

PENGUNAAN SPSS DALAM STATISTIK

AGUS TRI BASUKI

The SPSS logo is displayed in a large, bold, red font. The letters are stylized, with the 'S' and 'P' being particularly prominent. A small registered trademark symbol (®) is located at the top right of the 'S'.

PENERBIT : DANISA MEDIA

PENGUNAAN SPSS DALAM STAISTIK

Katalog Dalam Terbitan (KDT)

Agus Tri Basuki.; **PENGUNAAN SPSS DALAM STATISTIK**

- Yogyakarta : 2014

105 hal.; 17,5 X 24,5 cm

Edisi Pertama, Cetakan Pertama, 2014

Edisi Revisi, 2015

Hak Cipta 2015 pada Penulis

© Hak Cipta Dilindungi oleh Undang-Undang

Dilarang memperbanyak atau memindahkan sebagian atau seluruh isi buku ini dalam bentuk apapun, secara elektronis maupun mekanis, termasuk memfotokopi, merekam, atau dengan teknik perekaman lainnya, tanpa izin tertulis dari penerbit

Penulis : Agus Tri Basuki

Desain Cover : Yusuf Arifin

ISBN :

Penerbit :

Danisa Media

Banyumeneng, V/15 Banyuraden, Gamping, Sleman

Telp. (0274) 7447007

Email : danisamedia_yk@yahoo.com

DAFTAR ISI

Kata Pengantar

Daftar Isi

Bab 1	Mengenal SPSS	3
Bab 2	Statistik Deskriptip	14
Bab 3	Uji t Satu Sampel	24
Bab 4	Uji t Sampel Berpasangan	29
Bab 5	Analisis Varians	33
Bab 6	Uji Validitas dan Realibilitas	65
Bab 7	Uji Normalitas dan Outlier	75
Bab 8	Analisis Regresi	83
Bab 9	Asumsi Klasik	94

Daftar Pustaka



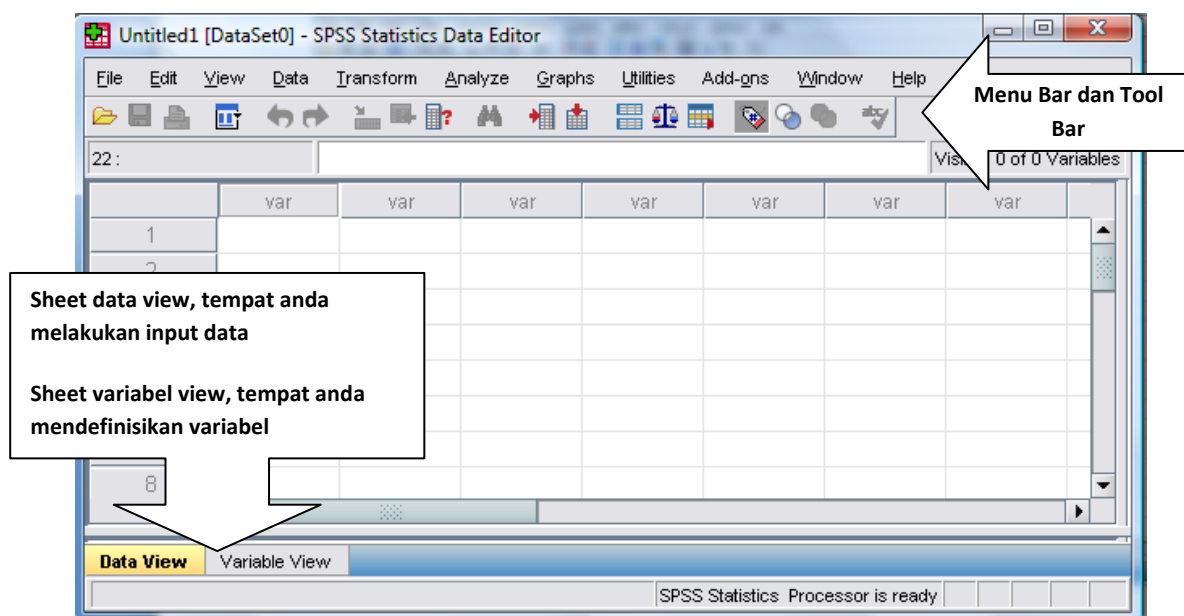
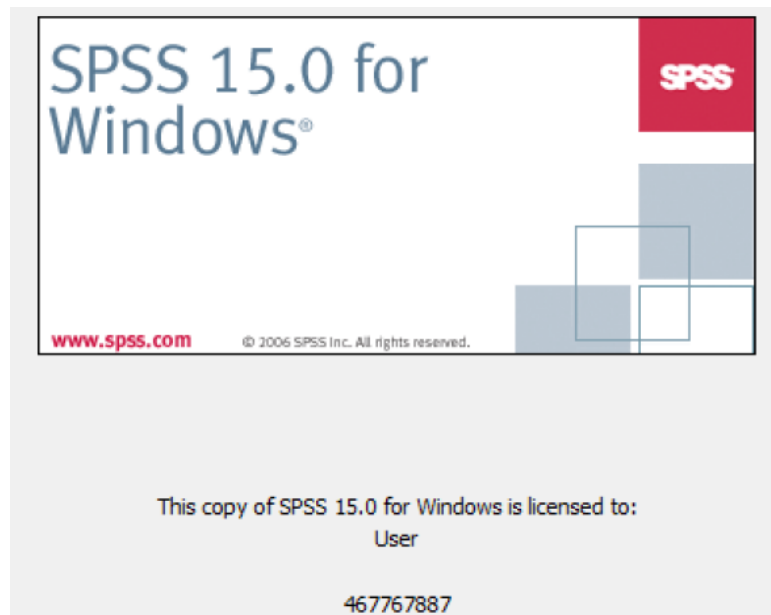
MENGENAL SPSS

MENGENAL SPSS (*Statistical Product and Service Solutions*)

SPSS adalah sebuah program aplikasi yang memiliki kemampuan analisis statistik cukup tinggi serta sistem manajemen data pada lingkungan grafis dengan menggunakan menu-menu deskriptif dan kotak-kotak dialog yang sederhana sehingga mudah untuk dipahami cara pengoperasiannya. SPSS banyak digunakan dalam berbagai riset pemasaran, pengendalian dan perbaikan mutu (*quality improvement*), serta riset-riset sains.

Pada awalnya SPSS dibuat untuk keperluan pengolahan data statistik untuk ilmu-ilmu sosial, sehingga kepanjangan SPSS itu sendiri adalah ***Statistical Package for the Social Sciences***. Sekarang kemampuan SPSS diperluas untuk melayani berbagai jenis pengguna (*user*), seperti untuk proses produksi di pabrik, riset ilmu sains dan lainnya. Dengan demikian, sekarang kepanjangan dari SPSS ***Statistical Product and Service Solutions***.

Langkah awal yang harus dilakukan untuk menganalisis data dengan menggunakan SPSS adalah melakukan **Input Data**. Pada saat kita membuka SPSS, maka akan muncul tampilan SPSS data Editor.



Setelah Data Editor aktif, berikut ini langkah-langkah untuk membuat data baru:

Latihan 1: Memasukkan Data ke dalam SPSS (1)

Berikut ini nilai PDRB tanpa migas tahun 2013 di beberapa kabupaten:

KABUPATEN	Pertanian	Industri	Jasa	Jumlah
Bodronoyo	5194485,32	5218350,93	5258136,50	15670972,75
Sulamanto	2762729,18	2997818,90	3178586,96	8939135,04
Ciganjur	2911111,03	2974994,19	3038805,78	8924911,01
Bandungan	4338609,02	4436742,73	4534050,73	13309402,49

KABUPATEN	Pertanian	Industri	Jasa	Jumlah
siGarut	3566959,43	3631799,79	3713444,29	10912203,51
Tasikmala	2886454,78	2920347,31	2970956,00	8777758,08
Ciamis	3236120,65	3305451,54	3416380,69	9957952,89
Kuning	2635535,23	2738240,12	2819521,59	8193296,93
Carubon	2580728,25	2671645,91	2742543,13	7994917,28
Malengka	2445604,05	2558835,15	2625150,66	7629589,86

Sumber : Data Hipotesis

Masukkan data di atas menggunakan SPSS. PROSES:

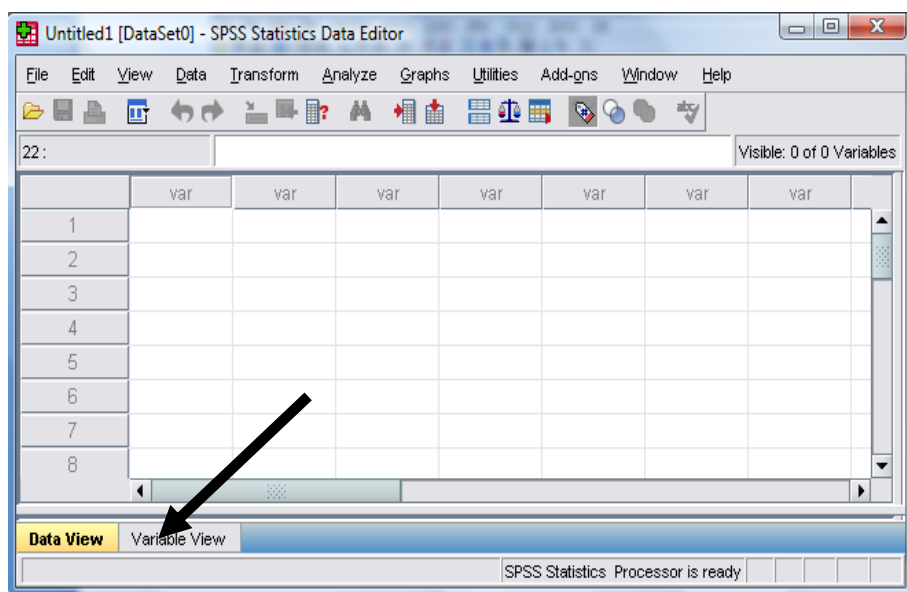
Input DATA BARU ke dalam SPSS mempunyai tahapan yang ber-urutan :

1. Membuka **VARIABLE VIEW** untuk mendefinisikan variabel
2. Membuka **DATA VIEW** untuk memasukkan data

Berikut proses lengkap pemasukan data.

A. Membuka VARIABLE VIEW untuk definisi variabel

Buka terlebih dahulu program SPSS, hingga tampak tampilan awal sebagai berikut.



Lihat pada bagian kiri bawah dari tampilan di atas. Lalu klik tab **VARIABLE VIEW** (lihat anak panah), hingga tampilan berubah menjadi:

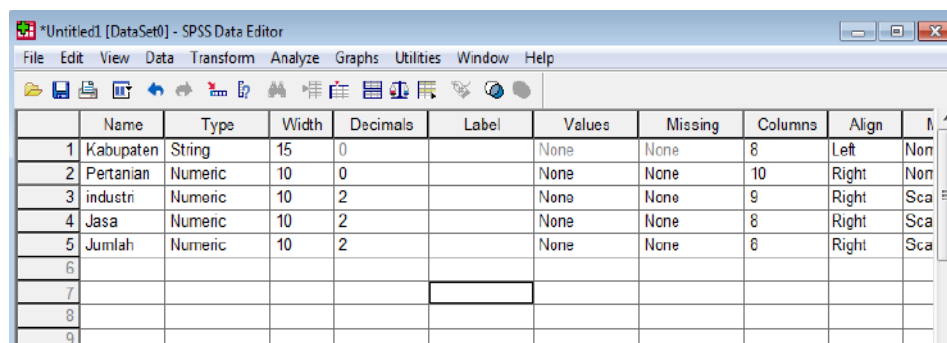
	Name	Type	Width	Decimals	Label	Values
1						
2						
3						
4						
5						

Perhatikan baris teratas yang berisi nama-nama untuk pendefinisian variabel, yang meliputi NAME, TYPE, dan sete-rusnya.

Soal di atas terdiri atas lima variabel, yakni Kabupaten, Pertanian, Industri, Jasa dan jumlah. Untuk itu, akan dilakukan input lima variabel di atas :

- Input variabel Kabupaten
 - Untuk NAME, isi Kabupaten.
 - Untuk TYPE, karena isi variabel NAMA adalah karakter, bukan angka, klik sekali pada sel TYPE di dekat isian Kabupaten (baris pertama dari TYPE), lalu klik sekali pada ikon  yang ada di bagian kanan. Tampak kotak dialog VARIABLE TYPE. Pilih STRING yang ada di paling bawah. Kemudian tekan OK untuk kembali ke VARIABLE VIEW.
 - Untuk WIDTH, isi angka **15**. Angka ini berarti maksimum dapat dimasukkan 15 karakter pada NAMA. Abaikan bagian yang lain, lalu masuk ke variabel kedua.
- Input variabel Pertanian
 - Untuk NAME, isi Pertanian.
 - Untuk TYPE, karena isi variabel Pertanian adalah angka, pilih **NUMERIC**. Abaikan bagian yang lain, lalu masuk ke variabel ketiga.
- Input variabel Industri
 - Untuk NAME, isi Industri.
 - Untuk TYPE, karena isi variabel Pertanian adalah angka, pilih **NUMERIC**. Abaikan bagian yang lain, lalu masuk ke variabel keempat.
- Input variabel Jasa
 - Untuk NAME, isi Jasa.
 - Untuk TYPE, karena isi variabel Pertanian adalah angka, pilih **NUMERIC**. Abaikan bagian yang lain, lalu masuk ke variabel kelima.
- Input variabel Jumlah
 - Untuk NAME, isi Jumlah.
 - Untuk TYPE, karena isi variabel Pertanian adalah angka, pilih **NUMERIC**.

Tampilan keseluruhan kelima variabel pada **VARIABLE VIEW**: (ditampilkan hanya sebagian)



	Name	Type	Width	Decimals	Label	Values	Missing	Columns	Align	M
1	Kabupaten	String	15	0		None	None	8	Left	Non
2	Pertanian	Numeric	10	0		None	None	10	Right	Non
3	industri	Numeric	10	2		None	None	9	Right	Sca
4	Jasa	Numeric	10	2		None	None	8	Right	Sca
5	Jumlah	Numeric	10	2		None	None	8	Right	Sca
6										
7										
8										
9										

Oleh karena pengisian lima variabel tersebut sudah selesai, klik ikon **DATA VIEW** yang ada di bagian kiri bawah untuk memasukkan data (tahap 2).

B. Memasukkan data ke dalam SPSS

Pada tampilan **DATA VIEW**, masukkan data seperti kita memasukkan data ke Excel atau sebuah tabel lainnya. Cara penempatan angka seperti tabel awal soal. Hanya perlu diperhatikan untuk tanda baca desimal, apakah menggunakan koma (,) ataukah titik (.). Pada contoh ini digunakan tanda baca 'koma', seperti angka 15,21.

Tampilan akhir input data:

12 :	Kabupaten	Pertanian	industri	Jasa	Jumlah	var	var	var	var
1	Bodronoyo	5194485	5218350.9	5258137	15670973				
2	Sulamanto	2762729	2997818.9	3178587	8939135				
3	Ciganjur	2911111	2974994.2	3038806	8924911				
4	Bandungan	4338609	4436742.7	4534051	13309402				
5	siGarut	3566959	3631799.8	3713444	10912204				
6	Tasikmala	2886455	2920347.3	2970956	8777758				
7	Cimis	3236121	3305451.5	3416381	9957953				
8	Kuning	2635535	2738240.1	2819522	8193297				
9	Carubon	2580728	2671645.9	2742543	7994917				
10	Malengka	2445604	2558835.2	2625151	7629590				
11									
12									
13									
14									
15									

C. Menyimpan data ke dalam SPSS

Setelah data diinput, data akan disimpan untuk digunakan pada waktu yang akan datang. Untuk itu:

- Buka menu **FILE** → **SAVE AS...**

Pada kotak dialog SAVE AS, simpan data dengan nama **PDRB tanpa migas Tahun 2013**. Lalu tempatkan pada folder tertentu dalam hard disk/tempat penyimpanan data Kita.

Latihan 2: Memasukan Data ke dalam SPSS (2)

Jika pada Latihan 1 tidak ada data kategori, pada Latihan 2 ini akan dimasukkan data kategori, sehingga akan digunakan VALUE.

Berikut data prestasi olahraga atletik sejumlah calon pegawai:

TINGGI BADAN	GENDER	PENDIDIKAN	LARI 100 M	LONCAT TINGGI	LOMPAT JAUH
170,50	PRIA	SMU	29,20	1,81	2,67
174,50	PRIA	SMU	28,10	1,45	2,45
168,40	WANITA	D3	38,40	1,12	2,15
165,70	WANITA	SARJANA	29,50	1,16	2,31
169,40	PRIA	SARJANA	38,40	1,37	2,51
152,80	WANITA	D3	45,70	1,32	2,16
171,60	PRIA	SARJANA	42,50	1,38	2,51
172,40	PRIA	SARJANA	45,90	1,31	2,41
156,50	WANITA	D4	35,40	1,26	1,81
154,20	WANITA	SMU	35,90	1,24	1,84

Data pertama adalah seorang calon pegawai pria berpendidikan SMU dengan tinggi badan 170,5 centimeter. Ia mampu berlari sejauh 100 meter dalam waktu 29,2 detik. Pada cabang loncat tinggi, ia mempunyai prestasi 1,81 meter. Sedang untuk lompat jauh ia mampu melompat sejauh 2,67 meter. Demikian seterusnya untuk arti data yang lain.

Pemasukan data di atas secara prinsip sama dengan Latihan 1. Perbedaan hanya pada saat memasukkan data GENDER dan PENDIDIKAN. Kedua jenis data tersebut adalah kategorikal karena isi data dapat diulang-ulang dan mempunyai batasan tertentu. Data gender hanya ada pria dan wanita, dan keduanya dapat diulang-ulang untuk sepuluh data tersebut. Demikian pula dengan data pendidikan dengan tiga kategori (SMU, Sarjana Muda, dan Sarjana).

PROSES:

Membuka VARIABLE VIEW untuk definisi variabel

Buka program SPSS, lalu klik sheet **VARIABLE VIEW**.

1. Input variabel TINGGI_BADAN
 - o Untuk NAME, isi TINGGI_BADAN.
Perhatikan adanya tanda baca garis bawah (_).
 - o Untuk TYPE, pilih NUMERIC karena data adalah angka.
 - o Untuk LABEL, isi Tinggi Badan (cm).

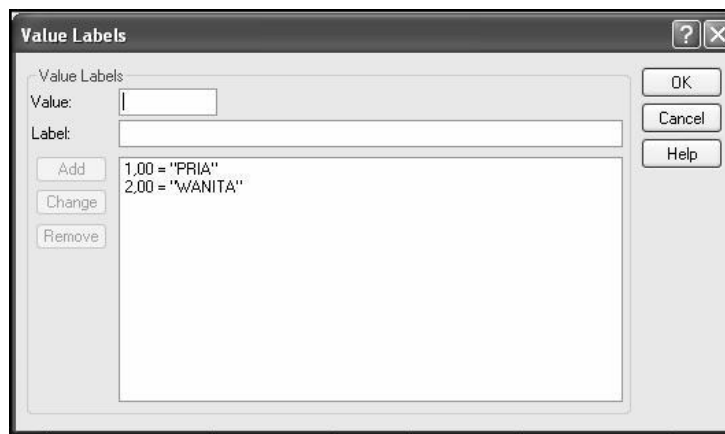
NB: tinggi badan akan diukur dalam cm. Abaikan bagian yang lain, lalu masuk ke variabel kedua.

2. Input variabel GENDER
 - o Untuk NAME, isi GENDER.
 - o Untuk TYPE, karena isi variabel GENDER adalah angka, akan digunakan proses kategorisasi dengan VALUE. Untuk itu, pilih NUMERIC.
 - o Untuk LABEL, kosongkan saja bagian ini.

- o Untuk VALUE, klik sekali pada sel VALUE, akan tampak ikon  di bagian kanan. Klik sekali pada ikon tersebut, akan tampak kotak dialog VALUE LABEL.
- o Pada bagian VALUE, masukan angka **1**.
- o Pada bagian LABEL, ketik **pria**.
- o Kemudian tekan tombol ADD untuk memasukan kode 1 yang adalah PRIA tersebut.

Ulangi proses di atas:

- o Pada bagian VALUE, masukan angka **2** Pada bagian LABEL, ketik **wanita**. Kemudian tekan tombol ADD untuk memasukkan kode 2 yang adalah WANITA tersebut. Oleh karena hanya ada dua kode, tekan OK. Tampilan akhir pemasukan dua kode gender tersebut.



Abaikan bagian yang lain, lalu masuk ke variabel ketiga.

3. Input variabel PENDIDIKAN

Variabel ini sama dengan GENDER, yakni data kategori; hanya di sini akan dimasukkan tiga kode.

- o Untuk NAME, isi PENDIDIKAN.
- o Untuk TYPE, pilih NUMERIC.
- o Untuk LABEL, kosongkan saja bagian ini.
- o Untuk VALUE, pada kotak dialog VALUE LABEL:
- o Pada bagian VALUE, masukkan angka **1**.
- o Pada bagian LABEL, ketik **SMU**. Kemudian tekan tombol ADD.
- o Pada bagian VALUE, masukkan angka **2**.
- o Pada bagian LABEL, ketik **D3**. Kemudian tekan tombol ADD.
- o Pada bagian VALUE, masukkan angka **3**.
- o Pada bagian LABEL, ketik **SARJANA**. Kemudian tekan tombol ADD. tekan OK.

4. Input variabel LARI_100_M
 - Untuk NAME, isi LARI_100_M.
 - Untuk TYPE, pilih NUMERIC.
 - Untuk LABEL, ketik lari 100 m (detik).

