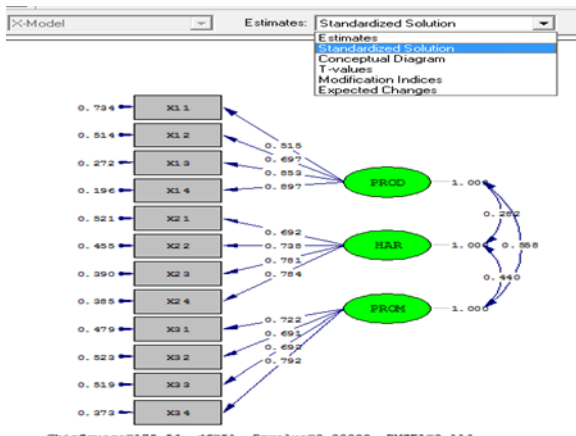
6. Kemudian klik SETUP kembali—Build Simplis Sintax seperti ini



7. Pada panel Simplis Output, non aktifkan set check admissibility to dan Number of decimal (0-8) in the output isikan dengan 3--OK. Terakhir Buat Persamaan seperti ini pada system file.



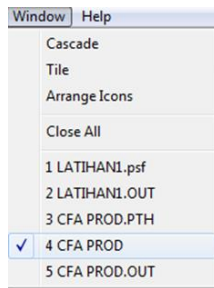
8. Akan Muncul Hasil gambar CFA



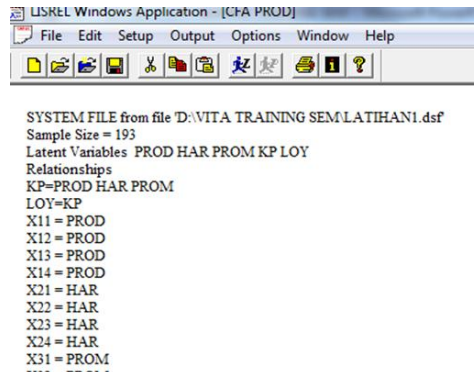
Ketentuannya adalah jika nilai loading factor lebih dari 0.5 maka indikator pada masing-masing kontruk adalah valid. Setelah ini dilanjutkan dengan uji model struktural.

Cara Membuat Pengujian Model Struktural

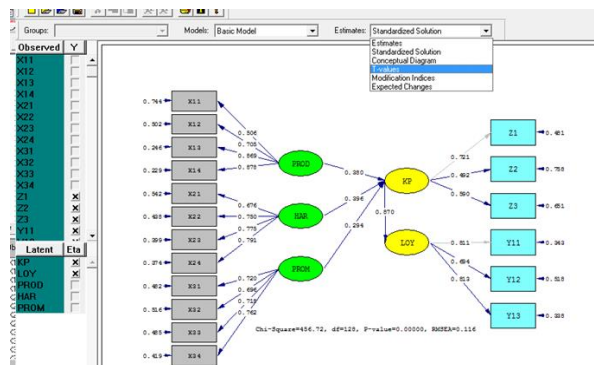
1. Pilih menu Windows lalu klik CFA Prod



2. Lengkapi semua syntaxnya seperti gambar berikut lalu save as menjadi model struktural lalu RUN (icon orang berlari).

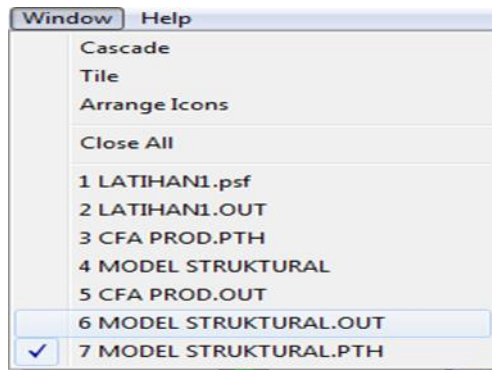


3. Berikut Hasil Gambar Model Struktural



Untuk Melihat Rincian Output SEM

1. Pilih meni windows dan klik Model structural.OUT



Maka akan muncul rincian output LISREL berikut.

DATE: 9/ 7/2021

TIME: 13:16

L I S R E L 8.70

BY

Karl G. Jöreskog & Dag Sörbom

This program is published exclusively by

Scientific Software International, Inc.

7383 N. Lincoln Avenue, Suite 100

Lincolnwood, IL 60712, U.S.A.

Phone: (800)247-6113, (847)675-0720, Fax: (847)675-2140

Copyright by Scientific Software International, Inc., 1981-2004

Use of this program is subject to the terms specified in the

Universal Copyright Convention.

Website: www.ssicentral.com

The following lines were read from file E:\VITA TRAINING SEM\Model Struktural1.SPJ:

Raw Data from file 'E:\VITA TRAINING SEM\LATIHAN1.psf'

Latent Variables PROD HAR PROM KP LOY

Relationships

KP=PROD HAR PROM

LOY=KP

X11 = PROD

X12 = PROD

X13 = PROD

X14 = PROD

X21 = HAR

X22 = HAR

X23 = HAR

X24 = HAR

X31 = PROM

X32 = PROM

X33 = PROM

X34 = PROM

Z1-Z3=KP

Y11-Y13=LOY

Path Diagram

End of Problem

Sample Size = 193

Covariance Matrix

	Z1	Z2	Z3	Y11	Y12	Y13
	-----	-----	-----	-----	-----	-----
Z1	0.26					
Z2	0.16	0.33				
Z3	0.13	0.10	0.22			
Y11	0.10	0.08	0.07	0.24		
Y12	0.13	0.09	0.08	0.14	0.26	
Y13	0.11	0.05	0.06	0.14	0.12	0.18
X11	0.06	0.05	0.06	0.10	0.07	0.06
X12	0.07	0.05	0.07	0.10	0.08	0.09
X13	0.10	0.04	0.04	0.09	0.08	0.10

X14	0.08	0.02	0.06	0.10	0.09	0.09
X21	0.07	0.04	0.09	0.10	0.07	0.07
X22	0.06	0.04	0.05	0.10	0.08	0.08
X23	0.06	0.07	0.07	0.17	0.09	0.10
X24	0.09	0.11	0.11	0.14	0.11	0.10
X31	0.09	0.07	0.06	0.07	0.11	0.06
X32	0.09	0.13	0.05	0.08	0.10	0.08
X33	0.14	0.10	0.11	0.10	0.12	0.08
X34	0.08	0.05	0.06	0.09	0.10	0.07

Covariance Matrix

	X11	X12	X13	X14	X21	X22
-----	-----	-----	-----	-----	-----	-----
X11	0.21					
X12	0.09	0.21				
X13	0.07	0.11	0.18			
X14	0.10	0.14	0.15	0.22		
X21	0.16	0.02	0.03	0.06	0.36	
X22	0.07	0.05	0.03	0.03	0.11	0.18
X23	0.12	0.04	0.02	0.04	0.23	0.17
X24	0.11	0.07	0.04	0.06	0.22	0.16
X31	0.09	0.05	0.09	0.09	0.05	0.02
X32	0.05	0.03	0.08	0.06	0.08	0.06
X33	0.13	0.08	0.10	0.11	0.11	0.05
X34	0.10	0.04	0.09	0.09	0.12	0.06

Covariance Matrix

	X23	X24	X31	X32	X33	X34
-----	-----	-----	-----	-----	-----	-----
X23	0.44					
X24	0.24	0.40				
X31	0.04	0.07	0.27			
X32	0.08	0.10	0.13	0.25		
X33	0.07	0.10	0.14	0.11	0.30	
X34	0.11	0.09	0.18	0.17	0.17	0.34

Number of Iterations = 31

LISREL Estimates (Maximum Likelihood)

Measurement Equations

Z1 = 0.37*KP, Errorvar.= 0.13 , R² = 0.51

(0.016)

7.88

Z2 = 0.28*KP, Errorvar.= 0.25 , R² = 0.23

(0.045)

(0.027)

6.09

9.25

Z3 = 0.27*KP, Errorvar.= 0.14 , R² = 0.34

(0.037)

(0.016)

7.28

8.88

Y11 = 0.40*LOY, Errorvar.= 0.082 , R² = 0.66

(0.012)

6.68

Y12 = 0.35*LOY, Errorvar.= 0.13 , R² = 0.49

(0.036)

(0.016)

9.78

8.31

Y13 = 0.35*LOY, Errorvar.= 0.061 , R² = 0.66

(0.030)

(0.0092)

11.55

6.61

X11 = 0.23*PROD, Errorvar.= 0.16 , R² = 0.26

(0.033)

(0.017)

7.12

9.40

X12 = 0.32*PROD, Errorvar.= 0.11 , R² = 0.50

(0.030)

(0.012)

10.75

8.61

X13 = 0.37*PROD, Errorvar.= 0.043 , R² = 0.76

(0.025)

(0.0075)

14.42

5.76

$$X14 = 0.41*PROD, \text{Errorvar.} = 0.051, R^2 = 0.77$$

(0.028)	(0.0094)
14.64	5.46

$$X21 = 0.40*HAR, \text{Errorvar.} = 0.20, R^2 = 0.45$$

(0.041)	(0.024)
9.83	8.39

$$X22 = 0.32*HAR, \text{Errorvar.} = 0.075, R^2 = 0.57$$

(0.028)	(0.010)
11.52	7.47

$$X23 = 0.52*HAR, \text{Errorvar.} = 0.18, R^2 = 0.60$$

(0.043)	(0.025)
11.93	7.15

$$X24 = 0.50*HAR, \text{Errorvar.} = 0.15, R^2 = 0.63$$

(0.041)	(0.022)
12.24	6.89

$$X31 = 0.38*PROM, \text{Errorvar.} = 0.13, R^2 = 0.52$$

(0.035)	(0.017)
10.65	7.73

$$X32 = 0.35*PROM, \text{Errorvar.} = 0.13, R^2 = 0.49$$

(0.034)	(0.016)
10.24	7.97

$$X33 = 0.39*PROM, \text{Errorvar.} = 0.14, R^2 = 0.51$$

(0.037)	(0.019)
10.57	7.78

$$X34 = 0.45*PROM, \text{Errorvar.} = 0.15, R^2 = 0.58$$

(0.039)	(0.020)
11.47	7.15

Structural Equations

$$KP = 0.38*PROD + 0.41*HAR + 0.29*PROM, \text{Errorvar.} = 0.28, R^2 = 0.72$$

(0.085)	(0.078)	(0.092)	(0.077)
4.48	5.23	3.19	3.64

$$\text{LOY} = 0.88 \cdot \text{KP}, \text{Errorvar.} = 0.23, R^2 = 0.77$$

(0.10)	(0.073)
8.77	3.10

Reduced Form Equations

$$\text{KP} = 0.38 \cdot \text{PROD} + 0.41 \cdot \text{HAR} + 0.29 \cdot \text{PROM}, \text{Errorvar.} = 0.28, R^2 = 0.72$$

(0.085)	(0.078)	(0.092)
4.48	5.23	3.19

$$\text{LOY} = 0.34 \cdot \text{PROD} + 0.36 \cdot \text{HAR} + 0.26 \cdot \text{PROM}, \text{Errorvar.} = 0.44, R^2 = 0.56$$

(0.075)	(0.068)	(0.081)
4.50	5.26	3.20

Correlation Matrix of Independent Variables

	PROD	HAR	PROM
PROD	1.00		
HAR	0.28 (0.08) 3.59	1.00	
PROM	0.57 (0.06) 9.31	0.43 (0.07) 5.81	1.00

Covariance Matrix of Latent Variables

	KP	LOY	PROD	HAR	PROM
KP	1.00				
LOY	0.88	1.00			
PROD	0.67	0.59	1.00		
HAR	0.64	0.57	0.28	1.00	
PROM	0.69	0.61	0.57	0.43	1.00

Goodness of Fit Statistics

Degrees of Freedom = 128

Minimum Fit Function Chi-Square = 478.35 (P = 0.0)

Normal Theory Weighted Least Squares Chi-Square = 458.28 (P = 0.0)

Estimated Non-centrality Parameter (NCP) = 330.28

90 Percent Confidence Interval for NCP = (268.71 ; 399.43)

Minimum Fit Function Value = 2.49

Population Discrepancy Function Value (F0) = 1.72

90 Percent Confidence Interval for F0 = (1.40 ; 2.08)

Root Mean Square Error of Approximation (RMSEA) = 0.12

90 Percent Confidence Interval for RMSEA = (0.10 ; 0.13)

P-Value for Test of Close Fit (RMSEA < 0.05) = 0.00

Expected Cross-Validation Index (ECVI) = 2.83

90 Percent Confidence Interval for ECVI = (2.51 ; 3.19)

ECVI for Saturated Model = 1.78

ECVI for Independence Model = 20.94

Chi-Square for Independence Model with 153 Degrees of Freedom =
3983.57

Independence AIC = 4019.57

Model AIC = 544.28

Saturated AIC = 342.00

Independence CAIC = 4096.30

Model CAIC = 727.57

Saturated CAIC = 1070.92

Normed Fit Index (NFI) = 0.88

Non-Normed Fit Index (NNFI) = 0.89

Parsimony Normed Fit Index (PNFI) = 0.74

Comparative Fit Index (CFI) = 0.91

Incremental Fit Index (IFI) = 0.91

Relative Fit Index (RFI) = 0.86

Critical N (CN) = 68.49

Root Mean Square Residual (RMR) = 0.025

Standardized RMR = 0.094

Goodness of Fit Index (GFI) = 0.79

Adjusted Goodness of Fit Index (AGFI) = 0.72

Parsimony Goodness of Fit Index (PGFI) = 0.59

The Modification Indices Suggest to Add the

Path to	from	Decrease in Chi-Square	New Estimate
Z2	LOY	8.2	-0.40
Z3	LOY	12.4	-0.40
X11	HAR	38.5	0.21
X11	PROM	20.5	0.19
X33	PROD	8.8	0.13
KP	LOY	54.0	-2.86
LOY	PROD	11.3	0.38
LOY	HAR	12.4	0.38

The Modification Indices Suggest to Add an Error Covariance

Between	and	Decrease in Chi-Square	New Estimate
LOY	KP	54.0	-0.65
Z2	Z1	23.8	0.07
Z3	Z1	14.1	0.05
Y11	Z1	13.5	-0.04
Y13	Z2	8.3	-0.03
X11	Y11	9.8	0.03
X11	Y13	8.0	-0.02
X13	Z1	12.2	0.03
X13	Z3	10.2	-0.02
X13	Y13	10.3	0.02
X13	X11	13.7	-0.03
X21	X11	30.6	0.08
X21	X12	10.4	-0.04
X22	X21	12.1	-0.04
X23	Y11	13.8	0.04
X32	Z2	23.4	0.07
X33	Z1	8.7	0.04
X33	Z3	9.8	0.04
X33	X11	14.8	0.05
X33	X32	8.3	-0.04
X34	Z2	7.9	-0.05

Time used: 0.094 Seconds

Persamaan yang Terbentuk

Structural equations

$$KP = 0.380 \cdot PROD + 0.396 \cdot HAR + 0.294 \cdot PROM, \text{Errorvar.} = 0.297, R^2 = 0.703$$

$$\begin{array}{cccc} (0.0860) & (0.0791) & (0.0940) & (0.0790) \\ 4.415 & 5.001 & 3.131 & 3.766 \end{array}$$

$$LOY = 0.870 \cdot KP, \text{Errorvar.} = 0.243, R^2 = 0.757$$

$$\begin{array}{cc} (0.0998) & (0.0747) \\ 8.725 & 3.249 \end{array}$$

2. Respesifikasi model

The Modification Indices Suggest to Add the				
Path to	from	Decrease in Chi-Square	New Estimate	
Z2	LOY	8.3	-0.38	
Z3	LOY	13.4	-0.39	
X11	HAR	39.3	0.21	
X11	PROM	21.0	0.19	
X33	PROD	8.9	0.13	
KP	LOY	56.4	-2.67	
LOY	PROD	13.3	0.40	
LOY	HAR	13.8	0.40	

Respesifikasi ini akan dilakukan untuk memperbaiki nilai Goodness of fit agar lebih baik. Jika muncul seperti itu dalam output maka yang harus dilakukan adalah membuat syntax dalam lembar SPJ dan buat bahasa syntax:

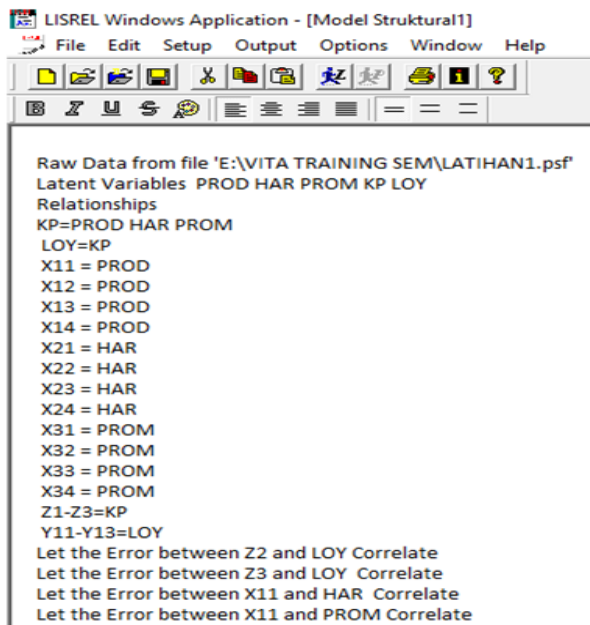
Let the Error between Z2 and LOY Correlate

Let the Error between Z3 and LOY Correlate

Let the Error between X11 and HAR Correlate

Let the Error between X11 and PROM Correlate, dst

Seperti tampilan berikut lalu klik RUN (ikon orang sedang berlari)



Setelah prosedur ini dilakukan, maka akan muncul nilai GOF akhir setelah dilakukan respesifikasi model. Prosedur ini merupakan langkah terakhir dalam analisis SEM.

DAFTAR PUSTAKA

- Arikunto, Suhaimi. (2010). *Prosedur Penelitian Suatu Pendekatan Praktik*. Jakarta: Rineka Cipta.
- Ghozali, Imam. (2005). *Structural Equation Modeling: Teori, Konsep, dan Aplikasi dengan Program LISREL*. Semarang: Badan Penerbit Undip.
- Ghozali, Imam. (2005). *Aplikasi Multivariate dengan Program SPSS*. Semarang: Badan Penerbit Universitas Diponegoro.
- Nasution, S. (2012). *Metode Research (Penelitian Ilmiah)*. Jakarta : PT Bumi Aksara.
- Ramadiani. (2005). *Analisis Pengukuran Keberhasilan Sistem Informasi Menggunakan Structural Equation Model dan LISREL*,. Thesis UGM.

ANALISIS REGRESI DATA PANEL

Pardomuan Robinson Sihombing, S.ST., M.Stat, C.PS

Badan Pusat Statistik

Salah satu aspek yang harus diperhatikan dalam melakukan pemodelan statistika adalah jenis data yang digunakan. Salah satu jenis data berdasarkan waktu pengumpulannya adalah data panel. Data panel merupakan gabungan data cross section dan time series (runtun/ deret waktu). Dengan kata lain, data panel merupakan data dari beberapa individu sama yang diamati dalam kurun waktu tertentu. Terkadang data panel disebut juga data longitudinal. 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.

Keuntungan dengan menggunakan data panel dalam pemodelan regresi maka akan menghasilkan degree of freedom yang lebih besar, sehingga dapat mengatasi masalah penghilangan variabel (omitted variabel). Selain itu juga dapat mengurangi bias dalam pengestimasian karena data cukup banyak. Hal lain yang dapat kita pelajari adalah terkait perilaku individu serta perubahannya yang bersifat dinamis (Gujarati, 2004).

Terdapat tiga pendekatan yang digunakan dalam model panel yaitu Common/ Pooled Effects, Fixed Effects dan Random Effects (Brooks, 2008). Penjelasan masing-masing ketiga pendekatan adalah sebagai berikut:

- **Model Common/ Pooled Effects (PEM)**

Pendekatan ini tidak memperhatikan dimensi individu maupun waktu, disebut juga Pooled Regression. Metode estimasinya menggunakan Ordinary Least Squares (OLS).

Model persamaan regresinya :

$$Y_{it} = \beta_0 + \beta_1 X_{1it} + \dots + \beta_p X_{pit} + u_{it}$$

- **Model Fixed Effects (FEM)**

Model ini mengasumsikan bahwa dalam berbagai kurun waktu, karakteristik masing-masing individu adalah berbeda. Perbedaan tersebut dicerminkan oleh nilai intersep pada model estimasi yang berbeda untuk setiap individu.

Model persamaan regresinya :

$$Y_{it} = \beta_{0i} + \beta_1 X_{1it} + \dots + \beta_p X_{pit} + u_{it}$$

Model di atas biasanya dituliskan dalam bentuk *dummy variabel* untuk menggantikan perbedaan intersep yang ada, sehingga dapat dituliskan sebagai berikut :

$$Y_{it} = \beta_0 + \alpha_2 D_{2i} + \alpha_3 D_{3i} + \dots + \alpha_N D_{Ni} + \beta_1 X_{1it} + \dots + \beta_p X_{pit} + u_{it}$$

Asumsi pada model PEM dan FEM adalah:

1. $u_{it} \sim N(0, \sigma^2)$, data berdistribusi normal
2. $E(u_{it}u_{is}) = E(u_{it}u_{jt}) = E(u_{it}u_{js}) = 0 (i \neq j; t \neq s)$, tidak terjadi masalah autokorelasi
3. $E(\varepsilon_i, \varepsilon_j) = \sigma^2$, varian residual homogen
4. Tidak ada multikolinieritas antar variabel bebas yang digunakan

- **Model Random Effects (REM)**

Model ini juga mengasumsikan bahwa dalam berbagai kurun waktu, karakteristik masing-masing individu adalah berbeda. Hanya saja, dalam REM perbedaan tersebut dicerminkan oleh *error* dari model.

Model persamaan regresinya :

$$Y_{it} = \beta_{0i} + \beta_1 X_{1it} + \dots + \beta_p X_{pit} + u_{it} \text{ Dimana : } \beta_{0i} = \beta_0 + \varepsilon_i$$

Sehingga modelnya dapat pula dituliskan sebagai berikut :

$$Y_{it} = \beta_0 + \beta_1 X_{1it} + \dots + \beta_p X_{pit} + (\varepsilon_i + u_{it})$$

$$\text{atau } Y_{it} = \beta_0 + \beta_1 X_{1it} + \dots + \beta_p X_{pit} + w_{it} \text{ dimana : } w_{it} = \varepsilon_i + u_{it}$$

$$\text{Asumsi : } \varepsilon_i \sim N(0, \sigma_\varepsilon^2) \quad u_{it} \sim N(0, \sigma_u^2)$$

$$E(\varepsilon_i u_{it}) = 0 \quad E(\varepsilon_i \varepsilon_j) = 0 \quad (i \neq j)$$

$$E(u_{it}u_{is}) = E(u_{it}u_{jt}) = E(u_{it}u_{js}) = 0 (i \neq j; t \neq s)$$

Akibat dari asumsi-asumsi di atas maka :

$$E(w_{it}) = 0$$

$$Var(w_{it}) = \sigma_{\varepsilon}^2 + \sigma_u^2$$

Jika $\sigma_{\varepsilon}^2 = 0$ maka REM dapat diestimasi dengan OLS (cukup menggunakan *Common Effects*), jika tidak maka diestimasi dengan *Generalized Least Squares* (GLS).

Dari ketiga model yang ada dipilih salah satu model terbaik yang akan diintrepretasikan. Pemilihan model terbaik dilakukan melalui uji Chow Likelihood Ratio, uji Lagrange Multiplier Breusch Pagan dan uji Hausman (Mouchart, 2004).

✚ Memilih antara model **Common Effects VS Fixed Effects**

Untuk memilih model mana yang lebih cocok antara *Common Effects* ataukah *Fixed Effects*, dapat digunakan Uji Chow (*Chow Test*) atau *Restricted F-Test* sebagai berikut

Ho : Model *Common Effects* lebih baik daripada *Fixed Effects*

H1 : Model *Fixed Effects* lebih baik daripada *Common Effects*

Tingkat signifikansi : α

Statistik Uji :

$$F_{obs} = \frac{(R_{UR}^2 - R_R^2)/(N - 1)}{(1 - R_{UR}^2)/(NT - k)}$$

Dimana : N = jumlah individu (dalam hal ini komoditi)

T = jumlah series (tahun), k = jumlah parameter, termasuk *intercept*

R_{UR}^2 = koefisien determinasi (R^2) dari model *unrestricted*/model *Fixed Effects*

R_R^2 = koefisien determinasi (R^2) dari model *restricted*/model *Common Effects*

Kriteria Pengambilan Keputusan : Tolak Ho jika $F_{obs} > F_{\alpha; (N-1), (NT-k)}$ atau jika $P - value \leq \alpha$

✚ Memilih antara model **Common Effects VS Random Effects**

Untuk memilih model mana yang lebih cocok antara *Common Effects* ataukah *Random Effects*, dapat digunakan Uji *Lagrange Multiplier* (LM Test), yaitu sebagai berikut :

Ho : $\sigma_{\varepsilon}^2 = 0$ (*intersep tidak bersifat random atau stochastic*)

H1 : $\sigma_{\varepsilon}^2 \neq 0$ (*intersep bersifat random atau stochastic*)

Tingkat signifikansi : α

Statistik Uji :

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

Kriteria Pengambilan Keputusan : Tolak H_0 jika $LM > \chi^2_{\alpha;1}$ atau jika $P - value \leq \alpha$

Memilih antara model **Fixed Effects VS Random Effects**

Untuk memilih model mana yang lebih cocok antara *Fixed Effects* ataukah *Random Effects*, dapat digunakan Uji *Hausman (Hausman's Test)*, yaitu sebagai berikut :

H_0 : Model *Random Effects* lebih baik daripada *Fixed Effects*

H_1 : Model *Fixed Effects* lebih baik daripada *Random Effects*

Tingkat signifikansi : α

Statistik Uji : $\chi^2_{obs} = (\hat{\beta} - \hat{\beta}_{GLS})' \hat{\Psi}^{-1} (\hat{\beta} - \hat{\beta}_{GLS})$

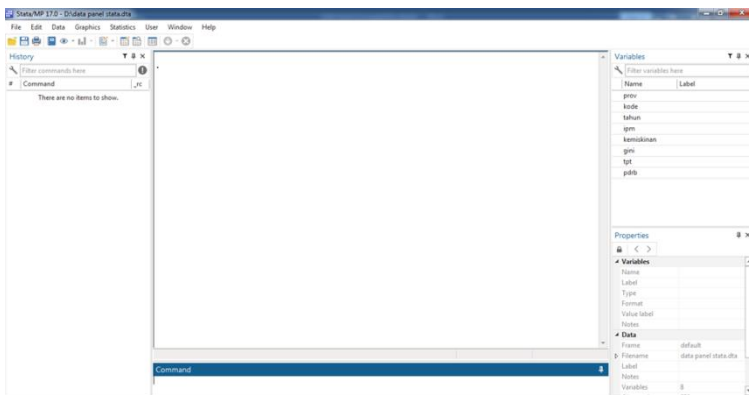
Kriteria Pengambilan Keputusan : Tolak H_0 jika $\chi^2_{obs} > \chi^2_{\alpha;p}$ atau jika $P - value \leq \alpha$

p = jumlah variabel bebas

Contoh kasus: Misalkan seorang peneliti tertarik melihat pengaruh Indeks Pembangunan Manusia (IPM), Gini Rasio dan Tingkat Pengangguran Terbuka (TPT) dan Pertumbuhan Ekonomi (PDRB) terhadap persentase kemiskinan di 34 Provinsi di Indonesia Tahun 2012-2019. Data bersumber dari Badan Pusat Statistik. Data data di akses pada CD yang disertakan dalam buku ini.

Langkah-langkah merun data panel

1. Buka data, dari menu open, pilih lokasi file data yang akan digunakan, dalam hal ini data panel.dta

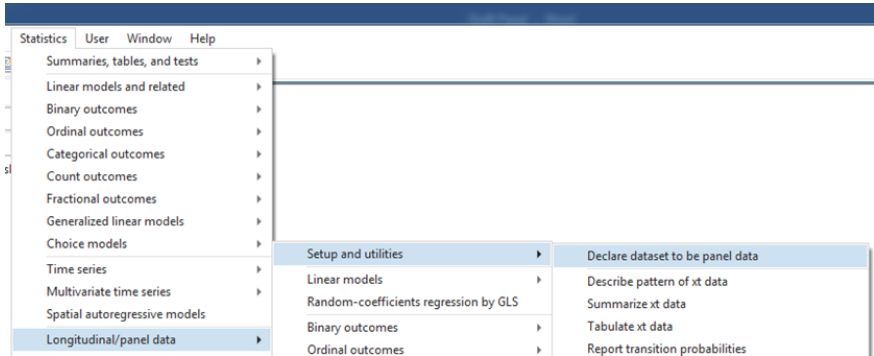


Pada sisi kanan layar akan terlihat data yang akan digunakan dalam model regresi data panel

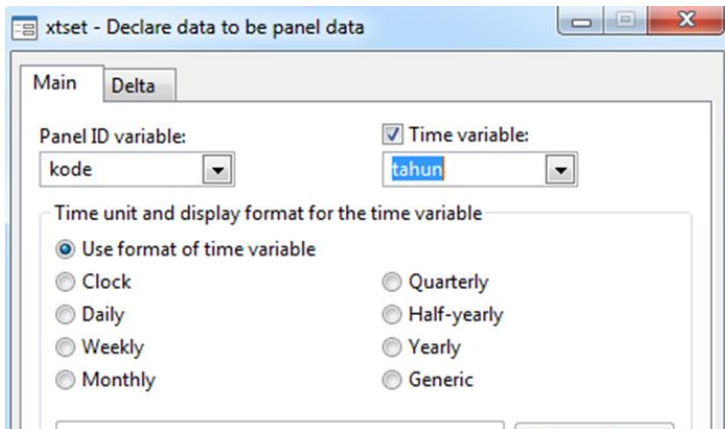
2. Deklarasikan sebagai data panel

a. menggunakan menu:

Statistics → Longitudinal/panel data → Setup and utilities → declare dataset to be panel data



Maka akan muncul jendela xtset



Pada Panel ID variable: klik dropdown dan pilih variabel cross section, misalnya “kode”

Klik centang time variable, dan klik dropdown dan pilih variabel time series, misalnya “tahun”

Lalu klik submit

b. dengan perintah pada kolom command

```
* xtset crosssection time series  
xtset no tahun
```

```
. xtset kode tahun
```

```
Panel variable: kode (strongly balanced)  
Time variable: tahun, 2012 to 2019  
Delta: 1 unit
```

Pada layar akan terlihat bahwa data yang digunakan adalah data panel, dengan periode dari tahun 2012 hingga 2019. Dapat juga

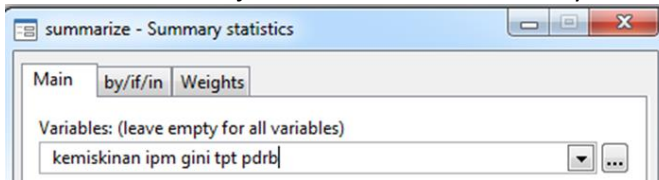
3. Selanjutnya adalah melakukan analisis deskriptif terhadap seluruh data ,

a. menggunakan menu:

Statistics → Summaries, tables, and tests → summary and descriptive statistics → summary statistics



Maka akan muncul jendela summarize-Summary statistics



Kata kunci variabel: klik dropdown dan pilih variabel penelitian misalnya “kemiskinan ipm gini tpt pdrb”. Lalu klik submit

b. dengan perintah pada kolom command

```
* sum y x1 x2 ... xp  
sum kemiskinan ipm gini tpt pdrb  
. sum kemiskinan ipm gini tpt pdrb
```

Variable	Obs	Mean	Std. dev.	Min	Max
kemiskinan	272	11.52143	6.14028	3.47	31.13
ipm	272	68.91699	4.293838	55.55	80.76
gini	272	.3666544	.038412	.27	.44
tpt	272	5.254522	1.95909	1.37	10.51
pdrb	272	5.497904	2.490634	-15.72	21.76

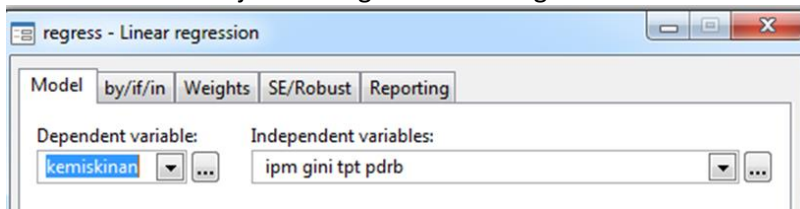
4. Selanjutnya adalah merun model PEM, dan menyimpan modelnya,

a. menggunakan menu:

Statistics → Linear modelas and related → Linear regression



Maka akan muncul jendela regress-Linear regression



Pada menu dependent variabel: klik dropdown dan pilih “kemiskinan”

Pada menu independent variable: klik droddown dan pilih “ipm gini tpt pdrb”

Lalu klik submit

b. dengan perintah pada kolom command,

* regress y x1 x2 ... xp

regress kemiskinan ipm gini tpt pdrb

estimate store pooled

. regress kemiskinan ipm gini tpt pdrb

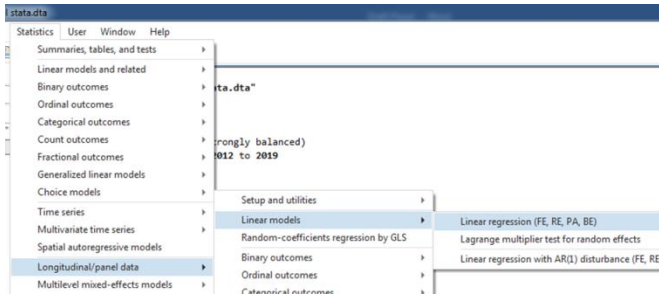
Source	SS	df	MS	Number of obs	=	272
Model	5488.31104	4	1372.07776	F(4, 267)	=	77.46
Residual	4729.21218	267	17.7124052	Prob > F	=	0.0000
				R-squared	=	0.5371
				Adj R-squared	=	0.5302
Total	10217.5232	271	37.7030377	Root MSE	=	4.2086

kemiskinan	Coefficient	Std. err.	t	P> t	[95% conf. interval]	
ipm	-.9473617	.0611537	-15.49	0.000	-1.067766	-.8269569
gini	50.62541	6.77846	7.47	0.000	37.27937	63.97144
tpt	-.0467209	.1351985	-0.35	0.730	-.3129117	.2194699
pdrb	-.382182	.1056091	-3.62	0.000	-.5901145	-.1742495
_cons	60.59541	4.806389	12.61	0.000	51.13217	70.05866

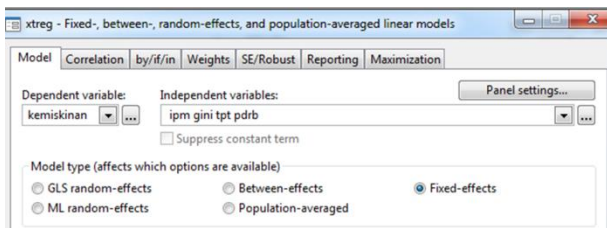
. estimate store pooled

5. Selanjutnya adalah merun model FEM, dan menyimpan modelnya, dengan perintah pada kolom command:

a. Statistics → longitudinal/panel data → linera model → Linera regression (FE, RE, PA, BE)



Maka akan muncul jendela : xtreg: Fixed-, between-, random-effect and poulation-averaged linier models



Pada menu dependent variabel: klik dropdown dan pilih “kemiskinan”

Pada menu independent variable: klik drodown dan pilih “ipm gini tpt pdrb”

Pada model type (affects which options are avaribel): pilih Fixed-effect
Lalu klik submit

b. dengan perintah pada kolom command:

* xtreg y x1 x2 ... xp, fe

xtreg kemiskinan ipm gini tpt pdrb, fe

estimate store fixed

```
. xtreg kemiskinan ipm gini tpt pdrb, fe

Fixed-effects (within) regression              Number of obs   =    272
Group variable: kode                          Number of groups =    34
                                             Obs per group:
                                             min   =      8
                                             avg   =     8.0
                                             max   =      8

R-squared:
  Within = 0.5142
  Between = 0.4282
  Overall = 0.4168

corr(u_i, Xb) = 0.4047                        F(4, 234)        =    61.92
                                              Prob > F         =    0.0000
```

		Coefficient	Std. err.	t	P> t	[95% conf. interval]
kemiskinan						
ipm		-.4466945	.0336599	-13.27	0.000	-.5130096 -.3803794
gini		-1.787561	2.594128	-0.69	0.491	-6.898392 3.32327
tpt		.0337096	.0523446	0.64	0.520	-.0684173 .1368365
pdrb		.0043947	.019136	0.23	0.819	-.0333061 .0420955
_cons		42.7604	2.943309	14.53	0.000	36.96163 48.55917
sigma_u		5.1529263				
sigma_e		.64366695				
rho		-.98463649				(fraction of variance due to u_i)

```

F test that all u_i=0: F(33, 234) = 338.81      Prob > F = 0.0000

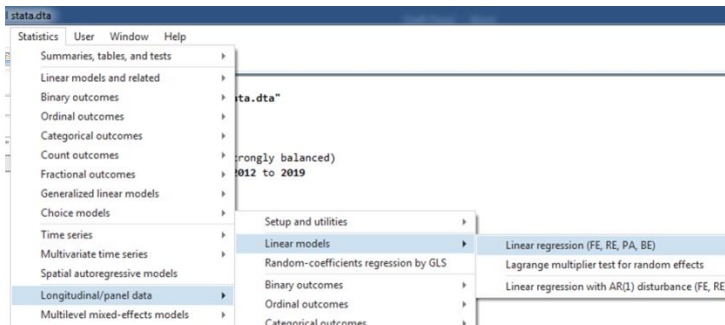
.
.
.      estimate store fixed

```

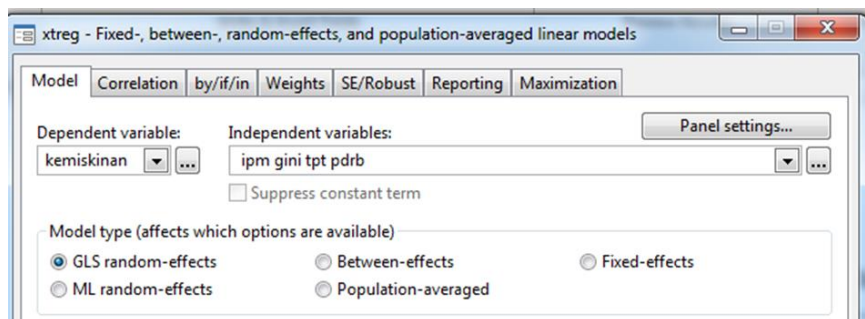
6. Selanjutnya adalah merun model REM, dan menyimpan modelnya,

a. menggunakan menu:

Statistics → longitudinal/panel data → linera model → Linera regression (FE, RE, PA, BE)



Maka akan muncul jendela : xtreg: Fixed-, between-, random-effect and poulation-averaged linier models



Pada menu dependent variabel: klik dropdown dan pilih “kemiskinan”

Pada menu independent variable: klik droddown dan pilih “ipm gini tpt pdrb”

Pada model type (affects which options are avaribel): pilih GLS Random-effect

Lalu klik submit

b. dengan perintah pada kolom command:

```
* xtreg y x1 x2 ... xp, re
xtreg kemiskinan ipm gini tpt pdrb, re
estimate store random
```

```
. xtreg kemiskinan ipm gini tpt pdrb, re
```

Random-effects GLS regression	Number of obs =	272
Group variable: kode	Number of groups =	34
R-squared:	Obs per group:	
Within = 0.5139	min =	8
Between = 0.4318	avg =	8.0
Overall = 0.4205	max =	8
corr(u_i, X) = 0 (assumed)	Wald chi2(4) =	251.02
	Prob > chi2 =	0.0000

	Coefficient	Std. err.	z	P> z	[95% conf. interval]
ipm	-.4634484	.0341921	-13.55	0.000	-.5304638 -.396433
gini	-1.477669	2.666031	-0.55	0.579	-6.702995 3.747656
tpt	.0192934	.0538344	0.36	0.720	-.0862201 .1248068
pdrb	.0007402	.0198925	0.04	0.970	-.0382484 .0397288
_cons	43.89725	3.06901	14.30	0.000	37.8821 49.9124
sigma_u	3.9578988				
sigma_e	.6436695				
rho	.97423346	(fraction of variance due to u_i)			

```

.
. estimate store random

```

7. Untuk melakukan pengujian pemilihan model pooled terhadap fixed dengan uji Chow dapat dilihat ada uji F, ada tabel model fix pada langkah 5

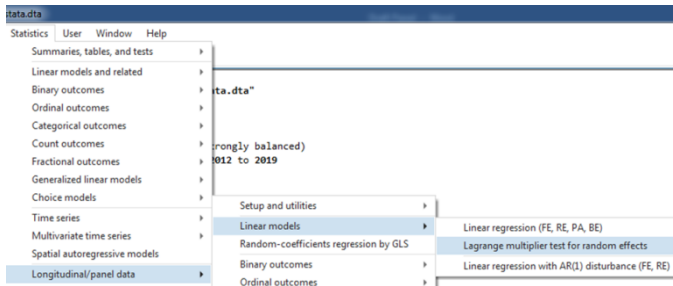
F test that all $u_i=0$: $F(33, 234) = 338.81$ Prob > F = 0.0000

Dari hasil di atas karena nilai Prob (F)=0.000 < $\alpha=0.05$ maka tolak H_0 dan dikatakan model fixed lebih baik dari model pooled.

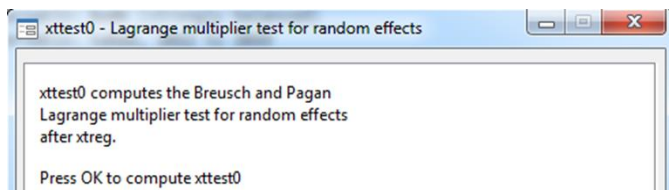
8. Untuk melakukan pengujian pemilihan model pooled terhadap random dengan uji LM BP

a. menggunakan menu

Statistics → longitudinal/panel data → linera model → Langrange multiplier test for random effect



Maka akan muncul jendela xttest0: Lagrange multiplier tests for random effect



Klik Submit

- b. dapat dilihat dengan command:
xttest0

```
. xttest0

Breusch and Pagan Lagrangian multiplier test for random effects

kemiskinan[kode,t] = Xb + u[kode] + e[kode,t]

Estimated results:

```

	Var	SD = sqrt(Var)
kemiskin~n	37.70304	6.14028
e	.4143071	.6436669
u	15.66496	3.957899

```

Test: Var(u) = 0
      chibar2(01) =    671.83
Prob > chibar2 =    0.0000

```

Dari hasil di atas karena nilai Prob (Chi2)=0.000 < α =0.05 maka tolak H_0 dan dikatakan model random lebih baik dari model pooled.

9. Untuk melakukan pengujian pemilihan model random terhadap fixed dengan uji Hausman dapat dilihat dengan command:

*hausman nama model fixed nama model random
hausman fixed random

```
. hausman fixed random
```

	Coefficients		(b-B)	sqrt(diag(V_b-V_B))
	(b)	(B)	Difference	Std. err.
	fixed	random		
ipm	-.4466945	-.4634484	.0167539	.
gini	-1.787561	-1.477669	-.3098914	.
tpt	.0337096	.0192934	.0144162	.
pdrb	.0043947	.0007402	.0036545	.

b = Consistent under H_0 and H_a ; obtained from xtreg.
B = Inconsistent under H_a , efficient under H_0 ; obtained from xtreg.

Test of H_0 : Difference in coefficients not systematic

```
chi2(4) = (b-B)'[(V_b-V_B)^(-1)](b-B)
        = -10.26
```

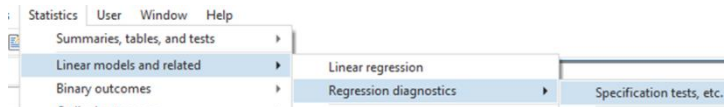
Warning: chi2 < 0 ==> model fitted on these data

Dari hasil di atas karena nilai Prob (χ^2) = 0.000 < α = 0.05 maka tolak H_0 dan dikatakan model fixed lebih baik dari model pooled.

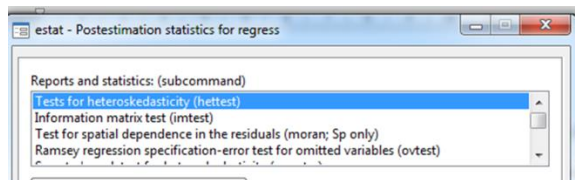
10. Langkah selanjutnya adalah melakukan pengujian asumsi klasik untuk model klasik:

a. menggunakan menu

Klik statistics → Linear models and relates → Regression diagnostic
→ Specification test, etc



Maka akan muncul menu: estat-Postestimation statistics for regress



Pilih asumsi klasik yang diinginkan seperti test for heteroskedasticity (hettest) dan lainnya secara bergantian.

b. atau menggunakan sintaks secara simultan

```
regress kemiskinan ipm gini tpt pdrb
*simpan residual untuk uji normalitas
predict resid, r
*uji multikol
vif
*uji normalitas
sktest resid
swilk resid
*uji hetero
hettest
*uji autokol
* xterial y x1 x2 ... xp
xtserial kemiskinan ipm gini tpt pdrb
```

Berikut output dalam uji asumsi klasik:

a. Uji Multikolinearitas

Suatu model mengalami pelanggaran asumsi multikolienaritas data jika nilai VIF variabel independennya lebih besar dari 10.

`. vif`

Variable	VIF	1/VIF
tpt	1.07	0.931656
pdrb	1.06	0.944683
ipm	1.05	0.947921
gini	1.04	0.964079
Mean VIF	1.06	

Dari hasil tabel di atas nilai VIF seluruh variabel independen < 10 , sehingga modelnya bebas dari asumsi multikolinearitas.

b. Normalitas Data

Suatu model dikatakan memiliki data yang berdistribusi normal jika nilai probabilita pengujian (baik skewnes kurtosi maupun pengujian shapiro wilk) lebih besar dari 0.05.

`. sktest resid`

Skewness and kurtosis tests for normality

Variable	Obs	Pr(skewness)	Pr(kurtosis)	Joint test	
				Adj chi2(2)	Prob>chi2
resid	272	0.0634	0.0113	9.00	0.0111

`. swilk resid`

Shapiro-Wilk W test for normal data

Variable	Obs	W	V	z	Prob>z
resid	272	0.97669	4.554	3.542	0.00020

Dari hasil di atas baik untuk pengujian sktest maupun swilk test nilai $\text{prob}(z) < \alpha = 0.05$ sehingga secara pengujian datanya dianggap belum berdistribusi normal. Akan tetapi berdasarkan teorema limit pusat (central limit theorema) jika datanya sudah cukup besar (lebih dari 30 data), dapat diasumsikan mendekati distribusi normal (Hildebrand, 2008)

c. Uji Heterokedastis

Suatu model dikatakan mengalami heterokedastis jika nilai probabilitas pengujian (Breusch Pagan) lebih kecil dari 0.05.

```
. hettest

Breusch-Pagan/Cook-Weisberg test for heteroskedasticity
Assumption: Normal error terms
Variable: Fitted values of kemiskinan

H0: Constant variance

      chi2(1) =    1.85
Prob > chi2 = 0.1736
```

Dari hasil di atas nilai prob (χ^2) > $\alpha=0.05$ sehingga secara pengujian modelnya bebas dari asumsi heterokedastis.

d. Uji Autorelasi

Suatu model dikatakan mengalami autokorelasi jika nilai probabilitas pengujian (wooldridge) lebih kecil dari 0.05.

```
. xtserial kemiskinan ipm gini tpt pdrb

Wooldridge test for autocorrelation in panel data
H0: no first order autocorrelation

      F( 1,    33) =    90.698
      Prob > F =    0.0000
```

Dari hasil di atas baik untuk pengujian wooldridge nilai prob(F) < $\alpha=0.05$ sehingga secara pengujian modelnya masih mengalami pelanggaran asumsi klasik.

11. Beberapa alternatif model yang dapat digunakan dalam mengatasi pelanggaran asumsi klasik:

- Jika terjadi pelanggaran asumsi normalitas, dan datanya sedikit, maka jika memungkinkan dapat melakukan transformasi datanya ke dalam bentuk logaritma natural. Cara lain yang dapat dilakukan adalah dengan mengeluarkan atau memberikan penimbang untuk data yang dianggap outlier/ pencilan dengan metode robust estimation. Cara lainnya, jika dependen tidak mengikuti distribusi normal tetapi mengikuti distribusi keluarga eksponensial maka dapat menggunakan model regresi linier terampat dengan data panel (Generalized Linier Model/ GLM).
- Jika terjadi pelanggaran asumsi heterokedastis maka dapat menggunakan metode estimasi general least square/ GLS (Greene, 2012). Adapun *command* yang digunakan adalah
* xtglm y x1 x2 ... xp

xtgls kemiskinan ipm gini tpt pdrb

- c. Jika terjadi pelanggaran asumsi autokorelasi maka dapat menambahkan lag data variabel dependen sebagai tambahan variabel independen di dalam model atau menggunakan model AR (autoregressive). Adapun *command* yang digunakan adalah:

* xtreg y x1 x2 ... xp l.y, fe

xtreg kemiskinan ipm gini tpt pdrb l.kemiskinan, fe

*atau

* xtregar y x1 x2 ... xp

xtregar kemiskinan ipm gini tpt pdrb, fe

xtregar kemiskinan ipm gini tpt pdrb, fe

- d. Jika terjadi pelanggaran asumsi autokorelasi dan heterokedastis secara bersamaan maka dapat menggunakan metode estimasi panel panel-corrected standard error/ PCSE (Hoechle, 2007). Adapun *command* yang digunakan adalah

* xtpcse y x1 x2 ... xp, corr(ar1)

xtpcse kemiskinan ipm gini tpt pdrb, corr(ar1)

- e. Jika terjadi pelanggaran multikolinearitas maka variabel salah satu variabel yang memiliki korelasi tinggi tidak disertakan dalam model atau digunakan reduksi variabel yang berkorelasi tinggi dengan analisis komponen utama (AKU)

12. Dari contoh di atas, pada model yang terbentuk hanya terjadi pelanggaran asumsi autokorelasi sehingga model FEM yang terpilih ditambahkan unsur AR seperti pada poin 9c. Berikut hasilnya:

```
. xtregar kemiskinan ipm gini tpt pdrb , fe
```

FE (within) regression with AR(1) disturbances	Number of obs	=	238
Group variable: kode	Number of groups	=	34
R-squared:	Obs per group:		
Within = 0.1757	min =		7
Between = 0.4577	avg =		7.0
Overall = 0.4422	max =		7
corr(u_i, Xb) = 0.4943	F(4,200)	=	10.66
	Prob > F	=	0.0000

	Coefficient	Std. err.	t	P> t	[95% conf. interval]
ipm	-.3432772	.0633969	-5.41	0.000	-.4682893 -.2182652
gini	1.488657	2.301498	0.65	0.518	-3.049657 6.026972
tpt	.0327196	.0460764	0.71	0.478	-.0581382 .1235775
pdrb	.0228186	.0136116	1.68	0.095	-.0040221 .0496593
_cons	34.18427	1.858418	18.39	0.000	30.51966 37.84888
rho_ar	.61259922				
sigma_u	5.2500758				
sigma_e	.49189448				
rho_fov	.99129804				(fraction of variance because of u_i)

F test that all u_i=0: F(33,200) = 78.31	Prob > F = 0.0000
--	-------------------

13. Langkah selanjutnya adalah melakukan intrepetasi terhada model final yang terbentuk

- a. Persamaan regres yang terbentuk adalah
kemiskinan = $34.184 - 0.3432 \text{ ipm} + 1.488 \text{ gini} + 0.032 \text{ tpt} + 0.033 \text{ pdrb}$
- b. Uji koefisien determinasi/ adjusted r square

```
R-squared:
  Within  = 0.1757
  Between = 0.4577
  Overall = 0.4422
```

Berdasarkan uji koefisien determinasi dengan nilai adjusted r square sebesar 0.4422 dapat diartikan variabel kemiskinan dapat dijelaskan oleh data ipm, gini, tpt dan pdrb sebesar 44.22 persen sisanya oleh variabel lain di luar model

- c. Uji simultan/ Uji F

```
F(4,200)      = 10.66
Prob > F      = 0.0000
```

Berdasarkan uji simultan dengan prob.(F) dalam model = $0.000 < \alpha = 0.05$ maka dapat dikatakan bahwa secara bersama-sama variabel ipm, gini, tpt dan pdrb mempengaruhi secara linier terhadap persentase kemiskinan, dengan kata lain modelnya sudah sesuai/ fit.

- d. Uji parsial/ uji hipotesis

kemiskinan	Coefficient	Std. err.	t	P> t
ipm	-.3432772	.0633969	-5.41	0.000
gini	1.488657	2.301498	0.65	0.518
tpt	.0327196	.0460764	0.71	0.478
pdrb	.0228186	.0136116	1.68	0.095
_cons	34.18427	1.858418	18.39	0.000

Berdasarkan uji parsial dengan $\text{prob.}(t)$ dalam model, hanya variabel IPM yang mempengaruhi kemiskinan dengan $\text{prob.}(t) = 0.000 < \alpha = 0.05$ sedangkan belum cukup variabel gini, tpt dn pdrb memengaruhi kemiskinan karena $\text{prob.}(t) = 0.000 > \alpha = 0.05$. Tanda koefisien IPM bernilai negatif artinya kenaikan nilai IPM akan menurunkan persentase kemiskinan. Sedangkan tanda koefisien Gini dan TPT positif artinya kenaikan nilai gini dan TPT akan menaikkan persentase kemiskinan.

DAFTAR PUSTAKA

- Brooks, C. (2008). *Introductory Econometrics for Finance 2nd Edition*. New York: Cambridge University Press.
- Greene, W. H. (2012). *Econometric Analysis. 7th edition*. Upper Saddle River, NJ: Prentice Hall.
- Gujarati, D. N. (2004). *Basic Econometrics 4th edition*. Mass: McGraw-Hill.
- Hildebrand, A. (2008). *Actuarial Statistics I*. Urbana-Champaign: University of Illinois.
- Hoechle, D. (2007). Robust standard errors for panel regressions with cross-sectional dependence. *Stata Journal*, 7, 281-312.
- Mouchart, M. (2004). *The Econometrics of Panel Data*. Brussel: Université Catholique de Louvain.

VECTOR AUTOREGRESSION (VAR) DENGAN EIEWS

Dr. Early Ridho Kismawadi, S.E.I., M.A
IAIN Langsa

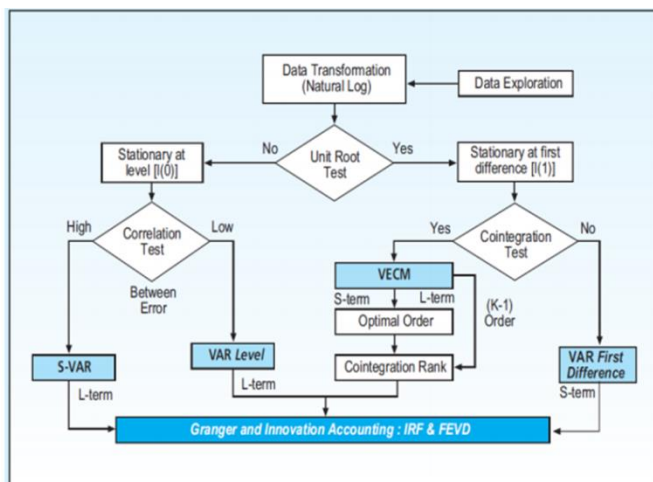
A. PEMBENTUKAN MODEL VAR

Sims (1980) mengembangkan Vector Autoregression (VAR) dalam makroekonometrika. Menurutnya, VAR adalah dinamika ad hoc model multivariat, memperlakukan set variabel simultan secara setara, di mana setiap variabel endogen diregresi pada lagnya sendiri dan lag semua variabel lain dalam sistem orde hingga. Tujuan dari pendekatan ini adalah untuk menguji respon dinamis dari sistem ke guncangan tanpa harus bergantung pada "pembatasan identifikasi yang luar biasa" yang melekat pada model struktural.

Sims (1992) dan peneliti lain mengikuti asumsi rekursif yang dibuat oleh Christiano et al. (1998a,b), yang mengatakan bahwa non-kebijakan variabel tidak bereaksi secara bersamaan terhadap variabel kebijakan, dan menempatkan variabel kebijakan terlebih dahulu sesuai dengan itu. Dengan demikian, diasumsikan bahwa keputusan kebijakan dibuat tanpa mempertimbangkan evolusi variabel ekonomi secara simultan. Jika kita ingin mengukur efek kontemporer dari variabel kebijakan pada variabel ekonomi, variabel kebijakan harus diurutkan terakhir. Jika korelasi di seluruh residual sangat kecil, posisi variabel dalam VAR tidak relevan. Jika semua variabel dalam VAR kita terintegrasi dengan orde 1 dan jika hubungan kointegrasi di antara mereka ada, kita dapat gunakan Vector Error Correction Model (VECM) untuk memperkirakan respon impuls dan fungsi dekomposisi varians.

Prosedur dalam pembentukan Model VAR yang merupakan model persamaan regresi yang menggunakan data time series yang berkaitan dengan stasioneritas dan kointegrasi antara variabel-variabel di dalamnya. Pembentukan model VAR dimulai dengan melakukan pengujian stasioneritas data, model VAR biasa (unrestricted VAR) diperoleh jika data telah stasioner pada tingkat level. Namun jika data tidak stasioner pada tingkat level tetapi stasioner pada proses diferensiasi (1st different) yang sama, maka harus dilakukan uji kointegrasi untuk mengetahui apakah data tersebut mempunyai hubungan dalam jangka panjang atau tidak.

Jika hal data stasioner pada proses diferensiasi namun tidak terkointegrasi, maka dapat dibentuk model VAR dengan data diferensiasi (VAR in difference). Namun apabila terdapat kointegrasi maka dibentuk Vector Error Correction Model (VECM), Dalam mencapai keseimbangan/ekuilibrium jangka panjang (kointegrasi), tentu akan terjadi beberapa deviasi/penyimpangan/diskrepansi. deviasi ekuilibrium model jangka panjang inilah yang harus dilakukan koreksi (Error Correction) secara tahap demi tahap melalui model series parsial penyesuaian jangka pendeknya. Vector Error Correction Model (VECM) merupakan model VAR yang terkektriksi (restricted VAR) karena kointegrasi menunjukkan hubungan jangka panjang antar variable dalam model VAR.



Gambar 11 Model VAR dan VECM (Basuki: 2016)

1. Model Vector Autoregressive (VAR)

Sebagaimana model-model ekonometrika deret waktu adalah model yang dibangun berdasarkan teori ekonomi yang ada. Dengan kata lain, teori ekonomi menjadi dasar dalam mengembangkan hubungan antarpeubah pada model. Model seperti ini disebut sebagai model struktural atau teoritis, dan estimasinya dapat memberikan informasi numeric sekaligus alat untuk menguji teori yang ada.

Namun sering kali teori ekonomi belum mampu menentukan spesifikasi yang tepat untuk model. Hal ini mungkin disebabkan teori ekonomi yang ada terlalu kompleks (rumit/majemuk) sehingga perlu dilakukan penyederhanaan dalam model atau sebaliknya fenomena yang ada terlalu kompleks sehingga

tidak cukup hanya dijelaskan dengan teori yang ada. Model VAR disebut sebagai model non-struktural atau model tidak teoritis (ateoritis).

Secara umum terdapat tiga bentuk model VAR

- 1) Unrestricted VAR, bentuk ini terkait dengan persoalan kointegrasi dan hubungan teoritis. Jika data yang digunakan dalam model VAR adalah data yang stasioner pada level, bentuk yang digunakan adalah Unrestricted VAR. Model VAR ini memiliki dua bentuk yaitu VAR in level dan VAR in difference. VAR in level digunakan jika data tidak stasioner pada level, sehingga datanya harus distasionerkan dulu sebelum menggunakan model VAR. VAR in difference digunakan jika data tidak stasioner dalam level dan tidak memiliki hubungan kointegrasi, maka estimasi VAR dapat dilakukan dalam bentuk data difference.
- 2) Restricted VAR atau disebut juga Vector Error Correction Model (VECM). Restriksi diberikan Karena data tidak stasioner namun terkointegrasi.
- 3) Struktural VAR merupakan bentuk VAR yang direstriksi berdasarkan hubungan teoritis yang kuat dan skema ordering (urutan) peta hubungan terhadap peubah-peubah yang digunakan dalam model VAR. Oleh karena itu S-VAR juga dikenal sebagai bentuk VAR yang teoritis (theoretical VAR)

Terdapat beberapa analisis penting dalam model VAR. empat diantaranya yang umum dilakukan adalah 1. Peramalan, 2. Impulse response, 3. Forecast error decomposition, dan 4. Uji kausalitas. Model VAR menganggap bahwa semua variabel ekonomi adalah saling tergantung dengan yang lain. Oleh karena itu, persamaan model VAR untuk penelitian ini dapat ditulis sebagai berikut:

$$Y_t = \alpha_1 + \delta_1 t + \phi_{11} Y_{t-1} + \dots + \phi_{1p} Y_{t-p} + \beta_{11} X_{t-q} + \epsilon_{1t}$$

Dan

$$X_t = \alpha_2 + \delta_2 t + \phi_{21} X_{t-1} + \dots + \phi_{2p} X_{t-p} + \beta_{21} Y_{t-1} + \dots + \beta_{2q} Y_{t-q} + \epsilon_{2t}$$

Kedua persamaan ini membentuk model VAR. VAR untuk k-variabel akan terdiri atas k-persamaan (yakni setiap persamaan dengan salah satu variable sebagai variable dependen, dan variable independen adalah laq dari seluruh variable yang lain, dan mungkin ditambah komponen *trend deterministic*).

2. Pengujian Stasioneritas (*Unit Root Test*)

Stasioneritas adalah hukum probabilitas yang mengharuskan dalam proses tersebut tidak berubah sepanjang waktu, dalam artian proses dalam keadaan setimbang secara statistik. Data dinyatakan stasioner jika nilai rata-rata dan

varian dari data tersebut tidak mengalami perubahan secara sistematis sepanjang waktu, atau rata-rata dan variannya konstan.

Dalam melakukan pengujian stasioneritas dapat dilakukan dengan melihat Grafik, Korelogram, namun demikian kedua metode itu mempunyai kelemahan, kedua metode tersebut memiliki kelemahan dalam objektivitas. Untuk itu diperlukan sebuah pengujian yang bersifat formal, untuk itu digunakan uji unit root, penggunaan uji unit root dikarenakan uji ini secara umum paling sering digunakan dalam pengujian stasioneritas. Uji ini disebut *Augmented Dickey-Fuller (ADF)*.

Pengujian ini melihat apakah data yang digunakan mengandung unit root atau tidak. Keberadaan unit root mengindikasikan bahwa data tersebut tidak stasioner. Untuk menguji keberadaan unit root dapat dilakukan pengujian *Augmented Dickey-Fuller Test/ADF Test* dan *Philips-Perron test*

$$Y_t = \rho Y_{t-1} + \mu_t \quad -1 \leq \rho \leq 1$$

Untuk mengestimasi persamaan di atas dengan OLS dan menguji hipotesis bahwa $\rho = 1$ dengan uji t biasa karena uji tersebut sangat bias pada kasus unit root, oleh karena itu kita manipulasi persamaan di atas sebagai berikut: Kurangi Y_{t-1} dari kedua sisi untuk memperoleh:

$$\begin{aligned} Y_t - Y_{t-1} &= \rho Y_{t-1} - Y_{t-1} + \mu_t \\ &= (\rho - 1) Y_{t-1} + \mu_t \end{aligned}$$

Yang dapat juga ditulis

$$\Delta Y_t = \delta Y_{t-1} + \mu_t$$

Untuk menguji hipotesis nol bahwa $\delta = 0$, hipotesis alternatif bahwa $\delta < 0$. Jika $\delta = 0$, kemudian $\delta < 0$, yaitu bahwa kita memiliki sebuah unit root, yang artinya data time series yang digunakan tidak stasioner.

Dalam menerapkan uji ADF yang pertama harus dilakukan adalah menentukan lag dari komponen differensi yang akan dimasukkan ke dalam model ($k=0$ adalah uji Dickey-Fuller/DF), sedangkan untuk uji ADF digunakan $k>0$).

3. Uji Stabilitas Model

Uji stabilitas dalam analisis VAR dilakukan untuk memastikan validitas impulse response dan variance decompositionnya. Untuk menguji stabil atau tidaknya estimasi VAR yang telah dibentuk maka dilakukan pengecekan kondisi

VAR stability berupa roots of characteristic polynomial. Suatu sistem VAR dikatakan stabil apabila seluruh roots-nya memiliki modulus lebih kecil dari satu.

4. Penentuan Panjang Selang (Lag) Optimal

Penentuan panjang selang (lag) optimal merupakan tahapan yang sangat penting dalam model VAR mengingat tujuan membangun model VAR adalah untuk melihat perilaku dan hubungan dari setiap variabel dalam model. Panjang lag selang (lag) Optimal digunakan untuk mengetahui lamanya periode ketergantungan suatu variabel terhadap variabel masa lalunya maupun variabel endogen lainnya

Untuk menentukan lag optimal, maka dalam penelitian ini digunakan beberapa kriteria informasi yang terdiri dari Akaike Information Criterion (AIC), Schwartz information Criterion (SIC) dan Hannan-Quinn (HQ). Penentuan lag optimal dengan menggunakan kriteria informasi tersebut diperoleh dengan memilih kriteria yang mempunyai nilai paling kecil di antara berbagai lag yang diajukan.

5. Uji Kausalitas

Setelah menentukan panjang lag optimal adalah melakukan uji kausalitas Granger, pengujian ini dimaksudkan untuk mengetahui apakah terdapat hubungan yang saling mempengaruhi antar variabel endogen sehingga spesifikasi model VAR menjadi tepat hal ini penting karena model VAR bersifat non struktural. pengujian kausalitas Granger juga digunakan untuk melihat pengaruh masa lalu terhadap kondisi sekarang sehingga uji ini memang dianggap tepat dipergunakan untuk data time series. Uji Kausalitas digunakan untuk pembentukan model VAR (VAR in difference).

Terdapat suatu aplikasi terkait dengan VAR, yakni Granger Causality Test, sebagai suatu uji sebab akibat, kita perlu membedakannya dengan sebab-akibat secara harfiah. Sebab-akibat secara Granger tidak memiliki arti fundamental, dalam artian kita dapat menelusuri alur logika mengapa suatu kejadian (X) akan menyebabkan kejadian lain (Y). Granger Causality adalah murni suatu konsep statistik. Dalam konsep ini X dikatakan menyebabkan Y jika realisasi X terjadi lebih dahulu dari pada Y dan realisasi Y tidak terjadi mendahului realisasi X.

Yang memiliki hubungan kausalitas adalah yang memiliki nilai probabilitas yang lebih kecil daripada $\alpha 0.05$ sehingga nanti H_0 akan ditolak yang berarti suatu variabel akan mempengaruhi variabel lain.

6. Uji Kointegrasi

Kointegrasi merupakan kombinasi hubungan linear dari variable-variable yang non stasioner, dan semua variable tersebut harus terintegrasi pada orde atau derajat yang sama. Variable-variable yang terintegrasi akan menunjukkan trend stochastic yang sama dan selanjutnya mempunyai arah pergerakan yang sama dalam jangka panjang. Uji kointegrasi merupakan kelanjutan dari unit root dan uji derajat integrasi. Untuk melakukan uji kointegrasi, pertama-tama peneliti perlu mengamati perilaku data ekonomi runtun waktu yang digunakan. Jika satu atau lebih variabel tidak stasioner dan tidak integrasi, maka variable tersebut tidak dapat berkointegrasi.

Dikarenakan data yang digunakan adalah data panel maka uji kointegrasi yang digunakan adalah Johansen Fisher Panel. Ide dasar uji unit root dalam data panel adalah pengembangan dari uji unit root dalam times series, yang dapat dijelaskan dalam model:

$$Y_{it} = \rho_t Y_{it} + X_{it} \delta_{it} + \varepsilon_{it}$$

$i = 1, 2, \dots, N$ (jumlah individu)

$t = 1, 2, \dots, T$ (jumlah periode individu)

Jika diasumsikan $\alpha = \rho - 1$ dengan lag p dan bervariasi antar cross section, maka uji hipotesisnya :

$H_0 : \alpha = 0$ (mempunyai akar unit)

$H_1 : \alpha < 0$ (tidak mempunyai akar unit)

Jika nilai $p_t = 1$ maka dikatakan bahwa variabel random Y mempunyai akar unit (unit root). Jika data panel mempunyai akar unit maka dikatakan data tersebut bergerak secara random (random walk) dan data yang mempunyai sifat random walk dikatakan data tidak stasioner. Oleh karena itu jika kita melakukan regresi Y_{it} pada lag $Y_{it} - 1$ dan mendapatkan nilai $p_t = 1$ maka data dikatakan tidak stasioner. Inilah ide dasar uji akar unit untuk mengetahui apakah data stasioner atau tidak. Formula uji unit root dengan dasar ADF adalah:

$$\Delta Y_{it} = \alpha Y_{it-1} + \sum_{j=1}^{p_t} \beta_{it} \Delta Y_{it-j} + X_{it} \delta + \varepsilon_{it}$$

Jika diasumsikan $\alpha = \rho - 1$ dengan lag p dan bervariasi antar cross section, maka uji hipotesisnya :

$H_0 : \alpha = 0$ (mempunyai akar unit)

$H_1 : \alpha < 0$ (tidak mempunyai akar unit)

Prosedur untuk menentukan apakah data stasioner atau tidak dengan cara membandingkan antara nilai statistik dengan nilai kritisnya. Jika nilai absolut statistik lebih besar dari nilai kritisnya, maka data yang diamati menunjukkan stasioner dan jika sebaliknya, nilai absolut statistik lebih kecil dari nilai kritisnya maka data tidak stasioner.

7. Impulse Response Function (IRF)

Analisis Impulse Response dapat melacak respon dari variabel endogen dalam model VAR yang dibangun karena adanya shocks atau guncangan atau perubahan di dalam variabel pengganggu, IRF menggambarkan bagaimana laju dari shock suatu variabel terhadap variabel-variabel yang lain sehingga melalui IRF ini, bisa diketahui lamanya pengaruh dari terjadinya suatu shock/goncangan suatu variabel terhadap variabel-variabel yang lain.

IRF digunakan untuk mengkonfirmasi apakah transmisi atau urutan proses variabel yang ditetapkan dalam teori dan penelitian empiris sebelumnya dapat dibuktikan dalam VECM yang diestimasi. Peneliti dapat mengatur transmisi variabel dengan menyusun (ordering) variabel dalam VEC. Urutan variabel yang disusun pada model VECM akan mempengaruhi beberapa hal, yaitu 1. Membentuk urutan rekurisif dari variabel pertama ke variabel terakhir. Dan 2. Membentuk variabel dependen dalam persamaan kointegrasi.

8. Variance Decomposition (VD)

Variance Decomposition (VD) ini akan memberikan keterangan tentang besarnya dan sampai berapa lama proporsi shock sebuah variabel terhadap variabel itu sendiri dan selanjutnya melihat besaran proporsi shock variabel lain terhadap variabel tersebut. Variance Decomposition menggambarkan kontribusi setiap variabel dalam model VAR karena adanya shock atau guncangan.

Variance decomposition adalah analisis persamaan VAR untuk melihat komponen-komponen pembentuk forecasting variance yang terjadi dekomposisi variance mengikuti struktur contemporaneous. Nilai dari Variance Decomposition atau Forecast Error Variance Decomposition (FEDV) selalu 100 persen, nilai FEDV lebih tinggi menjelaskan kontribusi varians satu variabel transmit terhadap variabel transmit lainnya lebih tinggi.

Langkah-langkah Analisis Data:

- 1) Cara Input Data
- 2) Uji Stasioner (unit root test)
- 3) Penentuan Panjang Selang (Lag) Optimal
- 4) Pengujian Stabilitas VAR

- 5) Uji Kausalitas Granger
- 6) Uji Kointegrasi
- 7) Analisis impulse response function
- 8) Analisis Variance decomposition

CONTOH KASUS :

Misalkan seorang peneliti ingin melihat dampak dari total pembiayaan bank syariah, pengeluaran pemerintah, investasi pemerintah dan jumlah uang beredar terhadap Produk Domestik Bruto (PDB) di Indonesia. Data-data yang berhasil dikumpulkan selama proses pengumpulan data adalah sebagai berikut:

DAMPAK PEMBIAYAAN BANK SYARIAH, PENGELUARAN PEMERINTAH, INVESTASI PEMERINTAH DAN JUMLAH UANG BEREDAR TERHADAP PERTUMBUHAN EKONOMI DI INDONESIA

Lampiran (Dalam Miliar)

TAHUN	X1 (PEMBIAYAAN)	X2 (PENGELUARAN PEMERINTAH)	X3 (INVESTASI)	X4 (PDB)	X5 (Jumlah Uang Beredar)
1994	Rp5,030.00	Rp300,505.20	Rp160,283.75	Rp840,854.90	Rp.
1995	Rp6,490.00	Rp400,176.56	Rp132,639.70	Rp783,077.90	Rp.2646480
1996	Rp7,232.00	Rp500,632.44	Rp87,653.10	Rp858,234.30	Rp.2102082
1997	Rp8,060.47	Rp607,128.86	Rp53,912.50	Rp967,040.30	Rp.2016023
1998	Rp9,834.71	Rp707,649.86	Rp97,405.65	Rp1,534,406.50	Rp.2123234
1999	Rp10,451.61	Rp985,730.69	Rp162,841.83	Rp1,427,633.50	Rp.2235144
2000	Rp11,530.00	Rp937,381.96	Rp101,662.00	Rp2,141,414.40	Rp.2227588
2001	Rp11,890.00	Rp1,042,117.23	Rp145,787.27	Rp3,446,851.90	Rp.2073859
2002	Rp12,232.00	Rp1,294,999.23	Rp176,594.76	Rp4,419,187.10	Rp.2066480
2003	Rp12,530.00	Rp376,505.20	Rp160,283.75	Rp1,840,854.90	Rp.2112082
2004	Rp13,490.00	Rp427,176.56	Rp132,639.70	Rp2,083,077.90	Rp.2116023
2005	Rp15,232.00	Rp509,632.44	Rp87,653.10	Rp2,458,234.30	Rp.2143234
2006	Rp21,060.47	Rp667,128.86	Rp53,912.50	Rp2,967,040.30	Rp.2231144
2007	Rp28,834.71	Rp757,649.86	Rp97,405.65	Rp3,534,406.50	Rp.2217588
2008	Rp39,451.61	Rp985,730.69	Rp162,841.83	Rp4,427,633.50	Rp.2236459
2009	Rp48,472.92	Rp937,381.96	Rp101,662.00	Rp5,141,414.40	Rp.2274954
2010	Rp70,241.44	Rp1,042,117.23	Rp145,787.27	Rp6,446,851.90	Rp.2308845
2011	Rp105,329.93	Rp1,294,999.23	Rp176,594.76	Rp7,419,187.10	Rp.2347806
2012	Rp151,058.52	Rp1,491,410.22	Rp221,082.30	Rp8,230,925.90	Rp.2073859

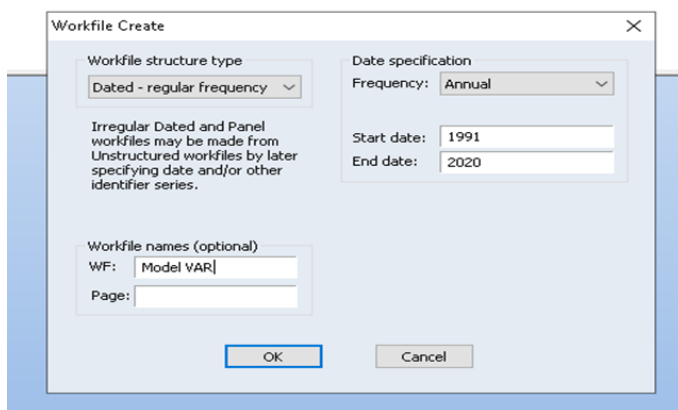
2013	Rp188,553.49	Rp1,650,563.73	Rp274,728.00	Rp9,087,276.50	Rp.2066480
2014	Rp204,334.91	Rp1,777,182.86	Rp330,944.52	Rp10,569,705.30	Rp.2112082
2015	Rp218,761.17	Rp1,806,515.20	Rp365,948.75	Rp11,526,332.80	Rp.2116023
2016	Rp254,669.56	Rp1,864,274.99	Rp391,015.35	Rp12,401,728.50	Rp.2143234
2017	Rp293,457.95	Rp2,007,351.80	Rp432,013.32	Rp13,589,825.70	Rp.2231144
2018	Rp329,277.47	Rp2,220,656.97	Rp392,725.99	Rp14,838,311.50	Rp.2217588
2019	Rp365,125.32	Rp2,461,112.08	Rp423,100.05	Rp15,833,943.40	Rp.2236459
2020	Rp375,125.32	Rp2,561,112.08	Rp433,100.05	Rp15,933,943.40	Rp.2336459

Cara Input Data



- 1) **Klik Create a new EViews workfile**
- 2) Pada Menu **Frequency** pilih **Annual** (Data Yang Kita Gunakan adalah **Data Tahunan**)
- 3) Pada menu **Start Date** tulis **tahun awal data penelitian**
- 4) Pada menu **End Date** tulis **tahun akhir data penelitian.**
- 5) **Klik OK**
- 6) **Klik Kanan** kemudian **Klik New Object**
- 7) Pada **Type of Object** pilih **Series**
- 8) Pada **Name For Object** Tulis Nama Variabel (tanpa Spasi dan Tanda Lainnya)
- 9) Pada Folder dengan Nama Variabel yang telah ditulisa sebelumnya **Klik Kanan** dan **Open**
- 10) **Klik Edit** Untuk Meninput Data
- 11) Untuk Membuka Data Yang Sudah Di Input

- **Blok Data** yang akan Dibuka
- Klik Open
- Kemudian **klik as Group**

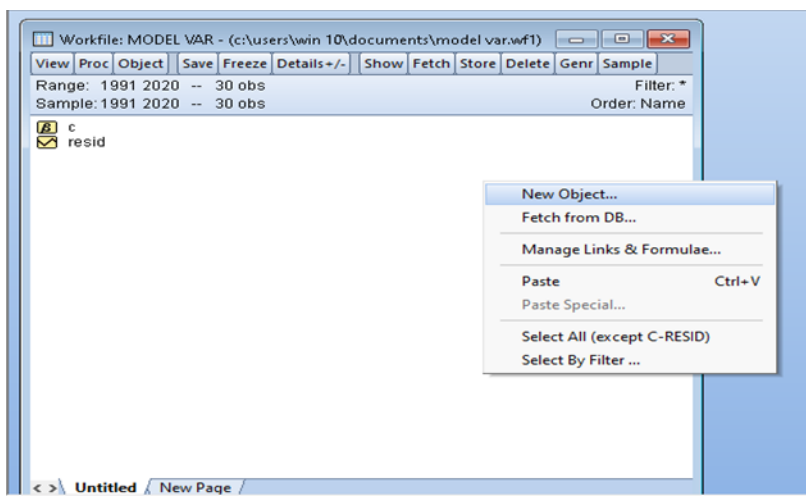


Pada **Menu Frequency** pilih **Annual** (Data Yang Kita Gunakan adalah **Data Tahunan**)

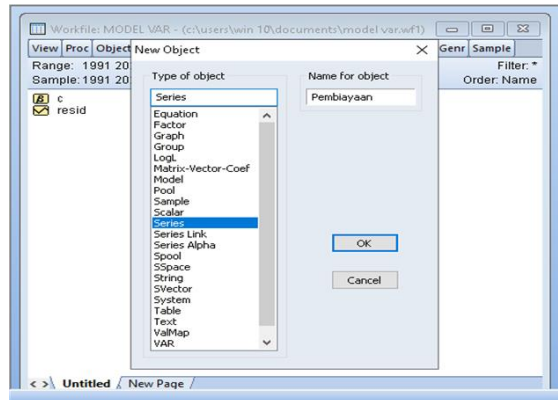
Pada menu **Start Date** tulis **tahun awal data penelitian**