NB: prestasi lari 100 meter akan diukur dalam detik. Abaikan bagian yang lain, lalu masuk ke variabel lima.

5. Input variabel LONCAT TINGGI
 - Untuk NAME, isi LONCAT_TINGGI.
 - Untuk TYPE, pilih NUMERIC.
 - Untuk LABEL, ketik **loncat tinggi (meter)**.

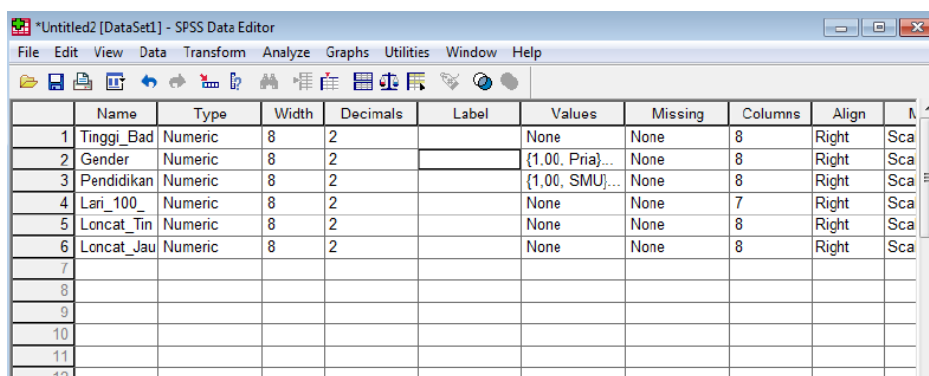
NB: prestasi loncat tinggi akan diukur dalam meter. Abaikan bagian yang lain, lalu masuk ke variabel enam.

6. Input variabel LOMPAT JAUH
 - Untuk NAME, isi **LOMPAT_JAUH**.
 - Untuk TYPE, pilih NUMERIC.
 - Untuk LABEL, ketik **lompat jauh (meter)**.

NB: prestasi lompat jauh akan diukur dalam meter. Oleh karena pemasukan variabel sudah selesai, klik sheet DATA VIEW yang ada di bagian kiri bawah.

Memasukkan data ke dalam SPSS

Pada tampilan **DATA VIEW**, masukan data seperti table soal di atas. Hanya saat memasukkan data GENDER dan PENDIDIKAN, gunakan *angka* dan bukan *kata*. Sebagai contoh, saat memasukkan gender pria, ketik angka 1 dan jangan kata 'pria'. Demikian pula untuk variabel pendidikan. Tampilan akhir input data:



	Name	Type	Width	Decimals	Label	Values	Missing	Columns	Align	Measure
1	Tinggi_Bad	Numeric	8	2		None	None	8	Right	Scale
2	Gender	Numeric	8	2		{1.00, Pria}...	None	8	Right	Scale
3	Pendidikan	Numeric	8	2		{1.00, SMU}...	None	8	Right	Scale
4	Lari_100	Numeric	8	2		None	None	7	Right	Scale
5	Loncat_Tin	Numeric	8	2		None	None	8	Right	Scale
6	Loncat_Jau	Numeric	8	2		None	None	8	Right	Scale
7										
8										
9										
10										
11										
12										

*Untitled2 [DataSet1] - SPSS Data Editor

File Edit View Data Transform Analyze Graphs Utilities Window Help

Visible: 6 of 6 Variables

	Tinggi_Badan	Gender	Pendidikan	Lari_100_M	Loncat_Tinggi	Loncat_Jauh	var	var	var	var
1	170,50	Pria	SMU	29,20	1,81	2,67				
2	174,50	Pria	SMU	28,10	1,45	2,45				
3	168,40	Wanita	D3	38,40	1,12	2,15				
4	165,70	Wanita	SARJANA	29,50	1,16	2,31				
5	169,40	Pria	SARJANA	38,40	1,37	2,51				
6	152,80	Wanita	D3	45,70	1,32	2,16				
7	171,60	Pria	SARJANA	42,50	1,38	2,51				
8	172,40	Pria	SARJANA	45,90	1,31	2,41				
9	156,50	Wanita	D3	35,40	1,26	1,81				
10	154,20	Wanita	SMU	35,90	1,24	1,84				
11										
12										
13										
14										
15										

Data View Variable View

SPSS Processor is ready

Catatan:

Jika pada tampilan data di atas, dari menu **VIEW** Kita non-aktifkan submenu **VALUE LABEL** (klik submenu tersebut hingga tanda ☒ hilang), maka tampilan akan berubah:

*Untitled2 [DataSet1] - SPSS Data Editor

File Edit View Data Transform Analyze Graphs Utilities Window Help

Visible: 6 of 6 Variables

	Tinggi_Badan	Gender	Pendidikan	Lari_100_M	Loncat_Tinggi	Loncat_Jauh	var	var	var	var
1	170,50	1,00	1,00	29,20	1,81	2,67				
2	174,50	1,00	1,00	28,10	1,45	2,45				
3	168,40	2,00	2,00	38,40	1,12	2,15				
4	165,70	2,00	3,00	29,50	1,16	2,31				
5	169,40	1,00	3,00	38,40	1,37	2,51				
6	152,80	2,00	2,00	45,70	1,32	2,16				
7	171,60	1,00	3,00	42,50	1,38	2,51				
8	172,40	1,00	3,00	45,90	1,31	2,41				
9	156,50	2,00	2,00	35,40	1,26	1,81				
10	154,20	2,00	1,00	35,90	1,24	1,84				
11										
12										
13										
14										
15										

Data View Variable View

Variabel GENDER dan PENDIDIKAN akan tampil dengan angka 1 dan 2 untuk GENDER, serta 1, 2, dan 3 untuk PENDIDIKAN. Angka-angka inilah yang seharusnya Kita masukkan untuk kedua variabel tersebut.

Untuk mengembalikan tampilan, aktifkan lagi submenu **VALUE LABEL** tersebut.

Menyimpan data ke dalam SPSS

Simpan data (lewat meui **FILE** → **SAVE AS**) dengan nama **PRESTASI OLAHRAGA**.

LATIHAN

1. Tabel berikut menunjukkan data perbandingan jumlah mobil dan sepeda motor yang terjual di Indonesia pada periode 2005 – 2011.

Tahun	MOBIL	SEPEDA MOTOR
2005	300.960	979.422
2006	299.630	1.650.770
2007	317.780	2.317.991
2008	354.360	2.820.000
2009	483.170	3.900.000
2010	533.840	5.090.000
2011	318.880	4.470.000

Masukkan data di atas ke dalam SPSS. Simpan data dengan nama **MOBIL dan MOTOR**. Petunjuk:

- Buat tiga variabel.
 - Semua tipe data adalah numerik; perhatikan untuk tidak menggunakan tanda baca 'titik (.)' saat input data.
 - Sesuai jenis data yang adalah bilangan bulat (*integer*), tetapkan desimal adalah 0.
2. Berikut tabel komposisi jenis sepeda motor yang terjual di Indonesia pada tahun 2012

TIPE SEPEDA MOTOR	PERSENTASE
Bebek	90,50%
Sport	6,20%
Bisnis	2,10%
Skuter	1,20%

Masukkan data di atas ke dalam SPSS. Simpan data dengan nama **KOMPOSISI MOTOR**. Petunjuk:

- Buat dua variabel.
 - Variabel pertama adalah karakter, sehingga dipilih tipe data STRING.
 - Jangan memasukkan tanda '%'. Sebagai contoh, angka 90,50% dimasukkan sebagai **90,50**.
3. P.T. SUKSES ABADI bergerak dalam bidang penjualan telepon seluler. Perusahaan tersebut mempunyai sejumlah tenaga penjualan (salesman) dengan data pada bulan Juni 2012 sebagai berikut.

Nama	gender	penddk	Masa Kerja	Nokia	Sony	Motorola	Samsung
RETNO	wanita	SMU	4	15	8	11	9
SUSI	wanita	SMU	3	20	5	14	6
SENTOT	pria	SMU	4	32	11	15	7
HUSIN	pria	SMU	4	14	14	12	8
HENDRO	pria	SMU	5	15	15	17	5
IIN	wanita	SARJANA	7	20	20	20	9
IAN	pria	SARJANA	3	20	7	21	7
RATNA	wanita	SMU	5	15	9	16	8
LILIK	pria	SMU	4	21	7	18	6
LUKITO	pria	SARJANA	6	27	14	19	11
BENNY	pria	SMU	5	21	15	14	10
ANITA	Wanita	SARJANA	2	9	18	16	10
LANY	Wanita	SMU	4	21	12	20	9
KURNIA	Pria	SMU	4	8	15	17	12
NANIK	Wanita	SMU	3	14	9	18	15

Masukkan data dan simpan dengan nama **PENJUALAN PONSEL**.

- Buat delapan variabel, dengan variabel pertama adalah karakter, sehingga tipe data adalah STRING.
- Variabel GENDER (2 kategori) dan PENDIDIKAN (3 kategori) adalah data kategori, sehingga pengisian lewat VALUES.
- Variabel keempat dalam satuan TAHUN MASA KERJA.
- Variabel lima sampai delapan adalah MERK PONSEL YANG TERJUAL (DALAM UNIT).
- Pada bagian LABEL (saat mendefinisikan variabel), isilah dengan bebas, namun dengan tetap memperhatikan nama variabel. Variabel PENGALAMAN KERJA, dapat ditulis **Pengalaman Kerja (tahun)**.

BAB 2

STATISTIK DESKRIPTIF

Pada Bab ini, akan dipelajari bagaimana menggambarkan/memaparkan suatu data dalam bentuk grafik maupun tabel.

Berikut adalah data mengenai Indeks Prestasi mahasiswa dari Fakultas Ekonomi, Fisipol dan Hukum sebagai berikut :

Tabel 3.1 Data IPK Mahasiswa

Fakultas	IPK	Fakultas	IPK	Fakultas	IPK
Ekonomi	3,3	Hukum	3,25	Ekonomi	3,3
Ekonomi	2,9	Hukum	3,25	Ekonomi	2,9
Ekonomi	3,4	Ekonomi	2,75	Ekonomi	3,1
Fisipol	3,8	Fisipol	2,5	Fisipol	3,1
Fisipol	2,75	Fisipol	2,5	Hukum	3,1
Ekonomi	3,25	Fisipol	3,25	Hukum	3,25
Fisipol	3,7	Hukum	3,85	Hukum	3,7
Hukum	3,8	Hukum	3,9	Fisipol	3,8
Ekonomi	3,5	Hukum	3,5	Fisipol	3,5
Fisipol	2,9	Ekonomi	2,9	Hukum	2,9

Kasus di atas akan dibuat dalam tabel frekuensi, baik berdasarkan Indeks Prestasi Mahasiswa maupun berdasarkan asal fakultas.

1. Input Data

- Menampilkan tampilan **VARIABLE VIEW** untuk memasukkan identitas variabel data sesuai dengan cara input masing-masing atribut pada

pembahasan sebelumnya. Sehingga akan menjadi tampilan **Variabel View** seperti pada (Gambar 3.1) dan Data View seperti pada (Gambar 3.2).

	Name	Type	Width	Decimals	Label	Values	Missing	Columns	Align
1	Fakultas	String	8	0	Asal Fakultas	(1, Ekonomi)...	None	8	Left
2	IPK	Numeric	8	2	Skor IPK	None	None	8	Right
3									
4									

(Gambar 3.1 Variabel View Data IPK)

	Fakultas	IPK
1	Ekonomi	3,30
2	Ekonomi	2,90
3	Ekonomi	3,40
4	Fisipol	3,80
5	Fisipol	2,75
6	Ekonomi	3,25
7	Fisipol	3,70
8	Hukum	3,80
9	Ekonomi	3,50
10	Fisipol	2,90
11	Hukum	3,25
12	Hukum	3,25
13	Ekonomi	2,75
14	Fisipol	2,50
15	Fisipol	2,50
16	Fisipol	3,25

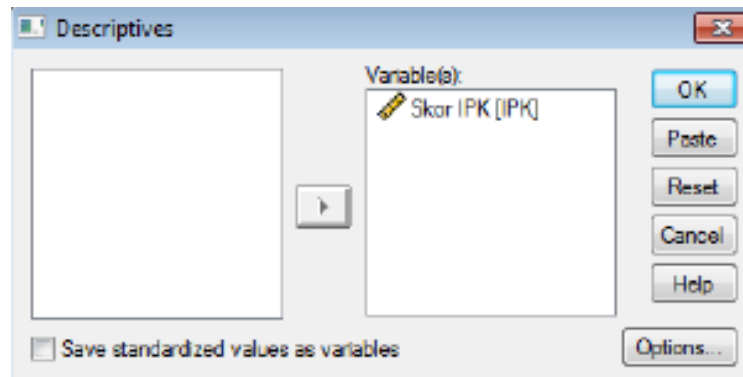
(Gambar 3.2 Data View Data IPK mahasiswa)

Note: Fakultas bertipe data string yang diubah menjadi numerik dengan pemberian **Value**.

2. STATISTIK DESKRIPTIF untuk IPK

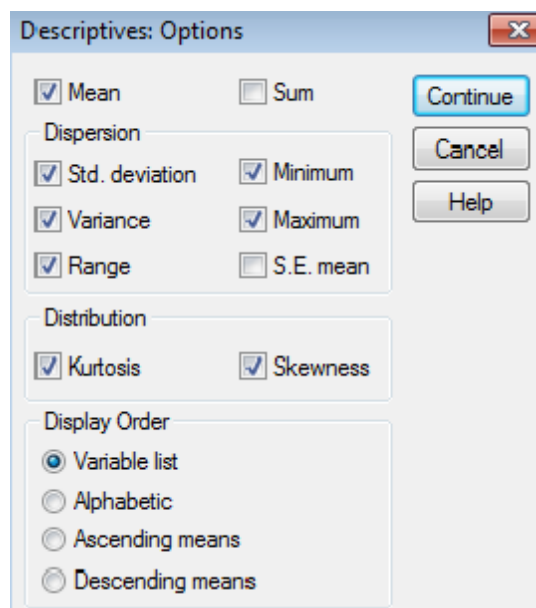
Oleh karena variabel **IPK** termasuk data kuantitatif, maka akan dibuat tabel frekuensi serta deskriptif statistik (meliputi Mean, Standart Deviasi, Range dan lainnya) untuk variabel tersebut. Selain itu akan dilengkapi dengan visualisasi data berupa Chart yang sesuai untuk data kuantitatif, yaitu Histogram atau Bar Chart.

- Dari menu utama SPSS, pilih menu **Analyze**, kemudian pilih submenu **Descriptive Statistics** → **Descriptives**, Maka akan keluar tampilan seperti **Gambar 3.5**.



(Gambar 3.3 Descriptives)

- Masukkan variabel **skor IPK** ke dalam kolom **Variable(s)**.
- Pilih **Options** maka akan tampil pada layar seperti **Gambar 2.4**. Dialog Box tersebut adalah untuk menampilkan karakteristik data apa saja yang ingin kita tampilkan. Beri tanda Cek List pada **Mean, Std deviation, Variance, Range, Minimum, dan Maximum**. Abaikan yang lain. Kemudian klik Continue untuk kembali pada Dialog Box **Descriptives**, kemudian pilih **OK**.



Gambar 3.4 Descriptives Options

- Output yang muncul adalah seperti pada **Gambar 3.5**.

Descriptive Statistics											
	N	Range	Minimum	Maximum	Mean	Std.	Variance	Skewness		Kurtosis	
	Statistic	Statistic	Statistic	Statistic	Statistic	Statistic	Statistic	Statistic	Std. Error	Statistic	Std. Error
Skor IPK	30	1,40	2,50	3,90	3,2533	,39891	,159	-,083	,427	-,803	,833
Valid N (listwise)	30										

Gambar 3.5 Descriptive Statistics

Analisis:

Berdasarkan **Gambar 2.5** didapatkan beberapa karakteristik data yaitu :

N = 30	Banyaknya data yang diolah adalah 30
Mean (Rata-rata) = 3,25	artinya besarnya IPK rata-rata berkisar diantara 3,25
Minimum = 2,5	Nilai Minimum IPK dari 30 mahasiswa tersebut adalah 2,5
Maximum = 3,9	Nilai Maksimum skor IPK dari 30 mahasiswa tersebut adalah 3,9
Range = 1,4	Merupakan selisih nilai Minimum dan Maksimum yaitu $3,90 - 2,50 = 1,4$
Variance = 0,159	Berkaitan erat dengan variasi data. Semakin besar nilai variance, maka berarti variasi data semakin tinggi
Standard Deviation = 0,398	Merupakan akar kuadrat dari Variance

Selain masih berkaitan dengan dengan variasi data, Penggunaan standard deviasi untuk memperkirakan dispersi rata-rata populasi (simpangan data). Untuk itu, dengan standard deviasi tertentu dan pada tingkat kepercayaan 95% (SPSS sebagian besar menggunakan angka ini sebagai standar), maka rata-rata tinggi badan populasi diperkirakan antara:

Rata-rata $\pm$ 2 Standart Deviasi

NB : Angka 2 digunakan, karena tingkat kepercayaan 95%.

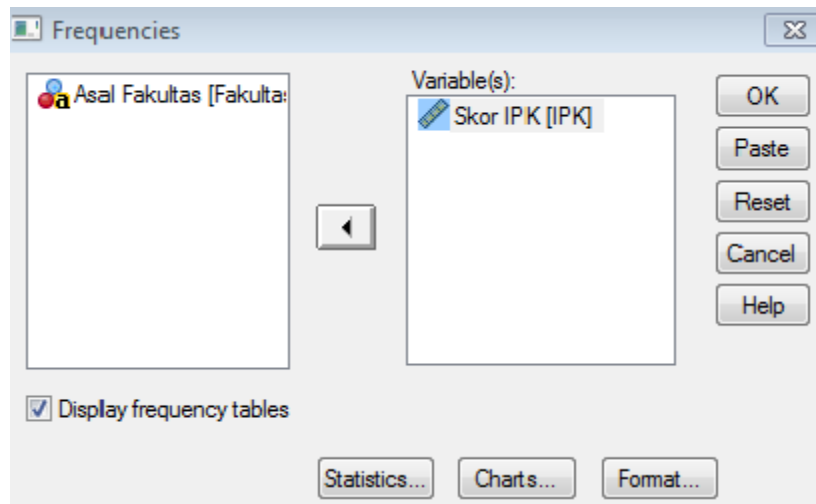
Maka : $3,25 \pm (2 \times 0,398)$
: 2,45 sampai 4,0.

Perhatikan kedua batas angka yang berbeda tipis dengan nilai minimum dan maksimum. Hal ini membuktikan sebaran data adalah baik.

3. TABEL FREKUENSI DATA INDEKS PRESTASI

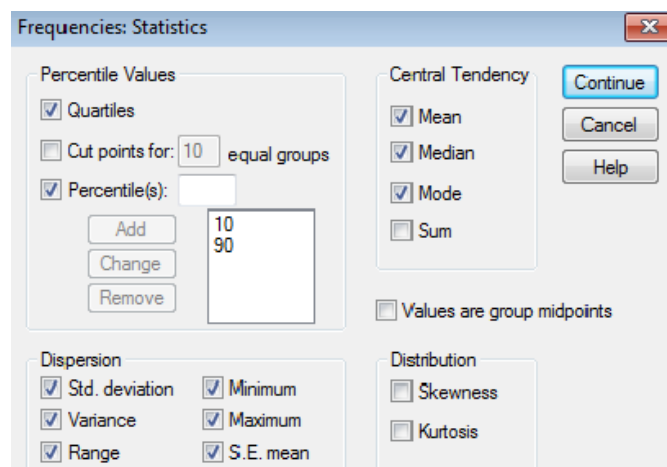
Selain dengan menggunakan menu **Descriptives**, informasi mengenai karakteristik data akan lebih tereksplorasi dengan menggunakan tabel frekuensi dan penyajian secara visual melalui grafik yang sesuai. Langkah-langkahnya adalah:

- Pilih menu **Analyze → Descriptive Statistics → Frequencies**, sehingga akan muncul Dialog Box sesuai Gambar 3.6



(Gambar 3.6 descriptive : Frequencies Skor IPK Mahasiswa)

- Masukkan variabel **Skor IPK** pada kolom **Variable(s)**. Un Chek **Display Frequency Tables**.
- Pilih **Statistics** sehingga muncul Dialog Box **Gambar 3.7**. Beri tanda Cek List pada **Quartiles**, **Percentiles** (isi dengan nilai 10 dan 90 sebagai contoh), **Mean**, **Median**, **Mode**, **Standard deviation**, **Variance**, **Range**, **Minimum**, **Maximum** dan **SE Mean**. Abaikan yang lain, lalu pilih **Continue** untuk kembali ke Dialog Box **Frequencies** → **OK**



Gambar 3.7 Statistics

- Output yang akan muncul adalah sesuai **Gambar 3.8**.

Statistics

Skor IPK		
N	Valid	30
	Missing	0
Mean		3,2533
Std. Error of Mean		,07283
Median		3,2500
Mode		2,90 ^a
Std. Deviation		,39891
Variance		,159
Range		1,40
Minimum		2,50
Maximum		3,90
Percentiles	10	2,7500
	25	2,9000
	50	3,2500
	75	3,5500
	90	3,8000

a. Multiple modes exist. The smallest value is shown

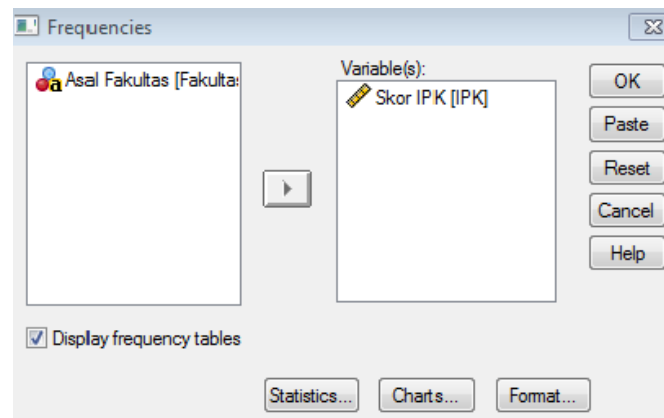
(Gambar 3.8 Statistics)

Analisis :

- **Mean, Std Deviation, Variance, Range, Minimum dan Maximum** telah dijelaskan sebelumnya.
- **Std Error of Mean** = 0,07283, digunakan untuk memperkirakan Rata-rata populasi berdasarkan 30 data sampel yang diolah.
Rata-Rata Populasi = Mean $\pm$ 2 Std Error of Mean. Sehingga berdasarkan data, didapatkan Rata-rata Populasi = $3,25 \pm 2 (0,07283)$
Maka Rata-rata Populasi bekisar antara 3,1-3,4
- **Median = 3,25**, merupakan nilai tengah data yang telah diurutkan (baik dari kecil ke besar maupun dari besar ke kecil). Sehingga Median = 3,25 berarti bahwa 50% berada di bawah (kurang dari) 3,25 dan 50% lainnya berada di atas (lebih dari) 3,25.
- **Percentiles 10, 25, 50, 75, dan 90** merupakan batasan-batasan yang menunjukkan proporsi sebaran data.
Percentile 10 = 2,75, artinya 10% data ($10\% \times 30 = 3$) atau ada 3 data terkecil bernilai kurang dari 2,75
Percentile 25 = 2,9, artinya 25% data ($25\% \times 30 = 8$) atau ada 8 data terkecil bernilai kurang dari 2,9..
 Demikian juga seterusnya untuk percentile 50, 75 dan 90.

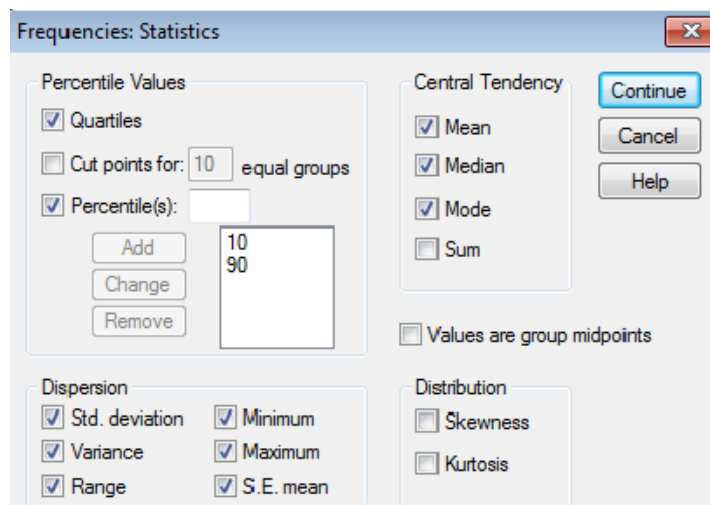
4. TABEL FREKUENSI DATA SKOR IPK

Langkah pembuatan Tabel Frekuensi pada Data **Skor IPK** Pilih Menu **analyze → Descriptive Statistics → Frequencies**, sehingga muncul **Dialog Box** sesuai **Gambar 3.9**.



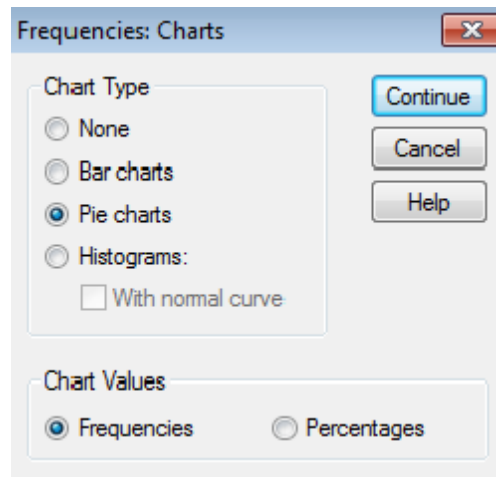
Gambar 3.9 Frequencies Data Mahasiswa Berdasarkan fakultas

- Masukkan Variabel **Skor IPK** pada kolom **Variable(s)**, Check List **Display Frequency Tables**.
- Pilih **Statistics**, Check semua item. **Continue**.



Gambar 3.10
Frequencies Statistics Data IPK Berdasarkan Fakultas

- Pilih **Charts**, kemudian pilih **Pie Chart → Continue → OK**



Gambar 3.11
Frequencies : Charts untuk Skor IPK

- Output yang muncul adalah sesuai **Gambar 3.12, Gambar 3.13, dan Gambar 3.14.**

Statistics

Asal Fakultas		
N	Valid	30
	Missing	0

Gambar 3.12 Output Statistics Data Jenis Bank)

Analisis

- **N Valid = 30**, Data yang diolah sebanyak 30
- **Missing = 0**, tidak ada data hilang

Asal Fakultas

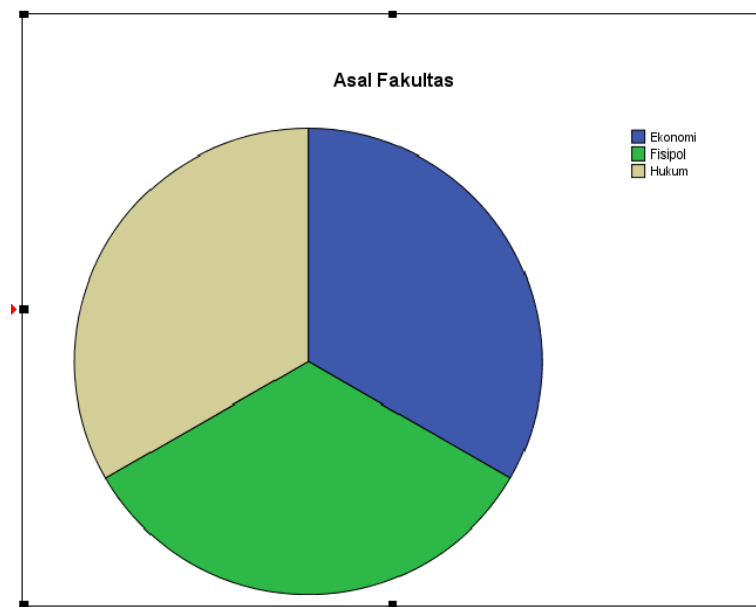
		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	Ekonomi	10	33,3	33,3	33,3
	Fisipol	10	33,3	33,3	66,7
	Hukum	10	33,3	33,3	100,0
	Total	30	100,0	100,0	

Gambar 3.13
Output Tabel Frekuensi Mahasiswa Berdasarkan Fakultas

Analisis

- Kolom **Frequency** menunjukkan banyaknya Jenis Fakultas pada data yang telah diolah. Pada tabel ditunjukkan bahwa terdapat 3 Fakultas.

- Kolom **Percent** berarti presentase jumlah masing-masing jenis Fakultas, yaitu 33.3% untuk Fakultas Ekonomi, 33,3 % untuk Fisipol dan 33,3 % untuk Fakultas Hukum, dapat disimpulkan bahwa dari sampel yang telah diambil rata-rata jumlah sampel adalah sama
- Kolom **Valid Percent** = Kolom **Percent**. Kolom **Cumulative Percent** merupakan jumlah kumulatif presentase.



(Gambar 3.14 Pie Chart Skor IPK)

Analisis

- **Pie Chart** menunjukkan bahwa proporsi ke tiga fakultas adalah sama

Latihan

1. Diketahui hasil nilai ujian STATISTIK EKONOMI mahasiswa ekonomi semester IV sebagai berikut :

65	44	46	95	55	39	55	89	48	34
34	60	40	40	60	89	85	70	80	62
50	55	67	48	49	45	45	50	89	98
65	70	77	70	59	52	55	49	35	30
80	65	81	60	70	76	78	65	65	88
75	58	55	76	48	70	70	85	64	77
30	30	30	55	95	67	90	68	61	70

Buatlah statistik deskriptif dari data diatas !

2. Diketahui berat badan mahasiswa Ilmu Ekonomi Universitas Sabar Menanti (USM) sebagai berikut:

51	43	55	45	45	53	46	43	54	70
58	45	66	57	46	50	49	55	55	36
56	57	53	58	58	54	48	43	63	58
55	48	50	68	65	41	42	50	64	60
44	46	52	54	56	58	60	62	64	66

Pertanyaan :

1. Buatlah Tabel Distribusi Frekuensinya !
2. Gambarkan ke dalam bentuk grafik



UJI T SATU SAMPEL

One sample t test merupakan teknik analisis untuk membandingkan satu variabel bebas. Teknik ini digunakan untuk menguji apakah nilai tertentu berbeda secara signifikan atau tidak dengan rata-rata sebuah sampel.

Uji t sebagai teknik pengujian hipotesis deskriptif memiliki tiga kriteria yaitu uji pihak kanan, kiri dan dua pihak.

Uji Pihak Kiri : dikatakan sebagai uji pihak kiri karena t tabel ditempatkan di bagian kiri Kurva

Uji Pihak Kanan : Dikatakan sebagai uji pihak kanan karena t tabel ditempatkan di bagian kanan kurva.

Uji dua pihak : dikatakan sebagai uji dua pihak karena t tabel dibagi dua dan diletakkan di bagian kanan dan kiri

Contoh Kasus

Contoh Rumusan Masalah : Bagaimana tingkat keberhasilan belajar siswa

Hipotesis kalimat :

1. Tingkat keberhasilan belajar siswa paling tinggi 70% dari yang diharapkan (uji pihak kiri / 1-tailed)
2. Tingkat keberhasilan belajar siswa paling rendah 70% dari yang diharapkan (uji pihak kanan / 1-tailed)
3. Tingkat keberhasilan belajar siswa tidak sama dengan 70% dari yang diharapkan (uji 2 pihak / 2-tailed)

Pengujian Hipotesis : Rumusan masalah Satu

Hipotesis kalimat

Ha : tingkat keberhasilan belajar siswa paling tinggi 70% dari yang diharapkan

Ho : tingkat keberhasilan belajar siswa paling rendah 70% dari yang diharapkan

Hipotesis statistik

Ha : $\mu_0 < 70\%$

Ho : $\mu_0 \geq 70\%$

Parameter uji :

Jika $-t_{\text{tabel}} \leq t_{\text{hitung}}$ maka H_0 diterima, dan H_a di tolak

Jika $-t_{\text{tabel}} > t_{\text{hitung}}$ maka H_0 ditolak, dan H_a diterima

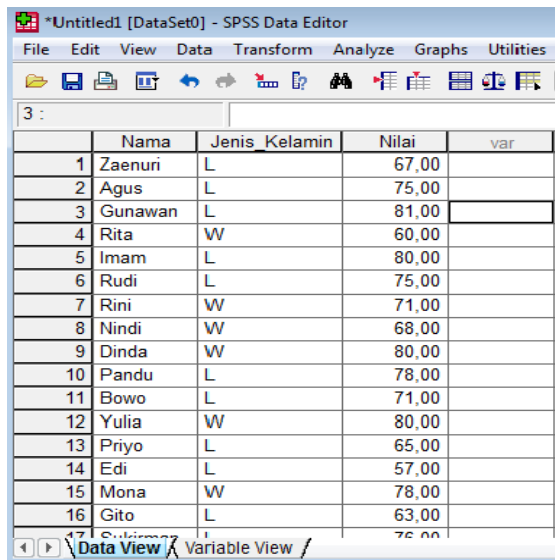
Penyelesaian Kasus 1 (uji t pihak kiri)

Data yang hasil ulangan matematika Universitas Gunung Kelud adalah sebagai berikut :

Tabel 4.1.
Hasil Ujian Matematika Universitas Gunung Kelud

No	Nama	Jenis_Kel	Nilai	No	Nama	Jenis_Kel	Nilai
1	Zaenuri	L	67	21	Nono	L	72
2	Agus	L	75	22	Rika	W	80
3	Gunawan	L	81	23	Tika	W	75
4	Rita	W	60	24	Tono	L	67
5	Imam	L	80	25	Toni	L	72
6	Rudi	L	75	26	Ika	W	79
7	Rini	W	71	27	Ian	L	80
8	Nindi	W	68	28	Lili	W	81
9	Dinda	W	80	29	Ari	L	75
10	Pandu	L	78	30	Aryani	W	71
11	Bowo	L	71	31	Tejo	L	74
12	Yulia	W	80	32	Tarjo	L	65
13	Priyo	L	65	33	Ngadiman	L	55
14	Edi	L	57	34	Ngadimin	L	70
15	Mona	W	78	35	Teno	L	72
16	Gito	L	63	36	Wuri	W	82
17	Sukirman	L	76	37	Wilian	L	67
18	Kirun	L	73	38	Ida	W	94
19	Maryati	W	63	39	Ita	W	60
20	Nani	W	65	40	Susi	W	79

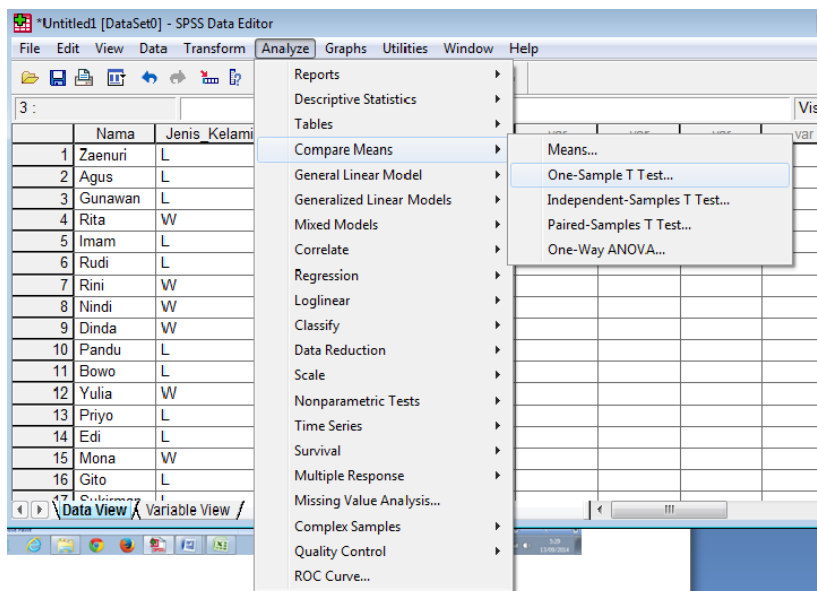
Masukan data diatas kedalam SPSS, sehingga diperoleh sebagai berikut :



*Untitled1 [DataSet0] - SPSS Data Editor

	Nama	Jenis Kelamin	Nilai	var
1	Zaenuri	L	67,00	
2	Agus	L	75,00	
3	Gunawan	L	81,00	
4	Rita	W	60,00	
5	Imam	L	80,00	
6	Rudi	L	75,00	
7	Rini	W	71,00	
8	Nindi	W	68,00	
9	Dinda	W	80,00	
10	Pandu	L	78,00	
11	Bowo	L	71,00	
12	Yulia	W	80,00	
13	Priyo	L	65,00	
14	Edi	L	57,00	
15	Mona	W	78,00	
16	Gito	L	63,00	

Klik Analyze – **Pilih Compare Means**, lalu pilih **One Sample T Test**, Masukkan variabel nilai ke dalam Test Variable Box, abaikan yang lain kemudian klik OK



Selanjutnya Uji Normalitas data :

Klik **Analyze**, Pilih **Non Parametrics Test** – pilih **1 Sampel K-S**, masukkan variabel nilai ke dalam **Test Variable List**, kemudian Klik OK

Hasil

One-Sample Statistics

	N	Mean	Std. Deviation	Std. Error Mean
Nilai	40	72,4000	7,99936	1,26481

One-Sample Test

	Test Value = 0					
	t	df	Sig. (2-tailed)	Mean Difference	95% Confidence Interval of the Difference	
					Lower	Upper
Nilai	57,242	39	,000	72,40000	69,8417	74,9583

One-Sample Kolmogorov-Smirnov Test

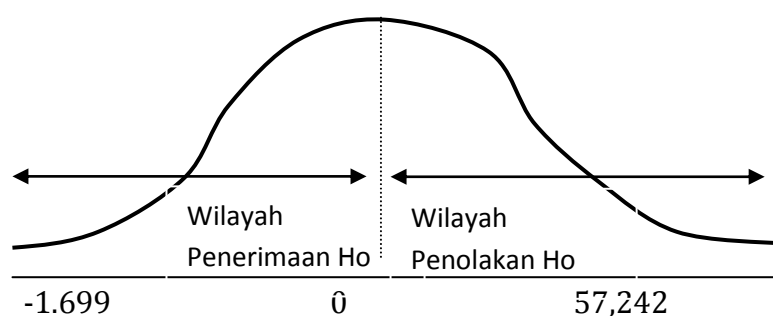
		Nilai
N		40
Normal Parameters(a,b)	Mean	72,4000
	Std. Deviation	7,99936
Most Extreme Differences	Absolute	,091
	Positive	,091
	Negative	-,083
Kolmogorov-Smirnov Z		,577
Asymp. Sig. (2-tailed)		,894

a Test distribution is Normal.

b Calculated from data.

Hasil uji di atas menunjukkan bahwa t hitung = 57,242. T tabel diperoleh dengan $df = 39$, sig 5% (1 tailed) = 1.699. Karena $-t$ tabel > dari t hitung ($57,242 > 1,699$), maka H_0 ditolak, artinya tingkat keberhasilan belajar siswa paling tinggi 70% terbukti, bahkan lebih dari yang diduga yaitu sebesar 72,4

Hasil uji normalitas data menunjukkan nilai Kol-Smirnov sebesar 0.600 dan Asymp. Sig tidak signifikan yaitu sebesar 0.577 (> 0.05), sehingga dapat disimpulkan data berdistribusi normal



Pengujian Hipotesis : Rumusan masalah Dua Hipotesis kalimat

H_a : tingkat keberhasilan belajar siswa paling tinggi 70% dari yang diharapkan

Ho : tingkat keberhasilan belajar siswa paling rendah 70% dari yang diharapkan

Hipotesis statistik

Ha : $\mu_0 < 70\%$

Ho : $\mu_0 > 70\%$

Parameter uji :

Jika $+ t_{\text{tabel}} > t_{\text{hitung}}$ maka Ho diterima, dan Ha di tolak

Jika $+ t_{\text{tabel}} < t_{\text{hitung}}$ maka Ho ditolak, dan Ha diterima

Penyelesaian Kasus 2 (uji t pihak kanan)

Data yang hasil ulangan matematika siswa sebanyak 40 mahasiswa sama seperti data di atas Klik **Analyze** – Pilih **Compare Means**, lalu pilih **One Sample T Test**, Masukkan variabel **nilai** ke dalam **Test Variable Box**, abaikan yang lain kemudian klik **OK**

Selanjutnya Uji Normalitas data : Klik **Analyze**, Pilih **Non Parametrics Test** – pilih **1 Sampel K-S**, masukkan variabel nilai ke dalam Test Variable List, kemudian Klik OK.

Masih menggunakan hasil analisis di atas, maka diperoleh t hitung sebesar 57,242, dan t tabel = 1.699. Karena $t_{\text{tabel}} < t_{\text{hitung}}$ ($1.699 < 57,242$), maka **Ho ditolak, dan Ha diterima**. Artinya Ha yaitu tingkat keberhasilan siswa paling tinggi 70% dari yang diharapkan diterima. Sedangkan Ho yang menyatakan bahwa keberhasilan belajar paling rendah 70% ditolak.



UJI T SAMPEL BERPASANGAN

Paired sample t test merupakan uji beda dua sampel berpasangan. Sampel berpasangan merupakan subjek yang sama namun mengalami perlakuan yang berbeda.

CONTOH KASUS

Akan diteliti mengenai perbedaan penjualan sepeda motor merk A disebuah Kabupaten sebelum dan sesudah kenaikan harga BBM. Data diambil dari 10 dealer.

Data yang diperoleh adalah sebagai berikut :

No	Sebelum	Sesudah
1	67	68
2	75	76
3	81	80
4	60	63
5	80	82
6	75	74
7	71	70
8	68	71
9	80	82
10	78	79

Masukan dalam SPSS

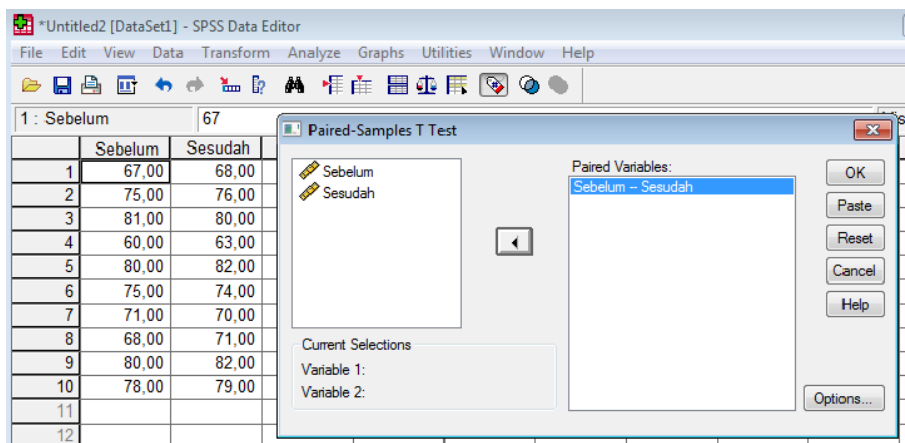
*Untitled2 [DataSet1] - SPSS Data Editor

	Sebelum	Sesudah	var	var	var	var
1	67,00	68,00				
2	75,00	76,00				
3	81,00	80,00				
4	60,00	63,00				
5	80,00	82,00				
6	75,00	74,00				
7	71,00	70,00				
8	68,00	71,00				
9	80,00	82,00				
10	78,00	79,00				
11						

PENYELESAIAN

Klik **ANALYZE > COMPARE MEANS > PAIRED SAMPLES t Test**

Masukkan jual_1 dan jual_2 pada kolom **“Paired variables”** seperti gambar di bawah ini



Abaikan yang lain, klik OK

HASIL

Paired Samples Statistics

		Mean	N	Std. Deviation	Std. Error Mean
Pair 1	Sebelum	73,5000	10	6,88396	2,17690
	Sesudah	74,5000	10	6,43342	2,03443

Paired Samples Correlations

		N	Correlation	Sig.
Pair 1	Sebelum & Sesudah	10	,975	,000

Bagian pertama. Paired Samples Statistic

Menunjukkan bahwa rata-rata penjualan pada sebelum dan sesudah kenaikan BBM. Sebelum kenaikan BBM rata-rata penjualan dari 10 dealer adalah sebanyak 73,4, sementara setelah kenaikan BBM jumlah penjualan rata-rata adalah sebesar 74,5 unit

Bagian Dua. Paired samples Correlation

Hasil uji menunjukkan bahwa korelasi antara dua variabel adalah sebesar 0.975 dengan sig sebesar 0.000. Hal ini menunjukkan bahwa korelasi antara dua rata-rata penjualan sebelum dan sesudah kenaikan adalah kuat dan signifikan.