Pada menu **End Data** tulis **tahun akhir data penelitian.**

Klik OK

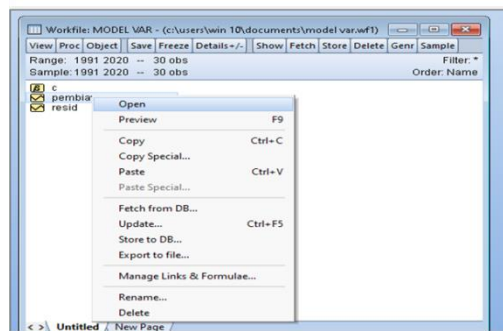
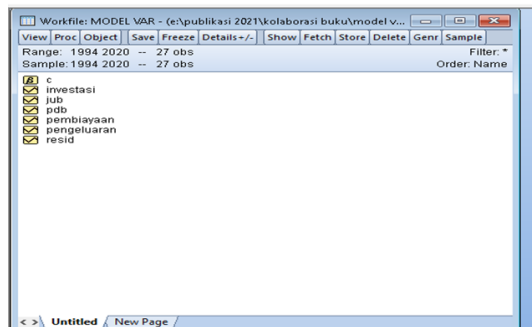


Klik Kanan kemudian **Klik New Object**



Pada **Type of Object** pilih Series

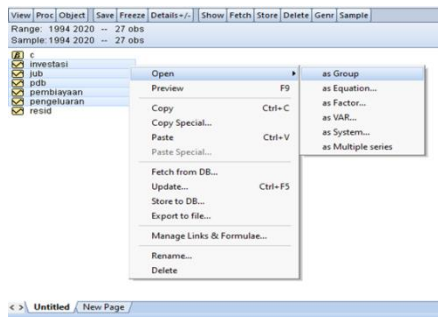
Pada **Name For Object** Tulis Nama Variabel (tanpa Spasi dan Tanda Lainnya)



Pada Folder dengan Nama Variabel yang telah ditulisa sebelumnya **Klik Kanan dan Open**

View	Proc	Object	Print	Name	Freeze	Default	Sort	Edit+/-	Smpl+/-	Adjust+/-	Label+/-	Wide+/-	Title	Sample	Genr
PEMBAYARAN															
Last updated: 08/31/21 - 23:45															
1991				NA											
1992				NA											
1993				NA											
1994				NA											
1995				NA											
1996				NA											
1997				NA											
1998				NA											
1999				NA											
2000				NA											
2001				NA											
2002				NA											
2003				NA											
2004				NA											
2005				NA											
2006				NA											
2007				NA											
2008				NA											
2009				NA											
2010				NA											
2011				NA											
2012				NA											
2013				NA											

Klik Edit Untuk Meninput Data



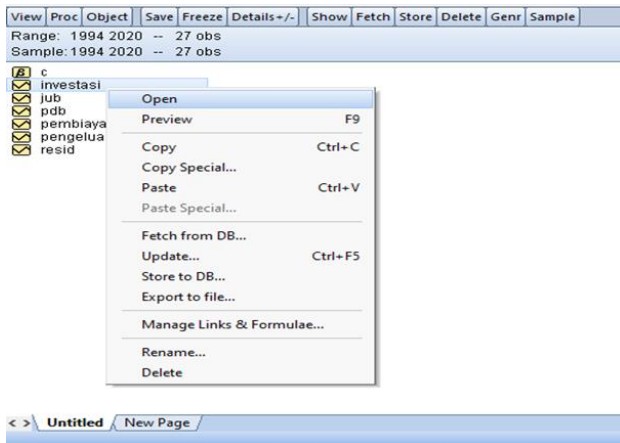
Untuk Membuka Data Yang Sudah Di Input

- Blok Data yang akan Dibuka
- Klik Open
- Kemudian klik as Group

View	Proc	Object	Print	Name	Freeze	Default	Sort	Edit+/-	Smpl+/-	Compare+/-	Transpose+/-	Title	Sample
		PDB		INVESTASI					JUB			PEMBAYARAN	PENGELUA...
1994		840854.9		160283.75					2053859			5030	300505.2
1995		783077.9		132639.7					2646480			6490	400176.56
1996		858234.3		87653.1					2102082			7232	500632.44
1997		967040.3		53912.5					2016023			8060.47	607128.86
1998		1534406.5		97405.65					2123234			9834.71	707649.86
1999		1427633.5		162841.83					2235144			10451.61	985730.69
2000		2141414.4		101662					2227588			11530	937381.96
2001		3446851.9		145787.27					2073859			11890	1042117.23
2002		4419187.1		176594.76					5066480			12232	1294999.23
2003		1840854.9		160283.75					2112082			12530	376505.2
2004		2083077.9		132639.7					2116023			13490	427176.56
2005		2458234.3		87653.1					2543234			15232	509632.44
2006		2967040.3		53912.5					2971144			21060.47	667128.86
2007		3534406.5		97405.65					3617588			28834.71	757649.86
2008		4427633.5		162841.83					4536459			39451.61	985730.69
2009		5141414.4		101662					2274954			48472.92	937381.96
2010		6446851.9		145787.27					2308845			70241.44	1042117.23
2011		7419187.1		176594.76					2347806			105329.93	1294999.23
2012		8230925.9		221082.3					2073859			151058.52	1491410.22
2013		9087276.5		274728					2086480			188553.49	1650563.73
2014		10569705.3		330944.52					2112082			204334.91	1777182.86
2015		11526332.8		365948.75					2116023			218761.17	1806515.2
2016		12401728.5		391015.35					2143234			254669.56	1864274.99
2017		13589825.7		432013.32					2231144			293457.95	2007351.8
2018		14838311.5		392725.99					2217588			329277.47	2220656.97
2019													

9. Uji Stasioner (unit root test)

Selanjutnya uji stationeritas variabel dengan cara **KLIK KANAN** salah satu Variabel Misal **INVESTASI** Lalu **Klik OPEN** pada jendela Pembiayaan pilih **VIEW** dan **klik GRAPH** maka akan muncul grafik seperti berikut:

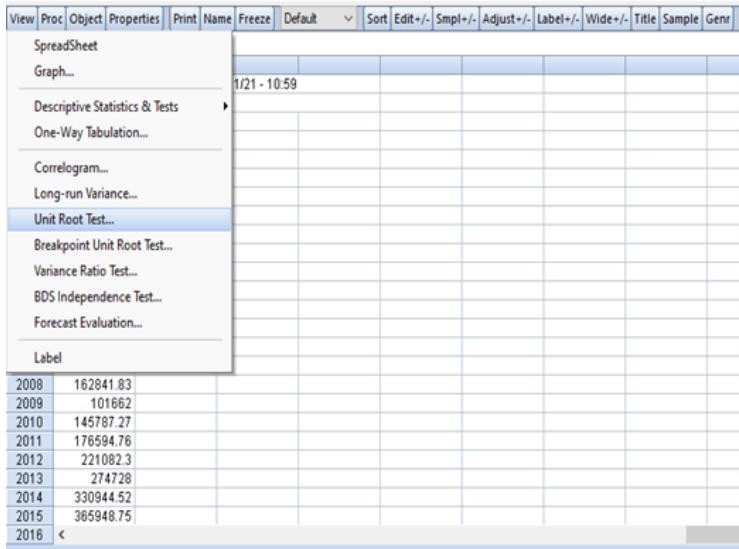


Klik Kanan dan Klik OPEN

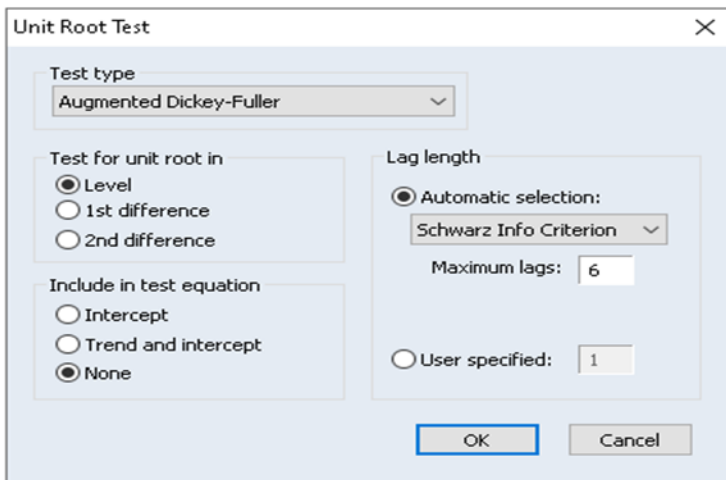
The screenshot shows the EViews software interface. At the top, there is a menu bar with options: View, Proc, Object, Properties, Print, Name, Freeze, Default, Sort, Edit+/-, Smpl+/-, Adjust+/-, Label+/-, Wide+/-, Title, Sample, Genr. Below the menu bar, a status bar indicates 'Last updated: 09/01/21 - 10:59'. The 'View' menu is open, and the 'Open' context menu is visible. The 'Open' context menu options are: Open, Preview (F9), Copy (Ctrl+C), Copy Special..., Paste (Ctrl+V), Paste Special..., Fetch from DB..., Update... (Ctrl+F5), Store to DB..., Export to file..., Manage Links & Formulae..., Rename..., and Delete. The 'Investasi' variable is selected, and the 'Open' context menu is open. The 'Open' context menu options are: Open, Preview (F9), Copy (Ctrl+C), Copy Special..., Paste (Ctrl+V), Paste Special..., Fetch from DB..., Update... (Ctrl+F5), Store to DB..., Export to file..., Manage Links & Formulae..., Rename..., and Delete.

Year	Investasi
1994	160283.75
1995	132638.7
1996	87653.1
1997	53912.5
1998	97405.65
1999	162841.83
2000	101662
2001	145787.27
2002	176594.76
2003	160283.75
2004	132638.7
2005	87653.1
2006	53912.5
2007	97405.65
2008	162841.83
2009	101662
2010	145787.27
2011	176594.76
2012	221082.3
2013	274728
2014	330944.52
2015	365948.75
2016	<
2017	160283.75

Klik VIEW dan Pilih GRAPH



Klik UNIT ROOT TEST



Pada Test Type Pilih Augmented Dickey-Fuller

Pada Test For Unit Root in Pilih LEVEL

Pada Include in test equation Pilih None

View	Proc	Object	Properties	Print	Name	Freeze	Sample	Genr	Sheet	Graph	Stats	Ident
------	------	--------	------------	-------	------	--------	--------	------	-------	-------	-------	-------

Null Hypothesis: INVESTASI has a unit root
 Exogenous: None
 Lag Length: 0 (Automatic - based on SIC, maxlag=6)

	t-Statistic	Prob.*
Augmented Dickey-Fuller test statistic	1.152592	0.9312
Test critical values: 1% level	-2.656915	
5% level	-1.954414	
10% level	-1.609329	

*MacKinnon (1996) one-sided p-values.

Berdasarkan output Augmented Dickey-Fuller Test (ADF) ternyata p-value > alpha = 0.05. maka terima H_0 yang artinya data mempunyai unit root (**data tidak stationer**) oleh karena data **tidak stationer Level** maka dilakukan pengulangan pengujian dengan memilih **1st difference**. Sama dengan langkah sebelumnya klik View lalu pilih unit root, dan pilih 1st difference, maka akan muncul tampilan sebagai berikut:

Unit Root Test

Test type
 Augmented Dickey-Fuller

Test for unit root in
☐ Level
☒ 1st difference
☐ 2nd difference

Include in test equation
☐ Intercept
☐ Trend and intercept
☒ None

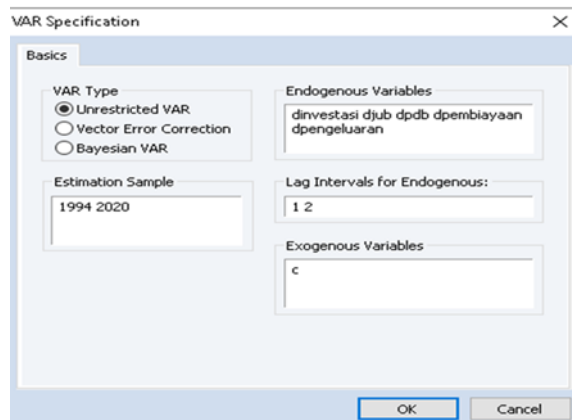
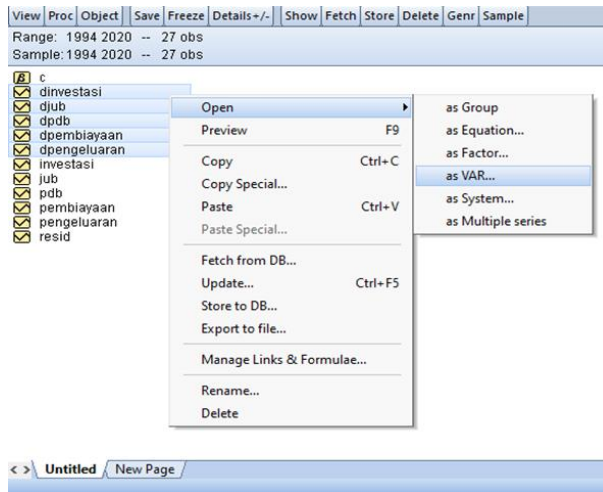
Lag length
☒ Automatic selection:
 Schwarz Info Criterion
 Maximum lags: 6
☐ User specified: 1

OK Cancel

Berdasarkan output Augmented Dickey-Fuller Test (ADF) ternyata p-value < alpha = 0.05. maka terima H_a yang artinya data tidak mempunyai unit root (data stationer) oleh karena data stationer 1st difference.

10. Penentuan Panjang Selang (Lag) Optimal

Langkah selanjutnya adalah menentukan lag optimum, sebelum menentukan lag optimum terlebih dahulu kita tentukan sampai lag seberapa model VAR stabil, hal ini dilakukan dengan cara **blok variabel** yang akan di uji **klik kanan**, pilih **OPEN**, kemudian **Klik as Group** kemudian klik **as VAR** maka akan muncul tampilan sebagai berikut:



pilih **Unrestricted VAR** karena model yang akan kita gunakan adalah VAR pada pilihan **VAR Type**. Kemudian ketik nama variabel yang akan digunakan pada kotak **Endogenous Variabel**. Isilah Lag Setinggi tingginya (bisa diulang berulang kali untuk memperoleh Lag Maksimal) pada kota **Lag Intervals.**, **Klik OK**

View	Proc	Object	Print	Name	Freeze	Estimate	Forecast	Stats	Impulse	Resids	Zoom
Vector Autoregression Estimates											
Date: 09/01/21 Time: 11:57											
Sample (adjusted): 1997 2020											
Included observations: 24 after adjustments											
Standard errors in () & t-statistics in []											
		DINVESTASI	DJUB	DPDB	DPBEMBIYAAAN	DPENGELU...					
DINVESTASI(-1)	-0.011590 (0.24353) [-0.04759]	2.403539 (5.26152) [0.45681]	4.124673 (3.82366) [1.07872]	-0.033718 (0.04536) [-0.74332]	1.010692 (0.92264) [1.09544]						
DINVESTASI(-2)	-0.182442 (0.22146) [-0.82382]	-13.52314 (4.78461) [-2.82638]	1.379041 (3.47706) [0.39661]	0.081222 (0.04125) [1.96904]	-2.133221 (0.83901) [-2.54255]						
DJUB(-1)	-0.005638 (0.01323) [-0.42605]	-0.697392 (0.28591) [-2.43917]	-0.660524 (0.20778) [-3.17897]	-0.004193 (0.00246) [-1.70090]	-0.182547 (0.05014) [-3.64100]						
DJUB(-2)	-0.015407 (0.01806) [-0.85331]	0.275429 (0.39009) [0.70606]	0.050843 (0.28349) [0.17935]	-0.003323 (0.00336) [-0.98815]	0.045709 (0.06841) [0.66822]						
DPDB(-1)	0.036112 (0.02238) [1.61388]	0.719319 (0.48343) [1.48796]	0.656457 (0.35132) [1.86856]	0.004974 (0.00417) [1.19338]	0.162205 (0.08477) [1.91342]						

Pada **Jendela VAR** diatas **Klik VIEW**, Kemudian **Klik Lag Struktire**, kemudian **Klik AR Root Table**.

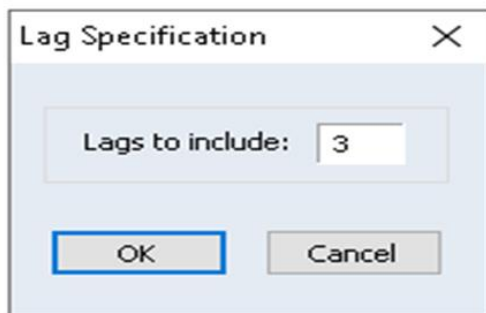
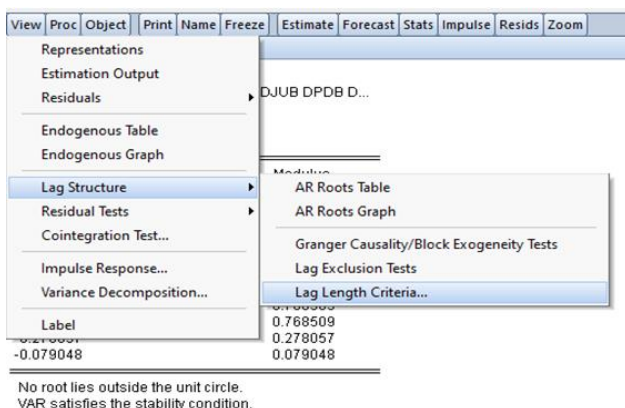
View	Proc	Object	Print	Name	Freeze	Estimate	Forecast	Stats	Impulse	Resids	Zoom																						
Representations																																	
Estimation Output																																	
Residuals																																	
Endogenous Table																																	
Endogenous Graph																																	
Lag Structure																																	
Residual Tests																																	
Cointegration Test...																																	
Impulse Response...																																	
Variance Decomposition...																																	
Label																																	
AR Roots Table																																	
AR Roots Graph																																	
Granger Causality/Block Exogeneity Tests																																	
Lag Exclusion Tests																																	
Lag Length Criteria...																																	
Roots of Characteristic Polynomial																																	
Endogenous variables: DINVESTASI DJUB DPDB D...																																	
Exogenous variables: C																																	
Lag specification: 1 2																																	
Date: 09/01/21 Time: 11:52																																	
<table><thead><tr><th>Root</th><th>Modulus</th></tr></thead><tbody><tr><td>-0.597865 - 0.598777i</td><td>0.846153</td></tr><tr><td>-0.597865 + 0.598777i</td><td>0.846153</td></tr><tr><td>0.489777 - 0.660776i</td><td>0.822500</td></tr><tr><td>0.489777 + 0.660776i</td><td>0.822500</td></tr><tr><td>0.775450 - 0.210053i</td><td>0.803396</td></tr><tr><td>0.775450 + 0.210053i</td><td>0.803396</td></tr><tr><td>-0.014111 - 0.768379i</td><td>0.768509</td></tr><tr><td>-0.014111 + 0.768379i</td><td>0.768509</td></tr><tr><td>-0.278057</td><td>0.278057</td></tr><tr><td>-0.079048</td><td>0.079048</td></tr></tbody></table>												Root	Modulus	-0.597865 - 0.598777i	0.846153	-0.597865 + 0.598777i	0.846153	0.489777 - 0.660776i	0.822500	0.489777 + 0.660776i	0.822500	0.775450 - 0.210053i	0.803396	0.775450 + 0.210053i	0.803396	-0.014111 - 0.768379i	0.768509	-0.014111 + 0.768379i	0.768509	-0.278057	0.278057	-0.079048	0.079048
Root	Modulus																																
-0.597865 - 0.598777i	0.846153																																
-0.597865 + 0.598777i	0.846153																																
0.489777 - 0.660776i	0.822500																																
0.489777 + 0.660776i	0.822500																																
0.775450 - 0.210053i	0.803396																																
0.775450 + 0.210053i	0.803396																																
-0.014111 - 0.768379i	0.768509																																
-0.014111 + 0.768379i	0.768509																																
-0.278057	0.278057																																
-0.079048	0.079048																																
No root lies outside the unit circle.																																	
VAR satisfies the stability condition.																																	

Maka akan Muncul tampilan seperti gambar dibawah ini:

View	Proc	Object	Print	Name	Freeze	Estimate	Forecast	Stats	Impulse	Resids	Zoom
Roots of Characteristic Polynomial											
Endogenous variables: DINVESTASI DJUB DPDB D...											
Exogenous variables: C											
Lag specification: 1 2											
Date: 09/01/21 Time: 11:52											
Root						Modulus					
-0.597865 - 0.598777i						0.846153					
-0.597865 + 0.598777i						0.846153					
0.489777 - 0.660776i						0.822500					
0.489777 + 0.660776i						0.822500					
0.775450 - 0.210053i						0.803396					
0.775450 + 0.210053i						0.803396					
-0.014111 - 0.768379i						0.768509					
-0.014111 + 0.768379i						0.768509					
-0.278057						0.278057					
-0.079048						0.079048					
No root lies outside the unit circle.											
VAR satisfies the stability condition.											

Stabilitas model dapat dilihat dari nilai modulus pada tabel AR Root. jika seluruh nilai AR Root nya dibawah satu, maka model tersebut stabil. Pada output diatas dapat dilihat bahwa sudah tidak ada lagi nilai modulus yang lebih dari 1 sehingga sampai lag 2 model masih stabil.

Selanjutnya pilih View kemudian klik Lag strukture dan kemudian klik Lag Length Criteria maka akan muncul tampilan sebagai berikut :

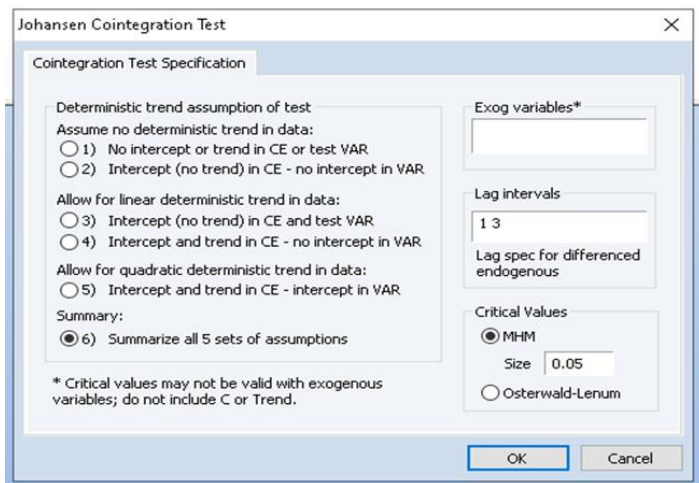
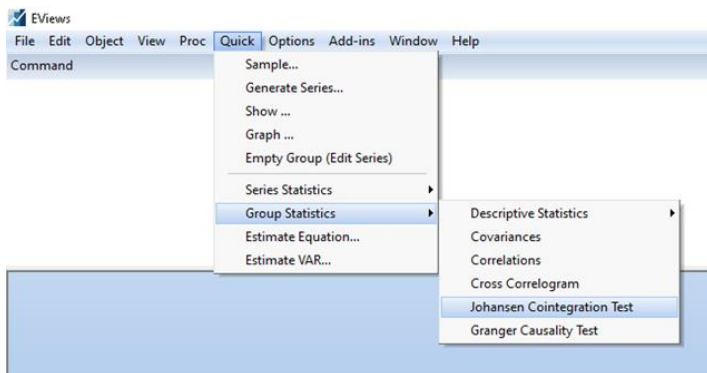


View	Proc	Object	Print	Name	Freeze	Estimate	Forecast	Stats	Impulse	Resids	Zoom
VAR Lag Order Selection Criteria											
Endogenous variables: DINVESTASI DJUB DPDB DPEMBIAAYAN DPENGELUARAN											
Exogenous variables: C											
Date: 09/03/21 Time: 14:36											
Sample: 1994 2020											
Included observations: 23											
Lag	LogL	LR	FPE	AIC	SC	HQ					
0	-1516.405	NA	1.96e+51	132.2961	132.5429	132.3582					
1	-1491.445	36.89828	2.10e+51	132.2995	133.7806	132.6720					
2	-1448.717	44.58524*	6.45e+50	130.7580	133.4733	131.4409					
3	-1387.264	37.40641	9.06e+49*	127.5881*	131.5377*	128.5814*					

* indicates lag order selected by the criterion
LR: sequential modified LR test statistic (each test at 5% level)
FPE: Final prediction error
AIC: Akaike information criterion
SC: Schwarz information criterion
HQ: Hannan-Quinn information criterion

Identifikasi model VAR dan VECM menggunakan nilai AIC, FPE, SC dan HQ. Pada output evies pilih lag yang paling banyak kode * . dan dari output diatas terlihat bahwa lg optimal 3. Semua variabel yang ada dalam model ini saling mempengaruhi satu sama lain, tidak hanya pada periode sekarang, namun variabel-variabel tersebut saling berkaitan sampai 3 periode sebelumnya.

B. UJI KOINTEGRASI



jika kita belum mengetahui apa asumsi data kita, maka kita pilih no 6 pada deterministic trend assumption. lalu isilah nilai lag optimal pada box lag intervals 1 3, dan klik ok.

Date: 09/03/21 Time: 14:44
 Sample: 1994 2020
 Included observations: 23
 Series: DINVESTASI DJUB DPDB DPEMBIAIYAN DPENGELUARAN
 Lags interval: 1 to 2

Selected (0.05 level*) Number of Cointegrating Relations by Model

Data Trend:	None	None	Linear	Linear	Quadratic
Test Type	No Intercept	Intercept	Intercept	Intercept	Intercept
	No Trend	No Trend	No Trend	Trend	Trend
Trace	4	3	3	3	3
Max-Eig	4	3	3	3	3

*Critical values based on MacKinnon-Haug-Michelis (1999)

Berdasarkan uji trace dan uji maximum eigenvalue maka terdapat satu persamaan kointegrasi. selanjutnya analisis dilakukan dengan VECM.

VAR Specification

Basics Cointegration VEC Restrictions

VAR Type

☐ Unrestricted VAR

☒ Vector Error Correction

☐ Bayesian VAR

Endogenous Variables

dinvestasi djub dpdb dpembiayaan
dpengeluaran

Estimation Sample

1994 2020

Lag Intervals for D(Endogenous):

1 2

Exogenous Variables

Do NOT include C or Trend in VEC's

OK Cancel

VAR Specification

Basics Cointegration VEC Restrictions

Rank

Number of cointegrating 1

Deterministic Trend Specification

No trend in data

☐ 1) No intercept or trend in CE or VAR

☐ 2) Intercept (no trend) in CE - no intercept in VAR

Linear trend in data

☐ 3) Intercept (no trend) in CE and VAR

☒ 4) Intercept and trend in CE - no trend in VAR

Quadratic trend in data

☐ 5) Intercept and trend in CE- linear trend in VAR

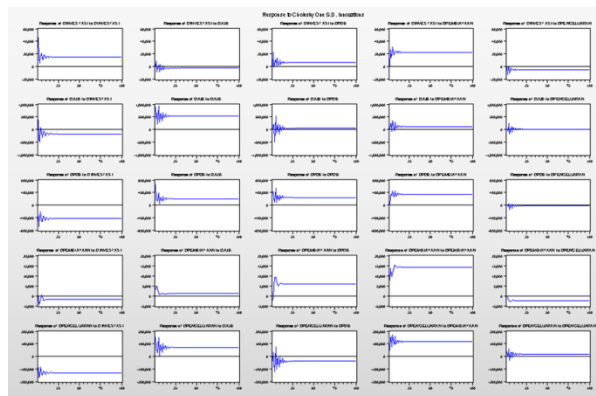
OK Cancel

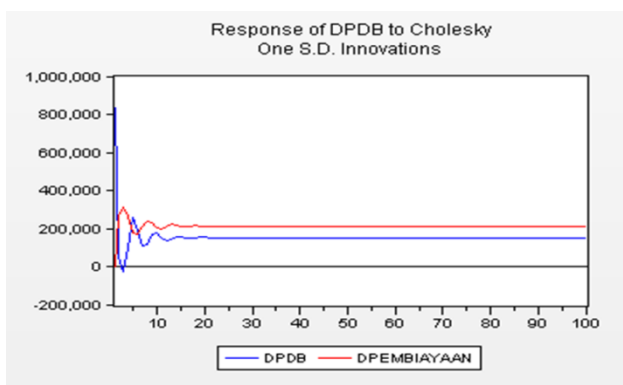
Vector Error Correction Estimates
Date: 09/03/21 Time: 14:50
Sample (adjusted): 1998 2020
Included observations: 23 after adjustments
Standard errors in () & t-statistics in []

Cointegrating Eq:	CointEq1
DINVESTASI(-1)	1.000000
DJUB(-1)	-0.002437 (0.00529) [-0.46027]
DPDB(-1)	0.143336 (0.00875) [16.3757]
DPEMBIAYAAN(-1)	-3.189372 (0.36033) [-8.62918]
DPENGELUARAN(-1)	-0.178243 (0.02985) [-5.97203]
@TREND(94)	-3638.633 (737.696) [-4.93243]
C	10257.51

Error Correction:	D(DINVEST...	D(DJUB)	D(DPDB)	D(DPEMBIA...	D(DPENGE...
CointEq1	-0.701210 (0.42166) [-1.66297]	-8.389631 (7.18410) [-1.16780]	-13.84067 (5.43564) [-2.54628]	-0.007816 (0.06914) [-0.11304]	-3.967274 (1.29339) [-3.06734]
D(DINVESTASI(-1))	-0.143331 (0.37658) [-0.38061]	3.072940 (6.41607) [0.47894]	8.611217 (4.85454) [1.77385]	-0.042375 (0.06175) [-0.68620]	2.519370 (1.15512) [2.18104]
D(DINVESTASI(-2))	0.221755 (0.45371) [0.48876]	-9.432620 (7.73023) [-1.22023]	9.494080 (5.84885) [1.62324]	0.086176 (0.07440) [1.15827]	0.705402 (1.39172) [0.50686]
D(DJUB(-1))	0.014450 (0.02104) [0.68692]	-0.913729 (0.35841) [-2.54937]	-0.389606 (0.27118) [-1.43669]	-0.002011 (0.00345) [-0.58307]	-0.068738 (0.06453) [-1.06526]
D(DJUB(-2))	-0.001105 (0.02190) [-0.05047]	0.077479 (0.37310) [0.20766]	-0.170600 (0.28229) [-0.60434]	-0.003220 (0.00359) [-0.89658]	0.049790 (0.06717) [0.74124]
D(DPDB(-1))	0.126331 (0.05708) [2.21331]	1.517601 (0.97247) [1.56056]	1.747026 (0.73579) [2.37434]	0.006308 (0.00936) [0.67395]	0.628427 (0.17508) [3.58937]

C. IMPULS RESPONSE





Terlihat baik dari grafik maupun tabel dibawah ini terlihat respon produk domestik bruto (PDB) terhadap guncangan variabel lainnya. Pada periode awal (periode 1) PDB merespon guncangan dari variabel Inveestasi secara negatif, hal ini dikarenakan investasi dalam jangka panjang jika tidak dikelola dengan baik malah akan memberikan beban terhadap PDB. dan PDB merespon guncangan dari jumlah uang beredar dan PDB itu sendiri. Namun demikian untuk varaibel pembiayaan dan pengeluaran belum merespon.

Namun dalam jangka panjang PDB merespon guncangan pembiayaan bank syariah, dan nilainya mengalami trend peningkatan sedangkan variabel lain direspon oleh PDB namun dengan kontribusi yang semakin menurun.hal ini menunjukkan bahwa dalam jangka panjang PDB didominasi oleh variabel pembiayaan bank syariah.

Response of DPDB:

Period	DINVESTASI	DJUB	DPDB	DPEMBIAYAAN	DPENGELUARAN
1	-369788.7	636439.0	410854.8	0.000000	0.000000
2	-673952.7	190765.8	428440.0	256897.1	-30876.96
3	-375665.8	168456.2	188539.7	332604.4	-38287.23
4	-210398.4	136223.8	82160.57	288467.3	38769.92
5	-521278.0	385936.5	531968.0	453813.7	-143306.4
6	-526091.5	36189.21	91175.19	255891.6	30040.77
7	-356901.3	240214.0	254749.8	316019.0	-552.1976
8	-369924.8	241297.4	305095.8	458131.8	-100402.3
9	-487035.2	120357.9	149658.3	280644.4	15669.64
10	-465409.7	273589.7	326596.3	336097.0	-35055.75
11	-379121.6	158073.0	205363.1	374664.7	-40526.10
12	-417177.0	192356.6	220344.9	323626.9	-15189.47
13	-460924.8	238425.3	297570.7	355175.3	-45457.93

14	-429413.2	161579.6	205887.5	333794.7	-18577.21
15	-409299.8	211698.9	246212.1	342754.3	-27433.83
16	-429660.6	207402.3	257332.3	358313.8	-40428.47
17	-442140.8	185401.7	229414.2	330484.3	-21097.50
18	-425800.7	209036.8	252441.2	343672.0	-30058.93
19	-419302.1	194569.4	239146.5	350887.4	-32739.90
20	-432646.5	198123.3	240852.6	339264.8	-27030.81
21	-433765.8	204018.0	250138.9	343898.7	-30679.08
22	-425004.3	193915.2	237615.7	343936.3	-28622.57
23	-426080.0	200881.5	243653.5	343754.4	-29457.23
24	-431164.8	200635.4	246091.4	344746.7	-30714.20
25	-429919.1	196717.9	240837.9	341844.1	-28164.40
26	-426716.0	200304.0	243944.9	344141.2	-29659.75
27	-427947.5	198931.0	243283.7	344777.8	-30140.23
28	-430042.5	198722.1	242874.0	342637.1	-28941.78
29	-428762.8	199631.4	243918.2	343674.5	-29540.19
30	-427601.2	198487.0	242514.4	344042.9	-29544.35
31	-428698.8	199309.1	243347.8	343558.2	-29452.80
32	-429214.5	199290.9	243665.6	343641.6	-29555.98
33	-428439.9	198685.6	242782.0	343527.7	-29317.21
34	-428231.4	199255.2	243303.0	343796.2	-29536.52
35	-428755.0	199110.6	243343.5	343755.1	-29562.95
36	-428828.6	198963.7	243128.5	343475.3	-29364.70
37	-428457.9	199136.3	243269.4	343709.3	-29490.52
38	-428469.9	199023.5	243164.4	343749.5	-29512.17
39	-428722.3	199090.6	243248.7	343604.3	-29453.26
40	-428663.0	199089.6	243267.8	343649.0	-29472.45
41	-428510.0	199015.9	243150.4	343677.5	-29465.88
42	-428576.3	199099.7	243242.3	343679.0	-29485.36
43	-428662.5	199076.3	243249.0	343657.7	-29478.76
44	-428607.3	199044.1	243195.2	343640.9	-29456.37
45	-428558.5	199079.6	243224.8	343680.0	-29480.58
46	-428602.6	199067.5	243223.6	343672.5	-29480.49
47	-428628.8	199066.6	243223.1	343647.3	-29466.34
48	-428594.1	199069.3	243223.0	343664.5	-29473.38
49	-428583.0	199062.8	243213.0	343670.0	-29475.50
50	-428608.1	199072.6	243226.5	343661.7	-29474.19
51	-428610.7	199068.1	243224.5	343660.2	-29472.71
52	-428594.1	199063.7	243215.5	343662.8	-29471.90
53	-428595.1	199070.3	243222.6	343666.3	-29475.26

54	-428605.7	199068.2	243223.0	343662.7	-29473.98
55	-428603.5	199066.6	243220.4	343660.5	-29471.95
56	-428596.7	199068.1	243220.9	343664.5	-29473.92
57	-428599.3	199067.6	243220.9	343664.2	-29474.11
58	-428603.4	199068.2	243222.1	343662.0	-29473.18
59	-428600.9	199067.6	243221.3	343662.8	-29473.28
60	-428598.7	199067.3	243220.4	343663.7	-29473.56
61	-428600.7	199068.3	243221.7	343663.4	-29473.75
62	-428601.8	199067.7	243221.5	343662.8	-29473.40
63	-428600.3	199067.5	243220.9	343663.0	-29473.29
64	-428599.8	199067.9	243221.2	343663.5	-29473.67
65	-428600.9	199067.8	243221.4	343663.2	-29473.57
66	-428601.1	199067.7	243221.3	343662.9	-29473.38
67	-428600.3	199067.7	243221.2	343663.2	-29473.50
68	-428600.3	199067.7	243221.2	343663.3	-29473.55
69	-428600.8	199067.8	243221.4	343663.1	-29473.50
70	-428600.7	199067.7	243221.3	343663.1	-29473.46
71	-428600.4	199067.7	243221.2	343663.2	-29473.49
72	-428600.5	199067.8	243221.3	343663.2	-29473.53
73	-428600.7	199067.8	243221.3	343663.1	-29473.49
74	-428600.6	199067.7	243221.2	343663.1	-29473.47
75	-428600.5	199067.8	243221.2	343663.2	-29473.51
76	-428600.6	199067.8	243221.3	343663.2	-29473.51
77	-428600.6	199067.8	243221.3	343663.1	-29473.49
78	-428600.6	199067.7	243221.2	343663.1	-29473.49
79	-428600.6	199067.7	243221.2	343663.2	-29473.50
80	-428600.6	199067.8	243221.3	343663.1	-29473.50
81	-428600.6	199067.7	243221.2	343663.1	-29473.49
82	-428600.6	199067.7	243221.2	343663.1	-29473.50
83	-428600.6	199067.8	243221.2	343663.2	-29473.50
84	-428600.6	199067.8	243221.2	343663.1	-29473.50
85	-428600.6	199067.7	243221.2	343663.1	-29473.50
86	-428600.6	199067.7	243221.2	343663.1	-29473.50
87	-428600.6	199067.7	243221.2	343663.1	-29473.50
88	-428600.6	199067.8	243221.2	343663.1	-29473.50
89	-428600.6	199067.7	243221.2	343663.1	-29473.50
90	-428600.6	199067.7	243221.2	343663.1	-29473.50
91	-428600.6	199067.7	243221.2	343663.1	-29473.50
92	-428600.6	199067.7	243221.2	343663.1	-29473.50
93	-428600.6	199067.7	243221.2	343663.1	-29473.50

94	-428600.6	199067.7	243221.2	343663.1	-29473.50
95	-428600.6	199067.7	243221.2	343663.1	-29473.50
96	-428600.6	199067.7	243221.2	343663.1	-29473.50
97	-428600.6	199067.7	243221.2	343663.1	-29473.50
98	-428600.6	199067.7	243221.2	343663.1	-29473.50
99	-428600.6	199067.7	243221.2	343663.1	-29473.50
100	-428600.6	199067.7	243221.2	343663.1	-29473.50

D. VARIANCE DECOMPOTITION

Tabel berikut ini merupakan output hasil dari variance decompotition, terlihat dari tabel dibawah bahwa pada periode awal (periode 1) varians Produk Domestik Bruto (PDB) dibentuk oleh kontribusi dari jumlah uang beredar sebesar 57 dan variabel lain hanya memberikan kontribusi yang lebih rendah. Namun dalam jangka panjang (periode 100) varians produk domestik bruto (PDB) didominasi oleh kontribusi investasi dan kontribusi pembiayaan bank syariah. Hal ini menunjukkan bahwa dalam jangka panjang apabila pemerintah ingin meningkatkan pertumbuhan ekonomi maka fokus yang harus dilakukan adalah meningkatkan investasi pemerintah dan meningkatkan pembiayaan bank syariah.