Hipotesis

Hipotesis yang diajukan adalah :

Ho : rata-rata penjualan adalah sama

H1 : rata-rata penjualan adalah berbeda

Hasil uji Hipotesis

Paired Samples Test									
		Paired Differences					t	df	Sig. (2-tailed)
		Mean	Std. Deviation	Std. Error Mean	95% Confidence Interval of the Difference				
					Lower	Upper			
Pair 1	Sebelum - Sesudah	-1,00000	1,56347	,49441	-2,11844	,11844	-2,023	9	,074

Nilai t hitung adalah sebesar -2,023 dengan sig 0.074. Karena sig > 0.05 maka dapat disimpulkan bahwa Ho diterima, artinya rata-rata penjualan sebelum dan sesudah kenaikan BBM adalah sama (tidak berbeda). dengan demikian dapat dinyatakan bahwa kenaikan harga BBM tidak mempengaruhi jumlah penjualan sepeda motor merek A di kabupaten tersebut.

Latihan

Apakah ada perbedaan kemampuan siswa sebelum dan sesudah diberitakan tambahan pelajaran ?

Responden	Nilai Ujian		Responden	Nilai Ujian	
	Sebelum Tambahan Les	Setelah Tambahan Les		Sebelum Tambahan Les	Setelah Tambahan Les
1	80	85	9	78	80
2	87	80	10	77	75
3	67	75	11	76	80
4	89	85	12	75	80
5	76	80	13	67	70
6	78	80	14	65	70
7	86	90	15	70	80
8	76	75	16	76	70



ANALISIS OF VARIANS

Setiap perusahaan perlu melakukan pengujian terhadap kumpulan hasil pengamatan mengenai suatu hal, misalnya hasil penjualan produk, hasil produksi produk, gaji pekerja di suatu perusahaan nilainya bervariasi antara satu dengan yang lainnya. Hal ini berhubungan dengan varian dan rata-rata yang banyak digunakan untuk membuat kesimpulan melalui penaksiran dan pengujian hipotesis mengenai parameter, maka dari itu dilakukan analisis varian yang ada dalam cabang ilmu statistika industri yaitu ANOVA. Penerapan ANOVA dalam dunia industri adalah untuk menguji rata-rata data hasil pengamatan yang dilakukan pada sebuah perusahaan ataupun industri.

Analisis varians (*analysis of variance*) atau ANOVA adalah suatu metode analisis statistika yang termasuk ke dalam cabang statistika inferensi. Uji dalam anova menggunakan uji F karena dipakai untuk pengujian lebih dari 2 sampel. Dalam praktik, analisis varians dapat merupakan uji hipotesis (lebih sering dipakai) maupun pendugaan (*estimation*, khususnya di bidang genetika terapan).

Anova (*Analysis of variances*) digunakan untuk melakukan **analisis komparasi multivariabel**. Teknik analisis komparatif dengan menggunakan tes “t” yakni dengan mencari perbedaan yang signifikan dari dua buah *mean* hanya efektif bila jumlah variabelnya dua. Untuk mengatasi hal tersebut ada teknik analisis komparatif yang lebih baik yaitu *Analysis of variances* yang disingkat anova.

Anova digunakan untuk **membandingkan rata-rata populasi** bukan ragam populasi. Jenis data yang tepat untuk anova adalah nominal dan ordinal pada variabel bebasnya, jika data pada variabel bebasnya dalam bentuk interval atau ratio maka harus diubah dulu dalam bentuk ordinal atau nominal. Sedangkan variabel terikatnya adalah data interval atau rasio.

Adapun asumsi dasar yang harus terpenuhi dalam analisis varian adalah :

1. Kenormalan
Distribusi data harus normal, agar data berdistribusi normal dapat ditempuh dengan cara memperbanyak jumlah sampel dalam kelompok.
2. Kesamaan variansi
Setiap kelompok hendaknya berasal dari populasi yang sama dengan variansi yang sama pula. Bila banyaknya sampel sama pada setiap kelompok maka kesamaan variansinya dapat diabaikan. Tapi bila banyak sampel pada masing-masing kelompok tidak sama maka kesamaan variansi populasi sangat diperlukan.
3. Pengamatan bebas
Sampel hendaknya diambil secara acak (*random*), sehingga setiap pengamatan merupakan informasi yang bebas.

Anova lebih akurat digunakan untuk sejumlah sampel yang sama pada setiap kelompoknya, misalnya masing-masing variabel setiap kelompok jumlah sampel atau responden sama-sama 250 orang.

Anova dapat digolongkan ke dalam beberapa kriteria, yaitu :

1. Klasifikasi 1 arah (*One Way ANOVA*)
Anova klasifikasi 1 arah merupakan ANOVA yang didasarkan pada pengamatan 1 kriteria atau satu faktor yang menimbulkan variasi.
2. Klasifikasi 2 arah (*Two Way ANOVA*)
ANOVA klasifikasi 2 arah merupakan ANOVA yang didasarkan pada pengamatan 2 kriteria atau 2 faktor yang menimbulkan variasi.
3. Klasifikasi banyak arah (*MANOVA*)
ANOVA banyak arah merupakan ANOVA yang didasarkan pada pengamatan banyak kriteria.

Anova Satu Arah (*One Way Anova*)

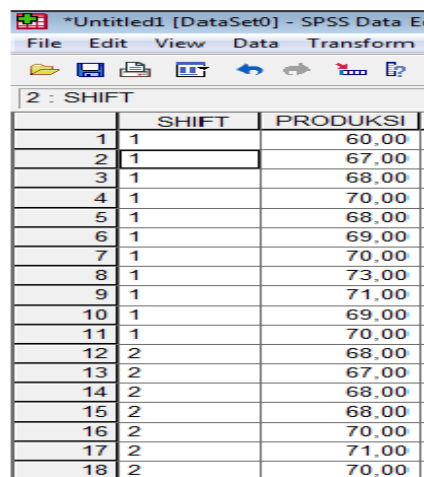
Anova satu arah (*one way anova*) digunakan apabila yang akan dianalisis terdiri dari **satu variabel terikat dan satu variabel bebas**. Interaksi suatu kebersamaan antar faktor dalam mempengaruhi variabel bebas, dengan sendirinya pengaruh faktor-faktor secara mandiri telah dihilangkan. Jika terdapat interaksi berarti efek faktor satu terhadap variabel terikat akan mempunyai garis yang tidak sejajar dengan efek faktor lain terhadap variabel terikat sejajar (saling berpotongan), maka antara faktor tidak mempunyai interaksi.

Pengolahan Data dengan Software

Dalam pengujian data ANOVA 1 arah dengan menggunakan software diperlukan software penunjang, yaitu program SPSS. Dalam pengujian kasus ANOVA 1 arah dengan menggunakan program SPSS, penyelesaian untuk pemecahan suatu masalah adalah sebagai berikut :

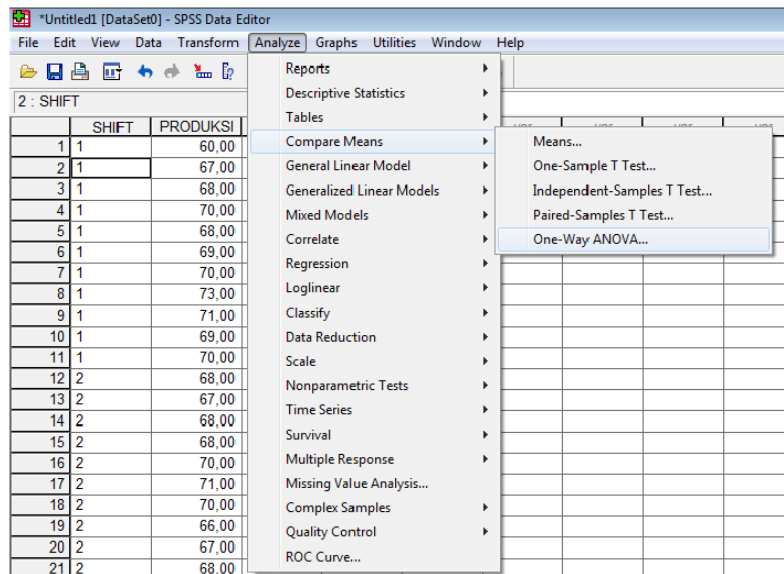
	A	B	C
1	PRODUKSI		
2	Shift Pagi	Shift Siang	Shift Malam
3	60	68	63
4	67	67	64
5	68	68	65
6	70	68	64
7	68	70	66
8	69	71	67
9	70	70	65
10	73	66	70
11	71	67	64
12	69	68	69
13	70	68	68

1. Memasukan data yang telah tersedia kedalam input data seperti gambar berikut. (terlebih dahulu isi bagian **Variabel View** seperti yang telah diajarkan pada penugasan sebelumnya) :

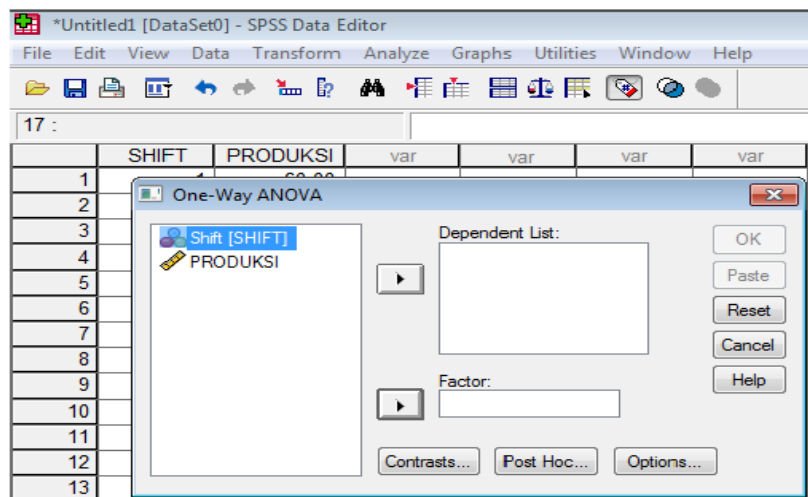


2 : SHIFT		
	SHIFT	PRODUKSI
1	1	60,00
2	1	67,00
3	1	68,00
4	1	70,00
5	1	68,00
6	1	69,00
7	1	70,00
8	1	73,00
9	1	71,00
10	1	69,00
11	1	70,00
12	2	68,00
13	2	67,00
14	2	68,00
15	2	68,00
16	2	70,00
17	2	71,00
18	2	70,00

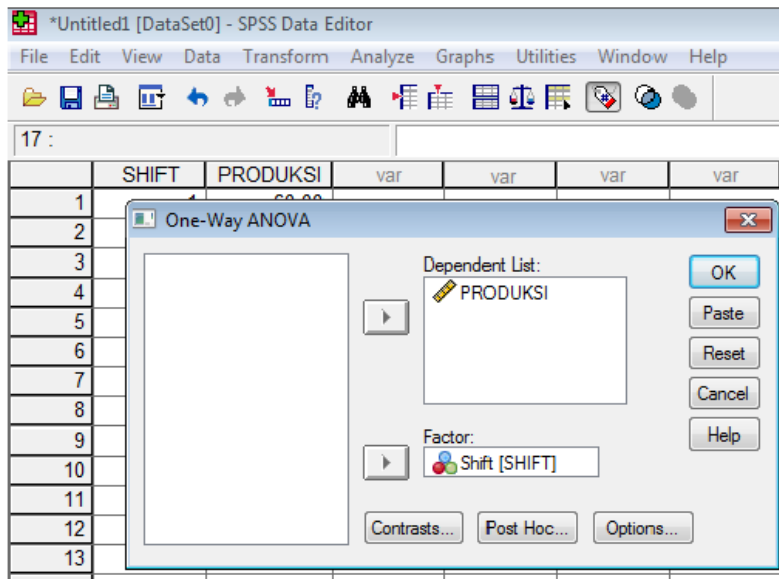
2. Melakukan setting analisis data sebagai berikut :
 - a. Pilih **analyze** pada menu file yang ada, pilih **compare mean** → **One Way Anova**



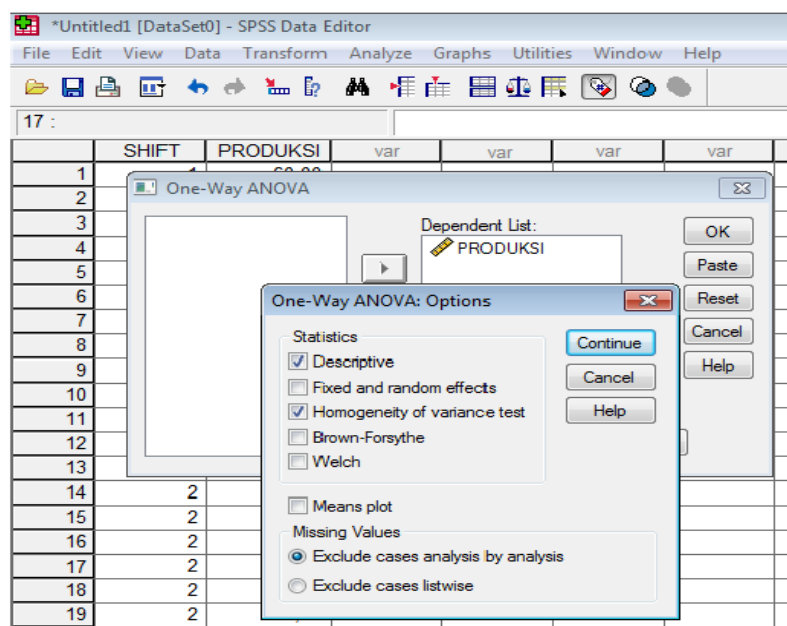
Setelah itu maka akan tampil gambar sebagai berikut :



- c. Pada Posisi **Dependent List** masukkan variabel yang menjadi variabel terikat. Dari data yang ada maka variabel terikatnya adalah variabel tingkat produksi, maka pilih tingkat penjualan.
- d. Pada Posisi **faktor** pilih variabel yang menjadi faktor penyebab terjadinya perubahan pada variabel terikat. Dalam hal ini adalah variabel shift. Sehingga akan berubah menjadi seperti ini :

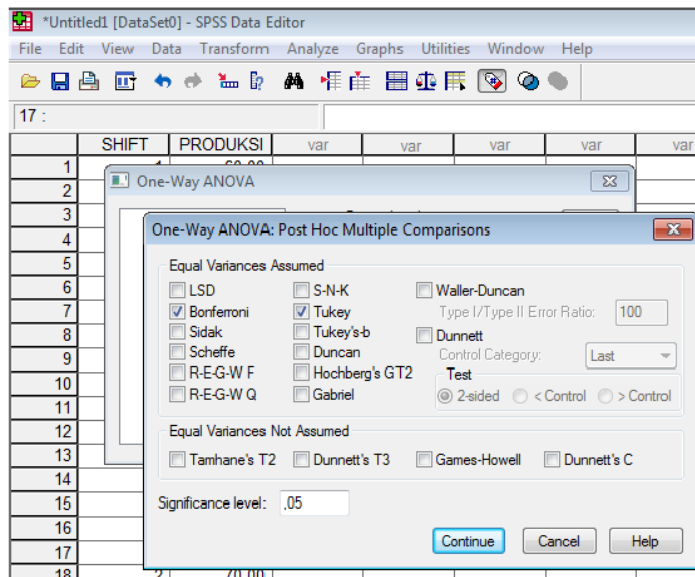


e. Klik tombol **options** dan klik pilihan yang diinginkan seperti berikut :



Untuk melihat keseragaman pada perhitungan statistik, maka dipilih **Descriptive** dan **Homogeneity-of-variance**. Untuk itu klik mouse pada pilihan tersebut. **Missing Value** adalah data yang hilang, karena data yang dianalisis tidak ada yang hilang, maka abaikan saja pilihan ini, kemudian klik **continue**.

Klik **post hoc** dan pilih jenis **post hoc** yang diinginkan.



Klik **Tukey** dan **Bonferroni** perhatikan **significance level** yang digunakan. Pada gambar diatas tertulis 0,05. Hal itu dikarenakan α sebesar 5%. Kemudian klik **Continue** jika pengisian dianggap selesai. Beberapa saat kemudian akan keluar tampilan *output* SPSS sebagai berikut :

Descriptives

PRODUKSI								
	N	Mean	Std. Deviation	Std. Error	95% Confidence Interval for Mean			
					Lower Bound	Upper Bound	Minimum	Maximum
1	11	68,6364	3,29462	,99337	66,4230	70,8497	60,00	73,00
2	11	68,2727	1,48936	,44906	67,2722	69,2733	66,00	71,00
3	11	65,9091	2,30020	,69354	64,3638	67,4544	63,00	70,00
Total	33	67,6061	2,69188	,46860	66,6516	68,5606	60,00	73,00

Test of Homogeneity of Variances

PRODUKSI			
Levene Statistic	df1	df2	Sig.
1,075	2	30	,354

Analisis Output :

1. **Output Descriptives**

Output Descriptives memuat hasil-hasil data statistic deskriptif seperti *mean*, standard deviasi, angka terendah dan tertinggi serta standard error. Pada bagian ini terlihat ringkasan statistik dari ketiga sampel.

2. *Output Test of Homogeneity of Variances*

Tes ini bertujuan untuk menguji berlaku tidaknya asumsi untuk Anova, yaitu apakah kelima sampel mempunyai varians yang sama. Untuk mengetahui apakah asumsi bahwa ketiga kelompok sampel yang ada mempunyai varian yang sama (homogen) dapat diterima. Untuk itu sebelumnya perlu dipersiapkan hipotesis tentang hal tersebut.

Adapun hipotesisnya adalah sebagai berikut :

H_0 = Ketiga variansi populasi adalah sama

H_1 = Ketiga variansi populasi adalah tidak sama

Dengan pengambilan Keputusan:

a) Jika signifikan > 0.05 maka H_0 diterima

b) Jika signifikan $< 0,05$ maka H_0 ditolak

Berdasarkan pada hasil yang diperoleh pada *test of homogeneity of variances*, dimana dihasilkan bahwa probabilitas atau signifikanya adalah 0,354 yang berarti lebih besar dari 0.05 maka dapat disimpulkan bahwa hipotesis nol (H_0) diterima, yang berarti asumsi bahwa ketiga varian populasi adalah sama (*homogeny*) dapat diterima.

3. *Output Anova*

Setelah kelima varians terbukti sama, baru dilakukan uji Anova untuk menguji apakah kelima sampel mempunyai rata-rata yang sama. Outpun Anova adalah akhir dari perhitungan yang digunakan sebagai penentuan analisis terhadap hipotesis yang akan diterima atau ditolak. Dalam hal ini hipotesis yang akan diuji adalah :

H_0 = Tidak ada perbedaan rata-rata hasil penjualan dengan menggunakan jenis kemasan yang berbeda. (Sama)

H_1 = Ada perbedaan rata-rata hasil penjualan dengan menggunakan jenis kemasan yang berbeda. (Tidak Sama)

Untuk menentukan H_0 atau H_a yang diterima maka ketentuan yang harus diikuti adalah sebagai berikut :

a) Jika $F_{hitung} > F_{tabel}$ maka H_0 ditolak

b) Jika $F_{hitung} < F_{tabel}$ maka H_0 diterima

c) Jika signifikan atau probabilitas > 0.05 , maka H_0 diterima

d) Jika signifikan atau probabilitas $< 0,05$, maka H_0 ditolak

ANOVA

PRODUKSI

	Sum of Squares	df	Mean Square	F	Sig.
Between Groups	48,242	2	24,121	3,941	,030
Within Groups	183,636	30	6,121		
Total	231,879	32			

Berdasarkan pada hasil yang diperoleh pada uji ANOVA, dimana dilihat bahwa F hitung = $F_{hitung} = 3,941$, yang berarti H_0 ditolak dan menerima H_a .

Sedangkan untuk nilai probabilitas dapat dilihat bahwa nilai probabilitas adalah $0,030 < 0,05$. Dengan demikian hipotesis nol (H_0) ditolak.

Hal ini menunjukkan bahwa ada perbedaan rata-rata hasil produksi dengan shift pagi, siang dan malam.