Variance Decomposition of DPDB:

Period	S.E.	DINVESTASI	DJUB	DPDB	DPEMBIAYAANDPENGLUARAN	
1	842970.9	19.24341	57.00178	23.75481	0.000000	0.000000
2	1204871.	40.70746	30.40862	24.27217	4.546082	0.065673
3	1329983.	41.38726	26.56089	21.92998	9.985090	0.136772
4	1386776.	40.36865	25.39487	20.52158	13.51094	0.203958
5	1689171.	36.73220	22.33653	23.74971	16.32434	0.857223
6	1790552.	41.32315	19.91960	21.39572	16.57049	0.791047
7	1885715.	40.83975	19.58254	21.11577	17.74872	0.713229
8	2015946.	39.10085	18.56686	20.76613	20.69407	0.872100
9	2101698.	41.34530	17.41062	19.61320	20.82293	0.807946
10	2220236.	41.44245	17.11961	19.73866	20.95038	0.748906
11	2298338.	41.39472	16.44889	19.21833	22.20811	0.729964
12	2376323.	41.80433	16.04222	18.83742	22.62911	0.686925
13	2476484.	41.95526	15.69771	18.78829	22.89257	0.666177
14	2549045.	42.43855	15.21854	18.38624	23.32256	0.634101
15	2624658.	42.46045	15.00489	18.22212	23.70352	0.609018
16	2704199.	42.52379	14.72339	18.07146	24.08529	0.596068
17	2775762.	42.89665	14.42014	17.83475	24.27695	0.571507
18	2848262.	42.97552	14.23400	17.72390	24.51267	0.553920
19	2916788.	43.04647	14.01802	17.57311	24.82160	0.540798

20	2984615.	43.21349	13.82877	17.43468	24.99836	0.524699
21	3052783.	43.32408	13.66471	17.33614	25.16344	0.511628
22	3116615.	43.42720	13.49783	17.21456	25.36108	0.499319
23	3180187.	43.50339	13.36259	17.12021	25.52568	0.488136
24	3243470.	43.58948	13.22888	17.03433	25.66908	0.478241
25	3304433.	43.68865	13.09967	16.94279	25.80087	0.468022
26	3364567.	43.74944	12.99002	16.86826	25.93306	0.459213
27	3423739.	43.81263	12.88249	16.79516	26.05850	0.451227
28	3481902.	43.88655	12.78143	16.72529	26.16354	0.443187
29	3539183.	43.94513	12.68922	16.66327	26.26645	0.435924
30	3595289.	43.99878	12.60105	16.60225	26.36875	0.429177
31	3650716.	44.05185	12.51938	16.54627	26.45975	0.422753
32	3705403.	44.10293	12.44184	16.49391	26.54459	0.416729
33	3759104.	44.15087	12.36826	16.44315	26.62673	0.410990
34	3812113.	44.19342	12.29988	16.39637	26.70469	0.405642
35	3864446.	44.23554	12.23448	16.35181	26.77758	0.400583
36	3916040.	44.27675	12.17236	16.30924	26.84594	0.395720
37	3966962.	44.31387	12.11385	16.26928	26.91185	0.391152
38	4017232.	44.34937	12.05802	16.23105	26.97474	0.386821
39	4066902.	44.38396	12.00493	16.19475	27.03369	0.382675
40	4115972.	44.41664	11.95436	16.16023	27.09003	0.378732
41	4164439.	44.44756	11.90610	16.12716	27.14420	0.374974
42	4212366.	44.47706	11.86012	16.09572	27.19570	0.371390
43	4259759.	44.50553	11.81609	16.06564	27.24477	0.367961
44	4306619.	44.53275	11.77396	16.03681	27.29180	0.364675
45	4352977.	44.55856	11.73368	16.00927	27.33695	0.361536
46	4398849.	44.58344	11.69503	15.98284	27.38017	0.358526
47	4444248.	44.60740	11.65795	15.95748	27.42154	0.355635
48	4489187.	44.63030	11.62236	15.93314	27.46134	0.352861
49	4533678.	44.65229	11.58815	15.90974	27.49962	0.350196
50	4577740.	44.67348	11.55526	15.88725	27.53637	0.347632
51	4621381.	44.69389	11.52360	15.86560	27.57174	0.345165
52	4664613.	44.71352	11.49311	15.84475	27.60583	0.342789
53	4707448.	44.73243	11.46373	15.82466	27.63869	0.340499
54	4749898.	44.75069	11.43538	15.80528	27.67036	0.338291
55	4791971.	44.76831	11.40803	15.78657	27.70093	0.336159
56	4833678.	44.78531	11.38163	15.76851	27.73045	0.334101
57	4875028.	44.80173	11.35611	15.75106	27.75898	0.332113
58	4916030.	44.81763	11.33144	15.73419	27.78655	0.330190
59	4956693.	44.83299	11.30758	15.71788	27.81322	0.328331
60	4997025.	44.84786	11.28448	15.70208	27.83904	0.326531
61	5037035.	44.86226	11.26212	15.68679	27.86404	0.324788
62	5076729.	44.87622	11.24045	15.67197	27.88826	0.323100
63	5116115.	44.88974	11.21945	15.65761	27.91174	0.321463

64	5155200.	44.90286	11.19908	15.64368	27.93451	0.319875
65	5193991.	44.91559	11.17932	15.63016	27.95660	0.318335
66	5232495.	44.92794	11.16013	15.61704	27.97804	0.316840
67	5270717.	44.93994	11.14150	15.60431	27.99887	0.315389
68	5308664.	44.95159	11.12341	15.59193	28.01910	0.313978
69	5346342.	44.96292	11.10582	15.57990	28.03876	0.312608
70	5383756.	44.97393	11.08872	15.56821	28.05788	0.311275
71	5420911.	44.98464	11.07208	15.55683	28.07647	0.309979
72	5457814.	44.99506	11.05589	15.54576	28.09456	0.308717
73	5494469.	45.00521	11.04014	15.53499	28.11218	0.307489
74	5530881.	45.01509	11.02480	15.52449	28.12933	0.306294
75	5567055.	45.02471	11.00985	15.51427	28.14603	0.305129
76	5602996.	45.03409	10.99529	15.50431	28.16231	0.303994
77	5638707.	45.04323	10.98109	15.49461	28.17818	0.302888
78	5674194.	45.05215	10.96725	15.48514	28.19365	0.301809
79	5709460.	45.06084	10.95375	15.47591	28.20874	0.300757
80	5744509.	45.06932	10.94058	15.46690	28.22346	0.299731
81	5779346.	45.07760	10.92773	15.45811	28.23783	0.298729
82	5813974.	45.08568	10.91518	15.44953	28.25186	0.297751
83	5848397.	45.09357	10.90292	15.44115	28.26556	0.296796
84	5882619.	45.10128	10.89095	15.43296	28.27894	0.295863
85	5916643.	45.10881	10.87925	15.42496	28.29202	0.294952
86	5950472.	45.11617	10.86782	15.41715	28.30479	0.294061
87	5984110.	45.12337	10.85665	15.40951	28.31728	0.293190
88	6017560.	45.13041	10.84572	15.40203	28.32950	0.292339
89	6050825.	45.13729	10.83504	15.39472	28.34145	0.291506
90	6083909.	45.14402	10.82458	15.38757	28.35313	0.290691
91	6116813.	45.15061	10.81435	15.38058	28.36457	0.289894
92	6149541.	45.15706	10.80434	15.37373	28.37576	0.289113
93	6182096.	45.16337	10.79453	15.36703	28.38672	0.288349
94	6214481.	45.16955	10.78493	15.36046	28.39745	0.287601
95	6246697.	45.17561	10.77553	15.35403	28.40796	0.286869
96	6278748.	45.18154	10.76632	15.34773	28.41826	0.286151
97	6310637.	45.18735	10.75729	15.34156	28.42835	0.285448
98	6342365.	45.19305	10.74845	15.33551	28.43823	0.284758
99	6373935.	45.19863	10.73978	15.32958	28.44792	0.284083
100	6405350.	45.20410	10.73128	15.32377	28.45743	0.283420

Dari tabel dibawah ini terlihat bahwa pada periode awal (periode 1) varians investasi terbentuk dari variance investasi itu sendiri sebesar 100, sedangkan variabel lain tidak memberikan kontribusi (0). Sedangkan jika kita lihat dalam jangka panjang terlihat bahwa variabel pembiayaan bank syariah

terus memberikan kontribusi dengan trend yang terus meningkat, sedangkan variabel lain dalam 100 periode kontribusinya terus menurun. Hal ini menunjukkan bahwa pembiayaan bank syariah dalam jangka panjang memberikan kontribusi yang terus meningkat terhadap investasi. Hal ini sesuai dengan berbagai penelitian yang pernah dilakukan, hal ini disebabkan oleh pembiayaan bank syariah dipergunakan untuk kegiatan yang bersifat real dan pada sektor yang produktif.

Variance Decomposition of DINVESTASI:						
Period	S.E.	DINVESTASI	DJUB	DPDB	DPEMBIAYAAN	DPENGELUARAN
1	45810.85	100.0000	0.000000	0.000000	0.000000	0.000000
2	61028.57	61.38892	1.777557	13.55481	18.81592	4.462788
3	63374.76	57.50326	3.517428	12.62201	22.15290	4.204408
4	68149.28	55.26727	3.086944	11.26845	26.68786	3.689483
5	79209.65	48.33130	2.330364	10.43352	34.37898	4.525840
6	82626.70	46.61500	3.210132	9.588405	36.31928	4.267185
7	87053.21	43.44491	2.974779	10.98914	38.32641	4.264764
8	92671.82	42.00768	3.057977	9.833735	40.90922	4.191390
9	96880.76	41.43603	2.973456	9.097048	42.46910	4.024367
10	101233.7	39.37293	2.725285	9.326206	44.45157	4.124008
11	104656.0	38.49623	2.838589	8.830979	45.79499	4.039210
12	108458.4	38.05070	2.683546	8.528226	46.76593	3.971601
13	112483.3	37.20482	2.533162	8.323806	47.93445	4.003756
14	115582.6	36.59437	2.516128	8.066672	48.86160	3.961236
15	118910.4	36.03542	2.410276	7.936297	49.67832	3.939693
16	122393.9	35.59927	2.347279	7.695822	50.43379	3.923842
17	125506.6	35.19144	2.289565	7.523388	51.09458	3.901026
18	128573.2	34.72727	2.218609	7.429328	51.72951	3.895281
19	131599.2	34.40143	2.181730	7.268516	52.27242	3.875902
20	134594.4	34.11180	2.129456	7.141491	52.75654	3.860706
21	137512.4	33.78466	2.080700	7.042423	53.23718	3.855038
22	140298.5	33.50837	2.047353	6.934914	53.66766	3.841700
23	143092.8	33.26656	2.007347	6.841580	54.05350	3.831009
24	145850.9	33.03556	1.972496	6.750440	54.41826	3.823237
25	148507.5	32.81545	1.941703	6.669607	54.75889	3.814355
26	151132.7	32.60924	1.910861	6.597677	55.07538	3.806849
27	153728.1	32.42671	1.884140	6.522875	55.36722	3.799061
28	156273.1	32.25236	1.857723	6.456438	55.64133	3.792150
29	158772.2	32.08255	1.832915	6.396467	55.90178	3.786292
30	161231.6	31.92858	1.810817	6.336970	56.14380	3.779836
31	163662.3	31.78488	1.788940	6.281758	56.37036	3.774067
32	166055.4	31.64645	1.768507	6.230145	56.58593	3.768973

33	168408.0	31.51578	1.749570	6.181323	56.78951	3.763814
34	170733.0	31.39324	1.731331	6.135166	56.98127	3.758993
35	173028.4	31.27754	1.714195	6.090898	57.16291	3.754463
36	175290.8	31.16685	1.697875	6.049422	57.33568	3.750176
37	177524.4	31.06146	1.682365	6.010147	57.49990	3.746121
38	179731.3	30.96211	1.667718	5.972319	57.65568	3.742179
39	181911.7	30.86730	1.653646	5.936496	57.80408	3.738483
40	184065.6	30.77642	1.640260	5.902459	57.94587	3.734988
41	186194.1	30.68993	1.627537	5.869853	58.08109	3.731594
42	188299.5	30.60746	1.615325	5.838692	58.21016	3.728368
43	190381.6	30.52844	1.603657	5.808890	58.33371	3.725299
44	192440.5	30.45269	1.592493	5.780405	58.45206	3.722350
45	194477.9	30.38018	1.581789	5.753090	58.56542	3.719522
46	196494.4	30.31070	1.571529	5.726844	58.67412	3.716807
47	198490.2	30.24393	1.561669	5.701691	58.77850	3.714208
48	200466.1	30.17974	1.552201	5.677533	58.87882	3.711711
49	202422.8	30.11809	1.543102	5.654277	58.97523	3.709303
50	204360.8	30.05876	1.534339	5.631906	59.06800	3.706990
51	206280.5	30.00160	1.525903	5.610372	59.15736	3.704765
52	208182.5	29.94652	1.517777	5.589620	59.24347	3.702618
53	210067.3	29.89343	1.509938	5.569605	59.32648	3.700547
54	211935.4	29.84220	1.502375	5.550290	59.40659	3.698551
55	213787.1	29.79272	1.495073	5.531649	59.48394	3.696623
56	215623.0	29.74492	1.488019	5.513638	59.55866	3.694760
57	217443.3	29.69873	1.481200	5.496223	59.63089	3.692959
58	219248.5	29.65405	1.474604	5.479383	59.70075	3.691218
59	221039.0	29.61080	1.468221	5.463087	59.76836	3.689532
60	222815.1	29.56893	1.462042	5.447307	59.83382	3.687900
61	224577.1	29.52838	1.456055	5.432020	59.89723	3.686319
62	226325.5	29.48907	1.450252	5.417204	59.95869	3.684787
63	228060.4	29.45094	1.444626	5.402838	60.01829	3.683302
64	229782.2	29.41396	1.439167	5.388900	60.07611	3.681860
65	231491.2	29.37807	1.433870	5.375372	60.13223	3.680461
66	233187.7	29.34322	1.428725	5.362237	60.18672	3.679103
67	234872.0	29.30936	1.423728	5.349477	60.23965	3.677783
68	236544.2	29.27646	1.418871	5.337076	60.29109	3.676501
69	238204.7	29.24447	1.414150	5.325020	60.34110	3.675254
70	239853.7	29.21336	1.409558	5.313295	60.38974	3.674041
71	241491.5	29.18309	1.405090	5.301886	60.43707	3.672862
72	243118.2	29.15363	1.400741	5.290782	60.48313	3.671713
73	244734.1	29.12494	1.396506	5.279970	60.52799	3.670595
74	246339.4	29.09700	1.392382	5.269439	60.57167	3.669506
75	247934.4	29.06978	1.388364	5.259178	60.61424	3.668445
76	249519.1	29.04324	1.384447	5.249177	60.65573	3.667411

77	251093.8	29.01737	1.380628	5.239426	60.69618	3.666402
78	252658.7	28.99213	1.376903	5.229916	60.73563	3.665419
79	254214.0	28.96752	1.373270	5.220638	60.77412	3.664459
80	255759.8	28.94349	1.369724	5.211584	60.81168	3.663523
81	257296.3	28.92004	1.366262	5.202745	60.84834	3.662609
82	258823.7	28.89714	1.362882	5.194114	60.88415	3.661716
83	260342.2	28.87477	1.359580	5.185684	60.91912	3.660844
84	261851.8	28.85292	1.356355	5.177448	60.95328	3.659993
85	263352.8	28.83156	1.353203	5.169399	60.98667	3.659160
86	264845.3	28.81069	1.350121	5.161531	61.01931	3.658347
87	266329.4	28.79028	1.347108	5.153838	61.05123	3.657551
88	267805.3	28.77031	1.344162	5.146314	61.08244	3.656773
89	269273.1	28.75078	1.341279	5.138954	61.11297	3.656012
90	270732.9	28.73167	1.338458	5.131752	61.14285	3.655267
91	272184.9	28.71297	1.335698	5.124703	61.17209	3.654538
92	273629.3	28.69466	1.332995	5.117802	61.20072	3.653824
93	275066.0	28.67674	1.330349	5.111046	61.22874	3.653126
94	276495.2	28.65918	1.327758	5.104429	61.25619	3.652441
95	277917.2	28.64198	1.325219	5.097946	61.28308	3.651771
96	279331.8	28.62513	1.322732	5.091595	61.30943	3.651114
97	280739.4	28.60861	1.320294	5.085371	61.33525	3.650470
98	282139.9	28.59243	1.317905	5.079270	61.36056	3.649839
99	283533.5	28.57656	1.315562	5.073289	61.38537	3.649221
100	284920.3	28.56099	1.313265	5.067424	61.40970	3.648614

DAFTAR PUSTAKA

- Bambang Juanda dan Junaidi. (2012). *Ekonometrika Deret Waktu Teori Dan Aplikasi*. Bogor: IPB Press.
- Christiano, Lawrence J., Martin Eichenbaum and Charles L. Evans (1998), "Modelling Money", NBER Working Paper 6371.
- Christiano, Lawrence J., Martin Eichenbaum and Charles L. Evans (1998), "Monetary Policy Shocks: What have we learned and to what end?", NBER Working Paper No. 6400.
- Damodar N Gujarati dan Dawn C Porter. (2013). *Dasar-dasar Ekonometrika, Buku 2*. Jakarta: Penerbit Salemba Empat. h 445-446
- Moch. Doddy Ariefianto, . (2012). *Ekonometrika esensi dan aplikasi dengan menggunakan Eviews*. Jakarta: Erlangga. h. 114
- Sims, C. A. (1980). Macroeconomics and reality. *Econometrica*, 48, 1–48.
- Sims, C. A. (1992). Interpreting the macroeconomic time series facts: The effects of monetary policy. *European Economic Review*, 36, 975–1011.

ANALISIS KOMPARASI

(Uji Independent Sample t Test)

Drs. Agung Pujiyanto, M.M

Dra. Awin Mulyati, M.M

Universitas 17 Agustus 1945 Surabaya

A. PENDAHULUAN

Penelitian (research) merupakan kegiatan ilmiah dalam rangka pemecahan suatu permasalahan. Hasil penelitian tidak pernah dimaksudkan sebagai suatu pemecahan masalah (solusi) langsung bagi permasalahan yang dihadapi, karena penelitian merupakan bagian saja dari upaya pemecahan masalah yang lebih besar. Fungsi penelitian adalah mencari penjelasan dan jawaban terhadap permasalahan serta memberikan alternatif bagi kemungkinan yang dapat digunakan untuk pemecahan masalah (saifudin Azwar).

Sering kali kesulitan awal yang dihadapi oleh para peneliti terutama untuk peneliti pemula adalah bagaimana menemukan dan memilih masalah. Dari mana masalah diperoleh serta bagaimana mengidentifikasi bahwa sesuatu itu sebenarnya adalah masalah. Yang jelas masalah sebenarnya merupakan bagian dari semua yang dialami oleh seseorang, dan masalah merupakan kebutuhan seseorang untuk dipecahkan. Seseorang melakukan kegiatan penelitian tentu dengan tujuan untuk mendapatkan jawaban atas masalah yang dihadapinya.

Secara garis besar permasalahan dalam penelitian dapat dikelompokkan dalam 3 jenis.

- a. Permasalahan untuk mengetahui status atau kedudukan fenomena yang diamati dan mendiskripsikan gejala atau fenomena.
Sehubungan dengan permasalahan jenis ini biasanya akan memunculkan jenis penelitian deskripsi. Yaitu penelitian yang tujuannya ingin mendapatkan gambaran tentang sebuah gejala atau fenomena.
- b. Permasalahan untuk mencari keterkaitan atau mencari hubungan antar fenomena .

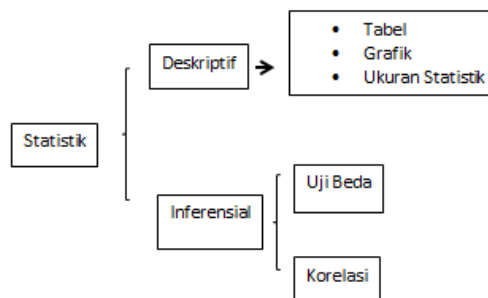
Permasalahan ini biasanya akan melahirkan jenis penelitian korelasi. Yaitu Penelitian yang bertujuan ingin melihat adanya keterkaitan antar fenomena, baik keterkaitannya searah ataupun dua arah (kausal)

c. Permasalahan untuk membandingkan atau komparasi antar fenomena .

Seringkali jenis permasalahan ini muncul ketika peneliti dihadapkan pada gejala atau fenomena yang nampak serupa ,dan penelitian ingin mencoba mengetahui adakah perbedaan dan kesamaan dari fenomena tersebut. Permasalahan ini biasanya melahirkan jenis penelitian komparasi.

Dalam Bab ini akan difokuskan pada bahasan point c (membandingkan fenomena atau jenis penelitian komparasi) yaitu bagaimana tahapan melakukan analisis uji komparasi atau uji beda. Analisis Uji beda merupakan bagian dari statistik induktif (infernsi).

Gambar : Lingkup Statistik



B. TAHAPAN PENELITIAN

Berikut akan disajikan contoh kasus dan tahapan analisisnya:

Kasus;

Mobilitas masyarakat yang semakin tinggi tentu menuntut pemenuhan keinginan dan kebutuhan yang semakin tinggi pula. Tidak terkecuali dalam pemenuhan akan kebutuhan makanan. Cara-cara instan dan cepat saji tentu menjadi pilihan bagi mereka. Tidak heran bila akhirnya bermunculan bergai tempat makan yang menyajikan makan siap saji. Mereka bersaing dengan menawarkan berbagai hidangan siap saji dengan harga bersaing dan semua menawarkan pelayanan yang cepat. Seorang peneliti mencoba ingin melihat adakah perbedaan diantara tempat makan tersebut. Peneliti mencoba mengambil Vareabel Harga yang ditawarkan, Cita rasa makanan yang disajikan dan tingkat layanan yang diberikan oleh masing-masing tempat makan siap saji. Untuk Kasus ini peneliti mengambil sampel 2 (dua) tempat makan siap saji

(tempat makan siap saji Bratang dan tempat makan siap saji Nginden) di Kota Surabaya.

a. Judul :

“Analisis Komparasi Cita Rasa makanan, Harga yang dan Tingkat Layanan”
(Studi kasus pada dua tempat makan siap saji Bratang dan Nginden di Kota Surabaya).

b. Rumusan Masalah:

Adakah perbedaan Cita Rasa Makanan yang disajikan, Harga yang Ditawarkan dan Tingkat Layanana yang diberikan antara tempat makan siap saji Bratang dan tempat makan siap saji Nginden di Kota Surabaya.

c. Tujuan Penelitian:

1. Untuk melihat adakah perbedaan Cita Rasa makanan yang disajikan, Harga yang ditawarkan dan Tingkat layanan yang diberikan oleh dua tempat makan siap saji tersebut.
2. Untuk mengetahui apakah perbedaan yang ditawarkan oleh kedua tempat makan siap saji tersebut sangat signifikan.

Berdasarkan pertanyaan penelitian yang telah terumuskan diatas, maka berikutnya dapat dibuat hipotesis penelitiannya. Hipotesis ini tentunya bersifat prediktif merupakan jawaban yang bersifat sementara dan akan diuji untuk pengambilan kesimpulan.

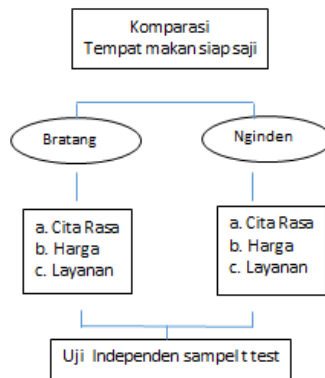
Mengacu pada rumusan masalah dan tujuan diatas maka hipotesis dapat dirumuskan kedalam bentuk Rumusan Hipotesis sbb:

Ho : Tidak terdapat perbedaan yang signifikan cita rasa makanan yang disajikan ,Harga yang ditawarkan dan tingkat layanan yang diberikan antara dua tempat makan siap saji tersebut.

Ha : Terdapat perbedaan yang signifikan cita rasa makanan yang disajikan, Harga yang ditawarkan dan Tingkat Layanan yang diberikan antara dua tempat makan siap saji tersebut.

Berikut adalah model kerangka Konseptual yang dapat digambarkan oleh peneliti:

Gambar: Kerangka Berfikir



Berdasarkan kerangka konseptual pada gambar diatas maka uji analisisnya akan dilakukan dengan menggunakan uji beda atau yang biasa disebut dengan uji independent sample t test, dimana uji ini digunakana bila ingin membandingkan dua objek penelitian berasal dari dua kelompok sampel yang berbeda. Berikut akan dijelaskan proses pengujian sampel independen pada kasus diatas.

a. Tabulasi Data

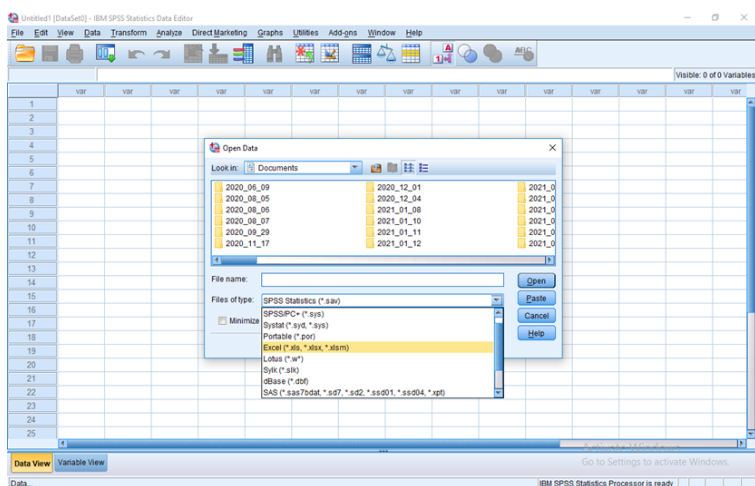
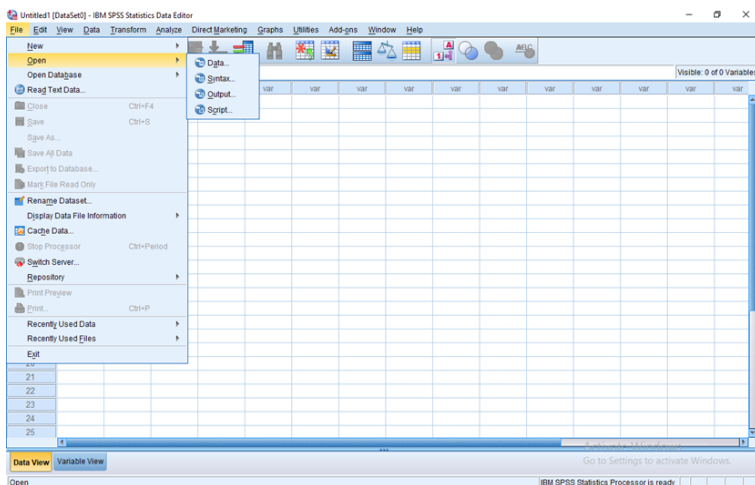
Dalam contoh kasus ini pengambilan data dengan instrumen berupa kuesionir dan skala yang digunakan adalah likert, mulai dari score 1 = sangat tidak setuju sampai dengan score 5=sangat setuju. Variabel Kualitas pelayanan menggunakan 5.indikator, Variabel Cita rasa menggunakan 3.indikator dan Variabel Harga ada 3 indikator. Peneliti menggunakan 50 responden yang pernah melakukan pembelian di dua tempat makan siap saji tersebut. Dibawah ini hasil tabulasi data dari 50 responden dalam format excel.

No		KUALITAS PELAYANAN					HARGA		CITA RASA		Jumlah		rata-rata	
1	2	3	4	5	6	7	8	9	10	11	12	13	14	15
1	Wanita	4	3	4	4	18	3,4	3	4	4	11	4,3	4	4
2	Wanita	4	3	3	4	17	3,4	4	3	4	11	3,7	4	4
3	Wanita	4	3	3	4	18	3,4	4	4	4	12	4	4	4
4	Wanita	4	3	3	4	17	3,4	3	3	3	9	3	3	3,7
5	Wanita	4	3	3	4	17	3,4	3	3	3	9	3	3	3,7
6	Wanita	4	4	4	4	20	4	4	4	4	12	4	4	4
7	Wanita	3	3	3	3	12	3	3	3	3	9	3	3	3
8	Wanita	4	3	3	4	17	3,4	4	4	4	12	4,3	4	4
9	Wanita	4	4	4	4	19	3,8	4	3	4	11	3,7	4	4
10	Wanita	4	4	4	4	19	3,8	4	4	4	12	4	3	3
11	Wanita	4	3	4	3	17	3,4	3	3	3	9	3	3	3,7
12	Wanita	4	3	3	3	17	3,4	3	3	3	9	3	3	3,7
13	Wanita	4	3	3	3	17	3,4	3	3	3	9	3	3	3,7
14	Wanita	4	4	4	4	21	4,2	3	3	3	9	3	3	3,7
15	Wanita	3	4	4	4	19	3,8	4	4	4	12	4	4	4
16	Wanita	4	3	3	3	17	3,4	3	3	3	9	3	3	3,7
17	Wanita	4	3	3	3	17	3,4	3	3	3	9	3	3	3,7
18	Wanita	4	3	3	3	17	3,4	3	3	3	9	3	3	3,7
19	Wanita	4	3	3	3	17	3,4	3	3	3	9	3	3	3,7
20	Wanita	4	3	3	3	17	3,4	3	3	3	9	3	3	3,7
21	Wanita	3	4	4	4	19	3,8	3	3	3	9	3	3	3,7
22	Wanita	4	4	4	4	20	4	4	3	4	11	3,7	4	4
23	Wanita	4	4	4	4	19	3,8	3	3	3	9	3	3	3,7
24	Wanita	4	3	3	3	17	3,4	3	3	3	9	3	3	3,7
25	Wanita	3	3	4	3	20	4	4	4	4	12	4	4	4
26	Wanita	3	3	3	3	12	3	3	3	3	9	3	3	3,7
27	Wanita	4	3	4	3	17	3,4	3	3	3	9	3	3	3,7

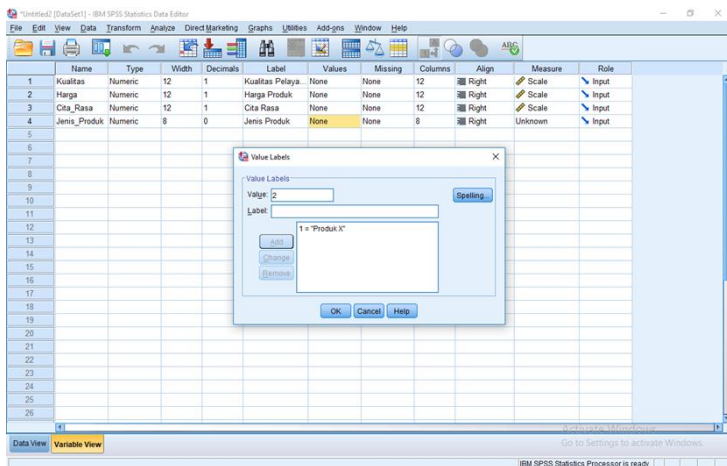
b. Memindah data microsoft excel ke SPSS

Sebelum melakukan pengolahan data, maka data dari MS Excel harus ditransfer kedalam SPSS. Berikut langkah transfer data MS excel ke SPSS.

- Buka SPSS dan klik File-open data, lalu masuk ke folder dimana file excel disimpan.
- Pilih File of type dari *.sav menjadi excel, agar data bentuk excel bisa ditampilkan
- Pilih data excel yang akan di transfer ke dalam SPSS
- Klik OK

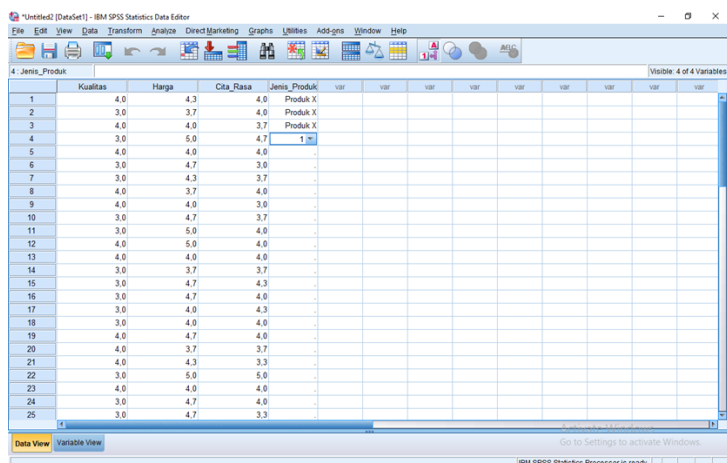


Selanjutnya data hasil transfer dari MS Excel, di edit sesuai dengan vareabel yang diperlukan (diolah dengan SPSS.)



Proses selanjutnya adalah menambahkan vareabel/kelompok sampelnya, sehingga pada kolom values diketikkan, sbb:

- Untuk tempat makan cepat saji Bratang (jenis produk X) diberi label = 1
- Untuk tempat makan cepat saji Nginden (jenis produk Y) diberi label = 2



Visible: 4 of 4 Variables

	Kualitas	Harga	Cita_Rasa	Jenis_Produk
13	4.0	4.0	4.0	Produk X
14	3.0	3.7	3.7	Produk X
15	3.0	4.7	4.3	Produk X
16	3.0	4.7	4.0	Produk X
17	3.0	4.0	4.3	Produk X
18	3.0	4.0	4.0	Produk X
19	4.0	4.7	4.0	Produk X
20	4.0	3.7	3.7	Produk X
21	4.0	4.3	3.3	Produk X
22	3.0	5.0	5.0	Produk X
23	4.0	4.0	4.0	Produk X
24	3.0	4.7	4.0	Produk X
25	3.0	4.7	3.3	Produk X
26	3.0	4.0	4.0	Produk Y
27	4.0	4.0	5.0	Produk Y
28	4.0	4.0	3.0	Produk Y
29	4.0	5.0	4.3	Produk Y
30	5.0	4.0	4.3	Produk Y
31	4.0	4.7	3.7	Produk Y
32	4.0	3.7	2.0	Produk Y
33	4.0	3.7	4.3	Produk Y
34	4.0	4.0	3.7	Produk Y
35	3.0	4.0	3.7	Produk Y
36	3.0	3.0	2.3	Produk Y
37	3.0	4.0	2.7	Produk Y

Data View Variable View

IBM SPSS Statistics Processor is ready

Pengolahan Data Dengan SPSS.

Setelah data masuk dalam bentuk format spss, maka pengolahan data dapat dilakukan dengan urutan sebagai berikut:

Menu Analyse → Compare Means → Independent Samples T test

Visible: 4 of 4 Variables

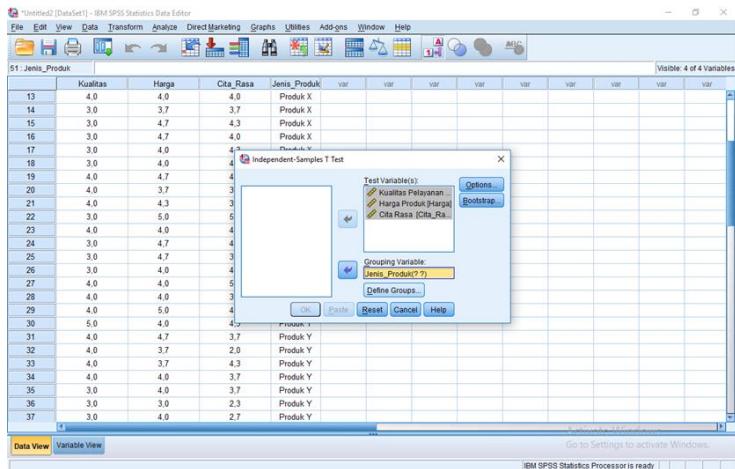
	Kualitas	Harga	Cita_Rasa	Jenis_Produk
13	4.0	4.0	4.0	Produk X
14	3.0	3.7	3.7	Produk X
15	3.0	4.7	4.3	Produk X
16	3.0	4.7	4.0	Produk X
17	3.0	4.0	4.3	Produk X
18	3.0	4.0	4.0	Produk X
19	4.0	4.7	4.0	Produk X
20	4.0	3.7	3.7	Produk X
21	4.0	4.3	3.3	Produk X
22	3.0	5.0	5.0	Produk X
23	4.0	4.0	4.0	Produk X
24	3.0	4.7	4.0	Produk X
25	3.0	4.7	3.3	Produk X
26	3.0	4.0	4.0	Produk Y
27	4.0	4.0	5.0	Produk Y
28	4.0	4.0	3.0	Produk Y
29	4.0	5.0	4.3	Produk Y
30	5.0	4.0	4.3	Produk Y
31	4.0	4.7	3.7	Produk Y
32	4.0	3.7	2.0	Produk Y
33	4.0	3.7	4.3	Produk Y
34	4.0	4.0	3.7	Produk Y
35	3.0	4.0	3.7	Produk Y
36	3.0	3.0	2.3	Produk Y
37	3.0	4.0	2.7	Produk Y

Data View Variable View

Independent-Samples T Test...

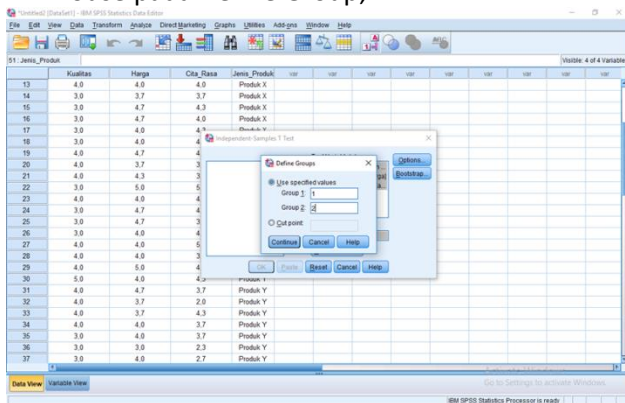
IBM SPSS Statistics Processor is ready

Setelah kotak Dialog independet Sample t test muncul. Masukkan Variabel **Kualitas Pelayanan, Harga dan Cita rasa** ke kotak test variable(s).



Variabel Jenis Produk masukkan ke Grouping Variabel, karena variabel pengelompokan ada pada variabel Jenis Produk (tempat makan siap saji Bratang dan Nginden)

Kemudian Klik mouse pada Define Group;



Untuk Group 1, isi dengan angka 1, (yang berarti Group 1 berisi Jenis tempat makan siap saji produk X)

Untuk Group 2, isi dengan angka 2, (yang berarti Group 2 berisi Jenis tempat makan siap saji produk Y.)

Kemudian tekan **Ok**

Berikut disajikan output hasil olah data dengan spss;

Output SPSS dan Analisis

➤ **Output bagian pertama (Group Statistics)**

Group Statistics					
Jenis Produk		N	Mean	Std. Deviation	Std. Error Mean
Kualitas Pelayanan	Produk X	25	3,440	,5066	,1013
	Produk Y	25	3,680	,5568	,1114
Harga Produk	Produk X	25	4,344	,4610	,0922
	Produk Y	25	4,100	,5236	,1047
Cita Rasa	Produk X	25	3,896	,4486	,0897
	Produk Y	25	3,624	,7944	,1589

Ket: Produk X = tempat makan siap saji Bratang

Produk Y = tempat makan siap saji Nginden

Pada bagian pertama terlihat ringkasan statistik dari kedua sampel. Untuk Kualitas Pelayanan tempat makan siap saji produk X memiliki nilai rata-rata 3,440 berada dibawah nilai rata-rata tempat makan siap saji produk Y, Bila dilihat dari Harga nilai rata-rata tempat makan siap saji produk X lebih tinggi bila dibandingkan dengan produk Y. Sedangkan bila dilihat dari Cita Rasa tempat makan siap saji produk X memiliki nilai rata-rata lebih tinggi bila dibandingkan dengan tempat makan siap saji Produk Y.

➤ **Output bagian kedua (Independent Sample Test)**

Independent Samples Test									
		Levene's Test for Equality of Variances		t-test for Equality of Means					
		F	Sig.	t	df	Sig. (2-tailed)	Mean Difference	Std. Error Difference	95% Confidence Interval of the Difference Lower Upper
Kualitas Pelayanan	Equal variances assumed	,004	,950	-1,594	48	,117	-,2400	,1506	-,5427 ,0627
	Equal variances not assumed			-1,594	47,579	,118	-,2400	,1506	-,5428 ,0628
Harga Produk	Equal variances assumed	,101	,752	1,749	48	,087	,2440	,1395	-,0365 ,5245
	Equal variances not assumed			1,749	47,243	,087	,2440	,1395	-,0367 ,5247
Cita Rasa	Equal variances assumed	8,781	,005	1,491	48	,143	,2720	,1825	-,0949 ,6389
	Equal variances not assumed			1,491	37,893	,144	,2720	,1825	-,0974 ,6414

Uji t dua kelompok sampel dilakukan dalam dua tahapan:

- Menguji apakah varian masing-masing dari dua populasi bisa dianggap sama
- Menguji untuk melihat ada tidaknya perbedaan rata-rata (populasi).

- **Variabel Kualitas Pelayanan**

- a. **Uji Varians;**

Pertama dilakukan pengujian apakah ada kesamaan varians pada dua tempat makan siap saji tersebut.