4. Output Tes Pos Hoc

Post Hoc dilakukan untuk mengetahui kelompok mana yang berbeda dan yang tidak berbeda. Hal ini dapat dilakukan bila F hitungnya menunjukkan ada perbedaan. Kalau F hitung menunjukkan tidak ada perbedaan, analisis sesudah anova tidak perlu dilakukan.

Multiple Comparisons

Dependent Variable: PRODUKSI

	(I) Shift	(J) Shift	Mean Difference (I-J)	Std. Error	Sig.	95% Confidence Interval	
						Lower Bound	Upper Bound
Tukey HSD	1	2	,36364	1,05496	,937	-2,2371	2,9644
		3	2,72727*	1,05496	,038	,1265	5,3280
	2	1	-,36364	1,05496	,937	-2,9644	2,2371
		3	2,36364	1,05496	,081	-,2371	4,9644
	3	1	-2,72727*	1,05496	,038	-5,3280	-,1265
		2	-2,36364	1,05496	,081	-4,9644	,2371
Bonferroni	1	2	,36364	1,05496	1,000	-2,3115	3,0388
		3	2,72727*	1,05496	,045	,0522	5,4024
	2	1	-,36364	1,05496	1,000	-3,0388	2,3115
		3	2,36364	1,05496	,098	-,3115	5,0388
	3	1	-2,72727*	1,05496	,045	-5,4024	-,0522
		2	-2,36364	1,05496	,098	-5,0388	,3115

*. The mean difference is significant at the .05 level.

PRODUKSI

	Shift	N	Subset for alpha = .05	
			1	2
Tukey HSD ^a	3	11	65,9091	
	2	11	68,2727	68,2727
	1	11		68,6364
	Sig.		,081	,937

Means for groups in homogeneous subsets are displayed.

a. Uses Harmonic Mean Sample Size = 11,000.

Dari tabel diatas dapat dilihat bahwa perbedaan *mean* Shift 1 dan Shift 2 adalah 0,3636 (rata-rata lebih kecil banyak 0,3636 poin dibanding shift 2). Angka tersebut berasal dari *mean* shift 1 adalah 68,6364 dan shift 2 adalah 68,2727 sehingga didapatkan 0,3636 (lihat *output* descriptive statistics). Perbedaan *mean* shift 1 dan shift 3 adalah 2,727 (shift 1 lebih besar 2,727 dari shift 3). Angka tersebut berasal dari *mean* shift 1 adalah 68,6364 dan shift 3 adalah 65,9 sehingga didapatkan 2,727. Untuk selanjutnya dapat dilihat gambar diatas untuk perbandingan shiftseterusnya.

Catatan :

Hasil uji signifikansi dengan mudah bisa dilihat pada *output* dengan ada atau tidak adanya tanda “*” pada kolom “*Mean Difference*”. Jika tanda * ada di angka *meandifference* maka perbedaan tersebut nyata atau signifikan. Jika tidak ada tanda *, maka perbedaan tidak signifikan.

Interpretasi :

- Shift yang paling baik untuk meningkatkan produksi adalah shift 1. Hal ini dapat dilihat dari jumlah rata-rata tertinggi pada shift 1. Sedangkan yang kurang baik dalam meningkatkan produksi adalah shift 3.
- Ada perbedaan tingkat produksi pada shift 1 dan shift 3, dan tidak ada perbedaan tingkat produksi pada shift 1 dan shift 2, shift 2 dan shift 3.
- Ada pengaruh yang signifikan antara produksi pada shift 1 dan shift 3.

Contoh 2 :

Uji anova satu arah akan digunakan untuk mengetahui adakah hubungan antara tingkat stress mahasiswa pada tiap kelompok Fakultas di Universitas Tugu Muda (UNTUMU). Tingkat stress diukur pada skala 1-10. Skala 1 hingga 3 menunjukkan mahasiswa cukup stress. Skala 4 sampai 6 menunjukkan mahasiswa dalam keadaan stress dan skala 7 keatas menunjukkan mahasiswa sangat stress. Pengamatan dilakukan pada waktu yang berbeda dengan menggunakan metode pengumpulan data yaitu kuisioner yang disebarakan pada 75 responden.

Tabel 6.1.
Tabel rekapitulasi tingkat stress mahasiswa tiap kelompok jurusan
yang ada di Fakultas dilingkungan UNTUMU

Pengamatan	Fakultas				
	Ekonomi	Hukum	ISIPOL	Teknik	Pertanian
1	4	4	1	4	1
	6	3	2	7	4
	2	2	3	9	5
	8	1	5	5	4
	8	8	2	4	7
2	2	9	1	2	8
	2	5	9	1	8
	3	3	8	1	7
	4	1	4	4	7
	5	5	7	7	7
3	6	7	5	9	5
	2	9	1	9	6
	1	6	3	2	7
	9	7	2	1	3
	8	3	5	4	4

1. Hipotesis

- H_0 : Semua rata – rata populasi fakultas sama, tidak ada hubungan antara tingkat stress dan fakultas di UNTUMU.
- H_1 : Tidak semua sama. beberapa atau semua rata – rata populasi fakultas sama, ada hubungan antara tingkat stress dan fakultas di UNTUMU.

2. Tingkat signifikansi

Dengan tingkat kepercayaan 95 persen maka tingkat signifikansi $(1 - \alpha) = 5$ persen atau sebesar 0,05.

3. Derajat kebebasan

Df jumlah kuadrat penyimpangan total = $N - 1$

Df jumlah kuadrat penyimpangan total = $75 - 1 = 74$

Df jumlah kuadrat dalam = $N - k$

Df jumlah kuadrat dalam = $75 - 5 = 70$

Df jumlah kuadrat antar kelompok = $k - 1$

Df jumlah kuadrat antar kelompok = $5 - 1 = 4$

4. Kriteria pengujian Untuk uji normalitas :

Signifikan atau probabilitas > 0.05 , maka data berdistribusi normal Signifikan atau probabilitas < 0.05 , maka data tidak berdistribusi normal

Untuk uji homogenitas :

Signifikan atau probabilitas > 0.05 , maka H_0 diterima Signifikan atau probabilitas < 0.05 , maka H_0 ditolak

Untuk uji ANOVA :

Jika signifikan atau probabilitas > 0.05 , maka H_0 diterima Jika signifikan atau probabilitas < 0.05 , maka H_0 ditolak

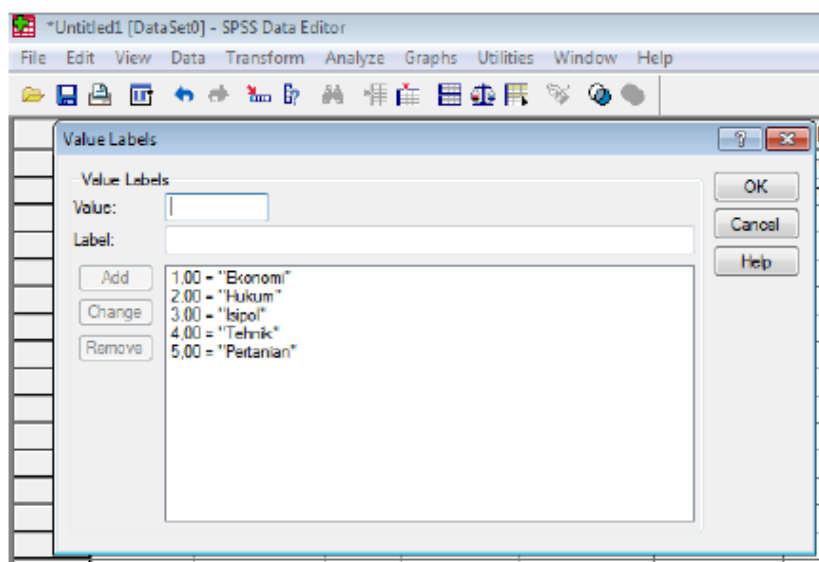
5. Pengolahan Data SPSS

a) Pengisian variabel

Pada kotak Name, sesuai kasus, ketik "stress" kemudian pada baris kedua ketik "**fakultas**" Pada Kotak Label variabel jurusan isi dengan "**tingkat stress**" dan pada kotak label variabel responden isi dengan "jurusan".

Klik Values dua kali untuk variabel "fakultas"		
o	Values : 1 ; Label : Ekonomi	Add
o	Values : 2 ; Label : Hukum	Add
o	Values : 3 ; Label : ISIPOL	Add
o	Values : 4 ; Label : Teknik	Add
o	Values : 5 ; Label : Pertanian	Add

Klik Ok



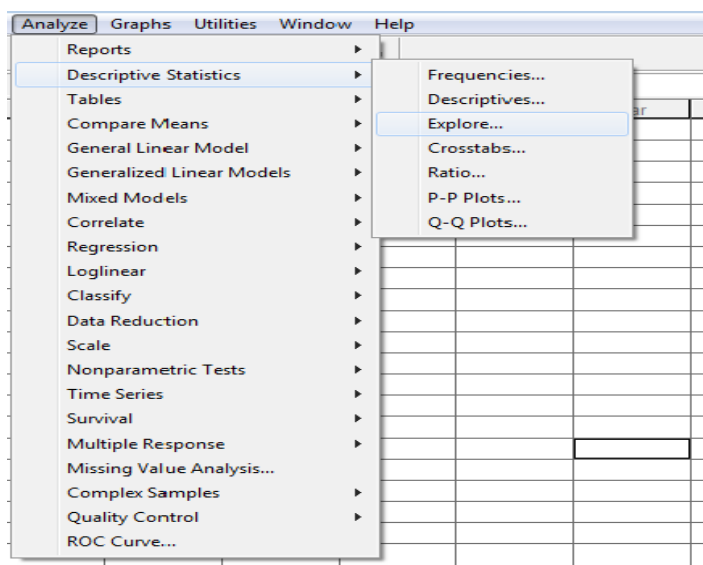
b. Pengisian DATA VIEW

Masukkan data mulai dari data ke-1 sampai dengan data ke-75.

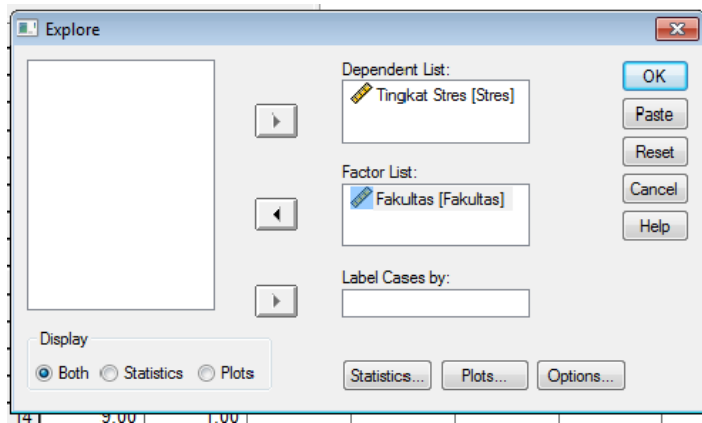
A	B	C
	Stress	Fakultas
	4	1
	6	1
	2	1
	8	1
	8	1
	2	1
	2	1
	3	1
	4	1
	5	1
	6	1
	2	1
	1	1
	9	1
	8	1
	4	2
	3	2
	2	2
	1	2
	8	2

c. Uji normalitas

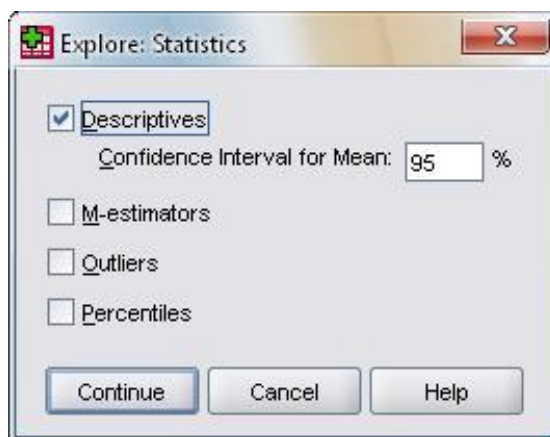
1) Menu **Analyze -> Descriptive statistics -> explore**



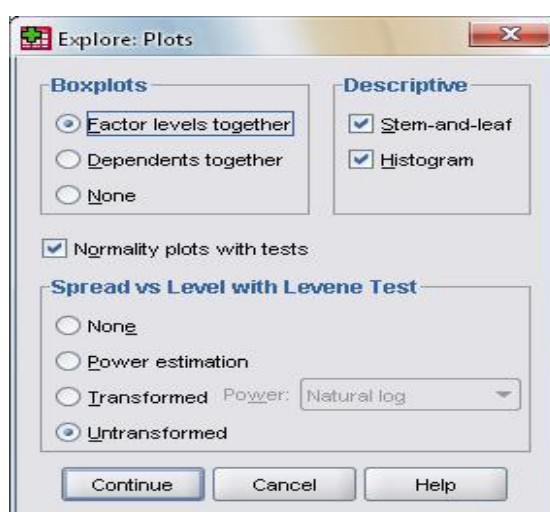
- 2) Masukkan variabel tingkat stress ke dependent list sebagai variabel terikat dan masukkan variabel jurusan ke faktor list sebagai variabel bebas, lalu klik Ok



- 3) Pada pilihan **Statistics**, isi *confidence interval for mean* dengan 95 % yang menandakan bahwa tingkat kepercayaan yang diambil sebesar 95 %. Lalu klik **continue**.



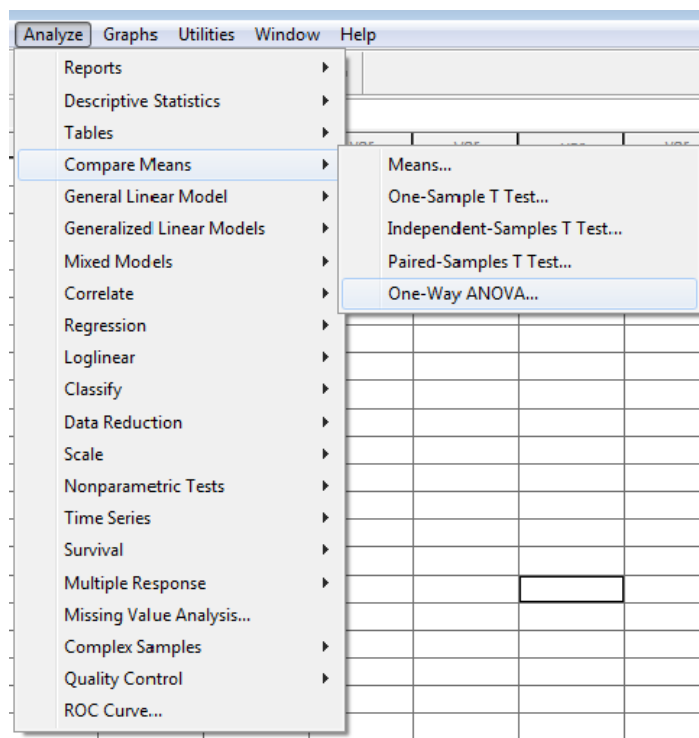
- 4) Pada pilihan **Plots**, tandai **normality plots with tests**, **histogram** pada **descriptive** dan **untransformed**. Lalu klik **continue**.



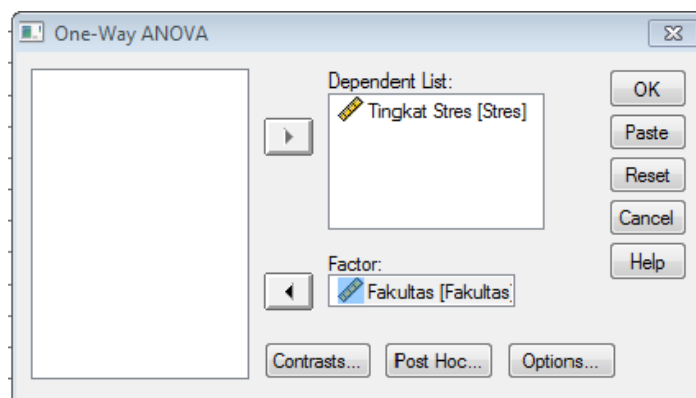
5) Klik Ok hingga muncul output SPSS.

d. Uji One Way ANOVA

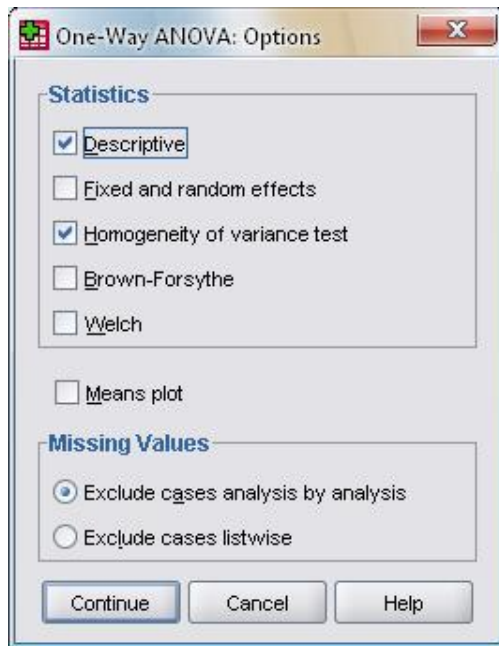
1) Menu **Analyze -> Compare means -> One way ANOVA**



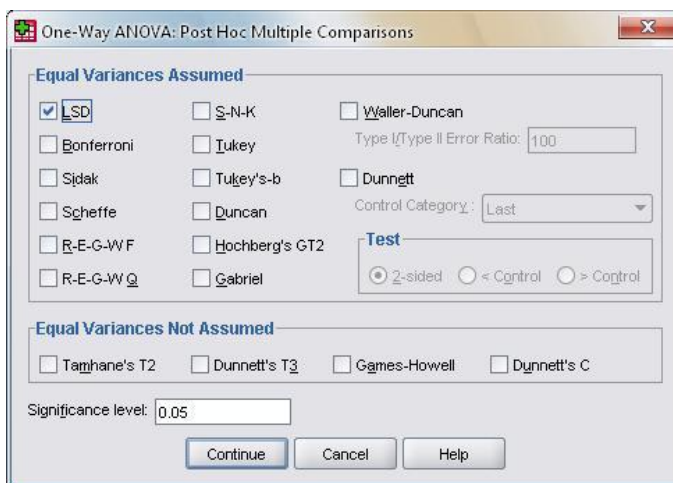
2) Masukkan variabel tingkat **tingkat stress** ke **dependent list** sebagai variabel terikat dan masukkan **variabel fakultas** ke **faktor** sebagai variabel bebas, lalu klik Ok.



3) Pada pilihan Options, tandai **descriptives** serta **homogeneity of variant tests** pada statistics. Lalu klik **continue**.



- 4) Pada pilihan **Post hoc**, tandai **LSD** pada equal variances assumed serta isi **significance level** berdasarkan tingkat signifikansi yang telah diberikan. Lalu klik continue.



- 5) Klik Ok hingga muncul output SPSS.

Hasil Output SPSS

a. Test of normality

Tests of Normality

	Fakultas	Kolmogorov-Smirnov(a)			Shapiro-Wilk		
		Statistic	df	Sig.	Statistic	df	Sig.
Tingkat Stres	Ekonomi	,173	15	,200(*)	,905	15	,113
	Hukum	,154	15	,200(*)	,938	15	,354
	Isipol	,165	15	,200(*)	,905	15	,112
	Teknik	,180	15	,200(*)	,886	15	,058
	Pertanian	,232	15	,030	,908	15	,125

* This is a lower bound of the true significance.

a Lilliefors Significance Correction

b. Test of homogeneity of variance

Test of Homogeneity of Variance

		Levene Statistic	df1	df2	Sig.
Tingkat Stres	Based on Mean	,729	4	70	,575
	Based on Median	,471	4	70	,757
	Based on Median and with adjusted df	,471	4	64,184	,757
	Based on trimmed mean	,722	4	70	,580

c. Anova

ANOVA

Tingkat Stres

	Sum of Squares	df	Mean Square	F	Sig.
Between Groups	21,413	4	5,353	,780	,542
Within Groups	480,133	70	6,859		
Total	501,547	74			

d. Post hoc

Multiple Comparisons

Dependent Variable: Tingkat Stres

	(I) Fakultas	(J) Fakultas	Mean Difference (I-J)	Std. Error	Sig.	95% Confidence Interval	
						Lower Bound	Upper Bound
Tukey HSD	Ekonomi	Hukum	-,20000	,95632	1,000	-2,8778	2,4778
		Isipol	,80000	,95632	,918	-1,8778	3,4778
		Tehnik	,06667	,95632	1,000	-2,6112	2,7445
		Pertanian	-,86667	,95632	,894	-3,5445	1,8112
	Hukum	Ekonomi	,20000	,95632	1,000	-2,4778	2,8778
		Isipol	1,00000	,95632	,833	-1,6778	3,6778
		Tehnik	,26667	,95632	,999	-2,4112	2,9445
		Pertanian	-,66667	,95632	,956	-3,3445	2,0112
	Isipol	Ekonomi	-,80000	,95632	,918	-3,4778	1,8778
		Hukum	-1,00000	,95632	,833	-3,6778	1,6778
		Tehnik	-,73333	,95632	,939	-3,4112	1,9445
		Pertanian	-1,66667	,95632	,415	-4,3445	1,0112
	Tehnik	Ekonomi	-,06667	,95632	1,000	-2,7445	2,6112
		Hukum	-,26667	,95632	,999	-2,9445	2,4112
		Isipol	,73333	,95632	,939	-1,9445	3,4112
		Pertanian	-,93333	,95632	,865	-3,6112	1,7445
	Pertanian	Ekonomi	,86667	,95632	,894	-1,8112	3,5445
		Hukum	,66667	,95632	,956	-2,0112	3,3445
		Isipol	1,66667	,95632	,415	-1,0112	4,3445
		Tehnik	,93333	,95632	,865	-1,7445	3,6112
Bonferroni	Ekonomi	Hukum	-,20000	,95632	1,000	-2,9721	2,5721
		Isipol	,80000	,95632	1,000	-1,9721	3,5721
		Tehnik	,06667	,95632	1,000	-2,7054	2,8388
		Pertanian	-,86667	,95632	1,000	-3,6388	1,9054
	Hukum	Ekonomi	,20000	,95632	1,000	-2,5721	2,9721
		Isipol	1,00000	,95632	1,000	-1,7721	3,7721
		Tehnik	,26667	,95632	1,000	-2,5054	3,0388
		Pertanian	-,66667	,95632	1,000	-3,4388	2,1054
	Isipol	Ekonomi	-,80000	,95632	1,000	-3,5721	1,9721
		Hukum	-1,00000	,95632	1,000	-3,7721	1,7721
		Tehnik	-,73333	,95632	1,000	-3,5054	2,0388
		Pertanian	-1,66667	,95632	,858	-4,4388	1,1054
	Tehnik	Ekonomi	-,06667	,95632	1,000	-2,8388	2,7054
		Hukum	-,26667	,95632	1,000	-3,0388	2,5054
		Isipol	,73333	,95632	1,000	-2,0388	3,5054
		Pertanian	-,93333	,95632	1,000	-3,7054	1,8388
	Pertanian	Ekonomi	,86667	,95632	1,000	-1,9054	3,6388
		Hukum	,66667	,95632	1,000	-2,1054	3,4388
		Isipol	1,66667	,95632	,858	-1,1054	4,4388
		Tehnik	,93333	,95632	1,000	-1,8388	3,7054

7. Analisis Hasil Output SPSS

a. *Test of normality*

Uji normalitas menunjukkan dari hasil keseluruhan tersebut dapat ditarik kesimpulan bahwa **signifikansi seluruh fakultas > 0,05 yang artinya distribusi data normal**. Maka data yang diambil dinyatakan tidak terjadi penyimpangan dan layak untuk dilakukan uji ANOVA.

b. *Test of homogeneity of variance*

Tes ini bertujuan untuk menguji berlaku tidaknya asumsi untuk ANOVA, yaitu apakah kelima kelompok sampel mempunyai variansi yang sama. Uji keseragaman variansi menunjukkan probabilitas atau signifikansi seluruh sampel adalah 0,58, yang berarti signifikansi = 0,58 > 0,05 maka sesuai dengan kriteria pengujian dapat disimpulkan bahwa **hipotesis nol (H_0) diterima, yang berarti asumsi bahwa kelima varian populasi adalah sama (homogen) dapat diterima**.

c. ANOVA

Setelah kelima varian terbukti sama, baru dilakukan uji ANOVA untuk menguji apakah kelima sampel mempunyai rata-rata yang sama. Uji ANOVA menunjukkan nilai probabilitas atau signifikansi adalah 0,542. Hal ini berarti signifikansi lebih besar dari 0.05 maka **H_0 juga diterima yang artinya ternyata tidak ada perbedaan rata-rata antara kelima kelompok fakultas yang diuji**. Maka tidak ada pengaruh tingkat stress terhadap kelompok **fakultas** yang ada di UNTUMU.

d. Post hoc

Post Hoc dilakukan untuk mengetahui kelompok mana yang berbeda dan yang tidak berbeda. Atau dapat dikatakan dalam kasus ini, kelompok jurusan mana yang memberikan pengaruh signifikan terhadap perbedaan tingkat stress. Uji post hoc merupakan uji kelanjutan dari uji ANOVA jika hasil yang diperoleh pada uji ANOVA adalah H_0 diterima atau terdapat perbedaan antara tiap kelompok. Namun karena uji ANOVA menunjukkan H_0 ditolak, maka otomatis uji post hoc menunjukkan tidak ada kelompok Fakultas di lingkungan UNTUMU yang memberikan pengaruh pada tingkat stress. Hal ini juga dapat dilihat pada tabel Post hoc yang tidak menunjukkan tanda (*) sebagai penanda bahwa terdapat kelompok yang signifikan.

8. Keputusan

Dari keseluruhan uji yang dilakukan maka dapat disimpulkan bahwa tidak terdapat pengaruh maupun perbedaan yang signifikan antara kelima fakultas yang ada di UNTUMU yang artinya tidak terdapat hubungan antara tingkat stress mahasiswa dengan kelompok Fakultas di Universitas Tugu Muda.

Anova Dua Arah (*Two Way Anova*)

ANOVA dua arah ini digunakan bila sumber keragaman yang terjadi tidak hanya karena satu faktor (perlakuan). Faktor lain yang mungkin menjadi sumber keragaman respon juga harus diperhatikan. Faktor lain ini bisa perlakuan lain atau faktor yang sudah terkondisi. Pertimbangan memasukkan faktor kedua sebagai sumber keragaman ini perlu bila faktor itu dikelompokkan (blok), sehingga keragaman antar kelompok sangat besar, tetapi kecil dalam kelompok sendiri.

Tujuan dan pengujian ANOVA 2 arah ini adalah untuk mengetahui apakah ada pengaruh dari berbagai kriteria yang diuji terhadap hasil yang diinginkan. Misal, seorang manajer teknik menguji apakah ada pengaruh antara jenis pelumas yang dipergunakan pada roda pendorong dengan kecepatan roda pendorong terhadap hasil penganyaman sebuah karung plastik pada mesin *circular*.

A. Pengolahan Menggunakan Software

Studi Kasus 1

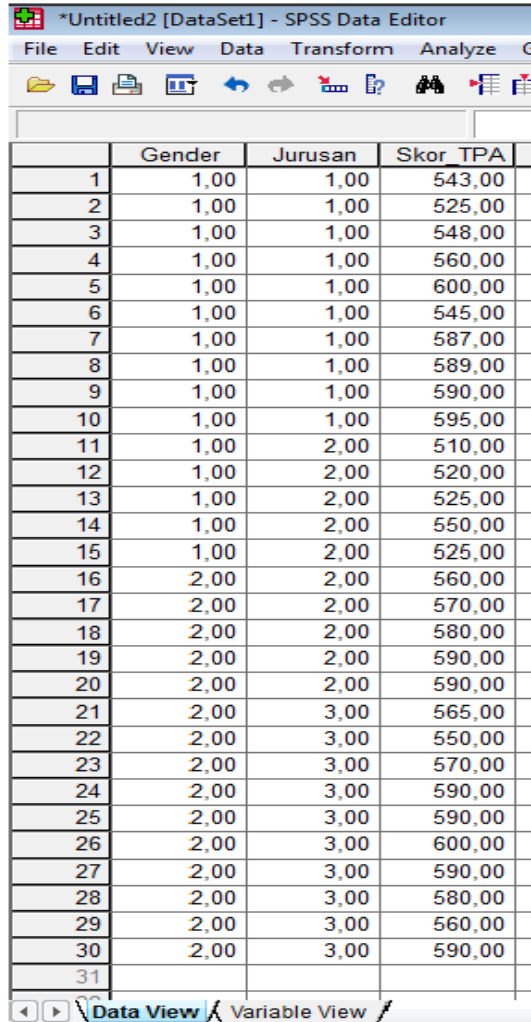
Ingin diketahui apakah jurusan dan gender mempengaruhi skor TPA mahasiswa. Didapat data sebagai berikut :

		Skor TPA	
		Laki-Laki	Perempuan
Jurusan	Ilmu Ekonomi	543	560
		525	570
		548	580
		560	590
		600	590
	Manajemen	545	565
		587	550
		589	570
		590	590
		595	590
	Akuntansi	510	600
		520	590
		525	580
		550	560
		525	590

Dalam pengujian kasus ANOVA 2 arah dengan menggunakan program SPSS untuk pemecahan suatu masalah adalah sebagai berikut:

1. Memasukan data ke SPSS

Hal yang perlu diperhatikan dalam pengisian variabel Name adalah “tidak boleh ada spasi dalam pengisiannya”.

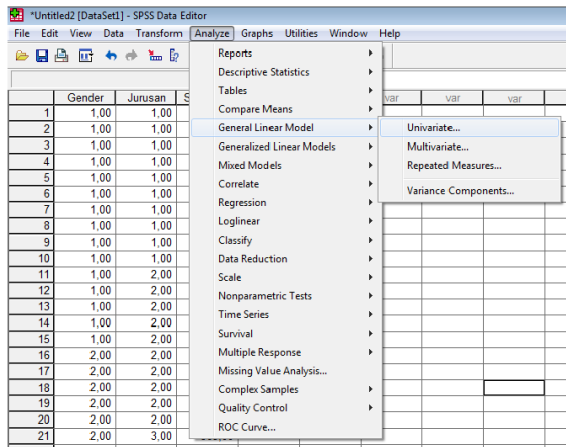


*Untitled2 [DataSet1] - SPSS Data Editor

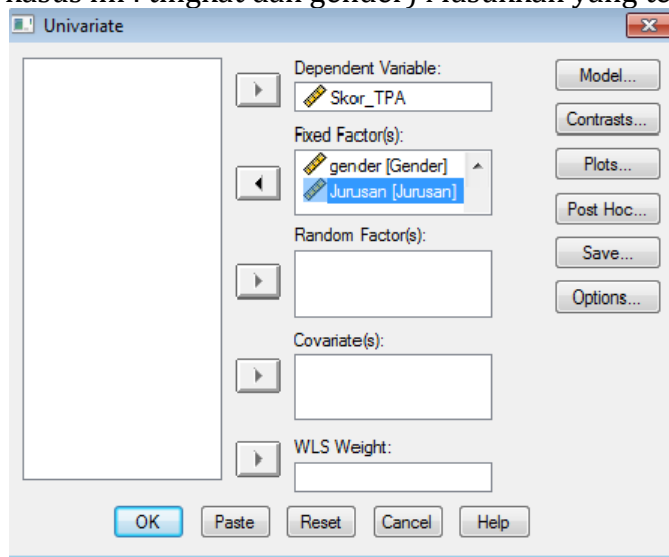
	Gender	Jurusan	Skor TPA
1	1,00	1,00	543,00
2	1,00	1,00	525,00
3	1,00	1,00	548,00
4	1,00	1,00	560,00
5	1,00	1,00	600,00
6	1,00	1,00	545,00
7	1,00	1,00	587,00
8	1,00	1,00	589,00
9	1,00	1,00	590,00
10	1,00	1,00	595,00
11	1,00	2,00	510,00
12	1,00	2,00	520,00
13	1,00	2,00	525,00
14	1,00	2,00	550,00
15	1,00	2,00	525,00
16	2,00	2,00	560,00
17	2,00	2,00	570,00
18	2,00	2,00	580,00
19	2,00	2,00	590,00
20	2,00	2,00	590,00
21	2,00	3,00	565,00
22	2,00	3,00	550,00
23	2,00	3,00	570,00
24	2,00	3,00	590,00
25	2,00	3,00	590,00
26	2,00	3,00	600,00
27	2,00	3,00	590,00
28	2,00	3,00	580,00
29	2,00	3,00	560,00
30	2,00	3,00	590,00
31			

2. Pengolahan data dengan SPSS Langkah-langkahnya :

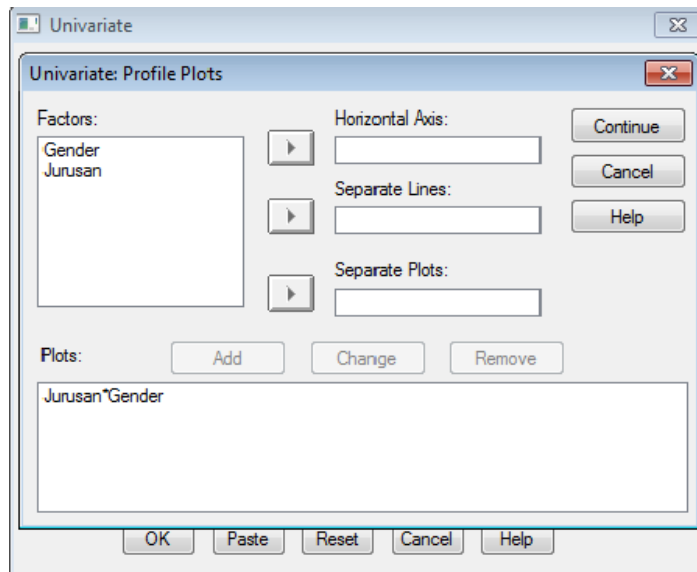
a. Pilih **Analyze** ----- **General Linear Model** ----- **Univariate**



- b. Kemudian lakukan pengisian terhadap : Kolom *Dependent Variable* masukan skor TPA, Kolom Faktor(s) Masukkan yang termasuk *Fixed Factor(s)* (dalam kasus ini : tingkat dan gender) Masukkan yang termasuk *Random Factor(s)*

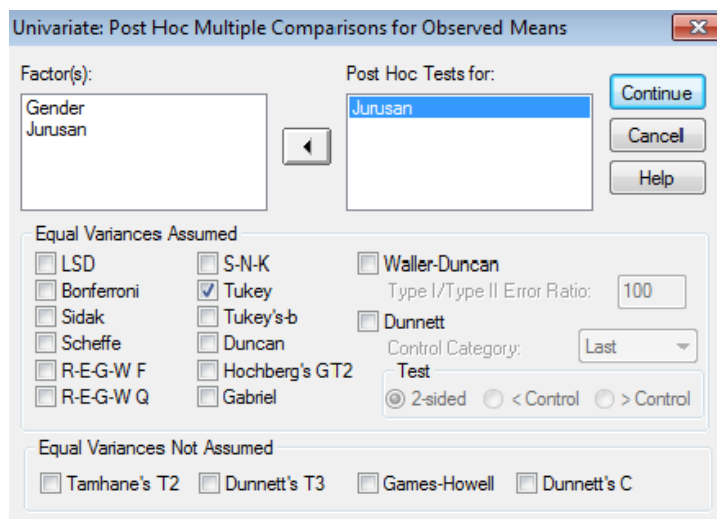


- c. Klik Plots
Horizontal Axis : ... (jurusan) Separate lines : ... (gender)

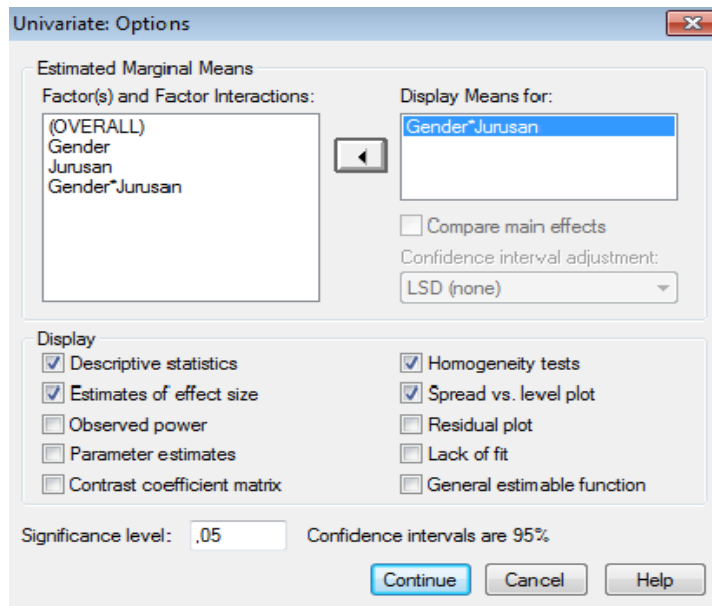


d. Klik Post Hoc

Masukan variabel yang akan di uji MCA ... (tingkat) ² Tukey



e. Options



f. Klik OK, diperoleh output :

Levene's Test of Equality of Error Variances^a

Dependent Variable: Skor TPA

F	df1	df2	Sig.
.586	5	24	.711

Tests the null hypothesis that the error variance of the dependent variable is equal across groups.

a. Design: Intercept+gender+jurusan+gender * jurusan

Tests of Between-Subjects Effects

Dependent Variable: Skor TPA

Source	Type III Sum of Squares	df	Mean Square	F	Sig.	Partial Eta Squared
Corrected Model	12321.767 ^a	5	2464.353	6.986	.000	.593
Intercept	9618605.633	1	9618605.633	27268.774	.000	.999
gender	4392.300	1	4392.300	12.452	.002	.342
jurusan	2444.067	2	1222.033	3.464	.048	.224
gender * jurusan	5485.400	2	2742.700	7.776	.002	.393
Error	8465.600	24	352.733			
Total	9639393.000	30				
Corrected Total	20787.367	29				

a. R Squared = .593 (Adjusted R Squared = .508)

Uji Interaksi

1. $H_0: \gamma_{ij}=0$ Tidak ada interaksi antara faktor jurusan dan gender
 $H_1: \gamma_{ij} \neq 0$ Ada interaksi antara jurusan dan gender
2. Tingkat Signifikasi $\alpha = 5\%$
3. Statistik Uji P-value = 0,02
(p_value diambil dari tabel dengan sig yang berasal dari source gender *jurusan)
4. Daerah Kritik
 H_0 ditolak jika $P\text{-value} < \alpha$
5. Kesimpulan
Karena $p\text{-value} (0,02) < \alpha (0,05)$ maka H_0 ditolak.
Jadi tidak ada interaksi antara faktor jurusan dengan faktor gender pada tingkat signifikasi 5%. Hal tersebut menyatakan bahwa uji efek untuk faktor gender dan jurusan bisa dilakukan.

Uji Efek faktor gender

1. $H_0: \alpha_1 = \alpha_2 = \dots = \alpha_i$ (Tidak ada efek faktor gender)
 H_1 : minimal ada satu $\alpha_i \neq 0$ (Ada efek faktor gender)
2. Tingkat Signifikasi $\alpha = 5\%$
3. Statistik Uji P-value = 0,002
(p_value diambil dari sig pada tabel dengan source gender)
4. Daerah Kritik
 H_0 ditolak jika $P\text{-value} < \alpha$
5. Kesimpulan
Karena $p\text{-value} (0,002) < \alpha (0,05)$ maka H_0 ditolak. Jadi ada efek faktor gender untuk data tersebut pada tingkat signifikasi 5% Karena faktor gender hanya terdiri dari 2 level faktor, sehingga tidak diperlukan uji MCA

Jurusan

1. $H_0: \alpha_1 = \alpha_2 = \dots = \alpha_i$ (Tidak ada efek faktor jurusan)
 H_1 : minimal ada satu $\alpha_i \neq 0$
2. Tingkat Signifikasi $\alpha = 5\%$
3. Statistik Uji P-value = 0,048 (p_value diambil dari tabel pada sig dengan source jurusan)
4. Daerah Kritik
 H_0 ditolak jika $P\text{-value} < \alpha$
5. Kesimpulan
Karena $p\text{-value} (0,048) < \alpha (0,05)$ maka H_0 ditolak.

Jadi ada efek faktor jurusan untuk data tersebut pada tingkat signifikasi 5% Karena faktor jurusan mempengaruhi SKOR secara signifikan, sehingga perlu dilakukan uji MCA Analisis perbandingan Ganda :

Multiple Comparisons

Dependent Variable: Skor TPA

Tukey HSD

(I) Jurusan	(J) Jurusan	Mean Difference (I-J)	Std. Error	Sig.	95% Confidence Interval	
					Lower Bound	Upper Bound
Ilmu_ekonomi	Manajemen	16,2000	8,96922	,187	-6,0876	38,4876
	Akuntansi	-10,3000	8,96922	,494	-32,5876	11,9876
Manajemen	Ilmu_ekonomi	-16,2000	8,96922	,187	-38,4876	6,0876
	Akuntansi	-26,5000*	8,96922	,017	-48,7876	-4,2124
Akuntansi	Ilmu_ekonomi	10,3000	8,96922	,494	-11,9876	32,5876
	Manajemen	26,5000*	8,96922	,017	4,2124	48,7876

Based on observed means.

*. The mean difference is significant at the ,05 level.

Skor_TPA

Tukey HSD^{a,b}

Jurusan	N	Subset	
		1	2
Manajemen	10	552,0000	568,2000 578,5000
Ilmu_ekonomi	10	568,2000	
Akuntansi	10		
Sig.		,187	,494

Means for groups in homogeneous subsets are displayed.

Based on Type III Sum of Squares

The error term is Mean Square(Error) = 402,235.

a. Uses Harmonic Mean Sample Size = 10,000.

b. Alpha = ,05.

Dapat juga disimpulkan bahwa terdapat perbedaan Skor TPA yang signifikan antara mahasiswa Akuntansi dan Manajemen. Sedangkan antara jurusan Akuntansi dan jurusan Ilmu Ekonomi serta Jurusan Ilmu ekonomi dengan jurusan manajemen menunjukkan tidak adanya perbedaan yang signifikan dalam hal Skor TPA.

Studi Kasus 2

Sebuah pabrik selama ini memperkerjakan karyawannya dalam 3 shift (satu shift terdiri atas sekelompok pekerja yang berlainan). Manajer pabrik tersebut ingin mengetahui apakah ada perbedaan produktifitas yang nyata diantara 3 kelompok kerja shift yang ada selama ini. Selama ini setiap kelompok kerja terdiri atas wanita semua atau pria semua, dan setelah kelompok pria bekerja dua hari berturut-turut, ganti kelompok wanita (tetap terbagi tiga kelompok) yang bekerja. Demikian seterusnya, dua hari untuk pria dan dua hari untuk wanita.