Hipotesis

Hipotesis untuk pengujian Varian:

- H_0 : Kedua varian populasi adalah identik (varian pada kelompok tempat makan siap saji Bratang dan Nginden adalah Sama)
- H_a : Kedua varian populasi adalah tidak identik (varian pada kelompok tempat makan siap saji Bratang dan Nginden adalah tidak sama)

Pengambilan Keputusan

Jika probabilitas $> 0,05$ maka H_0 diterima

Jika Probabilitas $< 0,05$ maka H_0 ditolak

Keputusan

Dari tabel output kedua diatas dapat dilihat bahwa F hitung kualitas pelayanan dengan nilai probabilitas $0,950 > 0,05$. Karena nilai probabilitas $> 0,05$ maka H_0 diterima, yang berarti kedua varian populasi identik.

- b. **Uji Beda Rata-rata;**

Hipotesis

Hipotesis untuk kasus ini:

- H_0 : Kedua rata-rata populasi identik (rata-rata kualitas pelayanan tempat makan siap saji Bratang dan tempat makan siap saji Nginden adalah sama)
- H_a : Kedua rata-rata populasi tidak identik (rata-rata kualitas pelayanan tempat makan siap saji Bratang dan tempat makan siap saji Nginden adalah tidak sama).

Keputusan

Terlihat bahwa t hitung untuk variabel Kualitas Pelayanan dengan equal variance assumed adalah $-1,594$ dengan probabilitas $0,117$. Karena $0,117 > 0,05$ maka H_0 diterima. Ini berarti rata-rata kualitas pelayanan tempat makan siap saji Bratang dan tempat makan siap saji Nginden adalah identik atau tidak berbeda.

Hasil Analisis juga menunjukkan keputusan yang sama untuk variabel Harga.

- **Variabel Cita Rasa makanan**

Tahapan uji untuk variabel Cita Rasa sama dengan proses diatas, pertama dilakukan pengujian apakah ada kesamaan varians pada dua tempat makan siap saji tersebut.

Hipotesis

Hipotesis untuk pengujian Varian:

- H_0 : Kedua varian populasi adalah identik (varian pada dua kelompok tempat makan siap saji tersebut sama)
- H_a : Kedua varian populasi adalah tidak identik (varian pada dua kelompok tempat makan siap saji tersebut adalah tidak sama)

Pengambilan Keputusan

Jika probabilitas $> 0,05$ maka H_0 diterima

Jika Probabilitas $< 0,05$ maka H_0 ditolak

Keputusan

Dari tabel output kedua diatas dapat dilihat bahwa F hitung Cita Rasa dengan nilai probabilitas $0,005 < 0,05$. Karena nilai probabilitas $< 0,05$ maka H_0 ditolak, yang berarti kedua varian populasi tidak identik.

Perbedaan yang nyata dari kedua varians membuat penggunaan varian untuk membandingkan rata-rata populasi dengan t test sebaiknya menggunakan dasar Equal variance not assumed.

Uji Beda Rata-rata

Hipotesisi

Hipotesis untuk kasus ini:

- H_0 : Kedua rata-rata populasi identik (rata-rata Cita rasa makanan di dua tempat makan siap saji tersebut adalah sama)
- H_a : Kedua rata-rata populasi tidak identik (rata-rata Cita Rasa makanan di dua tempat makan siap saji tersebut adalah tidak sama).

Keputusan

Terlihat bahwa t hitung untuk variabel Cita Rasa Makanan dengan equal variance not assumed adalah 1,491 dengan probabilitas 0,144. Karena $0,144 > 0,05$ maka H_0 diterima. Ini berarti rata-rata Cita Rasa makanan di dua tempat makan siap saji tersebut adalah identik atau tidak berbeda.

DAFTAR PUSTAKA

- Arikunto, S. (1990). *Prosedur Penelitian*. Yogyakarta: Rineka.
- Azwar, S. (1998). *Metode Penelitian, Edisi I*. Yogyakarta: Pustaka Pelajar.
- Rachbini, W. (2020). *Metode Riset Ekonomi dan Bisnis*. Jakarta: Indef.
- Santoso, S. (2015). *SPSS 20*. Jakarta: PT Elex Media Komputindo.

ANALISIS REGRESI SPASIAL DATA PANEL DENGAN STATA

Yuhana Astuti, Ph.D
Universitas Telkom

Analisis regresi spasial data panel merupakan pengembangan dari model data panel tradisional dimana pada analisis spasial mempertimbangkan faktor ruang dan waktu. Kajian tentang efek spasial pertama kali dipergunakan dalam ilmu geografi dengan pemikiran bahwa lokasi suatu wilayah akan mempengaruhi wilayah lainnya. Wilayah yang lokasinya terletak berdekatan akan memiliki pengaruh yang lebih kuat daripada yang terletak berjauhan. Selanjutnya mengikuti perkembangan waktu kajian tentang korelasi spasial diadopsi ke dalam berbagai penelitian di berbagai bidang ilmu lainnya seperti ilmu sosial.

A. KONSEP EKONOMETRIKA SPASIAL

Suatu model menggunakan analisis ekonometrika spasial karena adanya ketergantungan waktu, variabel yang terabaikan, heterogenitas spasial, basis eksternalitas dan ketidakpastian model (LeSage and Pace, 2010). Pemilihan model terbaik dalam model ekonometrika spasial dapat dilakukan dengan membandingkan tiga model spasial yang berkembang diantaranya yaitu:

a. General Spatial Model (GSM)

Model regresi spasial yang dikembangkan oleh Anselin (1988). Asumsi pada regresi spasial sama halnya dengan asumsi pada model regresi klasik yaitu kehomogenan, kenormalan, dan tidak ada otokorelasi dari galat, dimana pendugaan parameter pada model GSM diperoleh dengan metode penduga kemungkinan maksimum.

b. Spatial Autogressive (SAR)

Model otogresi spasial (SAR) biasa disebut juga dengan Spatial Lag Model (SLM) merupakan model regresi linier yang pada peubah responnya terdapat korelasi spasial yang mengkombinasikan regresi biasa dengan model regresi spasial lag pada peubah.

c. Spatial Errors Model (SEM)

Model galat spasial adalah model regresi linier yang pada bentuk galatnya terdapat korelasi spasial. Hal ini disebabkan oleh adanya peubah penjelas yang tidak dilibatkan dalam model regresi linier sehingga akan dihitung sebagai galat dan peubah tersebut berkorelasi dengan galat pada lokasi lain.

B. UJI EFEK SPASIAL

Uji efek spasial dilakukan melalui dua tahap yaitu pertama, melakukan uji Ketergantungan Spasial untuk mengetahui ada tidaknya pola interaksi spasial data dalam model dengan menggunakan Uji Pesaran. Kedua adalah melakukan Uji Keragaman Spasial untuk mendeteksi heterogenitas ragam spasial dalam model dengan menggunakan uji Breusch-Pagan.

C. MATRIKS PEMBOBOT SPASIAL

Keterkaitan wilayah berdasarkan jarak (distance) ditentukan oleh jarak antar wilayah. Semakin dekat jarak antar dua wilayah yang bertetangga maka akan memiliki hubungan yang lebih kuat atau semakin kuat interaksi yang terjadi. Matriks pembobot spasial (W) diperoleh berdasarkan informasi jarak antar suatu wilayah dengan wilayah yang lain dengan membakukan matriks pembobot (Arbia, 2006; Caroline, 2020). Adapun formulanya adalah sebagai berikut:

$$\begin{aligned} W_{ij}(d) &= 0 \text{ jika } i = j \\ W_{ij}(d) &= 1/d_{ij}^2 \text{ jika } i \neq j, \text{ dan } d_{ij} \leq D \\ W_{ij}(d) &= 0 \text{ jika } i \neq j, \text{ dan } d_{ij} > D \end{aligned}$$

Dimana d_{ij} adalah jarak euklidian antara titik tengah suatu wilayah, sedangkan D adalah batas jarak terjauh yang diasumsikan sudah tidak terjadi interaksi.

D. EVALUASI MODEL

Akaike's Information Criterion (AIC) adalah metode yang digunakan untuk memilih model terbaik. Model dikatakan baik jika memiliki nilai AIC yang kecil. Selain metode AIC, model terbaik dapat dipilih dengan melihat koefisien determinasi (R-square). Semakin besar nilai R-square maka model dikatakan semakin tepat menjelaskan peubah respon dengan kata lain apabila nilai R-square bernilai 1 maka ragam peubah respon mutlak dipengaruhi oleh peubah-peubah penjelas yang terdapat pada model.

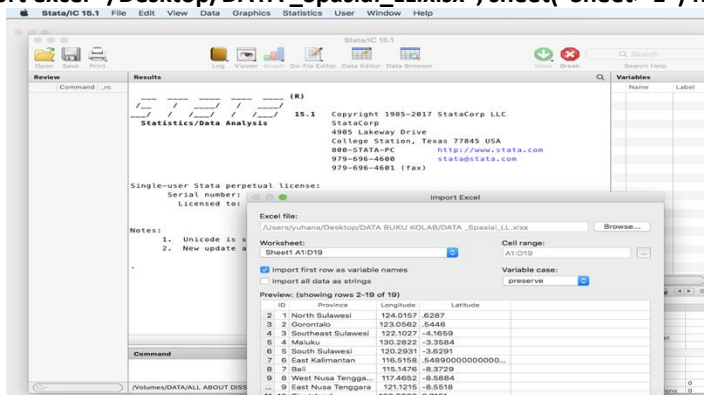
Contoh Kasus:

Indonesia adalah salah satu negara yang menjadi negara tujuan investasi modal asing (FDI). Data dari Badan Pusat Statistik menunjukkan bahwa jumlah realisasi FDI proyek pada periode 2009-2014 lebih banyak terkonsentrasi di provinsi Jakarta (44%), West Java (29%), dan Banten (13%), sedangkan untuk nilai realisasi FDI lebih terkonsentrasi di West Java (38%), Jakarta (23%), dan Banten (19%). Permasalahan tidak meratanya penyebaran investasi di setiap provinsi menjadi salah satu perhatian pemerintah Indonesia. Studi ini bertujuan menerapkan analisis panel spasial untuk mengetahui faktor-faktor yang mempengaruhi pemilihan suatu provinsi sebagai lokasi investasi. Data yang digunakan dapat dilihat pada Lampiran 1 dan Lampiran 2. Adapun langkah-langkah yang dilakukan sesuai tujuan studi diatas adalah sebagai berikut:

1. Import Data dari Excel

Untuk data yang akan di import dari excel ke STATA harus dipastikan bahwa data sudah mengikuti format untuk pemodelan regresi spasial data panel. Selanjutnya sebelum melakukan analisis regresi spasial maka terlebih dahulu dihitung bobot spasialnya yaitu dengan menggunakan data lokasi (spatial) yang berkaitan dengan suatu koodinat geografi yaitu lintang (latitude) dan bujur (longitude). Kedua data tersebut mewakili titik lokasi dari suatu provinsi yang dapat diperoleh dari google map, atau menggunakan software GIS atau GeoDa. Data latitude dan longitude yang digunakan terdapat pada data Lampiran 1. Untuk syntax (perintah) yang diketik dalam kolom *command* sebagai berikut:

. import excel "/Desktop/DATA _Spasial_LL.xlsx", sheet("Sheet1") firstrow



2. Menghitung Matrik Pembobot Spasial

Salah satu hal penting dalam analisis model spasial adalah menghitung matrik pembobot spasial. Dalam studi ini kita fokus pada perhitungan matrik bobot spasial berdasarkan jarak. Untuk hasil pembobotan spasial selanjutnya diberi nama `M_bobot`. Pembobotan spasial dilakukan dengan menggunakan perintah yang diketik dalam kolom *command* sebagai berikut:

```
.spmat idistance M_bobot Longitude Latitude, id(ID) normalize (row)
.spmat summarize M_bobot
```

```
Summary of spatial-weighting object M_bobot
```

Matrix	Description
Dimensions	18 x 18
Stored as	18 x 18
Values	
min	0
min>0	.0104889
mean	.0555556
max	.4206013

Hasil pembobotan spasial berdasarkan invers jarak menunjukkan dimensi matriks adalah 18 x 18 dengan nilai pembobot minimum adalah 0, nilai pembobot maksimum adalah 0,420 dan nilai rata-ratanya adalah 0,055. Data pembobot spasial ini akan otomatis tersimpan dan tidak akan hilang selama kita tidak menutup windows software STATA.

3. Analisis Regresi Panel Spasial

a. Import data panel

Untuk melakukan pengolahan analisis regresi spasial maka kembali terlebih dahulu dilakukan import data panel dari excel menggunakan data Lampiran 2. Data kemudian disimpan dalam format `.dta` dengan menggunakan perintah yang diketik dalam kolom *command* sebagai berikut:

```
.import excel "Desktop/Data_Spasial.xlsx", sheet("Sheet1") firstrow
```

b. Melakukan set data panel

Langkah selanjutnya adalah menetapkan data yang kita gunakan sebagai data panel dengan menggunakan perintah yang diketik dalam kolom *command* sebagai berikut:

. xtset ID Year

```
. xtset ID Year
      panel variable:  ID (strongly balanced)
      time variable:  Year, 2009 to 2014
      delta:          1 unit
```

Hasil menunjukkan bahwa data panel yang digunakan sudah sangat seimbang (strongly balanced) dengan waktu data analisis adalah dari periode 2009 sampai 2014.

c. Uji ketergantungan spasial

Mendeteksi ada tidaknya ketergantungan spasial pada model perlu terlebih dahulu harus dilakukan sebelum melakukan pengolahan regresi spasial data panel dengan melakukan uji Pesaran. Hipotesis uji ketergantungan spasialnya adalah:

$$H0: \delta_i^2 = \alpha_i$$

$$H1: \delta_i^2 \neq \alpha_i$$

Adapun perintah yang diketik dalam kolom *command* adalah sebagai berikut:

. xtreg CFDI GDRP Wage Road School CFDI Electricity

```
. xtreg CFDI GDRP Wage Road School CFDI Electricity
note: CFDI omitted because of collinearity

Random-effects GLS regression              Number of obs   =       108
Group variable: ID                        Number of groups  =        18

R-sq:                                     Obs per group:
      within = 0.3361                      min =           6
      between = 0.5121                     avg =          6.0
      overall = 0.4719                     max =           6

Wald chi2(5) =       56.92
corr(u_i, X) = 0 (assumed)                Prob > chi2      =    0.0000
```

	Coef.	Std. Err.	z	P> z	[95% Conf. Interval]	
GDRP	.0013352	.0002614	5.11	0.000	.0008229	.0018476
Wage	-.0093762	.0504535	-0.19	0.853	-.1082632	.0895108
Road	-43.68292	26.87798	-1.63	0.104	-96.3628	8.996954
School	4.199399	2.200547	1.91	0.056	-.1135938	8.512392
CFDI	0 (omitted)					
Electricity	7.168388	5.664947	1.27	0.206	-3.934704	18.27148
_cons	-280.141	110.4179	-2.54	0.011	-496.5561	-63.72593
sigma_u	153.39833					
sigma_e	86.45849					
rho	.75891627	(fraction of variance due to u_i)				

. xtcsd, pes abs

. xtcsd, pes abs

Pesaran's test of cross sectional independence = 3.740, Pr = 0.0002

Average absolute value of the off-diagonal elements = 0.672

Hasil uji Pesaran menunjukkan nilai probabilitas (Pr) adalah 0,0002. Karena nilai probabilitas $0,0002 < 0,05$ maka H_0 ditolak yang berarti tidak terdapat ketergantungan spasial di dalam model.

d. Uji Breush-Pagan

Selanjutnya dilakukan uji Breusch-Pagan untuk mendeteksi keheterogenan ragam spasial di dalam model. Hipotesis untuk pengujian keheterogenan ragam adalah:

$$H_0: \delta_i^2 = \alpha_i$$

$$H_1: \delta_i^2 \neq \alpha_i$$

Adapun perintah yang diketik dalam kolom *command* adalah sebagai berikut:

. regress CFDI GDRP Wage Road School Electricity

. hettest GDRP Wage Road School Electricity

. regress CFDI GDRP Wage Road School Electricity						
Source	SS	df	MS	Number of obs = 108		
Model	2377637.13	5	475527.426	F(5, 102)	= 18.89	
Residual	2567031.64	102	25166.9769	Prob > F	= 0.0000	
				R-squared	= 0.4808	
				Adj R-squared	= 0.4554	
Total	4944668.77	107	46211.8576	Root MSE	= 158.64	
CFDI	Coef.	Std. Err.	t	P> t	[95% Conf. Interval]	
GDRP	.0012011	.0001443	8.33	0.000	-.0009149	.0014872
Wage	.0108338	.0627526	0.17	0.863	-.1136357	.1353033
Road	-26.76986	35.55885	-0.75	0.453	-97.30067	43.76096
School	1.751891	2.386344	0.73	0.465	-2.981412	6.485194
Electricity	2.560081	6.823987	0.38	0.708	-10.97527	16.09543
_cons	-137.2361	125.2922	-1.10	0.276	-385.7526	111.2804
. hettest GDRP Wage Road School Electricity						
Breusch-Pagan / Cook-Weisberg test for heteroskedasticity						
Ho: Constant variance						
Variables: GDRP Wage Road School Electricity						
chi2(5)	= 301.41					
Prob > chi2	= 0.0000					

Jika digunakan taraf signifikansi sebesar 5% maka terlihat nilai probabilitas adalah 0,00. Karena nilai probabilitas < 0,05 maka H0 ditolak, yang berarti tidak terdapat heterogenitas ragam spasial di dalam model.

e. Pemodelan spasial data panel

Analisis efek spasial dapat dilakukan dengan menggunakan tiga model yaitu General Spatial Model (GSM), Spatial Autogressive Model (SAR), dan Spatial Error Model (SEM). Dari ketiga model ini nantinya akan dipilih satu model terbaik. Analisis menggunakan tiga model dapat kita lakukan sebagai berikut:

i. Model GSM

Pemodelan GSM dilakukan untuk mengetahui efek spasial pada nilai lag dan error termnya. Analisis pemodelan GSM panel dapat dilakukan dengan perintah yang diketik dalam kolom *command* sebagai berikut:

. xsmle CFDI GDRP Wage Road School Electricity, fe wmat(M_bobot) emat(M_bobot) mod(sac)

```

. xsmle CFDI GDRP Wage Road School Electricity, fe wmat(M_bobot) emat(M_bobot) mod(sac)
Iteration 0: Log-likelihood = -621.59282
Iteration 1: Log-likelihood = -619.92718
Iteration 2: Log-likelihood = -619.56602
Iteration 3: Log-likelihood = -619.56329
Iteration 4: Log-likelihood = -619.56329

SAC with spatial fixed-effects                                Number of obs =      108

Group variable: ID                                           Number of groups =     18

Time variable: Year                                           Panel length =        6

R-sq:   within = 0.4054
        between = 0.3817
        overall = 0.3464

Mean of fixed-effects = -2.8e+02

Log-likelihood = -619.5633

```

	CFDI	Coef.	Std. Err.	z	P> z	[95% Conf. Interval]
Main						
	GDRP	.0012222	.0003731	3.28	0.001	.0004908 .0019535
	Wage	-.0963289	.0461354	-2.09	0.037	-.1867526 -.0059052
	Road	-17.45393	25.94081	-0.67	0.501	-68.29698 33.38913
	School	3.714527	2.163233	1.72	0.086	-.5253321 7.954386
	Electricity	25.09998	8.733983	2.87	0.004	7.981693 42.21828
Spatial						
	rho	.4278548	.166291	2.57	0.010	.1019305 .7537792
	lambda	-.5404629	.376512	-1.44	0.151	-1.278413 .197487
Variance						
	sigma2_e	6390.389	779.8399	8.19	0.000	4861.931 7918.847

Berdasarkan hasil pemodelan dengan model GSM diperoleh nilai R-square sebesar 0,34 yang dapat diartikan bahwa GDRP, Wage, Road, School dan Electricity berpengaruh sebesar 34 % terhadap pemilihan lokasi investasi disuatu daerah sedangkan sisanya 66 % dipengaruhi oleh faktor lainnya diluar model. Hasil pemodelan GSM diperoleh nilai signifikansi rho=0,010 < alpa=0,05 yang berarti tidak terdapat efek spasial yang berpengaruh terhadap pemilihan lokasi investasi disuatu provinsi. Namun, nilai signifikansi lamda=

0,151 > $\alpha = 0,05$ menunjukkan terdapat efek spasial yang berpengaruh terhadap pemilihan lokasi investasi di suatu provinsi.

.estat ic

. estat ic						
Akaike's information criterion and Bayesian information criterion						
Model	Obs	ll(null)	ll(model)	df	AIC	BIC
.	108	.	-619.5633	8	1255.127	1276.584
Note: N=Obs used in calculating BIC; see [R] BIC note.						

Akaike's Information Criterion (AIC) untuk pemodelan GSM adalah sebesar 1255.127.

ii. Model SAR

Pemodelan SAR dilakukan untuk mengetahui efek spasial antara variabel bebas dengan variabel tidak bebasnya. Jika nilai signifikansi spasial (ρ) < 0.05 maka tidak terdapat efek spasial di dalam model. Analisis pemodelan SAR dapat dilakukan dengan perintah yang diketik dalam kolom *command* sebagai berikut:

.xsmle CFDI GDRP Wage Road School Electricity, wmat(M_bobot) mod(sar) hausman nolog

. xsmle CFDI GDRP Wage Road School Electricity, wmat(M_bobot) mod(sar) hausman nolog						
... estimating fixed-effects model to perform Hausman test						
SAR with random-effects			Number of obs =		108	
Group variable: ID			Number of groups =		18	
Time variable: Year			Panel length =		6	
R-sq:	within = 0.3514					
	between = 0.5442					
	overall = 0.5005					
Log-likelihood = -657.7276						
	CFDI	Coef.	Std. Err.	z	P> z	[95% Conf. Interval]
Main						
	GDRP	-.0012305	.0002383	5.16	0.000	-.0007634 .0016977
	Wage	-.0320322	.0503276	-0.64	0.524	-.1306724 .066608
	Road	-41.44895	25.87209	-1.60	0.109	-92.15732 9.259413
	School	3.01638	2.17405	1.39	0.165	-1.244679 7.277439
	Electricity	6.389802	5.367001	1.19	0.234	-4.129327 16.90893
	_cons	-202.688	111.0082	-1.83	0.068	-420.2602 14.88415
Spatial						
	rho	.3096784	.1677747	1.85	0.065	-.019154 .6385108
Variance						
	lgt_theta	-1.059027	.2545077	-4.16	0.000	-1.557853 -.5602011
	sigma2_e	7169.051	1082.064	6.63	0.000	5048.244 9289.858
Ho: difference in coeffs not systematic chi2(6) = 655.59 Prob>=chi2 = 0.0000						

Berdasarkan hasil perhitungan dengan SAR panel diperoleh nilai R-square sebesar 0,50 yang dapat diartikan bahwa GDRP, Wage, Road, School dan Electricity berpengaruh sebesar 50 % terhadap pemilihan lokasi investasi di suatu provinsi sedangkan sisanya 50 % dipengaruhi oleh faktor lainnya diluar model. Sedangkan nilai signifikansi $\rho=0,06 > \alpha=0,05$ menunjukkan terdapat efek spasial yang berpengaruh terhadap terpilihnya suatu provinsi sebagai lokasi investasi.

.estat ic

Akaike's information criterion and Bayesian information criterion						
Model	Obs	ll(null)	ll(model)	df	AIC	BIC
.	108	.	-657.7276	9	1333.455	1357.594
Note: N=Obs used in calculating BIC; see [R] BIC note.						

Akaike's Information Criterion (AIC) untuk pemodelan SAR panel adalah 1333.455.

iii. Model SEM

Pemodelan SEM dilakukan untuk mengetahui efek spasial pada nilai error termnya. Jika nilai signifikansi spasial (λ) $< 0,05$ maka tidak terdapat efek pembobotan spasial terhadap terpilihnya suatu provinsi sebagai lokasi investasi. Analisis pemodelan SEM panel dapat dilakukan dengan perintah yang diketik dalam kolom *command* sebagai berikut:

. xsmle CFDI GDRP Wage Road School Electricity, emat(M_bobot) mod(sem) hausman nolog

. xsmle CFDI GDRP Wage Road School Electricity, emat(M_bobot) mod(sem) hausman nolog ... estimating fixed-effects model to perform Hausman test						
SEM with random-effects			Number of obs =		108	
Group variable: ID			Number of groups =		18	
Time variable: Year			Panel length =		6	
R-sq:	within = 0.3348					
	between = 0.5178					
	overall = 0.4757					
Log-likelihood = -658.5825						
CFDI	Coef.	Std. Err.	z	P> z	[95% Conf. Interval]	
Main						
GDRP	.0012806	.0002389	5.36	0.000	.0008125	.0017488
Wage	-.0161348	.0521614	-0.31	0.757	-.1183693	.0860996
Road	-35.64278	26.86414	-1.33	0.185	-88.29552	17.00996
School	3.97162	2.229074	1.78	0.075	-.3972861	8.348525
Electricity	5.105025	5.546218	0.92	0.357	-5.765362	15.97541
_cons	-255.4988	122.0361	-2.09	0.036	-494.685	-16.31248
Spatial						
lambda	.2797261	.2247323	1.24	0.213	-.1607412	.7201934
Variance						
ln_phi	.8327153	.400448	2.08	0.038	.0478517	1.617579
sigma2_e	7334.098	1103.474	6.65	0.000	5171.33	9496.866
Ho: difference in coeffs not systematic chi2(6) = 27.42 Prob>=chi2 = 0.0001						

Berdasarkan hasil perhitungan dengan SEM panel diperoleh nilai R-square sebesar 0,47 yang dapat diartikan bahwa GDRP, Wage, Road, School dan Electricity berpengaruh sebesar 47 % terhadap pemilihan provinsi sebagai lokasi investasi dan sisanya 53 % dipengaruhi oleh faktor lainnya diluar model. Sedangkan nilai signifikansi lamda = 0,213 > $\alpha = 0,05$ menunjukkan terdapat efek spasial yang berpengaruh terhadap terpilihnya suatu provinsi sebagai lokasi investasi.

.estat ic

Akaike's information criterion and Bayesian information criterion						
Model	Obs	ll(null)	ll(model)	df	AIC	BIC
.	108	.	-658.5825	9	1335.165	1359.304
Note: N=Obs used in calculating BIC; see [R] BIC note.						

Akaike's Information Criterion (AIC) untuk pemodelan SEM adalah sebesar 1335.165.

4. Evaluasi Model

Pemilihan model terbaik dilakukan dengan menggunakan nilai Akaike's Information Criterion (AIC). Dengan membandingkan hasil perhitungan ketiga model diatas maka model GSM memiliki nilai AIC terkecil yaitu 1225.128. Studi ini menyimpulkan bahwa dengan mempertimbangkan efek spasial maka pemodelan GSM yang lebih baik digunakan untuk menganalisis faktor-faktor yang mempengaruhi pemilihan suatu provinsi sebagai lokasi investasi di Indonesia.

DAFTAR PUSTAKA

- Anselin, L.1988. *Model Validation in Spatial Econometrics: A Review and Evaluation of Alternative Approaches*. International Regional Science Review, 11 (3), 279-316.
- Arbia G.2006. *Spatial Econometrics: Statistical Foundations and Applications to Regional Convergence*. Berlin: Springer Publications.
- Caroline, E. 2020. *Aplikasi Ekonometrika Spasial dengan Software STATA*. Scopindo Media Pustaka.
- LeSage, J.P., & Pace, R.K. 2010. *Spatial Econometric Models*. Handbook of Applied Spatial Analysis.

ANALISA SPASIAL MENGGUNAKAN R STUDIO

Wahyu Catur Adinugroho, S.Hut., M.Si

Rinaldi Imanuddin, S.Hut., M.Sc

Pusat Penelitian dan Pengembangan Hutan, Kllhk

R adalah suatu bahasa pemrograman komputer berorientasi obyek yang penulisannya “case-sensitive”. R juga merupakan “open source software” sehingga dapat digunakan secara gratis. “R” mempunyai kemampuan digunakan untuk ekstrak atau pengambilan data, manajemen data (organizing), penampilan data (visualizing), pemodelan data (modeling), dan aplikasi data untuk berbagai analisis (performing) (Chamber, 2008). R dapat digunakan pada berbagai bidang penelitian dan jenis analisis termasuk analisis spasial. Pemanfaatannya berkembang pesat dan semakin populer karena gratis, dukungan fungsi yang powerful melalui package-package yang tersedia, memiliki kemampuan untuk membaca berbagai macam jenis data, memiliki kemampuan pemrosesan data yang cepat untuk pemrosesan “big data”, mempunyai kemudahan untuk membuat produk yang sama (reproducible) dari hasil analisis atau visualisasi, ditambah adanya dukungan komunitas ilmiah dan dokumentasi dokumen “help” yang memadai. Terlepas dari kelemahan dari R yang sering disampaikan yaitu sulit dipelajari karena pengguna harus menuliskan sebuah “script” untuk melakukan komputing, analisa dan visualisasi sehingga kurang “user friendly”. Untuk mengenal lebih lanjut dapat mengunjungi : <https://www.r-project.org/>.

Perangkat lunak RStudio IDE (Integrated Development Environment) mulai dikembangkan untuk memudahkan bagi pengguna dalam menggunakan R sehingga terlihat sedikit lebih “user friendly” meski pengguna tetap harus menuliskan dan menjalankan sebuah “script”.

Pada bab ini akan dijelaskan pengenalan RStudio dan aplikasinya dalam analisis spasial. Selain membuat dan memvisualisasikan peta, analisis data spasial juga berkaitan dengan pertanyaan yang tidak dapat langsung dijawab melalui visualisasi data itu sendiri. Pertanyaan-pertanyaan ini mengacu pada proses hipotetis yang menghasilkan data untuk diamati dan dianalisis lebih lanjut. Deskriptif dan Inferensi statistik untuk proses spasial seperti itu

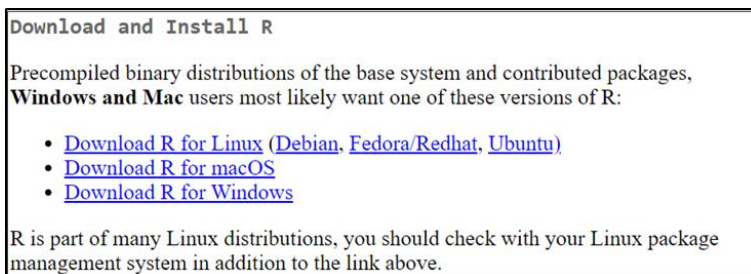
seringkali menantang, tetapi diperlukan ketika kita mencoba menarik kesimpulan tentang pertanyaan yang menarik minat kita.

A. INSTALASI

Untuk bisa menggunakan RStudio, ada dua hal yang harus diinstall. Pertama, software R-nya sendiri, lalu aplikasi RStudio yang merupakan interface yang user-friendly dan akan memudahkan pengguna dalam me-run kode-kode di dalam R.

Cara instalasi R

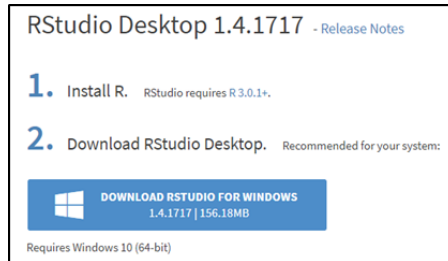
- Buka internet browser, lalu buka halaman <https://repo.bppt.go.id/cran/> untuk mengunduh R sesuai dengan sistem di komputer yang digunakan
- Klik link/tautan "install R for the first time"
- Klik "Download R for Windows" jika operating system komputer Anda menggunakan "windows", dan simpan file .exe tersebut di komputer Anda, lalu run dan ikuti petunjuk instalasi.



Gambar 12 Tampilan halaman akses download R

Cara instalasi RStudio

- Setelah R selesai diinstalasi, sekarang, Anda bisa mengunduh dan menginstall RStudio
- Buka <https://rstudio.com/products/rstudio/download/> dan pilih versi yang 'Free', lalu klik tombol "Download".
- Pilih instalasi sesuai sistem komputer yang Anda miliki, download file instalasi .exe, lalu run, dan ikuti petunjuk instalasi



Gambar 13 Tampilan halaman akses download RStudio

B. HALAMAN ANTARMUKA RSTUDIO

Rstudio memungkinkan kita untuk menyimpan perintah R ke file “script”, melihat variabel saat kita mendefinisikannya, dan melihat output dan visualisasi langsung pada “environment” yang tersedia. Rstudio didukung dengan “Menu bar” dan “Toolbar” untuk memudahkan pengguna. Halaman antarmuka Rstudio terdiri dari 4 jendela utama, yaitu :

1. Script

Script merupakan jendela tempat menuliskan “script” dan menjalankannya. Jendela ini pada saat awal membuka tidak muncul dan dapat diakses melalui tiga cara, yaitu melalui **Menu File > New File > R Script**, melalui toolbar **New** dan melalui shortcut **CTRL+SHIFT+N**.

2. Console

Console merupakan jendela tempat proses dijalankannya “script”, menampilkan hasil jika script berjalan dengan baik, serta menunjukkan kesalahan script melalui error code. Script dapat juga langsung ditulis dan dijalankan di Console tetapi tidak disarankan karena tidak dapat disimpan.

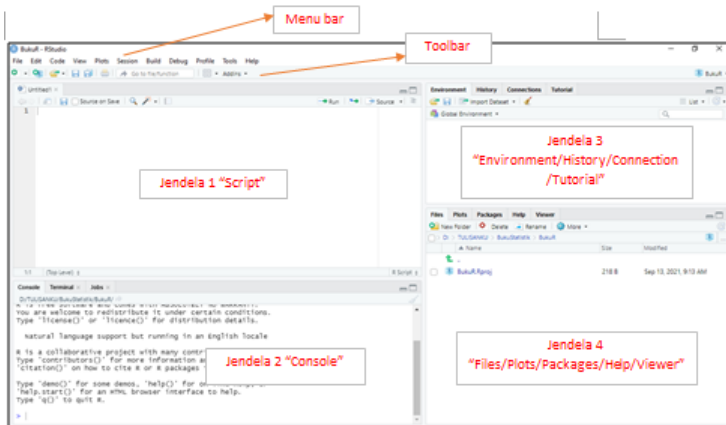
3. Environment/History/Connections/Tutorial

Environment menunjukkan daftar “object” yang tersedia beserta yang telah kita buat. Pada tab History, ditampilkan rangkaian proses-proses yang telah dijalan secara berurutan dan pada tab Tutorial terdapat tutorial penggunaan. Pada jendela ini juga terdapat toolbar yang membantu dalam “import dataset” untuk diproses lebih lanjut dalam “script”.

4. File/Plots/Packages/Help/Viewer

- File tab digunakan sebagai file explorer untuk melakukan navigasi file dari direktori.
- Plots tab tempat menampilkan visualisasi hasil berbentuk plot, grafik atau gambar.

- Packages tab menunjukkan daftar packages yang sudah terinstal pada library. Packages dapat digunakan dalam analisis dengan mengaktifkan packages dengan melakukan “checklist” pada packages yang akan digunakan. Pada jendela ini juga terdapat toolbar untuk melakukan install packages yang akan digunakan dalam analisis menggunakan RStudio.
- Help tab berisi dokumentasi packages, digunakan untuk mencari dan menampilkan bantuan untuk memahami fungsi atau packages.
- Viewer berfungsi untuk melihat data frame dan data lainnya.

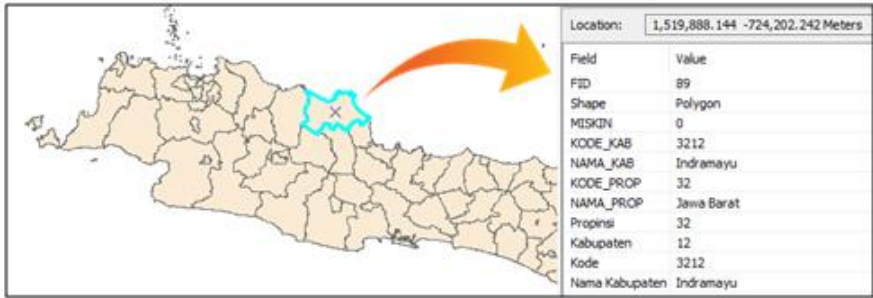


Gambar 14 Tampilan halaman jendela antarmuka Rstudio

C. TIPE DATA SPASIAL

Pemahaman terhadap data spasial dan kemampuan melakukan analisis spasial akan membantu kita menyingkap pengetahuan dan menjawab permasalahan kompleks dengan lebih menyeluruh. Data spasial mampu menunjukkan pola hubungan selanjutnya pola, tren, dan hubungan antara fenomena ini dapat digunakan untuk menjawab berbagai masalah yang berhubungan dengan ruang dan waktu.

Data spasial adalah sebuah data yang berorientasi geografis dan memiliki sistem koordinat tertentu sebagai dasar referensinya. Mempunyai dua bagian penting yaitu informasi lokasi (spasial) dan informasi deskriptif (atribut). Data spasial meliputi data geografis baik dalam format raster maupun vektor.



Gambar 15 Contoh data spasial

- Data vektor – titik, garis, dan wilayah (poligon)

Data vektor merupakan data yang digambarkan dengan titik atau titik yang terhubung pada diagram kartesian (koordinat) sehingga membentuk titik, garis dan poligon. Format file data vektor dapat berupa Shape file (.shp), Simple feature, Autocad DXF, GeoJSON. Untuk memahami data vektor dan memvisualisasikan dalam Rstudio dapat dilihat contoh berikut berdasarkan data pada package “sp”.

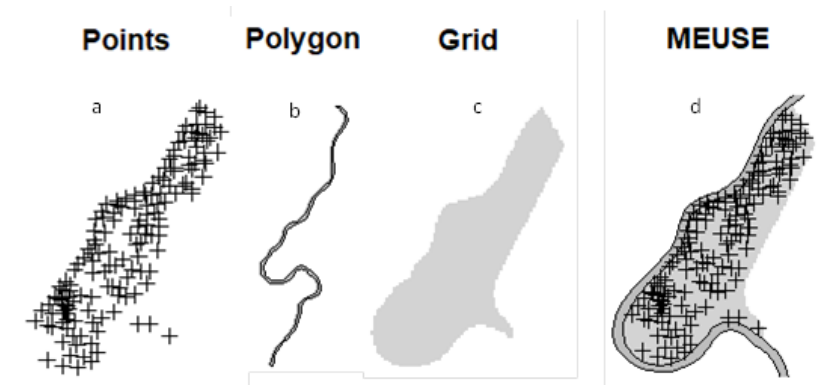
```
1 ##Tipe Data Spasial##
2
3 library(sp)
```

- Data point dan visualisasinya (Gambar 5a)

```
5 ##Point##
6 data("meuse")
7 coordinates(meuse) <- c("x", "y")
8 plot(meuse)
9 title("Points")
```

- Data Polygon dan visualisasinya (Gambar 5b)

```
11 ##Polygon##
12 data("meuse.riv")
13 meuse.lst <- list(Polygons(list(Polygon(meuse.riv)), "meuse.river"))
14 meuse.sr <- SpatialPolygons(meuse.lst)
15 plot(meuse.sr, col="grey")
16 title("polygon")
```

Gambar 16 Data spasial dan visualisai pada RStudio

- Data raster – data grid seperti citra satelit, data elevasi di seluruh permukaan, total curah hujan di seluruh permukaan selama periode waktu tertentu. Layer pada data raster ini berupa rasterlayer atau raster stack. Dalam model raster posisi obyek dipermukaan bumi digambarkan dalam ruang dan disebut sebagai sel. Representasi dari sel berada dalam suatu posisi yang berupa kolom dan baris, yang biasa diistilahkan sebagai pixel (picture element). Format file data raster dapat berupa GeoTIFF, ESRI Grid, IMG – Erdas Imagine Format, ECW (Enhanced Compression Wavelet) JPEG 2000.

- **Data Grid dan visualisasinya (Gambar 5c)**

```
18 ##Raster##
19 data("meuse.grid")
20 coordinates(meuse.grid) <- c("x", "y")
21 meuse.grid <- as(meuse.grid, "SpatialPixels")
22 image(meuse.grid, col="lightgrey")
23 title("Grid")
```

- **Visualisasi gabungan (Gambar 5d)**

```
25 ##Visualisasi Gabungan##
26 image(meuse.grid, col="lightgrey")
27 plot(meuse.sr, col="grey", add=TRUE)
28 plot(meuse, add=TRUE)
29 title("MEUSE")
```

D. CARA MEMPEROLEH DAN IMPORT DATA SPASIAL

Data spasial dapat diperoleh dari berbagai cara dan sumber, seperti : melakukan registrasi dan digitasi peta analog, survei dan pengukuran langsung dengan menyertakan data koordinat, data GPS, data tabular yang terdapat informasi koordinat, data penginderaan jauh. Salah satu keunggulan menggunakan R adalah kemampuan untuk langsung mengunduh atau mengakses data-data spasial yang disediakan oleh pihak eksternal dan megimportnya dalam Rstudio.

- Data batas administrasi

Data batas administrasi semua negara di dunia hingga level provinsi dapat diunduh dari portal data GADM (<https://gadm.org/>). Terdapat dua cara untuk akses data GADM dari Rstudio yaitu menggunakan package raster dan package GADMtools.

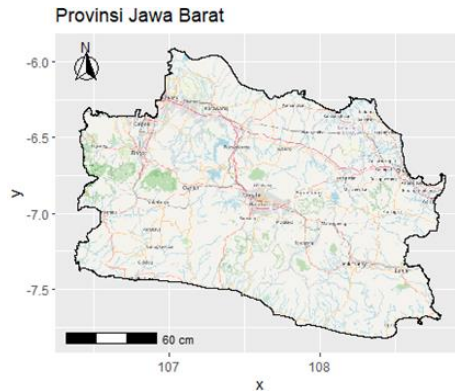
```
1 ##Download Data Spasial##
2 #Batas Administrasi#
3
4 #Menggunakan package "GADMTools"#
5 library(GADMTools)
6 library(sp)
7
8 #Batas negara#
9 IDN <- gadm_sp_loadCountries("IDN", basefile = "./")
10 plotmap(IDN)
11
12 #Batas Provinsi#
13 IDNProv <- gadm_sp_loadCountries(c("IDN"), level=1, basefile = "./")
14 plotmap(IDNProv)
15
16 #Batas Kabupaten#
17 IDNKab <- gadm_sp_loadCountries(c("IDN"), level=2, basefile = "./")
18 plotmap(IDNKab)
```

Untuk melakukan proses ini, harus terdapat koneksi internet dimana “script” akan menjalankan proses untuk menghubungkan ke <https://biogeo.ucdavis.edu/data/> dan mendownload data sesuai wilayah yang kita inginkan, jika berhasil data akan muncul pada jendela environment dan tersimpan pada folder proyek kita dengan format file .rds, tampilan pada Console adalah seperti berikut :

```
> IDNKab <- gadm_sp_loadCountries(c("IDN"), level=2, basefile = "./")
trying URL 'https://biogeo.ucdavis.edu/data/gadm3.6/Rsp/gadm36_IDN_2_sp.rds'
Content type 'application/x-rds' length 25914575 bytes (24.7 MB)
downloaded 24.7 MB
```

Untuk menampilkan hanya provinsi tertentu dapat menggunakan “script” berikut :

```
21 #Menampilkan wilayah tertentu dan membuat peta#
22 listNames(IDNProv, level=1)
23 Jabar = gadm.subset(IDNProv, regions=c("Jawa Barat"))
24 Jabar2 <- gadm_getBackground(Jabar, "Jabar", "osm")
25 gadm_plot (Jabar2, title = "Provinsi Jawa Barat"
26           ) %>% gadm_showNorth("t1") %>% gadm_showScale("b1")
```



Gambar 17 Plot view data spasial batas administrasi pada provinsi Jawa Barat

- Data elevasi

Data elevasi dari data citra SRTM dengan resolusi spasial 90 meter dapat didownload pada Rstudio menggunakan package raster. Untuk mendownload data ini diperlukan informasi lintang dan bujur dari AOI (Area of Interest). Untuk mengetahui koordinat lat/long dari AOI ini kita dapat memanfaatkan Google Maps.

```
50 #Data Elevasi SRTM#
51 library(raster)
52 library(sp)
53
54 #Download
55
56 srtm <- getData('SRTM', lon=106, lat=-6, download = TRUE)
57 srtm2 <-getData("SRTM", lon=112, lat=-6, download = TRUE)
58 plot(srtm)
59 plot(srtm2)
```

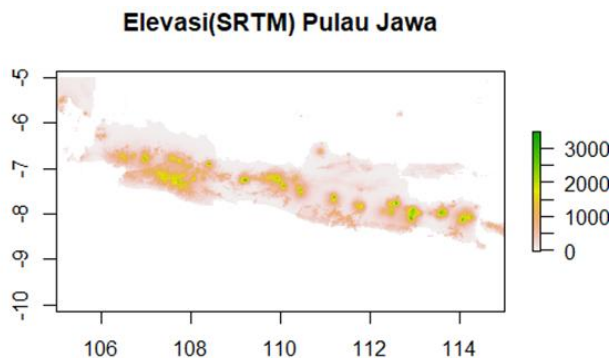
Untuk melakukan proses ini, harus terdapat koneksi internet dimana “script” akan menjalankan proses untuk menghubungkan ke <http://srtm.csi.cgiar.org/> dan mendownload data srtm sesuai wilayah yang kita inginkan, jika berhasil data akan muncul pada jendela environment dan

tersimpan pada folder proyek kita dengan format file .tif , tampilan pada Console adalah seperti berikut :

```
> srtm2 <-getData("SRTM", lon=112, lat=-6, download = TRUE)
trying URL 'http://srtm.csi.cgiar.org/wp-content/uploads/files/srtm_5x5/TIFF/srtm_59_14.zip'
Content type 'application/zip' length 9137457 bytes (8.7 MB)
downloaded 8.7 MB
```

Proses mosaik dapat dilakukan untuk menggabungkan sejumlah tile dari AOI yang sudah kita download dan kemudian dapat divisualisasikan menggunakan “script” berikut :

```
61 #Mosaic/merge srtm tiles
62 srtmmosaic <- mosaic(srtm, srtm2, fun=mean)
63 plot(srtmmosaic, main="Elevasi(SRTM) Pulau Jawa")
```



Gambar 18 Hasil download dan visualisasi data SRTM di Pulau Jawa

- Data iklim

Data iklim secara global tersedia pada <https://www.worldclim.org/>. WorldClim adalah database data cuaca dan iklim global dengan resolusi spasial tinggi. Terdapat 19 bioclimate variables dengan resolusi 2.5 minutes of degrees. Data ini dapat digunakan untuk pemetaan dan pemodelan spasial. Data dapat kita download pada Rstudio dengan packages “raster” dengan fungsi “getData”. Terdapat beberapa argumen yang diperlukan untuk mendownload data tersebut yaitu :

- Select Dataset: dataset yang digunakan adalah 'worldclim'
- Pilih variabel: Variabel yang akan didownload dapat berupa: 'tmin', 'tmax', 'prec' dan 'bio'

BIO1	= Annual Mean Temperature
BIO2	= Mean Diurnal Range (Mean of monthly (max temp – min temp))
BIO3	= Isothermality (BIO2/BIO7) (* 100)
BIO4	= Temperature Seasonality (standard deviation *100)
BIO5	= Max Temperature of Warmest Month
BIO6	= Min Temperature of Coldest Month
BIO7	= Temperature Annual Range (BIO5-BIO6)
BIO8	= Mean Temperature of Wettest Quarter
BIO9	= Mean Temperature of Driest Quarter
BIO10	= Mean Temperature of Warmest Quarter
BIO11	= Mean Temperature of Coldest Quarter
BIO12	= Annual Precipitation
BIO13	= Precipitation of Wettest Month
BIO14	= Precipitation of Driest Month
BIO15	= Precipitation Seasonality (Coefficient of Variation)
BIO16	= Precipitation of Wettest Quarter
BIO17	= Precipitation of Driest Quarter
BIO18	= Precipitation of Warmest Quarter
BIO19	= Precipitation of Coldest Quarter

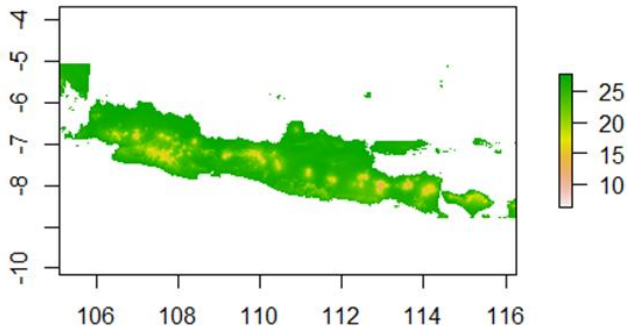
- Tentukan resolusi: 0.5, 2.5, 5, dan 10 (menit derajat). Jika dipilih res=0.5, diperlukan argumen informasi lat dan lon AOI.

```

66 #Data IKLIM#
67 library(raster)
68 library(sp)
69 library(sf)
70
71 pulaujawa <- st_read("D:/TULISANKU/BukuStatistik/BukuR/Data/pulaujawa.shp")
72 plot(pulaujawa)
73
74 # download iklim jawa
75 iklim_jawa <- getData('worldclim', var='bio',
76                       lat = -6, lon = 111,
77                       res=0.5, download = TRUE)
78
79 #clip pulau jawa
80 image_iklimjawa <- crop(iklim_jawa, pulaujawa)
81
82 #Data temperature
83 plot(image_iklimjawa$bio1_39/10, #data perlu dibagi 10 karena satuannya °C x 10
84      main="Annual Mean Temperature Jawa (°C)")

```

Annual Mean Temperature Jawa (°C)



Gambar 19 Hasil download dan visualisasi data suhu pada Worldclim

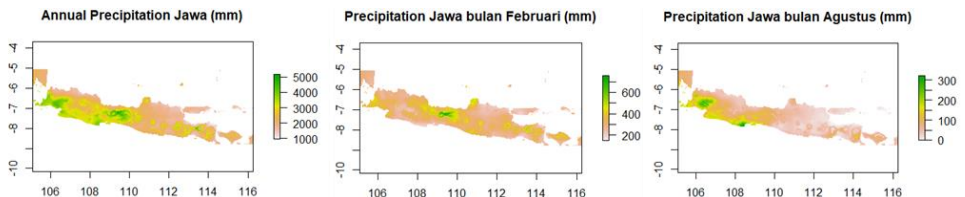
- Data curah hujan

Data curah hujan juga dapat di download pada <https://www.worldclim.org/>, dengan proses yang sama dilakukan dalam mendownload data temperatur. Data curah hujan dapat ditampilkan berdasarkan data rata-rata tahunan atau curah hujan bulanan.

```

87 #Data Curah Hujan#
88 library(raster)
89 library(sp)
90 library(sf)
91
92 #Data curah hujan tahunan
93 plot(image_iklimjawasbio12_39,
94      main="Annual Precipitation Jawa (mm)")
95
96 #Data curah hujan bulanan
97 CH_jawa <- getData('worldclim', var='prec',
98                  lat = -6, lon = 111,
99                  res=0.5, download = TRUE)
100
101 #Clip pulau jawa
102 image_CHJawa <- crop(CH_jawa, pulaujawa)
103 plot(image_CHJawa)
104
105 # ganti nama band dengan bulan untuk menampilkan CH bulanan
106 names(image_CHJawa) <- c("Jan", "Feb", "Mar", "Apr", "Mei", "Jun",
107                        "Jul", "Agu", "Sep", "Okt", "Nov", "Des")
108
109 plot(image_CHJawa$Feb,
110      main="Precipitation Jawa bulan Februari (mm)")
111 plot(image_CHJawa$Agu,
112      main="Precipitation Jawa bulan Agustus (mm)")

```



Gambar 20 Hasil download dan visualisasi data curah hujan pada Worldclim

E. APLIKASI ANALISA DATA SPASIAL DATA DEMOGRAFI

Analisa data spasial seringkali digunakan untuk melihat pengaruh tempat atau pengaruh spasial pada data demografi yang dianalisis. Adanya efek spasial merupakan hal yang lazim terjadi antara satu wilayah dengan wilayah yang lain. Pada beberapa kasus, peubah tak bebas yang diamati memiliki keterkaitan dengan hasil pengamatan di wilayah yang berbeda, terutama wilayah yang berdekatan. Adanya hubungan spasial dalam peubah tak bebas akan menyebabkan pendugaan menjadi tidak tepat karena asumsi keacakan galat dilanggar. Untuk mengatasi permasalahan di atas diperlukan suatu model regresi yang memasukkan hubungan spasial antar wilayah ke dalam model. Adanya informasi hubungan spasial antar wilayah menyebabkan perlu mengakomodir keragaman spasial ke dalam model, sehingga model yang digunakan adalah model regresi spasial. Sebagai studi kasus akan digunakan data kemiskinan penduduk di Jawa Tengah pada tahun 2020, data batas administrasi kabupaten/kota Jawa Tengah. Tahapan analisa yang dilakukan, yaitu :

- Membuat file new project sebagai direktori semua data input dan hasil analisis dari proses yang dilakukan
- Mengaktifkan packages-packages yang akan digunakan dalam analysis

```
1 #ANALISA DATA SPASIAL DEMOGRAFI#
2 library(readr)           #membaca file
3 library(rgdal)           #membaca file shp, fs readOGR
4 library(raster)          #eksplorasi data
5 library(spdep)           #pembobot spasial (fungsi poly2nb)
6 library(lmtest)          #regresi klasik
7 library(nortest)         #uji kenormalan galat
8 library(DescTools)       #uji kebebasan galat
9 library(spatialreg)     #model SAR
10 library(RColorBrewer)   #memberi warna peta
11 library(car)            #uji multikolinieritas
12 library(GWmodel)        #GWR Jarak Spasial
13 library(sf)
14 library(sp)
15 library(ggplot2)        #membuat grafik
```

- Import data yang akan digunakan yaitu data frame “kemiskinan200_jateng.csv” menggunakan import dataset pada environment, data shapefile batas administrasi “JawaTengah.shp” dengan fungsi “readOGR” pada packages “rgdal”

```
18 #Import Dataset, data "kemiskinan2020_jateng.csv"
19 view(Miskin2020_jateng)
20
21 jateng <- readOGR("D:/TULISANKU/BukuStatistik/BukuR/Data/JawaTengah.shp")
22 names(jateng)
23 plot(jateng)
```

- Menggabungkan data kemiskinan dan batas administrasi

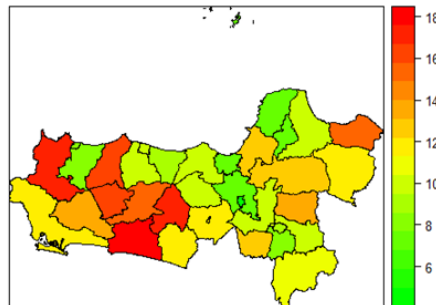
```
25 # menggabungkan data kemiskinan dan batas administrasi dengan kunci 'KODE_KAB'
26 jatengmiskin <- merge(jateng, Miskin2020_jateng, by='KODE_KAB')
27
28 # perhaps save as shapefile again
29 shapefile(jatengmiskin, "D:/TULISANKU/BukuStatistik/BukuR/Data/jatengmiskin.shp")
30 names(jatengmiskin)
31 plot(jatengmiskin)
```

- Eksplorasi data

Berikut adalah sebaran persentase penduduk miskin tahun 2020 di Jawa Tengah. Terdapat 3 pembagian warna yaitu hijau, kuning, dan merah yang didasarkan pada persentase penduduk miskin. Semakin wilayah berwarna merah, maka semakin tinggi persentase penduduk miskin di wilayah tersebut.

```
35 #Eksplorasi Data
36
37 colfunc <- colorRampPalette(c("green", "yellow", "red"))
38 color <- colfunc(16)
39 spplot(jatengmiskin, "PersenMiskin2020", col.regions=color,
40       main="Peta Sebaran Persentase Penduduk Miskin 2020")
```

Peta Sebaran Persentase Penduduk Miskin 2020



Gambar 21 Hasil eksplorasi data

- Autokorelasi spasial menggunakan Indeks Moran

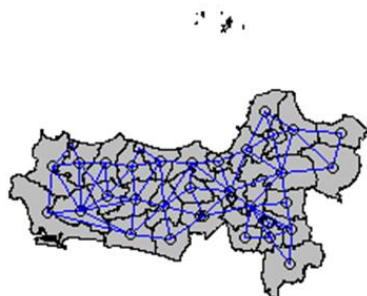
Autokorelasi spasial merupakan teknik dalam analisis spasial untuk mengukur kemiripan nilai atribut dalam suatu ruang (jarak, waktu dan area). Jika terdapat pola sistematis dalam nilai atribut tersebut, maka terdapat autokorelasi spasial. Adanya autokorelasi spasial mengindikasikan bahwa nilai atribut pada area tertentu terkait oleh nilai atribut tersebut pada area lain yang letaknya saling berdekatan (bertetangga). Ketetanggaan tersebut diharapkan dapat mencerminkan derajat ketergantungan area (spasial) yang tinggi apabila dibandingkan dengan area lain yang letaknya terpisah jauh. Hasil

analisis autokorelasi spasial persentase kemiskinan di Jawa Tengah pada tahun 2020, adalah sebagai berikut :

- Penyiapan dataset

```
42 #Autokorelasi Spasial
43 w_jateng <- poly2nb(jatengmiskin)
44 ww_jateng <- nb2listw(w_jateng)
45
46 ##visualisasi
47 plot(jatengmiskin, col="gray")
48 plot(w_jateng, coordinates(jatengmiskin), add = TRUE, col="blue")
```

Neighbours Jawa Tengah



- Moran Index Test

Analisis Moran Index dapat dilakukan dengan menggunakan random test atau juga Monte Carlo Simulation.

```
50 ##Moran Test
51 moran.test(jatengmiskin$PersenMiskin2020, ww_jateng, alternative="greater")
52
53 moran.mc(jatengmiskin$PersenMiskin2020, ww_jateng, nsim = 499)
```

Hasil test ini dapat dilihat pada jendela Console :

```
> moran.test(jatengmiskin$PersenMiskin2020, ww_jateng, alternative="greater")

Moran I test under randomisation

data: jatengmiskin$PersenMiskin2020
weights: ww_jateng

Moran I statistic standard deviate = 2.007, p-value = 0.02238
alternative hypothesis: greater
sample estimates:
Moran I statistic      Expectation      Variance
0.20439342          -0.02941176          0.01357124
```

```
> moran.mc(jatengmiskin$PersenMiskin2020, ww_jateng, nsim = 499)

Monte-Carlo simulation of Moran I

data:  jatengmiskin$PersenMiskin2020
weights: ww_jateng
number of simulations + 1: 500

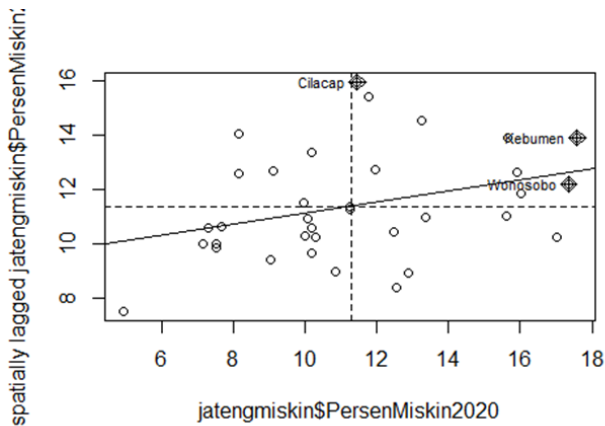
statistic = 0.20439, observed rank = 483, p-value = 0.034
alternative hypothesis: greater
```

Berdasarkan hasil tersebut diketahui nilai Index Moran sebesar 0.204 dengan hasil uji baik menggunakan random test atau menggunakan monte carlo simulation mempunyai p-value dibawah 0.05. Dari pengujian Indeks Moran diperoleh kesimpulan bahwa pada taraf signifikansi 5% dinyatakan terdapat autokorelasi spasial terhadap persen kemiskinan di Jawa Tengah pada tahun 2020. Nilai Indeks Moran sebesar 0,204 berada pada rentang $0 < I \leq 1$ dan menunjukkan adanya autokorelasi spasial positif. Berarti disimpulkan bahwa antar kabupaten satu dengan yang lainnya pada lokasi yang berdekatan memiliki kemiripan nilai atau mengindikasikan bahwa data berkelompok.

- Moran's Scatterplot

Moran's Scatterplot menunjukkan hubungan antara nilai amatan pada suatu lokasi (distandarisasi) dengan rata-rata nilai amatan dari lokasi-lokasi yang bertetangga dengan lokasi yang bersangkutan (Lee dan Wong, 2001). Scatterplot tersebut terdiri atas empat kuadran (Perobelli dan Haddad, 2003), yaitu: Kuadran I (High-High), menunjukkan lokasi yang mempunyai nilai amatan tinggi dikelilingi oleh lokasi yang mempunyai nilai amatan tinggi. Kuadran II (Low-High), menunjukkan lokasi yang mempunyai nilai amatan rendah dikelilingi oleh lokasi yang mempunyai nilai amatan tinggi. Kuadran III (Low-Low), menunjukkan lokasi yang mempunyai nilai amatan rendah dikelilingi oleh lokasi yang mempunyai nilai amatan rendah. Kuadran IV (High-Low), menunjukkan lokasi yang mempunyai nilai amatan tinggi dikelilingi oleh lokasi yang mempunyai nilai amatan rendah.

```
55 ##Moran's Scatterplot
56 moran.plot(jatengmiskin$PersenMiskin2020, ww_jateng,
57           labels=jatengmiskin$NAMA_KAB.x)
```



Gambar 22 Moran scatterplot

Gambar 22 merupakan Moran's scatterplot yang menunjukkan pola hubungan antara persentase penduduk miskin pada suatu kabupaten/kota di Jawa Tengah dengan kabupaten/kota lain. Kabupaten/Kota Wonosobo, Kebumen, Cilacap berada pada Kuadran 1 Kabupaten ini memiliki nilai persentase penduduk miskin tinggi dan berdekatan dengan kabupaten lain yang memiliki nilai persentase penduduk miskin tinggi juga.

- Analisis Regresi Spasial

Suatu analisis pemodelan regresi untuk mengetahui faktor-faktor kemiskinan dengan melibatkan pengaruh aspek spasial adalah sangat penting. Hal ini disebabkan aspek-aspek kemiskinan tidak hanya dijelaskan oleh peubah-peubah penjelas saja, namun aspek lokasi juga menentukan dimana pengamatan di suatu wilayah dipengaruhi oleh pengamatan di wilayah lain. Adanya efek spasial merupakan hal yang lazim terjadi antara satu wilayah dengan wilayah yang lain. Pada beberapa kasus, peubah tak bebas yang diamati memiliki keterkaitan dengan hasil pengamatan di wilayah yang berbeda, terutama wilayah yang berdekatan. Adanya hubungan spasial dalam peubah tak bebas akan menyebabkan pendugaan menjadi tidak tepat karena asumsi keacakan galat dilanggar. Untuk mengatasi permasalahan di atas diperlukan suatu model regresi yang memasukkan hubungan spasial antar wilayah ke dalam model. Adanya informasi hubungan spasial antar wilayah menyebabkan perlu mengakomodir keragaman spasial ke dalam model, sehingga model yang digunakan adalah model regresi spasial. Beberapa metode pada model spasial yang digunakan antara lain model lag

spasial/Spatial Autoregressive Model (SAR), model galat spasial/Spatial Error Model (SEM) dan model umum regresi spasial/General Spatial Model (GSM).

- **Uji Efek Spasial**

Identifikasi efek spasial ini bertujuan untuk mengetahui adanya heterogenitas spasial dan ketergantungan spasial. Kedua hal di atas dilakukan untuk menentukan pemodelan berikutnya, yaitu menentukan model spasial yang akan digunakan untuk memodelkan persentase kemiskinan. Lagrange Multiplier (LM) / Robust LM digunakan untuk mendeteksi ketergantungan spasial secara lebih spesifik yaitu ketergantungan spasial dalam lag, error, atau keduanya (lag dan error).

```
59 #Regresi Klasik
60 reg.klasik <- lm(PersenMiskin2020~X1+X2+X3+X4+X5+X6+X7+X8+X9,data=jatengmiskin)
61 err.regklasik <- residuals(reg.klasik)
62 summary(reg.klasik)
63
64 reg.klasik <- lm(PersenMiskin2020~X3+X4,data=jatengmiskin)
65 err.regklasik <- residuals(reg.klasik)
66 summary(reg.klasik)
67
68 #Efek Spasial
69 LM <- lm.LMtests(reg.klasik, nb2listw(w_jateng, style="w"),
70                  test=c("LMerr", "LMlag", "RLMerr", "RLMlag", "SARMA"))
```

Hasil analisis dapat dilihat pada jendela Console

```
> #Efek Spasial
> LM <- lm.LMtests(reg.klasik, nb2listw(w_jateng, style="w"),
+                 test=c("LMerr", "LMlag", "RLMerr", "RLMlag", "SARMA"))
> summary(LM)
Lagrange multiplier diagnostics for spatial dependence
data:
model: lm(formula = PersenMiskin2020 ~ X3 + X4, data = jatengmiskin)
weights: nb2listw(w_jateng, style = "w")

      statistic parameter p.value
LMerr    0.30959         1 0.57793
LMlag    0.10671         1 0.74392
RLMerr   4.19840         1 0.04046 *
RLMlag   3.99552         1 0.04562 *
SARMA    4.30511         2 0.11619
---
Signif. codes:  0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1
```

Berdasarkan hasil analisis Robust LM Diagnostic tersebut diketahui terdapat signifikansi ketergantungan spasial dalam error dan lag.

- **Model SAR**

Spatial Autoregressive Model (SAR) merupakan model spasial yang terjadi akibat adanya pengaruh spasial pada variabel dependen. Apabila data yang diperoleh menghasilkan dependensi lag maka data dimodelkan dengan SAR.

```

60 #Model Regresi Spasial
61 ##SAR
62 sar <- lagsarlm(PersenMiskin2020~X1+X2+X3+X4+X5+X6+X7+X8+X9,data=jatengmiskin,ww_jateng)
63 summary(sar)

```

Hasil analisis dapat dilihat pada jendela Console

```

> summary(sar)

Call:lagsarlm(formula = PersenMiskin2020 ~ X3 + X4, data = jatengmiskin,
               listw = ww_jateng)

Residuals:
    Min       1Q   Median       3Q      Max
-4.57333 -1.94953 -0.40699  2.21764  4.78946

Type: lag
Coefficients: (asymptotic standard errors)
              Estimate Std. Error z value Pr(>|z|)
(Intercept) 12.572915   2.718147   4.6255 3.736e-06
X3           0.224352   0.082379   2.7234 0.006461
X4          -0.829658   0.296560  -2.7976 0.005148

Rho: 0.067758, LR test value: 0.10579, p-value: 0.74499
Asymptotic standard error: 0.20643
z-value: 0.32823, p-value: 0.74274
wald statistic: 0.10774, p-value: 0.74274

Log likelihood: -83.02926 for lag model
ML residual variance (sigma squared): 6.7235, (sigma: 2.593)
Number of observations: 35
Number of parameters estimated: 5
AIC: 176.06, (AIC for lm: 174.16)
LM test for residual autocorrelation
test value: 3.9497, p-value: 0.046879

```

Model tidak signifikan

- **Model SEM**

Spatial Error Model (SEM) merupakan model spasial yang terjadi akibat adanya pengaruh spasial pada error. Apabila data yang diperoleh menghasilkan dependensi error maka data dimodelkan dengan SEM.

```

65 ##SEM
66 sem <- errorsarlm(PersenMiskin2020~X1+X2+X3+X4+X5+X6+X7+X8+X9,data=jatengmiskin,ww_jateng)
67 summary(sem)

```

Hasil analisis dapat dilihat pada jendela Console

```

> summary(sem)

Call:errorsarlm(formula = PersenMiskin2020 ~ X3 + X4, data = jatengmiskin,
                 listw = ww_jateng)

Residuals:
    Min       1Q   Median       3Q      Max
-4.24480 -1.99926 -0.11025  2.28057  4.65032

Type: error
Coefficients: (asymptotic standard errors)
              Estimate Std. Error z value Pr(>|z|)
(Intercept) 13.536150   1.287410 10.5142 < 2.2e-16
X3           0.263459   0.075002   3.5127 0.0004436
X4          -0.946642   0.286962  -3.2988 0.0009709

Lambda: -0.20274, LR test value: 0.48598, p-value: 0.48573
Asymptotic standard error: 0.23908
z-value: -0.64797, p-value: 0.39645
wald statistic: 0.71906, p-value: 0.39645

Log likelihood: -82.83916 for error model
ML residual variance (sigma squared): 6.5983, (sigma: 2.5687)
Number of observations: 35
Number of parameters estimated: 5
AIC: 175.68, (AIC for lm: 174.16)

```

Model tidak signifikan

- **Model SARMA/GSM**

Jika data menghasilkan dependensi lag dan dependensi error maka data dimodelkan dengan Spatial Autoregressive Moving Average (SARMA).

```
69 ##SARMA/GSM
70 gsm <- saccsar1m(PerserMiskin2020~X1+X2+X3+X4+X5+X6+X7+X8+X9, data=jatengmiskin, ww_jateng)
71 summary(gsm)
```

Hasil analisis dapat dilihat pada jendela Console

```
> summary(gsm)

Call: saccsar1m(formula = PerserMiskin2020 ~ X3 + X4, data = jatengmiskin, listw = ww_jateng)

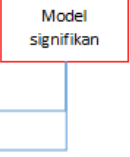
Residuals:
    Min       1Q   Median       3Q      Max
-4.3039 -1.9722 -0.5178  2.2465  4.0006

Type: sac
Coefficients: (asymptotic standard errors)
              Estimate Std. Error z value Pr(>|z|)
(Intercept)  8.670520   2.770771  3.1293 0.0017524
X3           0.202625   0.072618  2.7903 0.0052663
X4          -0.919778   0.248177 -3.7061 0.0002104

Rho: 0.44107
Asymptotic standard error: 0.22589
z-value: 1.9526, p-value: 0.050864
Lambda: -0.66897
Asymptotic standard error: 0.33673
z-value: -1.9867, p-value: 0.046958

LR test value: 3.2235, p-value: 0.19953

Log likelihood: -81.47039 for sac model
ML residual variance (sigma squared): 5.3138, (sigma: 2.3052)
Number of observations: 35
Number of parameters estimated: 6
AIC: 174.94, (AIC for 1m: 174.16)
```



Model terbaik adalah model yang memiliki nilai AIC terkecil. Sehingga jika dilihat dari nilai AIC, Model GSM adalah model terbaik dibandingkan Model SAR dan Model SEM karena mempunyai nilai AIC yang paling kecil. Model GSM memiliki signifikansi pada Rho dan Lambda sedangkan Rho pada Model SAR dan Lambda pada Model SEM tidak signifikan.

- **Uji Asumsi Klasik**

- Asumsi kenormalan galat

```
86 #Uji Asumsi Klasik
87 err.gsm<-residuals(gsm)
88
89 ##uji kenormalan galat
90 ad.test(err.gsm)
```

Hasil analisis dapat dilihat pada jendela Console, dimana diperoleh nilai p-value > 0.05 sehingga dilakukan penerimaan terhadap H0 yang berarti galat menyebar normal.

```
> ad.test(err.gsm)

Anderson-Darling normality test

data:  err.gsm
A = 0.61741, p-value = 0.09969
```

- Asumsi kehomogenan ragam

```
92 #uji kehomogenan ragam galat
93 bptest.Sarlm(gsm)
```

Hasil analisis dapat dilihat pada jendela Console, dimana diperoleh nilai p-value > 0.05 sehingga dilakukan penerimaan terhadap H₀ yang berarti ragam homogen.

```
> bptest.Sarlm(gsm)

studentized Breusch-Pagan test

data:
BP = 1.9304, df = 2, p-value = 0.3809
```

- Asumsi autokorelasi (dengan indeks moran)

```
95 #autokorelasi
96 moran.test(err.gsm, ww_jateng, randomisation=T, alternative="greater")
```

Hasil analisis dapat dilihat pada jendela Console, dimana diperoleh nilai p-value > 0.05 sehingga dilakukan penerimaan terhadap H₀ yang berarti tidak terdapat autokorelasi spasial.

```
Moran I test under randomisation

data:  err.gsm
weights: ww_jateng

Moran I statistic standard deviate = -0.0028313, p-value = 0.5011
alternative hypothesis: greater
sample estimates:
Moran I statistic      Expectation      Variance
-0.02974438          -0.02941176         0.01380170
```

Dari uji asumsi model SARMA/GSM terlihat bahwa dari uji kenormalan, kehomogenan ragam, dan autokorelasi, p-value > alpha (0.05) tidak terdapat cukup bukti tolak H₀. Artinya, galat menyebar normal, ragam galat homogen, dan tidak terdapat autokorelasi spasial pada model SARMA/GSM sehingga dapat dikatakan bahwa semua asumsi klasik terpenuhi.

F. APLIKASI ANALISA DATA SPASIAL DATA DEFORESTASI

Deforestasi didefinisikan sebagai hilangnya tutupan hutan menjadi non hutan. Pada analisis deforestasi ini akan digunakan data Global Forest Watch dari Hansen et al (2013). Dataset Hansen <https://data.globalforestwatch.org/documents/14228e6347c44f5691572169e9e107ad/explore> ini terdiri atas beberapa data yaitu: Tree canopy cover for year 2000 (treecover2000), Global forest cover gain 2000–2012 (gain), Year of gross forest cover loss event (lossyear), Data mask (datamask), Circa year 2000 Landsat 7 cloud-free image composite (first), Circa year 2019 Landsat cloud-free image composite (last).

Proses yang akan dilakukan dalam analisis ini adalah menggunakan package “gfcanalysis”. Tahapan yang dilakukan dalam analisis, yaitu : menentukan AOI (Area of Interest), mendownload data hansen sesuai dengan AOI, penentuan treshold, perhitungan luas area tutupan hutan, loss, gain.

- Penentuan AOI (Area of Interest)

Batas administrasi seringkali digunakan sebagai AOI. Cara memperoleh data ini sudah dijelaskan pada Bab sebelumnya. Pada studi kasus ini digunakan batas wilayah Provinsi Sumatera Selatan sebagai AOI.

```
1 ##Install dan aktifkan Packages yang dibutuhkan
2 if (!require(gfcanalysis)) install.packages('gfcanalysis')
3 if (!require(rgdal)) install.packages('rgdal')
4 if (!require(raster)) install.packages('raster')
5 if (!require(sp)) install.packages('sp')
6 if (!require(sf)) install.packages('sf')
7
8 library(gfcanalysis)
9 library(rgdal)
10 library(raster)
11 library(sp)
12 library(sf)
13
14 ##Batas administrasi tingkat kota/ kabupaten dari GADM
15 idn1 <- getData("GADM", country='IDN', level=1, download = TRUE) #level provinsi
16
17 ##plot data beserta judulnya
18 plot(idn1, main = "Provinsi di Indonesia")
19 view(idn1)
20
21 # memilih Provinsi Sumatera Selatan dan menyimpannya sebagai objek aoi data shp
22 aoi_sumsel <- idn1[idn1$NAME_1 == "Sumatera Selatan",]
23
24 plot(aoi_sumsel, main="Provinsi Sumatera Selatan")
```

- Download data

Hal yang perlu diperhatikan dalam download adalah menentukan folder dari hasil download akan disimpan serta mengetahui jumlah tile yang terdapat pada AOI sehingga dapat memperkirakan seberapa besar data yang akan didownload.


```

26 ##Download data
27 #Menentukan Output Folder
28 output_folder <- "D:/TULISANKU/BukuStatistik/BukuR/Data"
29
30 # menghitung jumlah tile yang dibutuhkan
31 tiles <- calc_gfc_tiles(aoi_sumsel)
32 print(length(tiles))
33
34 plot(tiles)
35 plot(aoi_sumsel, add=TRUE, lty=2, col="#00ff0050")
36 download_tiles(tiles, output_folder)

```

- Threshold

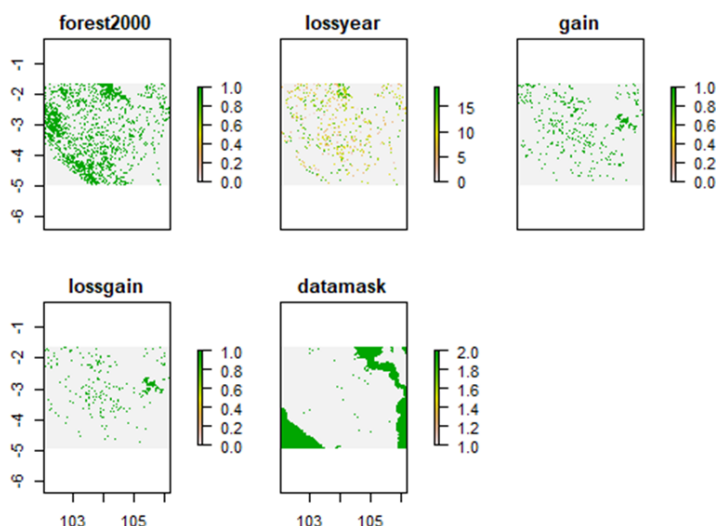
Hutan dalam analisis ini didefinisikan sebagai piksel dengan tutupan pohon lebih dari nilai threshold yang digunakan. Threshold 90 akan digunakan dalam analisis, jadi nilai piksel sama dengan atau lebih dari 90 (berarti memiliki 90% tutupan pohon) kita definisikan sebagai hutan.

```

38 ##Proses thresholding
39 #Penentuan threshold
40 threshold <- 90
41
42 #extract data
43 gfc_extract <- extract_gfc(aoi_sumsel, output_folder,
44                           filename="sumsel_GFC_extract.tif", overwrite = TRUE)
45
46 #Hasil extract di threshold
47 gfc_thresholded <- threshold_gfc(gfc_extract, forest_threshold = threshold,
48                                filename="sumsel_GFC_extract_thresholded.tif", overwrite = TRUE)
49
50 #Membaca hasil thresholding
51 plot(gfc_thresholded)

```

Hasil proses thresholding



- Perhitungan

Selanjutnya akan dilakukan perhitungan luas area berdasarkan layer yang diinginkan dan menyimpannya dalam format CSV pada folder yang telah ditentukan.

```
53 ##Menentukan luas
54 gfc_stats <- gfc_stats(aoi_sumsel, gfc_thresholded)
55 gfc_stats
56
57 ##Menyimpan hasil perhitungan ke format csv
58 write.csv(gfc_stats$loss_table, file='sumsel_lossstable.csv', row.names=FALSE)
59 write.csv(gfc_stats$gain_table, file='sumsel_gaintable.csv', row.names=FALSE)
```

Pada jendela Console akan ditampilkan nilai hasil perhitungan dalam satuan hektar, hasil perhitungan ini juga akan tersimpan pada folder output yang kita telah tentukan.