Tabel 6.2.
Berikut hasil pengamatan (angka dalam unit)

Hari	Shift 1	Shift 2	Shift 3	Gender
1	38	45	45	Pria
2	36	48	48	Pria
3	39	42	42	Wanita
4	34	46	46	Pria
5	35	41	41	Pria
6	32	45	45	Wanita
7	39	48	48	Pria
8	34	47	47	Pria
9	32	42	42	Wanita
10	36	41	41	Pria
11	33	39	39	Pria
12	39	33	33	Wanita

Nb : pada baris 1, di hari pertama kelompok shift 1 berproduksi 38 unit, kelompok shift 2 berproduksi 45 unit, kelompok shift 3 berproduksi 45 unit, dengan catatan semua anggota kelompok pria. Demikian untuk data yang lain.

Dalam pengujian kasus ANOVA 2 arah dengan menggunakan program SPSS ver 15.0, penyelesaian untuk pemecahan suatu masalah adalah sebagai berikut:

1. Memasukkan Data ke SPSS

Tabel pada kasus di atas harus kita dirubah dalam format berikut ini jika akan digunakan dalam uji ANOVA dengan SPSS

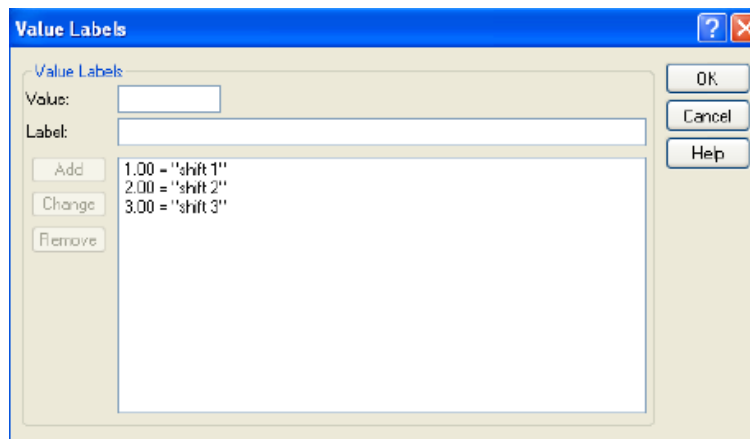
Produk	Shift	Gender
38	Satu	Pria
36	Satu	Pria
39	Satu	Wanita
34	Satu	Pria
35	Satu	Pria
32	Satu	Wanita
39	Satu	Pria
34	Satu	Pria
32	Satu	Wanita
36	Satu	Pria
33	Satu	Pria

Produk	Shift	Gender
39	Satu	Wanita
45	Dua	Pria
48	Dua	Pria
42	Dua	Wanita
46	Dua	Pria
41	Dua	Pria
45	Dua	Wanita
48	Dua	Pria
47	Dua	Pria
42	Dua	Wanita
41	Dua	Pria
39	Dua	Pria
33	Dua	Wanita
45	Tiga	Pria
48	Tiga	Pria
42	Tiga	Wanita
46	Tiga	Pria
41	Tiga	Pria
45	Tiga	Wanita
48	Tiga	Pria
47	Tiga	Pria
42	Tiga	Wanita
41	Tiga	Pria
39	Tiga	Pria
33	Tiga	Wanita

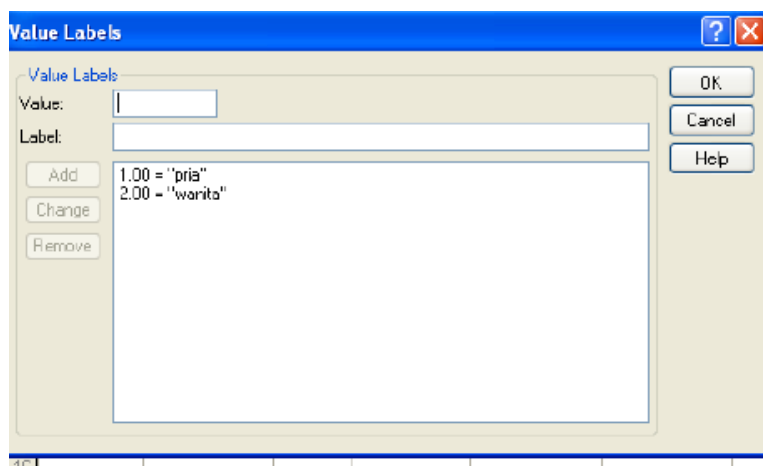
Langkah-langkah :

- a. Dari menu utama **file**, pilih menu **new**, lalu klik **Data**. Kemudian klik pada sheet tab **Variabel View**.
 Pengisian variable PRODUK
 o Name, sesuai kasus, ketik Produk
 Pengisian Variabel SHIFT
 o Name sesuai kasus ketik Shift
 Values, pilihan ini untuk proses pemberian kode, dengan isian :

Kode	Label
1	Satu
2	Dua
3	Tiga



Pengisian Variabel Gender

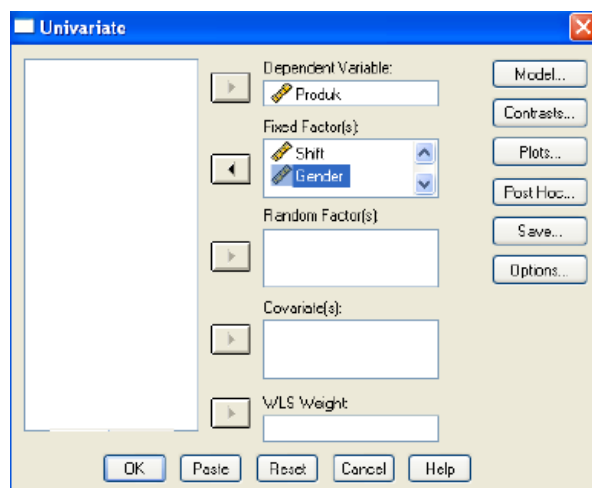
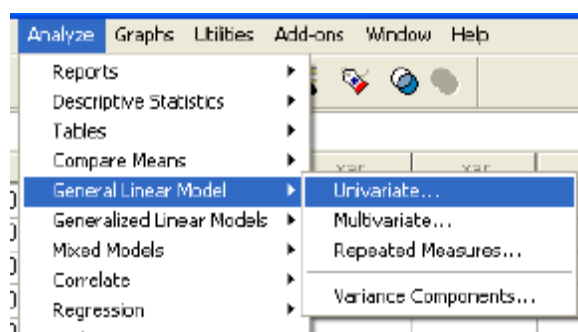


- b. Abaikan bagian yang lain kemudian tekan CTRL+T untuk pindah ke **DATA VIEW**
- c. Mengisi Data
 1. Isikan data sesuai data pada table
 2. Aktifkan value label dengan menu View kemudian klik Value Label

	Produk	Shift	Gender
1	38.00	1.00	1.00
2	36.00	1.00	1.00
3	39.00	1.00	2.00
4	34.00	1.00	1.00
5	35.00	1.00	1.00
6	32.00	1.00	2.00
7	39.00	1.00	1.00
8	34.00	1.00	1.00
9	32.00	1.00	2.00
10	36.00	1.00	1.00
11	33.00	1.00	1.00
12	39.00	1.00	2.00
13	45.00	2.00	1.00
14	48.00	2.00	1.00
15	42.00	2.00	2.00
16	46.00	2.00	1.00

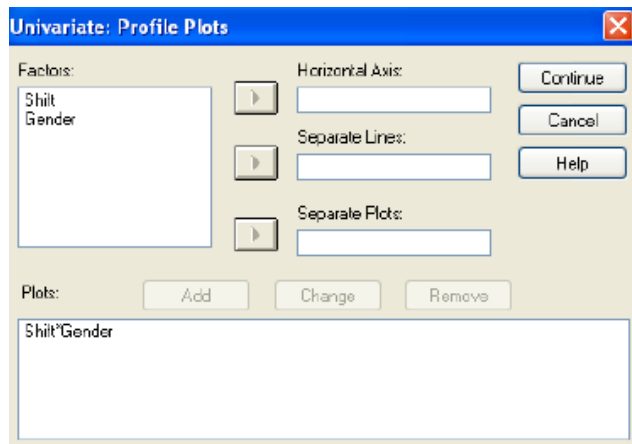
d. Pengolahan Data SPSS

1. Pilih menu **Analyze**, pilih **General-Linear Model**, ketik **Univariate**. Untuk pengisiannya sesuaikan dengan gambar dibawah ini :



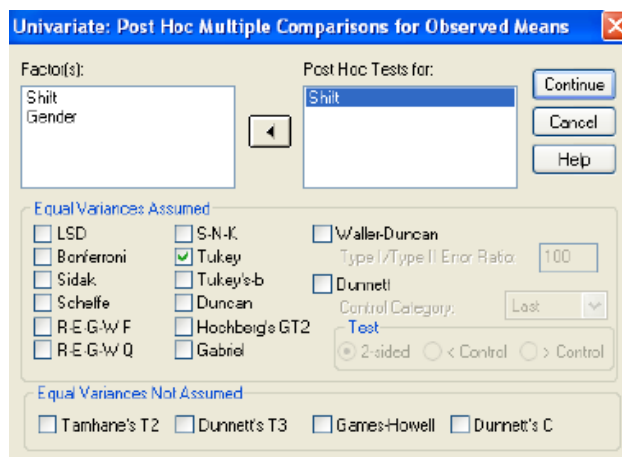
Klik Plots

o Horizontal Axis : ... (Shift) o Separate lines : ... (gender) o Add; Shift*Gender



Klik Post Hoc

Masukan variabel yang akan di uji MCA ... (tingkat) o Tukey



Options

Univariate: Options

Estimated Marginal Means
 Factor(s) and Factor Interactions: (OVERALL), Shift, Gender, Shift*Gender
 Display Means for: Shift*Gender
☐ Compare main effects
 Confidence interval adjustment: LSD (none)

Display
☒ Descriptive statistics
☒ Estimates of effect size
☐ Observed power
☐ Parameter estimates
☐ Contrast coefficient matrix
☒ Homogeneity tests
☒ Spread vs. level plot
☐ Residual plot
☐ Lack of fit
☐ General estimable function

Significance level: .05 Confidence intervals are 95%

Continue Cancel Help

Klik Ok

Tests of Between-Subjects Effects

Dependent Variable: Produk

Source	Type III Sum of Squares	df	Mean Square	F	Sig.	Partial Eta Squared
Corrected Model	530.125 ^a	5	106.025	7.635	.000	.560
Intercept	51574.014	1	51574.014	3713.700	.000	.992
Shift	336.111	2	168.056	12.101	.000	.447
Gender	55.125	1	55.125	3.969	.056	.117
Shift * Gender	25.000	2	12.500	.900	.417	.057
Error	416.625	30	13.887			
Total	60239.000	36				
Corrected Total	946.750	35				

a. R Squared = .560 (Adjusted R Squared = .487)

Estimated Marginal Means

Shift * Gender

Dependent Variable: Produk

Shift	Gender	Mean	Std. Error	95% Confidence Interval	
				Lower Bound	Upper Bound
shift 1	pria	35.625	1.318	32.934	38.316
	wanita	35.500	1.863	31.695	39.305
shift 2	pria	44.375	1.318	41.684	47.066
	wanita	40.500	1.863	36.695	44.305
shift 3	pria	44.375	1.318	41.684	47.066
	wanita	40.500	1.863	36.695	44.305

Post Hoc Tests

Multiple Comparisons

Dependent Variable: Produk

Tukey HSD

(I) Shift	(J) Shift	Mean Difference (I-J)	Std. Error	Sig.	95% Confidence Interval	
					Lower Bound	Upper Bound
shift 1	shift 2	-7.5000*	1.52138	.000	-11.2506	-3.7494
	shift 3	-7.5000*	1.52138	.000	-11.2506	-3.7494
shift 2	shift 1	7.5000*	1.52138	.000	3.7494	11.2506
	shift 3	.0000	1.52138	1.000	-3.7506	3.7506
shift 3	shift 1	7.5000*	1.52138	.000	3.7494	11.2506
	shift 2	.0000	1.52138	1.000	-3.7506	3.7506

Based on observed means.

*. The mean difference is significant at the .05 level.

Homogeneous Subsets

Produk

Tukey HSD^{a,b}

Shift	N	Subset	
		1	2
shift 1	12	35.5833	
shift 2	12		43.0833
shift 3	12		43.0833
Sig.		1.000	1.000

Means for groups in homogeneous subsets are displayed.

Based on Type III Sum of Squares

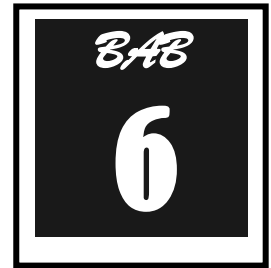
The error term is Mean Square(Error) = 13.887.

a. Uses Harmonic Mean Sample Size = 12.000.

b. Alpha = .05.

Analisis

Berdasarkan output diatas, tampak bahwa shift 2 dan shift 3 tidak terdapat perbedaan produksi yang signifikan, tetapi memiliki perbedaan yang signifikan apabila dibandingkan dengan shift 1.



UJI VALIDITAS DAN REALIBILITAS

Dalam penelitian, data mempunyai kedudukan yang paling tinggi, karena data merupakan penggambaran variabel yang diteliti dan berfungsi sebagai alat pembuktian hipotesis. Benar tidaknya data, sangat menentukan bermutu tidaknya hasil penelitian. Sedang benar tidaknya data, tergantung dari baik tidaknya instrumen pengumpulan data. Pengujian instrumen biasanya terdiri dari uji validitas dan reliabilitas.

A. Definisi Validitas dan Reliabilitas

Validitas adalah tingkat keandalan dan kesahihan alat ukur yang digunakan. Instrumen dikatakan valid berarti menunjukkan alat ukur yang dipergunakan untuk mendapatkan data itu valid atau dapat digunakan untuk mengukur apa yang seharusnya diukur (Sugiyono, 2004:137). Dengan demikian, instrumen yang valid merupakan instrumen yang benar-benar tepat untuk mengukur apa yang hendak diukur.

Penggaris dinyatakan valid jika digunakan untuk mengukur panjang, namun tidak valid jika digunakan untuk mengukur berat. Artinya, penggaris memang tepat digunakan untuk mengukur panjang, namun menjadi tidak valid jika penggaris digunakan untuk mengukur berat.

Uji reliabilitas berguna untuk menetapkan apakah instrumen yang dalam hal ini kuesioner dapat digunakan lebih dari satu kali, paling tidak oleh responden yang sama akan menghasilkan data yang konsisten. Dengan kata lain, reliabilitas instrumen mencirikan tingkat konsistensi. Banyak rumus yang dapat digunakan untuk mengukur reliabilitas diantaranya adalah rumus **Spearman Brown** :

$$r_{11} = \frac{2r_b}{1 + r_b}$$

Keterangan :

R_{11} adalah nilai reliabilitas

R_b adalah nilai koefisien korelasi

Nilai koefisien reliabilitas yang baik adalah diatas 0,7 (cukup baik), di atas 0,8 (baik). Pengukuran validitas dan reliabilitas mutlak dilakukan, karena jika instrument yang digunakan sudah tidak valid dan reliable maka dipastikan hasil penelitiannya pun tidak akan valid dan reliable. Sugiyono (2007: 137) menjelaskan perbedaan antara penelitian yang valid dan reliable dengan instrument yang valid dan reliable sebagai berikut :

Penelitian yang valid artinya bila terdapat kesamaan antara data yang terkumpul dengan data yang sesungguhnya terjadi pada objek yang diteliti. Artinya, jika objek berwarna merah, sedangkan data yang terkumpul berwarna putih maka hasil penelitian tidak valid. Sedangkan penelitian yang reliable bila terdapat kesamaan data dalam waktu yang berbeda. Kalau dalam objek kemarin berwarna merah, maka sekarang dan besok tetap berwarna merah.

Ada beberapa jenis validitas yang digunakan untuk menguji ketepatan ukuran, diantaranya validitas isi (*content validity*) dan validitas konsep (*concept validity*).

Validitas Isi

Validitas isi atau *content validity* memastikan bahwa pengukuran memasukkan sekumpulan item yang memadai dan mewakili yang mengungkap konsep. Semakin item skala mencerminkan kawasan atau keseluruhan konsep yang diukur, semakin besar validitas isi. Atau dengan kata lain, validitas isi merupakan fungsi seberapa baik dimensi dan elemen sebuah konsep yang telah digambarkan.

Validitas muka (*face validity*) dianggap sebagai indeks validitas isi yang paling dasar dan sangat minimum. Validitas isi menunjukkan bahwa item-item yang dimaksudkan untuk mengukur sebuah konsep, memberikan kesan mampu mengungkap konsep yang hendak di ukur.

Validitas Konsep

Validitas konsep atau *concept validity* menunjukkan seberapa baik hasil yang diperoleh dari pengukuran cocok dengan teori yang mendasari desain test. Hal ini dapat dinilai dari validitas konvergen dan validitas diskriminan.

Validitas konvergen terpenuhi jika skor yang diperoleh dengan dua instrument berbeda yang mengukur konsep yang sama menunjukkan korelasi yang tinggi.

Validitas diskriminan terpenuhi jika berdasarkan teori, dua variabel diprediksi tidak berkorelasi, dan skor yang diperoleh dengan mengukurnya benar-benar secara empiris membuktikan hal tersebut.

Secara umum, Sekaran (2006) membagi beberapa istilah validitas sebagai berikut :

1. Validitas isi yaitu apakah pengukuran benar-benar mengukur konsep ?
2. Validitas muka yaitu apakah para ahli mengesahkan bahwa instrument mengukur apa yang seharusnya diukur
3. Validitas berdasarkan criteria yaitu apakah pengukuran membedakan cara yang membantu memprediksi criteria variabel

4. Validitas konkuren yaitu apakah pengukuran membedakan cara yang membantu memprediksi criteria saat ini ?
5. Validitas prediktif yaitu apakah pengukuran membedakan individual dalam membantu memprediksi di masa depan ?
6. Validitas Konsep yaitu apakah instrument menyediakan konsep sebagai teori ?
7. Validitas konvergen yaitu apakah dua instrument mengukur konsep dengan korelasi yang tinggi ?
8. Validitas diskriminan yaitu apakah pengukuran memiliki korelasi rendah dengan variabel yang diperkirakan tidak ada hubungannya dengan variabel tersebut ?

Akan di uji validitas dan reliabilitas variabel **kepuasan kerja**. Variabel ini berjumlah 5 indikator yang diadaptasi dari Intrinsic factor dari teori dua factor Herzberg meliputi **pekerjaan itu sendiri, keberhasilan yang diraih, kesempatan bertumbuh, kemajuan dalam karier dan pengakuan orang lain**. Skala yang digunakan adalah skala Likert 1 – 5 dengan jumlah sampel sebanyak 36. Setelah angket ditabulasi maka diperoleh data sbb :

No	X1	X2	X3	X4	X5	No	X1	X2	X3	X4	X5
1	4	4	4	4	4	19	3	3	3	3	3
2	3	3	3	3	3	20	3	3	3	3	3
3	4	4	4	4	4	21	4	3	4	3	3
4	4	4	4	3	4	22	4	3	3	3	3
5	3	3	3	3	3	23	4	4	4	4	4
6	3	3	4	3	3	24	3	3	4	4	3
7	4	4	4	4	4	25	4	3	4	4	3
8	5	4	4	4	4	26	4	4	4	4	4
9	3	4	3	3	4	27	4	4	4	4	4
10	4	4	4	4	4	28	4	4	4	4	4
11	4	4	4	4	4	29	2	2	3	2	2
12	3	3	3	3	3	30	3	3	3	3	3
13	4	4	4	3	4	31	4	4	4	4	2
14	4	4	4	4	4	32	4	4	4	3	3
15	3	3	3	4	3	33	4	4	4	4	4
16	3	4	3	4	4	34	2	3	4	4	3
17	3	3	3	3	3	35	3	3	3	4	2
18	4	4	4	4	4	36	4	3	3	3	4

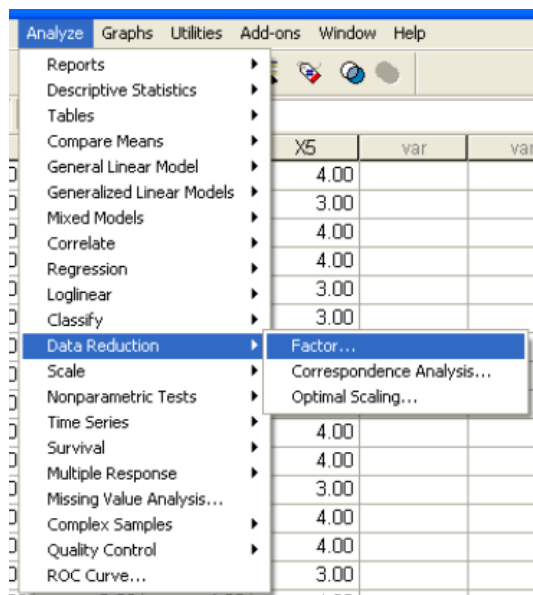
B. PENYELESAIAN

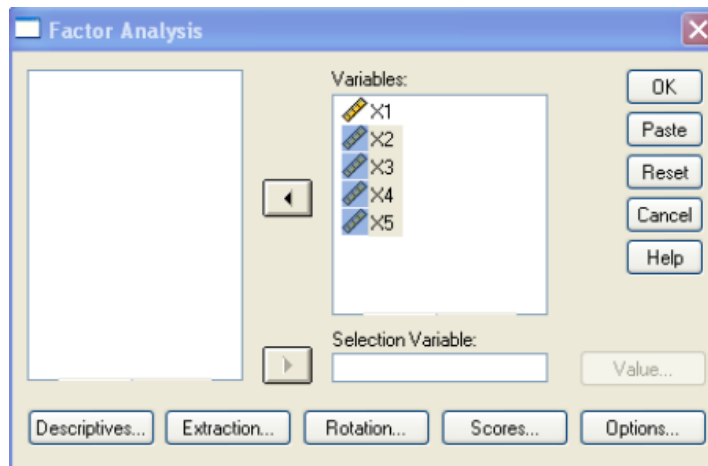
	X1	X2	X3	X4	X5
1	4.00	4.00	4.00	4.00	4.00
2	3.00	3.00	3.00	3.00	3.00
3	4.00	4.00	4.00	4.00	4.00
4	4.00	4.00	4.00	3.00	4.00
5	3.00	3.00	3.00	3.00	3.00
6	3.00	3.00	4.00	3.00	3.00
7	4.00	4.00	4.00	4.00	4.00
8	5.00	4.00	4.00	4.00	4.00
9	3.00	4.00	3.00	3.00	4.00
10	4.00	4.00	4.00	4.00	4.00
11	4.00	4.00	4.00	4.00	4.00
12	3.00	3.00	3.00	3.00	3.00
13	4.00	4.00	4.00	3.00	4.00
14	4.00	4.00	4.00	4.00	4.00
15	3.00	3.00	3.00	4.00	3.00