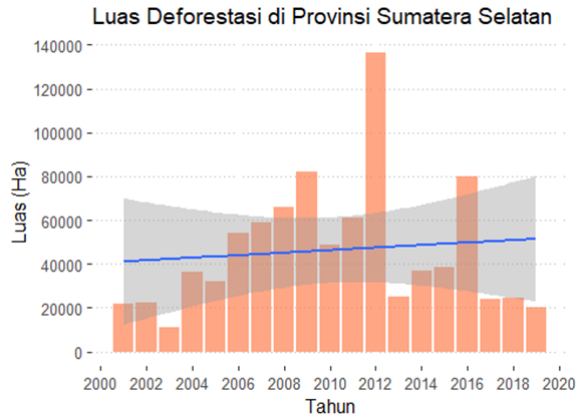
```
> gfc_stats
$loss_table
  year aoi cover loss
1 2000 AOI 1 2110992 NA
2 2001 AOI 1 2089157 21835.43
3 2002 AOI 1 2066665 22491.79
4 2003 AOI 1 2055542 11122.66
5 2004 AOI 1 2019005 36536.93
6 2005 AOI 1 1987046 31959.64
7 2006 AOI 1 1932633 54412.78
8 2007 AOI 1 1873715 58917.92
9 2008 AOI 1 1807598 66116.52
10 2009 AOI 1 1725358 82240.40
11 2010 AOI 1 1676473 48884.83
12 2011 AOI 1 1615206 61267.55
13 2012 AOI 1 1478853 136352.53
14 2013 AOI 1 1453914 24938.59
15 2014 AOI 1 1416901 37013.01
16 2015 AOI 1 1378614 38287.33
17 2016 AOI 1 1298702 79911.69
18 2017 AOI 1 1274757 23945.08
19 2018 AOI 1 1250283 24474.04
20 2019 AOI 1 1230265 20017.96

$gain_table
  period aoi gain lossgain
1 2000-2019 AOI 1 674585.7 493774.2
```

- Menyajikan dalam Grafik

```
61 ##Menyajikan hasil dalam grafik
62 library(ggplot2)
63 library(ggpubr)
64 ggplot(data = gfc_stats$loss_table, mapping = aes(x = year, y = loss)) +
65   geom_col(fill = "coral", alpha = 0.7) +
66   geom_smooth(method=lm)+
67   theme_pubclean()+
68   labs(title = "Luas Deforestasi di Provinsi Sumatera Selatan",
69        x = "Tahun",
70        y = "Luas (Ha)") +
71   scale_x_continuous(breaks = scales::breaks_width(2)) +
72   scale_y_continuous(breaks = scales::breaks_width(20000))
```

Hasil visualisasi grafik luas dan trend deforestasi di Provinsi Sumatera Selatan



- Melakukan analisis Statistik Deskriptif

```

74 ##Menyajikan Descriptive Statistic
75 if (!require(pastecs)) install.packages('pastecs')
76 library(pastecs)
77 stat.desc(gfc_stats$loss_table$loss, norm = TRUE)

```

Hasil statistik deskriptif dapat dilihat pada jendela Console

```

> stat.desc(gfc_stats$loss_table$loss, norm = TRUE)

```

	nbr.val	nbr.null	nbr.na	min	max	range	sum
1.	900000e+01	0.000000e+00	1.000000e+00	1.112266e+04	1.363525e+05	1.252299e+05	8.807267e+05
median		mean	SE.mean	CI.mean.0.95	var	std.dev	coef.var
3.	701301e+04	4.635404e+04	6.915672e+03	1.452929e+04	9.087038e+08	3.014471e+04	6.503148e-01
skewness		skew.2SE	kurtosis	kurt.2SE	normtest.w	normtest.p	
1.	343783e+00	1.282807e+00	1.603612e+00	7.905253e-01	8.585655e-01	9.276331e-03	

Berdasarkan hasil analisis statistik deskriptif tersebut diatas, diketahui bahwa rata-rata tingkat deforestasi di Provinsi Sumatera Selatan pada periode tahun 2000 – 2019 adalah 46354.04 ± 14529.29 Ha, pada tingkat keyakinan 95% dengan nilai SE of mean sebesar 6915.67 Ha.

DAFTAR PUSTAKA

- [BPS] Badan Pusat Statistik. 2021. Persentase penduduk miskin menurut kabupaten-kota. Retrieved from <https://www.bps.go.id/indicator/23/621/1/persentase-penduduk-miskin-menurut-kabupaten-kota.html>
- Bekti RD. Autokorelasi Spasial Untuk Identifikasi Pola Hubungan Kemiskinan Di Jawa Timur. ComTech Vol.3 No. 1 Juni 2012: 217-227
- Djuraidah A dan Wigena AH. Regresi Spasial untuk Menentukan Faktor-faktor Kemiskinan di Provinsi Jawa Timur. Statistika, Vol. 12 No. 1, 1 – 8 Mei 2012
- Hansen, M. C., P. V. Potapov, R. Moore, M. Hancher, S. A. Turubanova, A. Tyukavina, D. Thau, S. V. Stehman, S. J. Goetz, T. R. Loveland, A. Kommareddy, A. Egorov, L. Chini, C. O. Justice, and J. R. G. Townshend. 2013. High-Resolution Global Maps of 21st-Century Forest Cover Change. Science 342, (15 November): 850–853. Data available on-line from: <http://earthenginepartners.appspot.com/science-2013-global-forest>
- Hijmans RJ & Ghosh A. 2021. Spatial Data Analysis with R. Retrieved from <https://www.rspatial.org/raster/analysis/analysis.pdf>
- Islam S. 2018. Hands-On Geospatial Analysis with R and QGIS. Packt Publishing Ltd.
- Lee J & Wong D.W.S. 2001, Statistic for Spatial Data, John Wiley & Sons, Inc, New York
- Package 'gfcanalysis' Version 1.6.0 June 16, 2020. CRAN Repository.
- Package 'raster' Version 3.4-13 June 18, 2021. CRAN Repository.
- Package 'rgdal' Version 1.5-25 September 8, 2021. CRAN Repository.
- Package 'spatialreg' Version 1.1-8 May 3, 2021. CRAN Repository.
- Perobelli F & Haddad EA, 2003. "Brazilian Interregional Trade (1985-1996): an Exploratory Spatial Data Analysis," Anais do XXXI Encontro Nacional de Economia [Proceedings of the 31st Brazilian Economics Meeting]
- Spark C. 2013. Spatial Analysis in R: Part 1 Getting data from the ACS into R and Exploratory Spatial Data Analysis. Spatial Demography 2013 1 (1):131-139

ONE-WAY ANALYSIS OF VARIANCE (ANOVA) DENGAN SPSS

Kristia, M.B.A
Universitas Sanata Dharma

ANOVA merupakan analisis data yang digunakan ketika peneliti ingin menemukan perbedaan rata-rata pada tiga kelompok responden atau lebih, seperti pertanyaan-pertanyaan penelitian sebagai berikut:

- Apakah ketiga segmen pengguna social media (pengguna berat, medium, dan jarang mengakses social media) memiliki perbedaan yang signifikan dalam hal pengeluaran belanja online setiap bulannya?
- Apakah lima segmen demografis yang berbeda memiliki perbedaan sikap (attitude) yang signifikan terhadap Brand X?

Seperti pengujian parametrik lainnya, uji normalitas merupakan prasyarat sebelum uji ANOVA dapat dilakukan. Apabila data tidak berdistribusi normal dan sampel berjumlah kurang dari 30, maka uji non-parametrik seperti pengujian Kruskal-Wallis dapat dilakukan. Bila peneliti ingin membuat perbandingan rata-rata pada tiga atau lebih kelompok data, dengan mempertimbangan dua variabel independen, maka pengolahan data tersebut menggunakan two-way ANOVA.

One-way ANOVA merupakan bentuk ANOVA paling sederhana yang membandingkan 1 faktor independen variabel dengan tiga atau lebih dari tiga kelompok yang berbeda. Bentuk hipotesis penelitian pada pengujian one-way ANOVA adalah:

$$H_0 = \mu_1 = \mu_2 = \mu_3$$

Hipotesis nul tersebut menyatakan bahwa tidak ada perbedaan rata-rata yang signifikan pada ketiga kelompok yang diteliti.

H_1 : Setidaknya dua μ_1, μ_2, μ_3 memiliki perbedaan rata-rata yang signifikan

Hipotesis 1 menyatakan bahwa setidaknya rata-rata pada dua kelompok data yang diteliti mengalami perbedaan yang signifikan.

Berikut merupakan langkah pengolahan data ANOVA pada sebuah kasus:

Perusahaan yang bergerak di bidang kerajinan anyaman eceng gondok kuliner memiliki 4 cabang rumah produksi yaitu cabang A, B, C, dan D. Manajer

perusahaan tersebut ingin mengetahui apakah ada perbedaan rata rata produksi kerajinan yang signifikan pada keempat cabang rumah produksi yang ada. Perbedaan produksi kerajinan tersebut agak susah diobservasi secara langsung karena jumlah unit yang diproduksi pada masing masing cabang sering fluktuatif. Untuk mengetahui ada atau tidak adanya rata rata produksi antar masing masing rumah produksi, manajer perusahaan mendelegasikan supervisor pada tiap cabang untuk mencatat jumlah unit kerajinan yang berhasil diselesaikan para pekerja selama 35 hari kerja. Berikut merupakan tabel catatan jumlah unit anyaman eceng gondok yang mampu diproduksi pada masing masing cabang.

Hari	Cabang A	Cabang B	Cabang C	Cabang D
1	102	97	120	111
2	100	95	121	112
3	97	93	123	107
4	100	93	117	113
5	102	97	117	111
6	100	96	119	107
7	100	94	120	107
8	102	94	120	108
9	98	93	122	113
10	102	95	123	110
11	102	96	123	109
12	103	95	119	110
13	98	94	123	109
14	101	95	119	108
15	101	97	122	113
16	99	96	117	107
17	99	95	123	112
18	98	93	117	111
19	103	93	121	111
20	97	95	122	109
21	99	96	121	109
22	101	94	122	109
23	98	95	118	111
24	97	97	118	107
25	98	94	123	108
26	101	97	123	108
27	99	96	120	110
28	101	93	117	110
29	103	95	121	112
30	103	94	117	113
31	100	96	119	112
32	103	93	117	110
33	97	97	120	113
34	99	95	118	108
35	97	95	118	112

Keterangan:

Pada hari pertama para pekerja pada Cabang A mampu menyelesaikan 102 unit produk anyaman, Cabang B memproduksi 120 produk anyaman, Cabang C memproduksi 120 produk anyaman, dan Cabang D memproduksi 111 produk anyaman.

Penyelesaian:

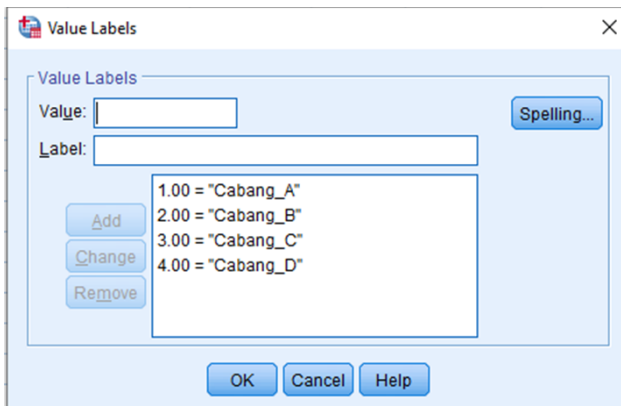
Pada kasus produksi kerajinan eceng gondok ini terdapat empat kelompok yang bebas atau tidak saling mempengaruhi satu dengan yang lain yaitu cabang rumah produksi yang berada pada lokasi yang berbeda dan juga memiliki pekerja yang berbeda beda pula. Pada kasus ini jumlah hari pengamatan hasil produksi adalah 35 hari kerja dan jumlah rumah produksi yang akan dibandingkan berjumlah lebih dari dua, maka pengujian yang digunakan untuk mengolah datanya adalah menggunakan uji One-Way ANOVA. Berikut merupakan langkah pengolahan data untuk menjawab pertanyaan yang diajukan oleh Manajer perusahaan.

1. Proses input data ke SPSS

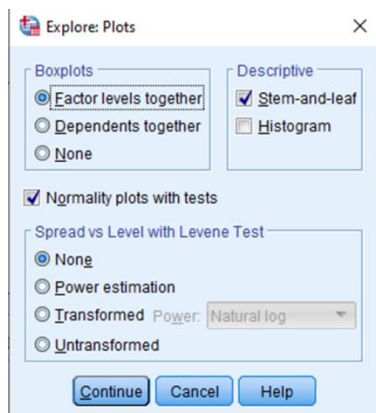
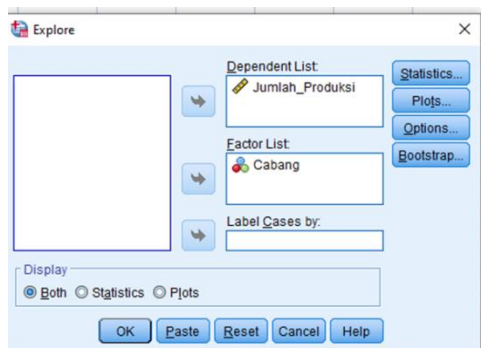
Ketika akan input data pada SPSS, tabel pada kasus sebelumnya harus diubah dalam format sebagai berikut:

Cabang	Unit	Cabang	Unit	Cabang	Unit	Cabang	Unit
1	102	2	97	3	120	4	111
1	100	2	95	3	121	4	112
1	97	2	93	3	123	4	107
1	100	2	93	3	117	4	113
1	102	2	97	3	117	4	111
1	100	2	96	3	119	4	107
1	100	2	94	3	120	4	107
1	102	2	94	3	120	4	108
1	98	2	93	3	122	4	113
1	102	2	95	3	123	4	110
1	102	2	96	3	123	4	109
1	103	2	95	3	119	4	110
1	98	2	94	3	123	4	109
1	101	2	95	3	119	4	108
1	101	2	97	3	122	4	113
1	99	2	96	3	117	4	107
1	99	2	95	3	123	4	112
1	98	2	93	3	117	4	111
1	103	2	93	3	121	4	111
1	97	2	95	3	122	4	109
1	99	2	96	3	121	4	109
1	101	2	94	3	122	4	109
1	98	2	95	3	118	4	111
1	97	2	97	3	118	4	107
1	98	2	94	3	123	4	108
1	101	2	97	3	123	4	108
1	99	2	96	3	120	4	110
1	101	2	93	3	117	4	110
1	103	2	95	3	121	4	112
1	103	2	94	3	117	4	113
1	100	2	96	3	119	4	112
1	103	2	93	3	117	4	110
1	97	2	97	3	120	4	113
1	99	2	95	3	118	4	108
1	97	2	95	3	118	4	112

2. Sebelum memasukkan data ke SPSS, kita perlu menyiapkan tampilan data menjadi dua variabel, yaitu variabel cabang dan variabel unit kerajinan eceng gondok yang mampu diproduksi pada masing masing cabang per harinya. Jumlah observasi data adalah sebanyak 35 hari kerja pada masing masing cabang rumah produksi, hanya penulisannya ditampilkan memanjang secara vertical (ke bawah). Kode cabang restoran 1 merepresentasikan Cabang A, 2 untuk Cabang B, 3 Untuk Cabang C, dan 4 untuk Cabang D. Pada setiap pengolahan data One-Way ANOVA, data akan ditampilkan sebagai dua variabel seperti contoh yang telah ditampilkan. Bila data telah disiapkan sedemikian rupa, maka langkah berikutnya adalah memasukkan data ke SPSS.
3. Buka software SPSS, lalu pada menu utama klik **File > New > Data**, lalu lihat pada bagian bawah pilih sheet tab sebelah kanan yaitu **Variable View**.
 - Pada kolom Name, sesuai kasus ketik **Kepuasan_Pelanggan**.
 - Masih pada sheet tab Variable View, di bawah Kepuasan_Pelanggan, sesuai kasus ketikkan **Cabang_Restoran**. Pada bagian sebelah kanan, pada kolom values ketikkan Value: 1, Label: Cabang_A dst sehingga terlihat seperti pada Gambar sebagai berikut.



- Klik sheet tab sebelah kiri, **Data View** untuk mengisi kolom Kepuasan_Pelanggan, lakukan copy paste data kepuasan pelanggan dan cabang restoran dari Excel (total data 140).
- Lakukan uji Normalitas dengan klik **Analyze > Descriptive Statistics > Explore**. Pada **factor list** masukkan Cabang_Restoran dan pada dependent list masukkan Kepuasan_Pelanggan.



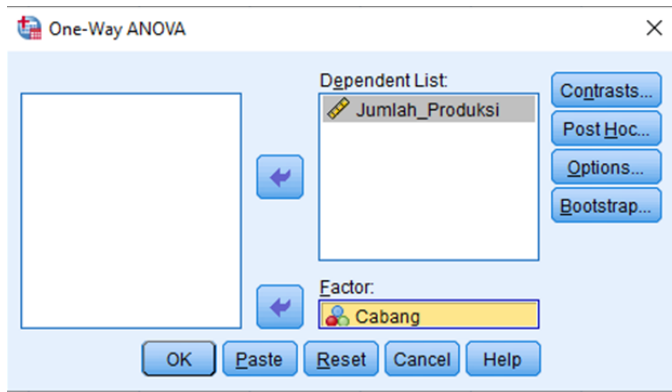
- Klik Plot dan centang Normality Plots with Tests > Continue > Ok. Lalu akan muncul dokumen output, lihat **Tabel Test of Normality**. Nilai Sig pada bagian Kolmogorov-Smirnov terlihat semua cabang bernilai > 0.05 yang berarti data berdistribusi normal, sehingga dapat melanjutkan ke proses Uji Homogenitas dan One Way ANOVA.

Tests of Normality

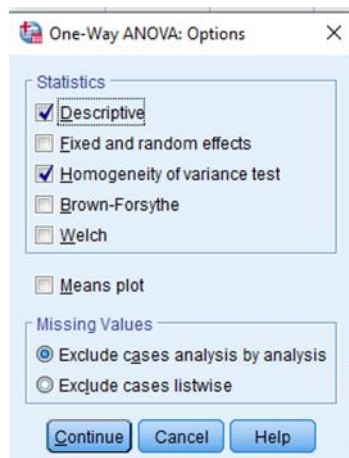
		Kolmogorov-Smirnov ^a			Shapiro-Wilk		
	Cabang	Statistic	df	Sig.	Statistic	df	Sig.
Jumlah_Produksi	Cabang_A	.124	35	.195	.920	35	.015
	Cabang_B	.145	35	.060	.900	35	.004
	Cabang_C	.132	35	.131	.892	35	.002
	Cabang_D	.124	35	.195	.920	35	.015

a. Lilliefors Significance Correction

- Langkah selanjutnya adalah klik menu Analyze > Compare-Means > One-Way ANOVA. Lalu pindahkan Kepuasan_Pelanggan pada Dependent List dan Cabang_Restoran pada Factor.



- Pada sebelah kanan klik Options, lalu centang **Descriptive** dan **Homogeneity of variance test** untuk menampilkan uji Homogenitas, lalu klik **Continue**.



- Pada bagian Post-Hoc centang Bonferroni dan Tukey untuk analisis lanjutan dan keseragaman, lalu tekan Continue dan OK, maka output siap untuk dianalisis.

4. Interpretasi output pengujian One-Way ANOVA

Pada tabel Descriptive, peneliti dapat melihat nilai ringkasan statistik dari keempat kelompok Cabang rumah produksi yang dibandingkan. Contoh analisis pada Cabang_A, rata-rata produksi kerajinan eceng gondok yang dapat diproduksi adalah 100 unit dengan minimum produksi 100 unit dan kemampuan produksi terbanyak (maksimum) pada angka 103 unit. Dengan tingkat kepercayaan 95% atau nilai signifikansi sebesar 5%, rata-rata produksi

kerajinan eceng gondok ada pada rentang 99 (pembulatan dari 99.3029) hingga 101 unit (pembulatan dari 100.6971). Dengan demikian analisis serupa dapat diterapkan pada ketiga cabang lainnya.

Descriptives								
Jumlah_Produksi								
	N	Mean	Std. Deviation	Std. Error	95% Confidence Interval for Mean		Minimum	Maximum
Cabang_A	35	100.0000	2.02920	.34300	Lower Bound	Upper Bound		
Cabang_B	35	94.9429	1.37076	.23170	99.3029	100.6971	97.00	103.00
Cabang_C	35	120.0000	2.20960	.37349	94.4720	95.4137	93.00	97.00
Cabang_D	35	110.0000	2.02920	.34300	119.2410	120.7590	117.00	123.00
Total	140	106.2357	9.84090	.83171	109.3029	110.6971	107.00	113.00
					104.5913	107.8801	93.00	123.00

Interpretasi uji Homogenitas

Hasil uji homogenitas dapat dilihat pada tabel berjudul Test of Homogeneity of Variances, seperti gambar sebagai berikut. Uji homogenitas ini dilakukan untuk menerima atau menolak Hipotesis 0 yang diteliti. Adapun bentuk hipotesis pada kasus ini adalah:

H_0 = Tidak terdapat perbedaan varians hasil produksi kerajinan eceng gondok yang signifikan pada keempat rumah produksi

H_1 = Terdapat perbedaan varians hasil produksi kerajinan eceng gondok yang signifikan pada keempat rumah produksi.

Adapun dasar pengambilan keputusan terkait hipotesis adalah sebagai berikut:

Jika nilai sig > 0.05, maka H_0 diterima, H_1 ditolak

Jika nilai sig < 0.05, maka H_0 ditolak, H_1 diterima

Test of Homogeneity of Variances					
		Levene Statistic	df1	df2	Sig.
Jumlah_Produksi	Based on Mean	4.128	3	136	.008
	Based on Median	4.262	3	136	.007
	Based on Median and with adjusted df	4.262	3	130.516	.007
	Based on trimmed mean	4.113	3	136	.008

Dapat kita lihat pada tabel bahwa nilai Sig. bernilai 0.08 yaitu lebih besar daripada 0.05, sehingga dapat disimpulkan bahwa **variens** hasil produksi kerajinan eceng gondok pada keempat rumah produksi tidak terdapat perbedaan yang signifikan. Dengan demikian persyaratan uji Homogenitas pada kasus ini telah dapat terpenuhi.

Interpretasi uji ANOVA

Setelah persyaratan uji Homogenitas terpenuhi, langkah selanjutnya adalah melakukan interpretasi terhadap hasil One-Way ANOVA, apakah hipotesis 0 pada penelitian ini dapat diterima atau ditolak. Berikut merupakan hipotesis pada penelitian ini:

H_0 = Tidak terdapat perbedaan **rata-rata** hasil produksi kerajinan eceng gondok yang signifikan pada keempat rumah produksi

H_1 = Terdapat perbedaan **rata-rata** hasil produksi kerajinan eceng gondok yang signifikan pada keempat rumah produksi

Adapun dasar pengambilan keputusan terkait hipotesis adalah sebagai berikut:

Jika nilai $\text{sig} > 0.05$, maka H_0 diterima, H_1 ditolak

Jika nilai $\text{sig} < 0.05$, maka H_0 ditolak, H_1 diterima

ANOVA

Jumlah_Produksi

	Sum of Squares	df	Mean Square	F	Sig.
Between Groups	12951.336	3	4317.112	1151.488	.000
Within Groups	509.886	136	3.749		
Total	13461.221	139			

Nilai Sig. dari tabel ANOVA menunjukkan nilai 0.000, yang nilainya lebih kecil dibandingkan 0.05. Sehingga dapat disimpulkan bahwa H_0 ditolak dan H_1 diterima, yaitu terdapat perbedaan rata-rata hasil produksi kerajinan eceng gondok yang signifikan pada keempat rumah produksi.

Interpretasi Post Hoc Test

Interpretasi Post Hoc Test ini digunakan untuk mengetahui kelompok observasi mana saja yang berbeda dan mana yang tidak memiliki perbedaan signifikan, terutama apabila hasil interpretasi pada Tabel ANOVA menunjukkan terdapat perbedaan rata-rata yang signifikan. Pada kasus ini telah ditunjukkan dari hasil ANOVA bahwa terdapat perbedaan rata-rata yang signifikan terhadap hasil produksi kerajinan pada keempat rumah produksi, lalu untuk melihat cabang cabang mana saja yang memiliki rata-rata produksi yang berbeda maka peneliti dapat meninjau tabel output Post Hoc Test, yang terdiri dari tabel Tukey Test dan Boferroni Test.

Multiple Comparisons

Dependent Variable: Jumlah_Produksi

	(I) Cabang	(J) Cabang	Mean Difference (I-J)	Std. Error	Sig.	95% ... Lower Bound
Tukey HSD	Cabang_A	Cabang_B	5.05714 [*]	.46286	.000	3.8532
		Cabang_C	-20.00000 [*]	.46286	.000	-21.2039
		Cabang_D	-10.00000 [*]	.46286	.000	-11.2039
	Cabang_B	Cabang_A	-5.05714 [*]	.46286	.000	-6.2611
		Cabang_C	-25.05714 [*]	.46286	.000	-26.2611
		Cabang_D	-15.05714 [*]	.46286	.000	-16.2611
	Cabang_C	Cabang_A	20.00000 [*]	.46286	.000	18.7961
		Cabang_B	25.05714 [*]	.46286	.000	23.8532
		Cabang_D	10.00000 [*]	.46286	.000	8.7961
	Cabang_D	Cabang_A	10.00000 [*]	.46286	.000	8.7961
		Cabang_B	15.05714 [*]	.46286	.000	13.8532
		Cabang_C	-10.00000 [*]	.46286	.000	-11.2039
Bonferroni	Cabang_A	Cabang_B	5.05714 [*]	.46286	.000	3.8179
		Cabang_C	-20.00000 [*]	.46286	.000	-21.2393
		Cabang_D	-10.00000 [*]	.46286	.000	-11.2393
	Cabang_B	Cabang_A	-5.05714 [*]	.46286	.000	-6.2964
		Cabang_C	-25.05714 [*]	.46286	.000	-26.2964
		Cabang_D	-15.05714 [*]	.46286	.000	-16.2964
	Cabang_C	Cabang_A	20.00000 [*]	.46286	.000	18.7607
		Cabang_B	25.05714 [*]	.46286	.000	23.8179
		Cabang_D	10.00000 [*]	.46286	.000	8.7607
	Cabang_D	Cabang_A	10.00000 [*]	.46286	.000	8.7607
		Cabang_B	15.05714 [*]	.46286	.000	13.8179
		Cabang_C	-10.00000 [*]	.46286	.000	-11.2393

Pada baris pertama pada hasil Tukey HSD dapat terlihat perbedaan mean pada Cabang_A dengan Cabang_B. Pada kolom Mean Difference atau perbedaan rata rata hasil yang ditunjukkan adalah 5.05714, yang merupakan hasil dari pengurangan mean hasil produksi Cabang_A (100) dikurangi mean hasil produksi Cabang_B (94.9429). Nilai mean hasil produksi kerajinan masing masing cabang dapat dilihat pada Tabel Descriptive Statistics yang telah ditampilkan sebelumnya. Berikut merupakan dasar penerimaan atau penolakan hipotesis terkait perbedaan mean (rata-rata) antara Restoran_A dan Restoran_B.

Jika nilai probabilitas (sig) > 0.05, maka H_0 diterima, H_1 ditolak

Jika nilai probabilitas (sig) < 0.05, maka H_0 ditolak, H_1 diterima

Dapat terlihat bahwa nilai sig pada perbandingan cabang A dengan cabang B adalah kurang dari 0.05 maka kesimpulan yang dapat diambil adalah terdapat perbedaan rata rata unit hasil produksi kerajinan yang signifikan pada rumah produksi Cabang_A dengan Cabang_B. Cara mudah lain untuk mengetahui apakah ada perbedaan rata rata yang signifikan antar Cabang yang dibandingkan adalah dengan melihat ada atau tidak adanya tanda

bintang (*) pada kolom Mean Difference. Jika terdapat tanda bintang setelah nilai mean, maka dapat diinterpretasikan bahwa terdapat perbedaan rata rata yang signifikan antara dua kelompok yang dibandingkan. Pada kasus ini semua cabang memiliki perbedaan yang signifikan terhadap rata rata hasil produksi kerajinan di masing masing cabangnya, karena dapat kita lihat tanda bintangnya terdapat pada semua nilai mean.

Interpretasi Homogeneous Subset

Pada tabel Homogeneous subset peneliti dapat menggali informasi mengenai kelompok mana saja yang tidak memiliki perbedaan rata-rata yang signifikan.

Homogeneous Subsets

Jumlah_Produksi						
			Subset for alpha = 0.05			
	Cabang	N	1	2	3	4
Tukey HSD ^a	Cabang_B	35	94.9429			
	Cabang_A	35		100.0000		
	Cabang_D	35			110.0000	
	Cabang_C	35				120.0000
	Sig.		1.000	1.000	1.000	1.000

Means for groups in homogeneous subsets are displayed.

a. Uses Harmonic Mean Sample Size = 35.000.

Pada hasil olah data kasus ini dapat terlihat ada 4 kelompok subset yang berbeda dan tidak ada cabang cabang yang berada pada satu kelompok subset yang sama. Hal ini menunjukkan bahwa tidak ada cabang yang memiliki rata rata hasil produksi kerajinan eceng gondok yang serupa.

DAFTAR PUSTAKA

- Ghozali, Imam. (2001). *Aplikasi Analisis Multivariate Dengan Program SPSS*. Semarang: UNDIP.
- Priyatno, Duwi. (2008). *Mandiri Belajar SPSS untuk Analisis Data & Uji Statistik*. Yogyakarta: MediaKom.
- PS, Djarwanto. (1997). *Statistik Nonparametrik*. Yogyakarta: BPFE.
- Santoso, Singgih. (2002). *SPSS Statistik Multivariat*. Jakarta: PT Elek Media Komputindo.
- Sudjana. (1996). *Teknik Analisis Regresi dan Korelasi Bagi Para Peneliti*. Bandung: Tarsito.

REGRESI LOGISTIK DENGAN STATA

Dr. Dra. Ani Nuraini, M.M
Universitas Respati Indonesia

A. REGRESI LOGISTIK

Analisis regresi merupakan suatu metode yang sangat umum digunakan dalam berbagai penelitian bidang ekonomi, keuangan, sosial, psikologi yang pada intinya adalah untuk mengetahui gambaran pengaruh atau hubungan antara variabel terikat dengan satu atau lebih variabel bebas (prediktor) (Hair, Joseph F.; Babin, Barry J.; Anderson, Rolph E.; Black, 2018). Variabel bebas dapat dalam bentuk kontinu, rasio, interval atau bahkan dalam bentuk kategorikal, begitu juga dengan variabel terikatnya (Ekananda, 2019). Regresi logistik merupakan salah satu metode analisis yang variabel terikatnya berbentuk kategorikal atau diskret atau non-metrik, sedangkan pada variabel bebasnya dapat menggunakan dalam bentuk diskret (non-metrik) atau metrik seperti pada persamaan di bawah ini.

$$Y_1 = X_1 + X_2 + X_3 + \cdots + X_n$$

(binary nonmetric) (nonmetric and metric)

Pada regresi linier variabel yang digunakan adalah dalam bentuk kontinu baik untuk variabel terikat maupun variabel bebas (Hosmer et al., 2013), seperti persamaan di bawah ini yang telah banyak digunakan pada penelitian-penelitian umumnya.

$$Y_i = \beta_{0i} + \beta_1 X_i + \beta_2 X_i + \beta_3 X_i + \cdots \dots \dots + \epsilon_i$$

Sebaiknya perlu dipahami terlebih dahulu mengenai model regresi logistik agar sesuai dengan tujuan analisis yang akan dilakukan, serta perbedaannya dengan model regresi linier berganda. Perbedaan model regresi logistik dengan model regresi linier adalah variabel hasil dalam regresi logistik bersifat biner atau dikotomis. Perbedaan antara regresi logistik dan regresi linier, analisis diskriminan ini tercermin baik dalam bentuk model maupun asumsinya.

Regresi logistik memiliki keuntungan karena tidak terlalu terpengaruh daripada analisis diskriminan ketika asumsi dasar yang mendasari inferensi statistik, seperti pada uji asumsi klasik normalitas variabel dan heteroskedastisitas yang hasilnya bergantung pada biner tidak terpenuhi. Selanjutnya regresi logistik juga dapat mengakomodasi variabel independen nonmetrik melalui pengkodean variabel dummy, seperti halnya regresi. Regresi logistik mempunyai kemampuan regresi terbatas, karena hanya untuk prediksi yang tergantung pada ukuran dua kelompok. Meskipun dapat diperluas ke ukuran dependen multi-kategori, situasi di mana tiga atau lebih kelompok membentuk ukuran dependen, analisis diskriminan lebih cocok dalam banyak kasus dalam situasi multi-kelompok ini.

Meskipun ada perbedaan yang sering kali mendukung regresi logistik daripada analisis diskriminan, banyak kesamaan juga ada di antara kedua metode tersebut. Ketika asumsi dasar dari kedua metode terpenuhi, mereka masing-masing memberikan perbandingan. Dalam regresi logistik terdapat dua model yaitu binary dan multinomial, binary ketika memprediksi dua kategorikal sedangkan multinomial apabila memprediksi lebih dari dua kategorikal.

Ciri-ciri yang menunjukkan bahwa dalam analisis regresi ketika variabel hasil dikotomis atau binary adalah sebagai berikut :

1. Model untuk mean kondisional dari persamaan regresi logistik binary harus dibatasi antara nol dan satu.
2. Distribusi binomial, bukan normal, menggambarkan distribusi kesalahan dan merupakan distribusi statistik yang menjadi dasar analisis.
3. Prinsip-prinsip pada analisis yang menggunakan regresi linier juga akan memberikan gambaran dalam regresi logistik.

B. CONTOH HASIL PENELITIAN

Hasil Penelitian dengan judul : Analisis Pengaruh Tata Kelola dan R&D Terhadap Probabilitas Default dengan CR sebagai Mediasi (Nuraini et al., 2020).

Rumusan hipotesis sebagai berikut :

H_{1a} : Audit Komite pada tata kelola berpengaruh terhadap Probabilitas Default

H_{1b} : Kepemilikan Manajerial pada tata kelola berpengaruh terhadap Probabilitas Default

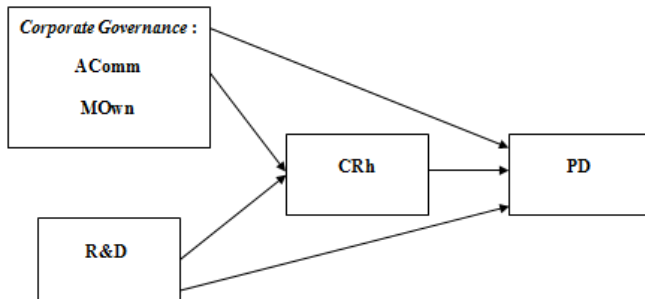
H_{2a} : Audit Komite pada tata kelola berpengaruh terhadap Probabilitas Default yang di mediasi oleh Current Ratio.

H_{2b} : Kepemilikan Manajerial pada tata kelola berpengaruh terhadap Probabilitas Default yang di mediasi oleh Current Ratio

H₃ : R&D berpengaruh terhadap Probabilitas Default

H₃ : R&D berpengaruh terhadap Probabilitas Default yang dimediasi oleh Current Ratio

Hipotesis tersebut di atas dapat di gambarkan dalam rerangka konsep di bawah ini :



Gambar 23 Rerangka Konsep

Penelitian ini menganalisis pengaruh antar variabel dengan menggunakan variabel mediasi, yang dilakukan melalui dua tahap regresi, yang pertama adalah regresi linier berganda dan yang kedua adalah regresi logistik. Data yang digunakan adalah pada perusahaan manufaktur di Indonesia yang memiliki data pengeluaran R&D pada periode 2009 -2018 dan terdaftar di Bursa Efek Indonesia yang terdiri dari 19 perusahaan, sehingga menghasilkan 190 data observasi.

C. REGRESI LINIER BERGANDA PANEL

Hasil estimasi pada regresi data panel akan menghasilkan intersep dan slope koefisien yang berbeda pada setiap perusahaan dan periode waktu (Widarjono, 2016). Struktur model yang dihasilkan terdiri common effect, fixed effect dan random effect, dengan menggunakan dua metode estimator yang terdiri dari Least square dummy variabel dan generalized least square (Ekananda, 2019), dengan model utama sebagai berikut:

$$Y_{it} = \beta_{0it} + \beta_1 X_{it} + \beta_2 X_{it} + \beta_3 X_{it} + \dots + \epsilon_{it} \quad (1)$$

Pada model *fixed effect* ditunjukkan dengan keberagaman individu yang ditunjukkan dengan intercept yang berbeda sehingga $\beta_{01t} \neq \beta_{02t} \neq \beta_{03t} \neq \dots \neq \beta_{0it}$ dan $\beta_{1t} = \beta_{2t} = \beta_{3t} = \dots = \beta_{3t}$, sedangkan pada model *random effect* adalah model yang memperhatikan adanya keberagaman dari variabel independen menurut individu dan memperhatikan dampak yang berbeda untuk setiap

individu yaitu $\beta_{01t} \neq \beta_{02t} \neq \beta_{03t} \neq \dots \neq \beta_{0i}$ dan $\beta_{1t} \neq \beta_{2t} \neq \beta_{3t} \neq \dots \neq \beta_{3t}$ (Ekananda, 2019).

D. REGRESI PANEL LOGIT

Model logistik merupakan model dimana variabel dependen mengalami transformasi menjadi probabilitas, yaitu $X\beta$ menjadi suatu nilai 0 sampai dengan 1. Bentuk persamaan logistik adalah sebagai berikut (Latan, 2014) :

$$\frac{P_i}{1-P_i} = e^{x_i\beta} \quad (2)$$

Sehingga :

$$\ln \left[\frac{P_i}{1-P_i} \right] = z = X\beta = b_1X_1 + b_2X_2 + \dots \quad (3)$$

Hasil estimasi regresi logistik menghasilkan model persamaan dengan koefisien regresi yang mempresentasikan prediksi logit, sedangkan *odds ratio* pada hasil regresi logistik merupakan ukuran *effect size* di dalam regresi logistik, yang artinya jika nilai di atas 1.0 menunjukkan pengaruh positif dan nilai di bawah 1.0 menunjukkan pengaruh negatif.

E. MODEL EMPIRIS

Penelitian ini menguji seberapa efektif tata kelola (*AComm & MOwn*) dan *R&D* dalam mempengaruhi probabilitas *default* (PD) dengan mediasi *CR*, sehingga dikembangkan model sebagai berikut :

Tahap I persamaan regresi data panel :

$$CR = \beta_0 + \beta_1 ACommR + \beta_2 MOwn + \beta_3 RD \dots + \epsilon_{it} \quad (4)$$

Tahap II persamaan regresi logistik :

$$\frac{P_i}{1-P_i} = PD = \beta_0 + \beta_1 CRh + \beta_2 ACommR + \beta_3 MOwn + \beta_4 RD \dots + \epsilon_{it} \quad (5)$$

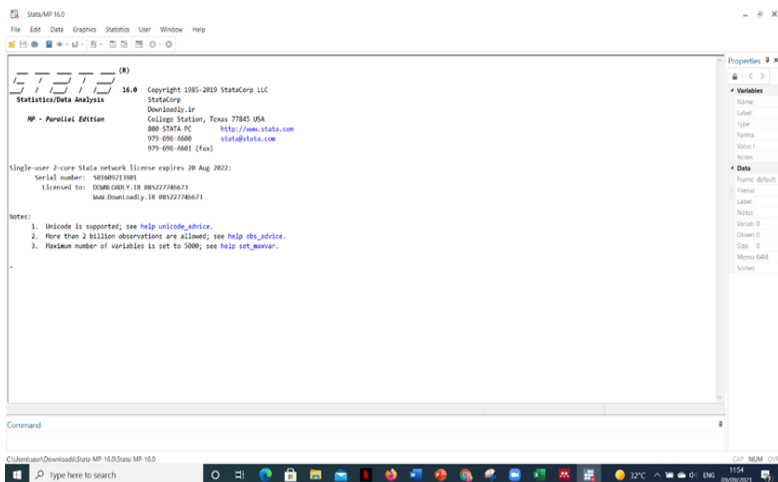
Persamaan (4) tersebut di atas adalah regresi data panel untuk mendapatkan hasil estimasi *CRh* dari variabel independennya yaitu komite audit, kepemilikan manajerial dan *R&D*, yang selanjutnya *CRh* digunakan sebagai mediasi pada persamaan (5), sehingga akan menghasilkan nilai PD pada masing-masing variabel independen baik pengaruh langsung maupun tidak langsung.

F. METODE PENGUKURAN DAN PREDIKSI

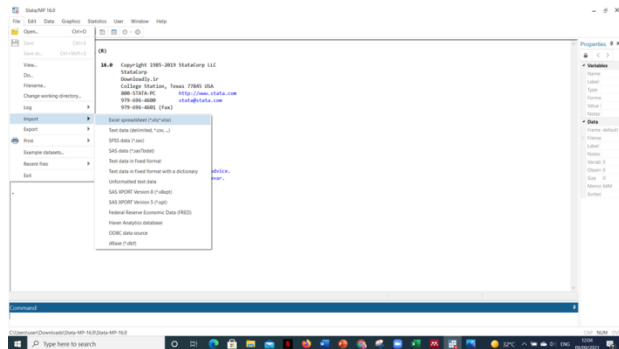
Pengukuran dan evaluasi prediksi model data panel menggunakan pemilihan model terbaik dengan menggunakan uji Chow, uji Hausman dan uji Lagrange Multiplier pada regresi tahap pertama, kemudian menentukan level of signifikan dengan nilai $\text{prob} < 0.05$. Model regresi logistik pada tahap kedua menggunakan nilai Pseudo R2 untuk mengukur fit model, yang berupa penjelasan dari varian variabel independen variabel dependen (Latan, 2014). Uji fit model yang lain yang digunakan Hosmer dan Lemeshow test yang melakukan pengelompokan berdasarkan pada nilai dari estimasi probabilitas. Uji ini secara proporsional direduksi di dalam nilai absolut dari log-likelihood yang diukur dan juga diukur dari berapa banyak fit yang buruk yang berpengaruh terhadap hasil estimasi dari prediktor variabel. Nilai R2L yang baik ditunjukkan dengan nilai $\text{Prob} > 0.05$ yang artinya model fit atau tidak terjadi misspesifikasi, dengan catatan sampel < 400 (Hosmer et al., 2013), uji yang sama bisa dilakukan dengan uji Pearson. Kurva ROC yang memberikan evaluasi pada batasan titik cut off, membandingkan antara sensitivitas versus 1-spesifikasi, secara ekuivalen ROC merupakan plot dari rata-rata positif yang benar (sensitivitas) dan dari rata-rata positif yang salah (1-spesifikasi) (Hosmer et al., 2013).

G. LANGKAH-LANGKAH REGRESI LOGISTIK (TAHAP II PENELITIAN)

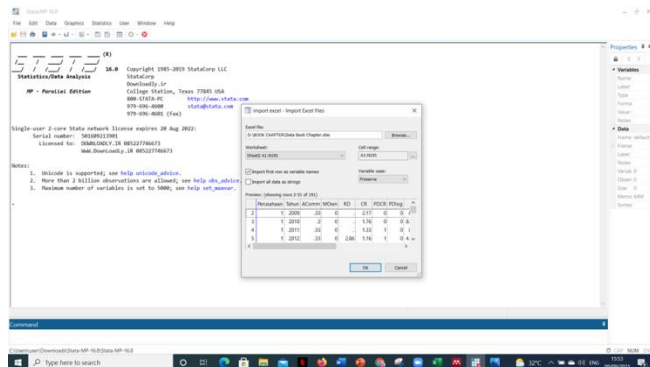
1. Membuka Program STATA



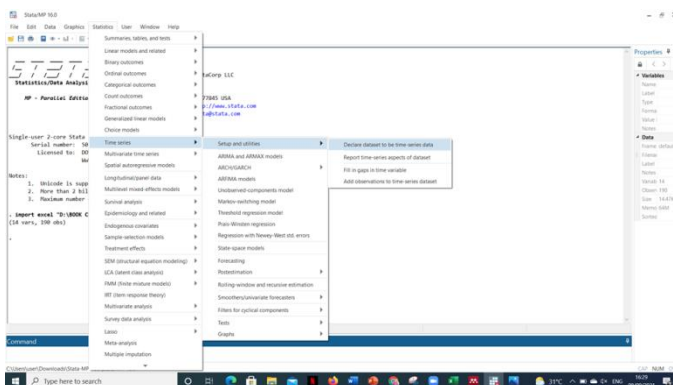
2. Import File dari Excel



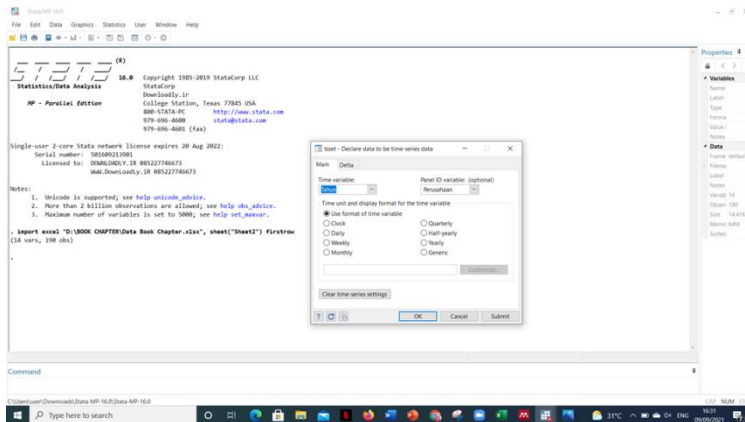
3. Memilih dan Menentukan File dan Sheet pada Excel



4. Menentukan Data Set (Data Time Series / Cross Section)



5. Pilih Time Variabel (Tahun) dan Panel ID Variabel (Perusahaan)



Klik OK, muncul :

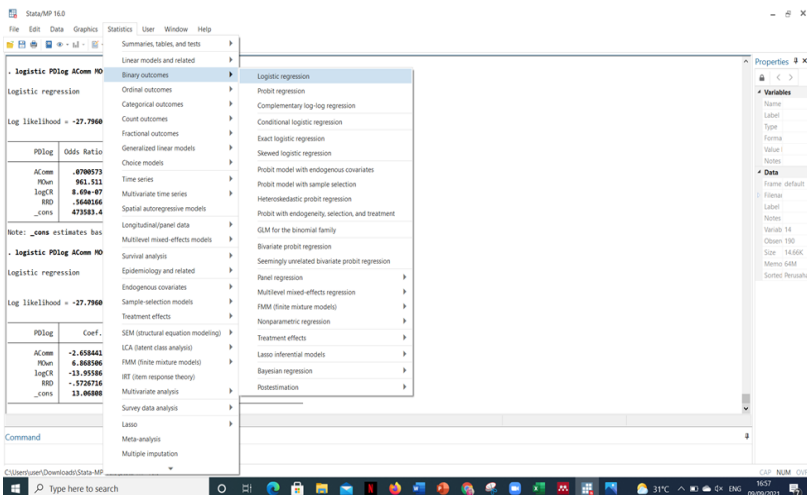
. tsset Perusahaan Tahun

panel variable: **Perusahaan (strongly balanced)**

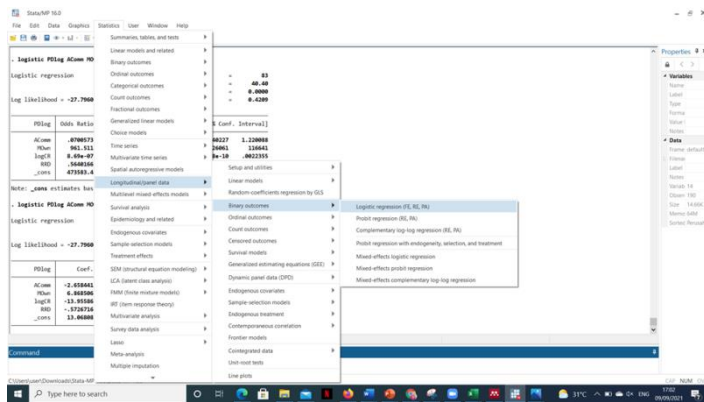
time variable: **Tahun, 2009 to 2018**

delta: **1 unit**

6. Regresi logistik dengan Common Effect

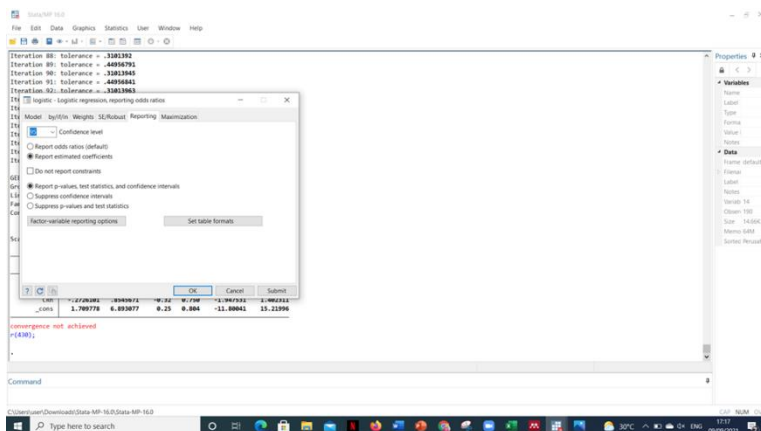


Regresi logistik dengan memilih model Common Effect, maka langsung klik Statistics → Binary outcomes → Logistic regression



Apabila memilih Random Effect atau Fixed Effect klik Statistics → Longitudinal/panel data → Binary outcomes → Logistic regression (FE/RE/PA)

Dalam contoh ini yang di jalankan adalah yang Common Effect, klik Statistics → Binary outcomes → Logistic regression → Reporting pilih Report estimated coefficients yang hasilnya dalam bentuk koefisien regresi sebagai berikut :



. logistic PDlog AComm MOwn CRh R&D, coef

```

Logistic regression                Number of obs =      83
                                LR chi2(4)  =    40.40
                                Prob > chi2  =    0.0000

```

```

Log likelihood = -27.796063        Pseudo R2    =    0.4209

```

PDlog	Coef.	Std. Err.	z	P> z	[95% Conf. Interval]	
-----+-----						
AComm	-2.658441	1.457866	-1.82	0.068	-5.515806	.1989231
MOwn	6.868506	2.448183	2.81	0.005	2.070156	11.66686
CRh	-13.95586	4.006487	-3.48	0.000	-21.80843	-6.103295
R&D	-.5726716	.1457592	-3.93	0.000	-.8583544	-.2869889
_cons	13.06808	3.827994	3.41	0.001	5.565354	20.57081

Hasil perhitungan estimasi sebagai hasil estimasi regresi logistik, hasil menunjukkan pada variabel tata kelola AComm tidak signifikan dengan prob. $\alpha > 0.05$ sedangkan pada MOwn berpengaruh positif signifikan dengan prob. $\alpha < 0.05$. Variabel R&D berpengaruh negatif signifikan terhadap probabilitas default yang ditunjukkan oleh nilai prob. $\alpha < 0.05$. Variabel CRh sebagai mediasi berpengaruh negatif signifikan dalam mempengaruhi probabilitas default yang ditunjukkan oleh nilai prob. $\alpha < 0.05$, yang artinya variabel AComm, MOwn dan R&D melalui CRh berpengaruh negatif signifikan dalam mempengaruhi probabilitas default.

Dalam contoh ini yang di jalankan adalah yang Common Effect, klik Statistics → Binary outcomes → Logistic regression → Reporting pilih Report odds ratios yang hasilnya sebagai berikut :

. logistic PDlog AComm MOwn CRh R&D

```

Logistic regression                Number of obs =      83
                                LR chi2(4)  =    40.40
                                Prob > chi2  =    0.0000

```

```

Log likelihood = -27.796063        Pseudo R2    =    0.4209

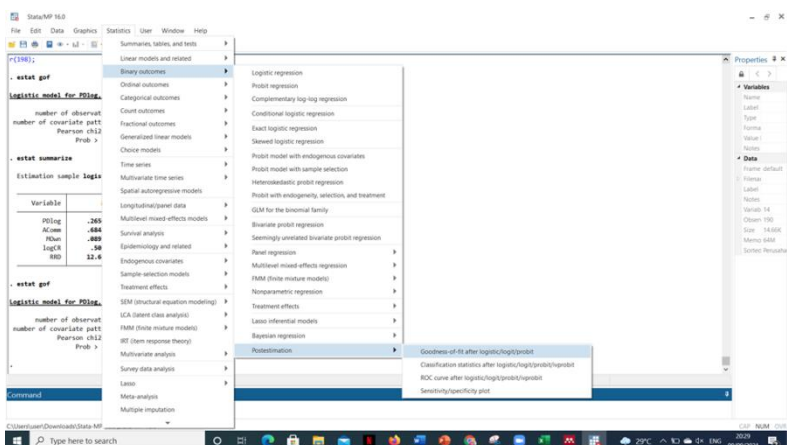
```

	PDlog	Odds Ratio	Std. Err.	z	P> z	[95% Conf. Interval]
AComm		.0700573	.1021342	-1.82	0.068	.0040227 1.220088
MOwn		961.511	2353.955	2.81	0.005	7.926061 116641
CRh		8.69e-07	3.48e-06	-3.48	0.000	3.38e-10 .0022355
R&D		.5640166	.0822106	-3.93	0.000	.423859 .7505201
_cons		473583.4	1812874	3.41	0.001	261.2176 8.59e+08

Note: _cons estimates baseline odds.

Odds rasio menunjukkan peluang default terhadap non-default pada variabel Acomm tidak signifikan, pada MOwn yang semakin naik berpengaruh signifikan terhadap meningkatnya peluang default terhadap non-default sebesar 961.511 kali, R&D yang semakin naik berpengaruh signifikan terhadap menurunnya peluang default terhadap non default sebesar 0.5640166 kali, CRh yang semakin meningkat berpengaruh signifikan terhadap menurunnya peluang default terhadap non-default sebesar 8.69e-07 kali yang artinya semua variabel yang terdiri dari AComm, MOwn dan R&D berpotensi menurunkan probabilitas default terhadap non-default dengan mediasi CRh yang meningkat.

Uji model pada regresi logistik goodness of fit ditunjukkan oleh :



Logistic model for PDlog, goodness-of-fit test

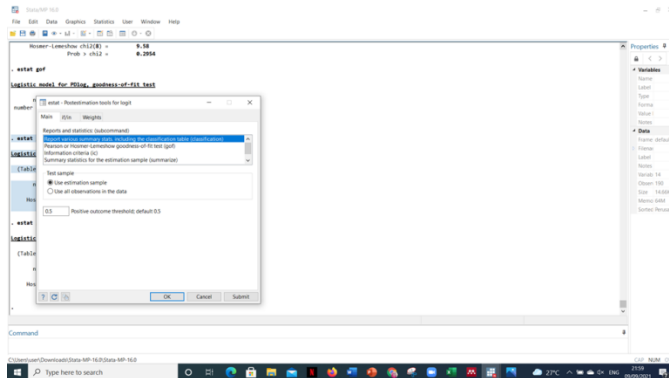
```
. estat gof, group(19)
```

number of observations	=	83
number of groups	=	19
Hosmer-Lemeshow chi2(17)	=	8.81
Prob > chi2	=	0.9460

Hasil uji Hosmer-Lemeshow dimana nilai R_L^2 yang baik ditunjukkan dengan nilai prob. > 0.05 , dimana nilai prob. $0.9460 > 0.05$ hasil ini menunjukkan bahwa model fit dan tidak terjadi misspesifikasi (Latan, 2014).

Untuk mengetahui tingkat akurasi model dalam membuat klasifikasi, maka dapat dilihat pada tabel di bawah ini sebagai hasil olah data :





. estat classification

Logistic model for PDlog

----- True -----			
Classified	D	~D	Total
-----+-----			
+	15	4	19
-	7	57	64
-----+-----			
Total	22	61	83

Classified + if predicted $\Pr(D) \geq .5$

True D defined as PDlog != 0

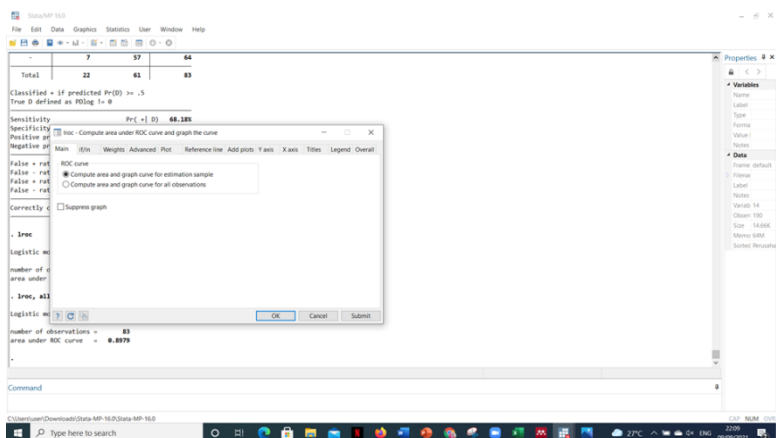
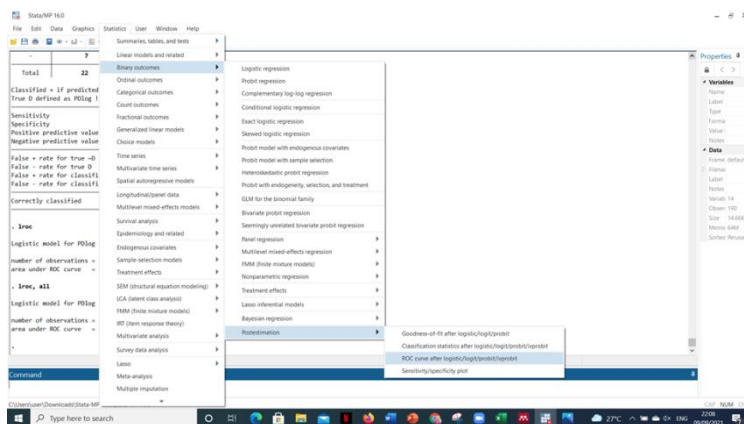
Sensitivity	$\Pr(+ D)$	68.18%
Specificity	$\Pr(- \sim D)$	93.44%
Positive predictive value	$\Pr(D +)$	78.95%
Negative predictive value	$\Pr(\sim D -)$	89.06%

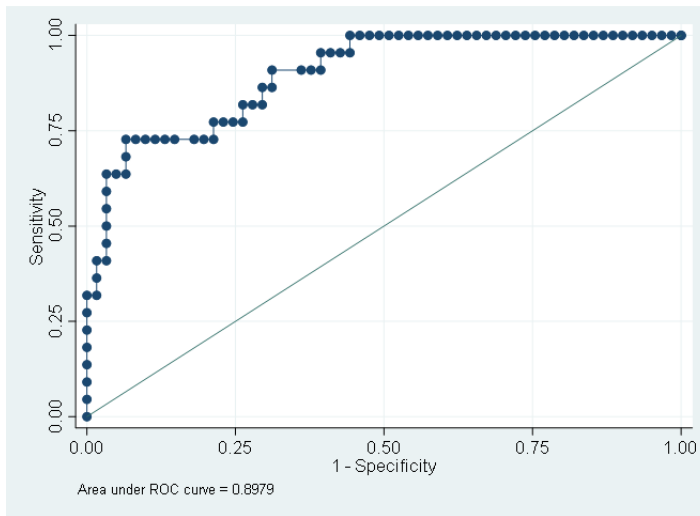
False + rate for true ~D	$\Pr(+ \sim D)$	6.56%
False - rate for true D	$\Pr(- D)$	31.82%
False + rate for classified +	$\Pr(\sim D +)$	21.05%
False - rate for classified -	$\Pr(D -)$	10.94%

Correctly classified	86.75%
----------------------	--------

Keterangan : D dimana $Y = 1$ dan $\sim D$ dimana $Y = 0$ untuk data observasi

Berdasarkan tabel klasifikasi model dapat dilihat bahwa ada 15 benar default, klasifikasi ini dilakukan dengan akurat pada model probabilitas tidak kurang dari 0.5 yang merupakan nilai cut-off, dan 57 benar tidak default, sehingga model yang diklasifikasi dengan benar berjumlah $15 + 57 = 72$ dari 83 atau sebesar 86.75%. Pada tabel 5. juga memberi informasi tentang probabilitas sensitivitas yang merupakan presentasi dari observasi dengan nilai probabilitas ≥ 0.5 yaitu 15 dari 22 atau sebesar 68.18%. Selanjutnya kurva ROC pada gambar 2. di bawah ini menjelaskan kemampuan model dalam mengukur diskriminasi likelihood antara subjek, dengan angka di bawah ROC = 0.8979, sesuai rule of thumb jika $0.8 \leq \text{ROC} \leq 0.9$ maka disimpulkan terdapat diskriminasi yang luar biasa (Hosmer et al., 2013).





Gambar 24 Kurva ROC

. Iroc

Logistic model for PDlo

number of observations = 83

area under ROC curve = 0.8979

DAFTAR PUSTAKA

- Ekananda, M. (2019). *Basic Econometrics* (2nd ed.). Mitra Wacana Media.
- Hair, Joseph F.; Babin, Barry J.; Anderson, Rolph E.; Black, W. C. (2018). *Multivariate Data Analysis* (Eight). CENGAGE.
- Hosmer, D. W., Lemeshow, S., & Sturdivant, R. X. (2013). Applied Logistic Regression: Third Edition. In *Applied Logistic Regression: Third Edition*. <https://doi.org/10.1002/9781118548387>
- Latan, H. (2014). *Aplikasi Analisis Data Statistik untuk Ilmu Sosial Sains dengan STATA* (Kesatu). ALFABETA. www.cvalfabeta.com
- Nuraini, A., Leon, F. M., & Usman, B. (2020). Analysis of the Effect of Governance and Research and Development on Probability of Default. *International Journal of Science and Society*, 2(4), 492–506. <https://doi.org/https://doi.org/10.200609/ij soc.v2i4.233>
- Widarjono, A. (2016). *Ekonometrika Pengantar dan Aplikasinya* (Keempat). UPP STIM YKPN.

STRUCTURAL EQUATIONAL MODELLING (SEM) DENGAN AMOS

M. Tirtana Siregar, S.TP., M.T
Politeknik App Jakarta

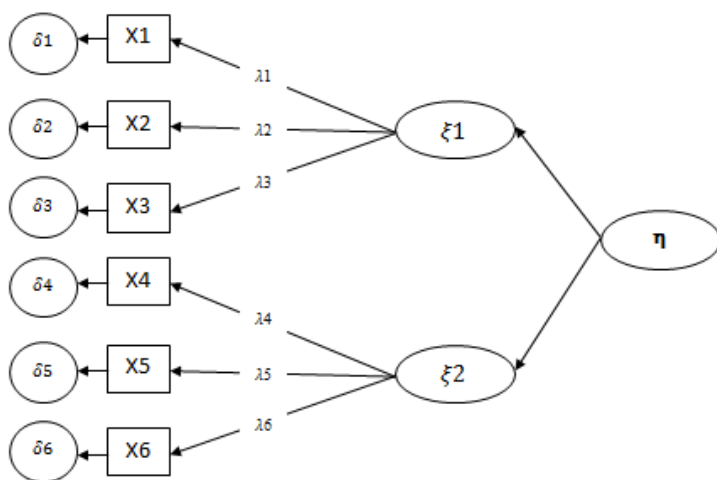
A. PENGERTIAN CFA (CONFIRMATORY FACTOR ANALYSIS)

CFA adalah salah satu metode analisis faktor yang digunakan ketika peneliti telah memiliki pengetahuan mengenai struktur suatu faktor laten. Struktur tersebut diperoleh berdasarkan kajian teoritis, hasil penelitian mengenai hubungan antara variabel yang diobservasi dengan variabel laten. CFA dibedakan menjadi First-Order CFA dan Second-Order CFA.

1. Second-Order Confirmatory Factor Analysis

Pada Second CFA untuk variabel laten tidak dapat diukur langsung melalui variabel-variabel indikatornya. Namun memiliki beberapa indikator dimana indikator tersebut tidak dapat diukur secara langsung, serta memerlukan beberapa indikator lagi.

Menurut Hair et al (2010), Second order merupakan CFA dari konstruk yang memiliki beberapa dimensi konstruk yang diukur oleh indikator. Analisis second order CFA merupakan model pengukuran yang terdiri dari dua tingkat. Tingkat pertama adalah sebuah Analisis Faktor Konfirmatori yang menunjukkan hubungan antara variabel-variabel observasi sebagai indikator-indikator dari variabel laten terkait. Tingkat kedua adalah sebuah Analisis Faktor Konfirmatori yang menunjukkan hubungan antara variabel-variabel laten pada tingkat pertama sebagai indikator-indikator dari sebuah variabel laten pada tingkat kedua. Variabel laten tingkat kedua disimbolkan dengan η sedangkan variabel tingkat pertama disimbolkan dengan lebih detailnya dapat dilihat pada ilustrasi berikut :



Gambar 25 second order CFA

Second order confirmatory factor analysis adalah bentuk model pengukuran dalam SEM yang terdiri dari 2 tingkat yang menunjukkan hubungan antara variabel – variabel laten pada tingkat pertama sebagai indikator-indikator dari sebuah variabel laten tingkat kedua. Notasi Matematik dari Hybrid Model secara umum (Struktural) sebagai berikut :

$$\eta = \beta\eta + \Gamma\xi + \zeta$$

Model Pengukuran sebagai berikut :

$$y = \Lambda y \eta + \varepsilon$$

$$x = \Lambda x \xi + \delta$$

Menurut Bollen (1989), hubungan antara first, second dan higher order factors dapat ditunjukkan melalui persamaan (1). Komponen $\Gamma\xi$ pada persamaan (1) tidak dibutuhkan jika higher order factors sebagai bagian dari η dengan koefisien masing-masing di β . Sebagai komponen $\beta\eta$ pada persamaan (1) dihapus jika hanya diperlukan second order factors dan tidak ada faktor tingkat pertama yang mempunyai efek langsung satu sama lain ($x=0$) . First order factors loading dari η pada y yaitu Λy .

B. STUDI KASUS SECOND ORDER

Seorang peneliti ingin mengukur tingkat kepuasan pelanggan terhadap produk tertentu dengan menggunakan indikator persepsi pelanggan terhadap kualitas produk, harga, dan kemudahan memperoleh produk. Dengan

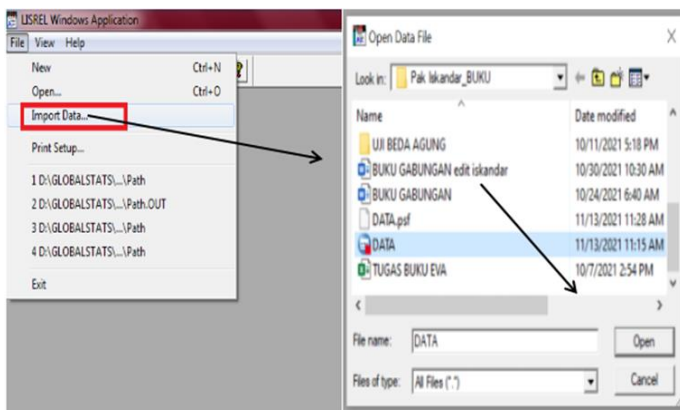
menggunakan CFA maka peneliti mampu melihat indikator apa saja yang berkontribusi besar dalam menggambarkan kepuasan pelanggan.

Variabel dalam Penelitian

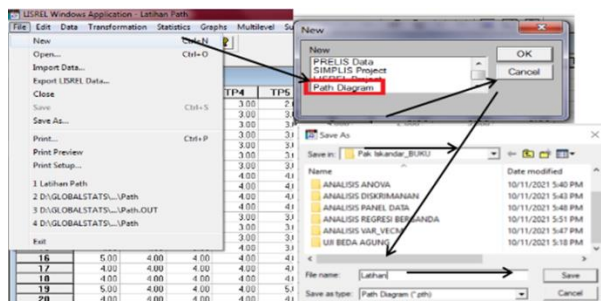
Variabel laten	Indikator
Kepuasan pelanggan	Kualitas (X1)
	Kebermanfaatan (X2)
	Kemudahan (X3)

a. CFA Second Order Menggunakan Lisrel

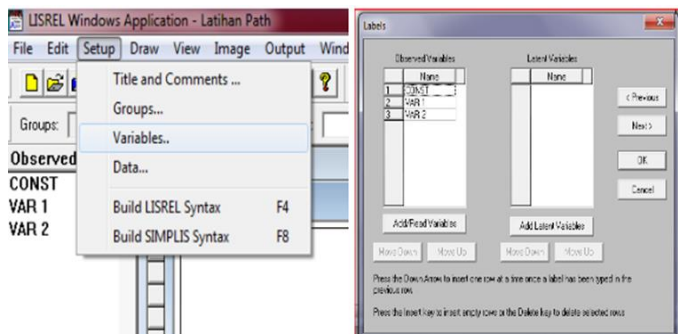
- Siapkan data dalam worksheet SPSS/ excel (.sav)
- Buka Lisrel
- Klik **File** lalu pilih **Import Data**, lalu Ubah “**File of Type**” menjadi “**SPSS Data File(*.sav)**” seperti pada gambar di bawah ini. kemudian carilah file dengan nama “**DATA**”, kemudian klik open



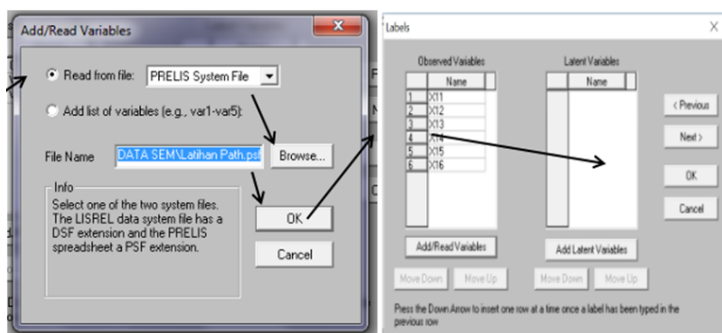
- Lalu klik menu **File**, kemudian pilih **File**, klik **New** lalu pilih **Path Diagram**, kemudian klik oke, setelah itu beri nama “**Latihan**” kemudian **Save As**.



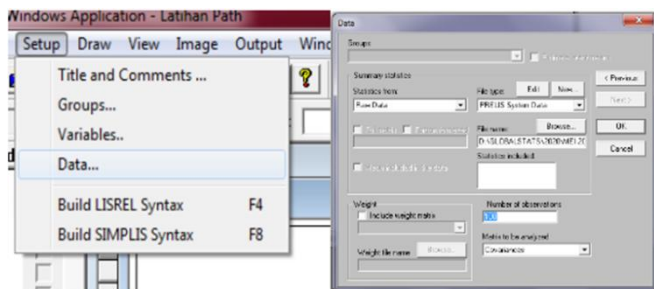
- e) Selanjutnya **Klik Setup**, pilih **Variables**, maka akan keluar kotak **Label**



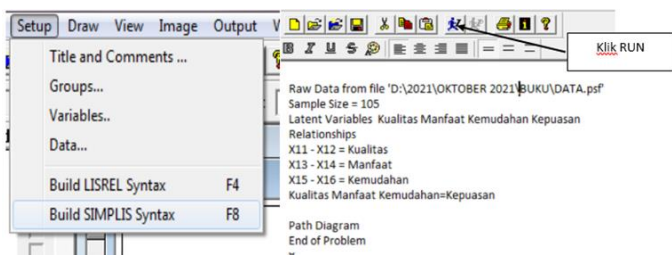
- f) Klik **Add/Read Variables**, pada **Read from file** pilih **PRELIS System File**, klik **Browse**, pilih data dengan ekstensi **.psf** yang telah disimpan sebelumnya, lalu klik OK. Kemudian klik **Add Latent Variables** untuk memberi nama laten variabel, kemudian setelah selesai maka klik OK



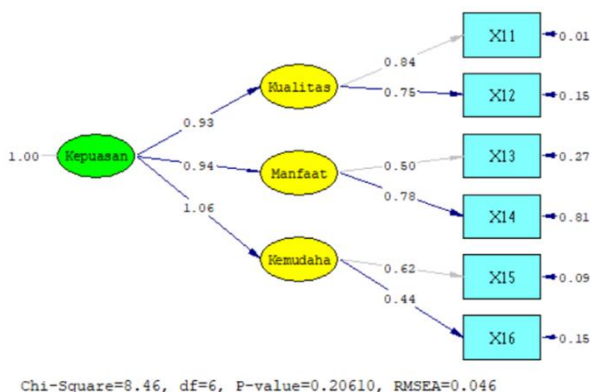
- g) Klik **Setup**, pilih **Data**, Kemudian masukkan Number of Observation dengan angka 105 (angka ini adalah jumlah observasi data) lalu klik OK.



- h) Klik Setup, pilih **Build SIMPLIS Syntax**, maka akan muncul kotak, kemudia lengkapi kotak SIntax seperti dibawah, lalu klik gambar Run



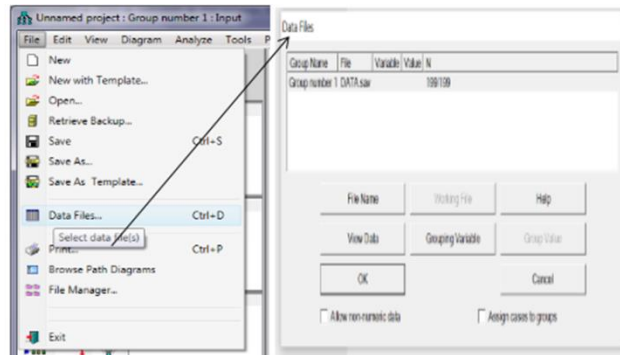
Dengan hasil Output sebagai berikut:



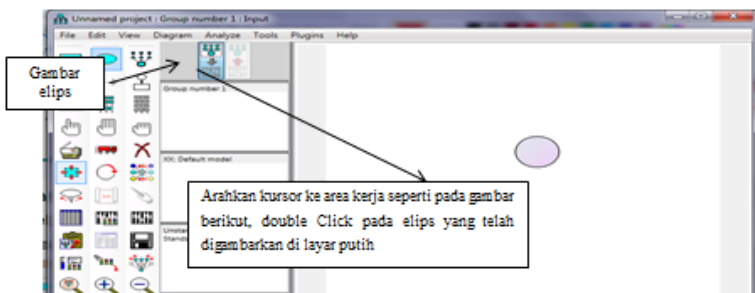
Interpretasi: Model pengukuran dihitung menggunakan teknik CFA (Confirmatory Factor Analysis), dimaksudkan untuk melihat keeratan hubungan (loading factor) dari indikator pembentuk dengan variabel latennya. Loading factor mensyaratkan di atas 0,5. Jika ada nilai yg kurang dari 0,5 maka indikator dihapus dan tidak digunakan dalam olah data selanjutnya. Nilai loading factor indikator X11 adalah 0.84, X120 adalah 0.75, X13 adalah 0.50, X14 adalah 0.78, X15 adalah 0.62 dan X16 adalah 0.44. Karena nilai loading factor untuk indikator X16 kurang dari 0,5 makan indikator X16 dibuang.

b. CFA Second Order Menggunakan Amos

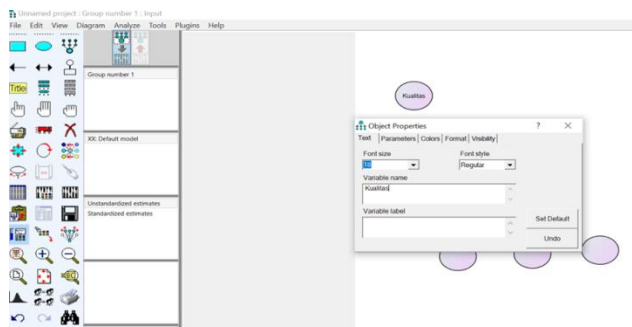
- Siapkan data dalam worksheet SPSS/ excel (.sav)
- Buka AMOS, masukan data dalam AMOS, klik **File-Data File**, lalu akan muncul kotak seperti berikut,



- c) Klik file name, pilih data yang telah disimpan dalam SPSS/excel (.csv), lalu klik open, setelah itu klik OK. Carilah file yang sudah disiapkan, kemudian klik **Open** lalu akan muncul tampilan seperti gambar berikut. Setelah itu klik **OK**
- d) Untuk membuat variabel laten, klik gambar elips



- e) setelah di double klik pada elips, mak akan muncul tampilan seperti berikut, beri nama variabel name/label



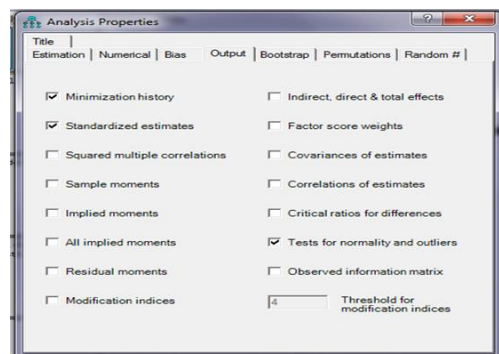
- f) Untuk membuat variabel indikator, klik gambar seperti dibawah



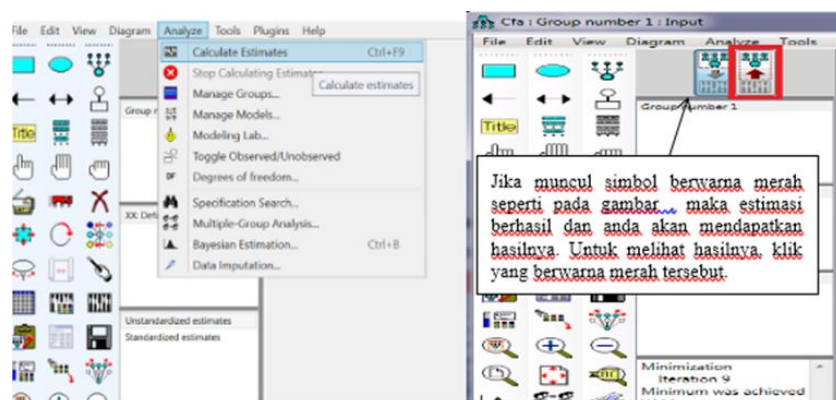
- g) Dengan Hasil gambar seperti berikut: (memberi nama indicator dan error bisa seperti point e). kemudian beri anak panah atau hubungan.



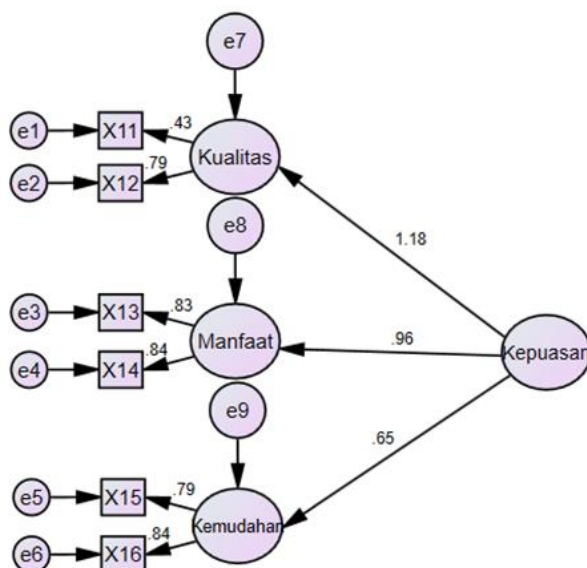
- h) Simpan gambar model, klik **File-Save As**. Pilih **View-Analys Properties**. Akan muncul kotak dialog sebagai berikut lalu pilih **Output** dan centang seperti pada gambar. Lalu close



- i) Untuk menjalankan program pilih **Analyze-calculate estimate**



- j) Setelah diklik gambar berwarna merah, maka akan muncul gambar dibawah ini, untuk melihat hasil estimate loading faktor, klik **View-Text Ouput**



Estimates (Group number 1 - Default model)

Standardized Regression Weights: (Group number 1 - Default model)

			Estimate
Kualitas	<---	Kepuasan	1.176
Manfaat	<---	Kepuasan	.964
Kemudahan	<---	Kepuasan	.649
X12	<---	Kualitas	.792
X11	<---	Kualitas	.435
X14	<---	Manfaat	.841
X13	<---	Manfaat	.828
X16	<---	Kemudahan	.843
X15	<---	Kemudahan	.792

DAFTAR PUSTAKA

- Bollen, K.A. (1989). *Structural Equation with Laten Variables*. New York: John Willey and Sons.
- Santoso, S. 2002. *Analisis SEM menggunakan AMOS*. Jakarta: Gramedia
- Hair, Joseph F, William C. B., Barry J.B., dan Rolph E.A. 1998. *Multivariate Data Analysis (Seventh Edition)*. New Jersey: Prentice Hall.
- Johnson, N. And Wichern, D. 2002. *Applied Multivariate Statistical Analysis, 5th Edition*. New Jersey: Prentice Hall, Englewood Cliffs.
- Effendi, R., Rahadhini, M. D., & Suddin, A. (2016). Analisis Pengaruh Word Of Mouth, Lokasi, Dan Kualitas Pelayanan Terhadap Keputusan Pembelian (Survei Konsumen Pada Warung Soto Seger Mbok Giyem Di Boyolali). *Jurnal ekonomi dan kewirausahaan* , Vol.16, No. 4.

STATISTIK MULTIVARIAT DALAM RISET

Sebuah buku yang merupakan hasil dari gabungan para peneliti pakar di Universitas Indonesia yang akan memberikan citra bagi kelangsungan kehidupan keilmuan karya ilmiah ditanah air dalam memfasilitasi para generasi muda dalam karya ilmiahnya maupun untuk memperkaya khazanah keilmuan para pelajar Universitas di Indonesia, buku ini dilengkapi cara cara praktis terkait dengan pengolahan data menggunakan SPSS, AMOS, Lisrel, R, Stata, Smart PLS, dan Eviews yang merupakan program komputer yang dipakai untuk analisis statistika.

Oleh sebab itu buku ini hadir dihadapan sidang pembaca sebagai bagian dari upaya diskusi sekaligus dalam rangka melengkapi khazanah keilmuan dibidang statistika, sehingga buku ini sangat cocok untuk dijadikan bahan acuan bagi kalangan intelektual dilingkungan perguruan tinggi ataupun praktisi yang berkecimpung langsung dibidang statistika.