Tahap 1. Analisis Faktor

Klik **Analyze – Data Reduction – Factor**

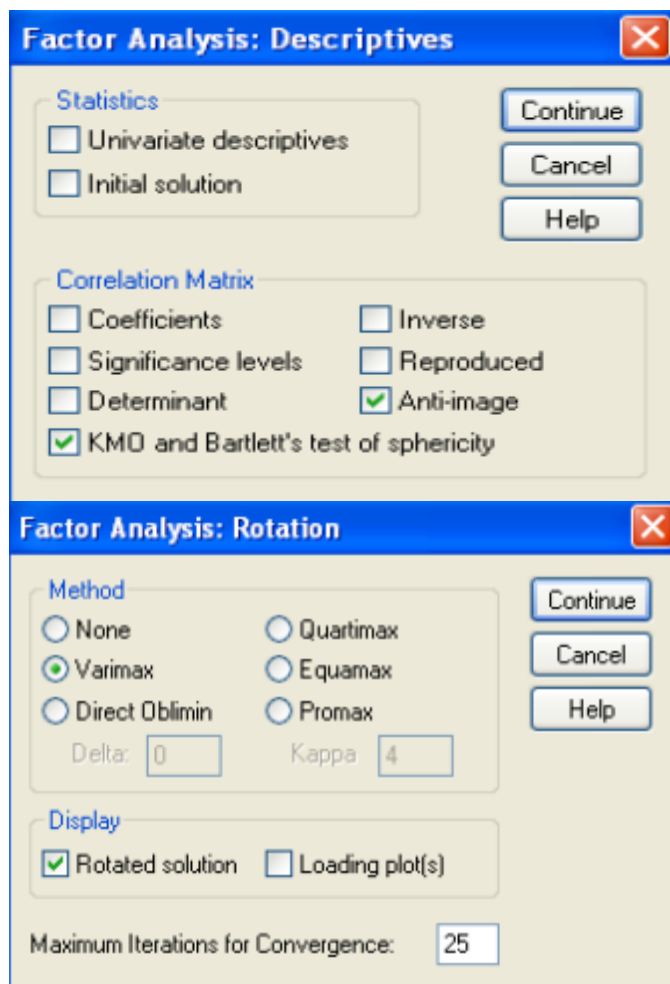
Masukkan seluruh pertanyaan ke box “Variables”





Klik **Descriptive** – Aktifkan **KMO and Bartlett's Test of Specirity** dan **Anti-Image**

Klik **Rotation** : Aktifkan **Varimax**



Hasil Analisis Faktor

KMO and Bartlett's Test

Kaiser-Meyer-Olkin Measure of Sampling Adequacy.		.804
Bartlett's Test of Sphericity	Approx. Chi-Square	85.478
	df	10
	Sig.	.000

Anti-image Matrices

		X1	X2	X3	X4	X5
Anti-image Covariance	X1	.448	-.118	-.149	.036	-.065
	X2	-.118	.264	-.075	-.130	-.190
	X3	-.149	-.075	.506	-.165	.038
	X4	.036	-.130	-.165	.586	.026
	X5	-.065	-.190	.038	.026	.431
Anti-image Correlation	X1	.851 ^a	-.342	-.313	.070	-.148
	X2	-.342	.754 ^a	-.205	-.330	-.562
	X3	-.313	-.205	.838 ^a	-.303	.082
	X4	.070	-.330	-.303	.828 ^a	.052
	X5	-.148	-.562	.082	.052	.782 ^a

a. Measures of Sampling Adequacy(MSA)

Nilai KMO sebesar 0.840 menandakan bahwa instrumen valid karena sudah memenuhi batas 0.50 ($0.840 > 0.50$). Korelasi anti image menghasilkan korelasi yang cukup tinggi untuk masing-masing item, yaitu 0.851 (X1), 0.754 (X2), 0.838 (X3), 0.828 (X4) dan 0.782 (X5). Dapat dinyatakan bahwa 5 item yang digunakan untuk mengukur konstruk kepuasan instrinsik memenuhi kriteria sebagai pembentuk konstruk.

Total Variance Explained

Component	Extraction Sums of Squared Loadings		
	Total	% of Variance	Cumulative %
1	3.280	65.604	65.604

Extraction Method: Principal Component Analysis.

Output ketiga adalah Total variance Explained menunjukkan bahwa dari 5 item yang digunakan, hasil ekstraksi SPSS menjadi 1 faktor dengan kemampuan menjelaskan konstak sebesar 65,604% .

Component Matrix

	Component 1
X1	.826
X2	.914
X3	.790
X4	.722
X5	.786

Extraction Method: Principal Component Analysis.

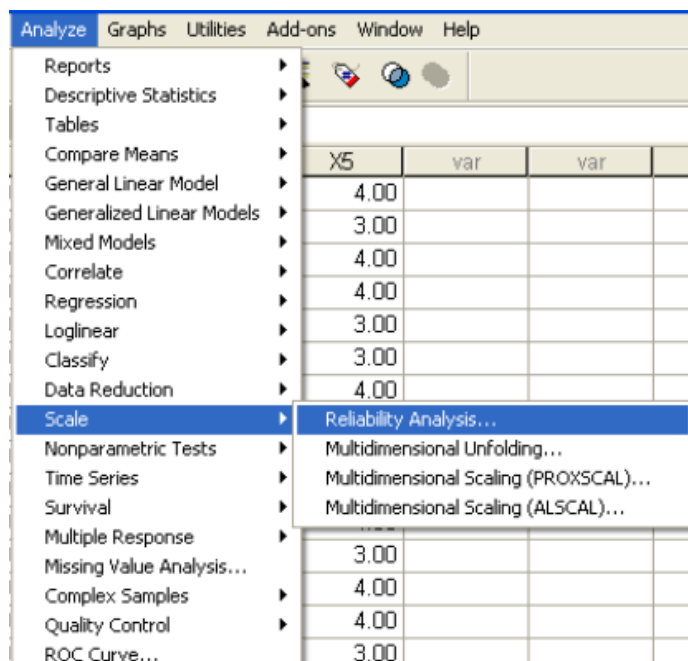
a. 1 components extracted.

Dengan melihat component matrix terlihat bahwa seluruh item meliputi pekerjaan itu sendiri (x1), keberhasilan yang diraih (x2), kesempatan bertumbuh (x3), kemajuan dalam karier (x4) dan pengakuan orang lain (x5) memiliki loading faktor yang besar yaitu di atas 0.50. Dengan demikian dapat dibuktikan bahwa 5 item valid.

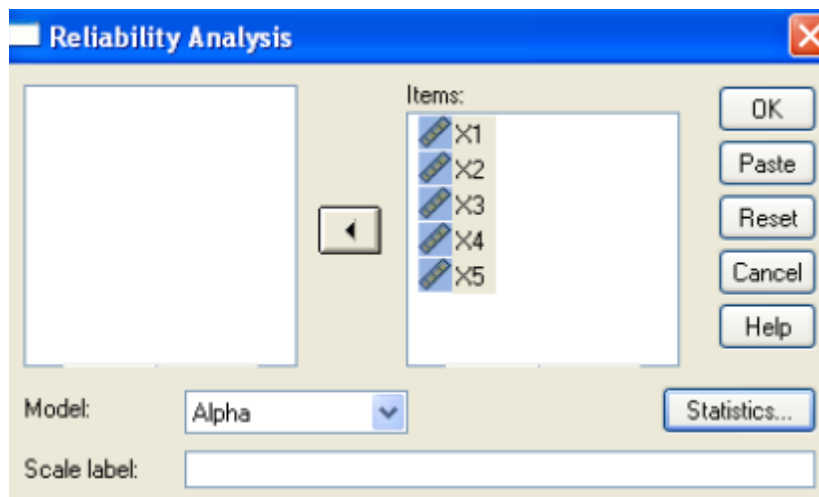
Tahap 2

Pilih **Analyze > Scale > Reliability Analysis**

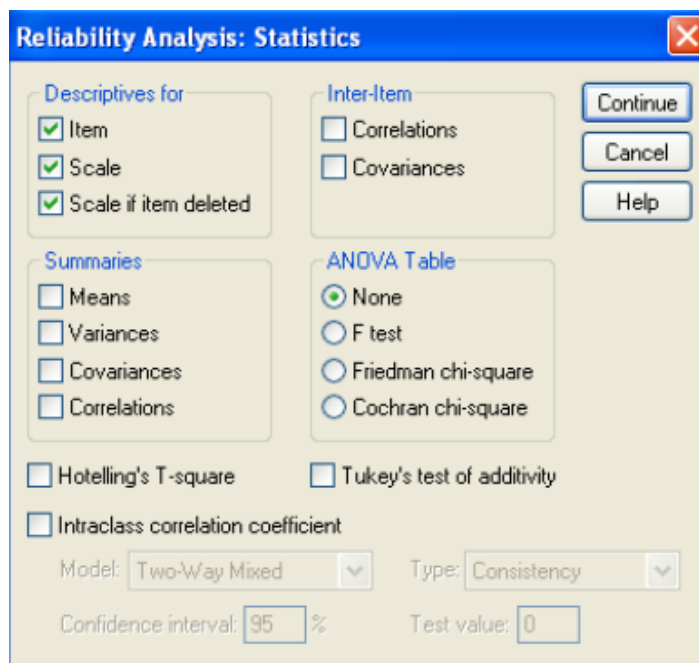
Masukkan semua variabel (item 1 s/d 5) ke kotak items



Klik **Reliability Analysis**, lalu masukan variabel **X1, X2, X3, X4** dan **X5** ke kotak items



Klik **Kotak Statistics**, lalu tandai **ITEM**, **SCALE**, dan **SCALE IF ITEM DELETED** pada kotak **DESCRIPTIVES FOR > Continue**



Klik OK

Maka akan tampil output sebagai berikut :

Reliability Statistics

Cronbach's Alpha	N of Items
.863	5

Item Statistics

	Mean	Std. Deviation	N
X1	3.5556	.65222	36
X2	3.5000	.56061	36
X3	3.6111	.49441	36
X4	3.5278	.55990	36
X5	3.4167	.64918	36

Item-Total Statistics

	Scale Mean if Item Deleted	Scale Variance if Item Deleted	Corrected Item-Total Correlation	Cronbach's Alpha if Item Deleted
X1	14.0556	3.425	.707	.830
X2	14.1111	3.473	.848	.794
X3	14.0000	4.000	.665	.842
X4	14.0833	3.964	.575	.861
X5	14.1944	3.533	.658	.844

D. INTERPRETASI

Reliabilitas

Sekaran (dalam Zulganef, 2006) yang menyatakan bahwa suatu instrumen penelitian mengindikasikan memiliki reliabilitas yang memadai jika koefisien alpha Cronbach lebih besar atau sama dengan 0,70. Sementara hasil uji menunjukkan koef cronbach alpha sebesar 0.863, dengan demikian dapat disimpulkan bahwa variabel ini adalah reliabel.

Analisis Item

Dalam prosedur kontruksi atau penyusunan test, sebelum melakukan estimasi terhadap reliabilitas dan validitas, dilakukan terlebih dahulu prosedur aitem yaitu dengan menguji karakteristik masing-masing item yang akan menjadi bagian test yang bersangkutan. Aitem-aitem yang tidak memenuhi persyaratan tidak boleh diikutkan sebagai bagian dari test. Pengujian reliabilitas dan validitas hanya layak dilakukan terhadap kumpulan aitem-aitem yang telah dianalisis dan diuji.

Beberapa teknik seleksi yang biasanya dipertimbangkan dalam prosedur seleksi adalah koefisien korelasi item-total, indeks reliabilitas item, dan indeks validitas item. Pada tes yang dirancang untuk mengungkap abilitas kognitif dengan format

item pilihan ganda, masih ada karakteristik item yang seharusnya juga dianalisis seperti tingkat kesukaran item dan efektivitas distraktor.

Salah satu parameter fungsi pengukuran item yang sangat penting adalah statistic yang memperlihatkan kesesuaian antara fungsi item dengan fungsi tes secara keseluruhan yang dikenal dengan istilah konsistensi item-total. Dasar kerja yang digunakan dalam analisis item dalam hal ini adalah memilih item-item yang fungsi ukurnya sesuai dengan fungsi ukur test seperti dikehendaki penyusunnya. Dengan kata lain adalah memilih item yang mengukur hal yang sama dengan apa yang diukur oleh tes secara keseluruhan.

Pengujian keselarasan fungsi item dengan fungsi ukur tes dilakukan dengan menghitung koefisien korelasi antara distribusi skor pada setiap item dengan distribusi skor total tes itu sendiri. Prosedur ini akan menghasilkan koefisien korelasi item total (r_{it}) yang juga dikenal dengan sebutan parameter daya beda item.

Tentang Cronbach Alpha

Cronbach's alpha is a measure of internal consistency, that is, how closely related a set of items are as a group. A "high" value of alpha is often used (along with substantive arguments and possibly other statistical measures) as evidence that the items measure an underlying (or latent) construct. However, a high alpha does not imply that the measure is unidimensional. If, in addition to measuring internal consistency, you wish to provide evidence that the scale in question is unidimensional, additional analyses can be performed. Exploratory factor analysis is one method of checking dimensionality. Technically speaking, Cronbach's alpha is not a statistical test – it is a coefficient of reliability (or consistency).

Didasarkan pada penjelasan di atas, maka penggunaan cronbach alpha bukanlah satu-satunya pedoman untuk menyatakan instrumen yang digunakan sudah reliabel. Untuk mengecek unidimensional pertanyaan diperlukan analisis tambahan yaitu ekplanatory factor analysis.

Teknik Yang Lebih Akurat Untuk Mengukur Validitas dan Reliabilitas

Untuk teknik yang lebih akurat untuk menguji validitas dan reliabilitas adalah analisis faktor konfirmatory. Menurut **Joreskog dan Sorbom (1993)**, CFA digunakan untuk menguji *"theoretical or hypotesical concepts, or construct, or variables, which are not directly measurable or observable"*.

Penjelasan Hair, dkk (2006) mengenai CFA adalah :

"CFA is way of testing how well measured variables represent a smaller number of construct...CFA is used to provide a confirmatory test of our measurement theory. A Measurement theory specifies how measured variables logically and systematically represent construct involved in a theoretical model. In Order words, measurement theory specifies a series relationships that suggest how variables represent a latent construct that is non measured directly" (dalam Kusnendi, 2008:97).



NORMALITAS DAN OUTLIER

NORMALITAS DATA

Pola sebaran data sangat penting diperhitungkan untuk menentukan jenis analisis statistik yang digunakan. Data dikatakan menyebar normal jika populasi data memenuhi kriteria:

68.27% data berada di sekitar $\text{Mean} \pm 1\sigma$ (standard deviasi)

95.45% data berada di sekitar $\text{Mean} \pm 2\sigma$ (standard deviasi)

99.73% data berada di sekitar $\text{Mean} \pm 3\sigma$ (standard deviasi)

Dan sisanya di luar range tersebut.

Metode statistika yang mengharuskan terpenuhinya asumsi normalitas disebut **Statistika Parametrik**, Sedangkan metode statistika yang digunakan untuk data tidak berdistribusi normal disebut **Statistika Nonparametrik**.

Hipotesis yang menandakan asumsi normalitas adalah:

H_0 : Data menyebar normal

H_1 : Data tidak menyebar normal

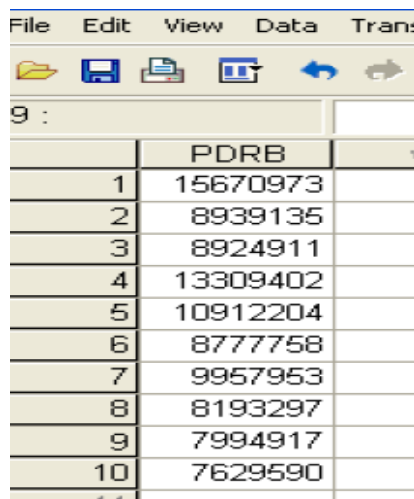
Cara menguji Normalitas data dapat dilakukan secara visual maupun uji yang relevan. Secara visual, uji normalitas dilakukan dengan:

Tabel 8.1.
PDRB dari 10 kabupaten/kota di propinsi HORE

KABUPATEN	Pertanian	Industri	Jasa	PDRB
Bodronoyo	5194485.32	5218350.93	5258136.5	15670972.75
Sulamanto	2762729.18	2997818.9	3178586.96	8939135.04
Ciganjur	2911111.03	2974994.19	3038805.78	8924911.01
Bandungan	4338609.02	4436742.73	4534050.73	13309402.49
siGarut	3566959.43	3631799.79	3713444.29	10912203.51
Tasikmala	2886454.78	2920347.31	2970956	8777758.08
Ciamis	3236120.65	3305451.54	3416380.69	9957952.89
Kuning	2635535.23	2738240.12	2819521.59	8193296.93
Carubon	2580728.25	2671645.91	2742543.13	7994917.28
Malengka	2445604.05	2558835.15	2625150.66	7629589.86

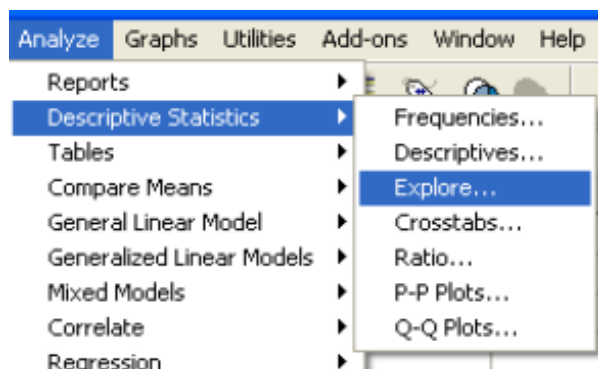
Sumber : data hipotesis

1. Tuliskan data PDRB kedalam SPSS, sehingga didapatkan hasil sbb:

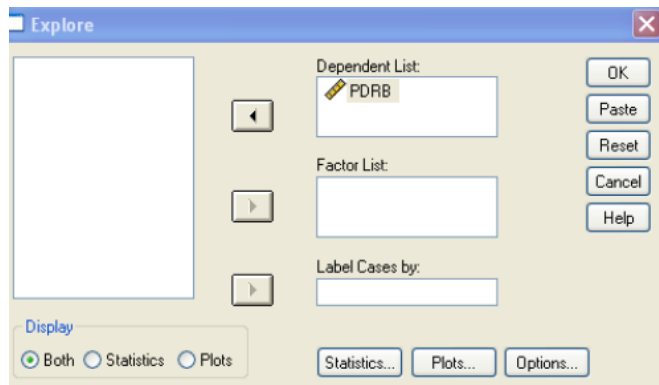


	PDRB
1	15670973
2	8939135
3	8924911
4	13309402
5	10912204
6	8777758
7	9957953
8	8193297
9	7994917
10	7629590

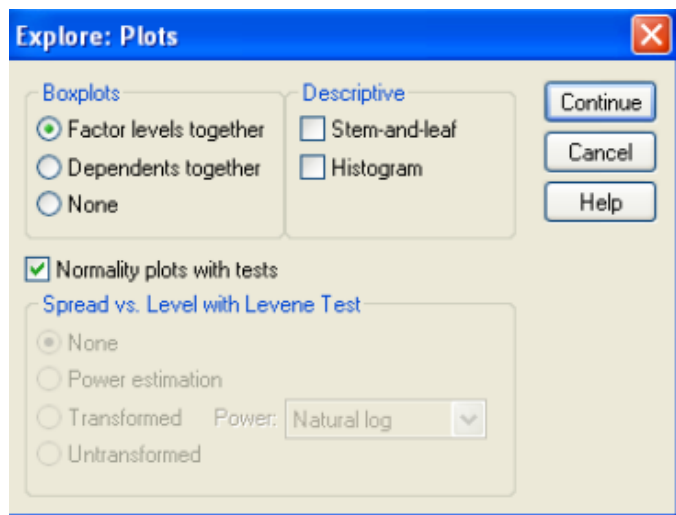
2. Pada menu utama SPSS pilih **Analyze** → **Descriptives statistics** → **Explore** sehingga muncul Dialog Box seperti pada gambar dibawah ini.



Masukan **PDRB** ke dependend list



3. Isi kolom **Dependent List** dengan variabel **PDRB** Pada **Display** pilih **Plots**, Kemudian Klik **Plots**, sehingga muncul Dialog Box seperti dibawah ini.



4. Pada Menu **Boxplots**, pilih **Factors levels together**, kemudian cek list pada **Normality plots with tests**. Pilih **Continue** → **OK**
5. Selanjutnya akan muncul output seperti ini

Tests of Normality

	Kolmogorov-Smirnov ^a			Shapiro-Wilk		
	Statistic	df	Sig.	Statistic	df	Sig.
PDRB	.263	10	.048	.829	10	.033

a. Lilliefors Significance Correction

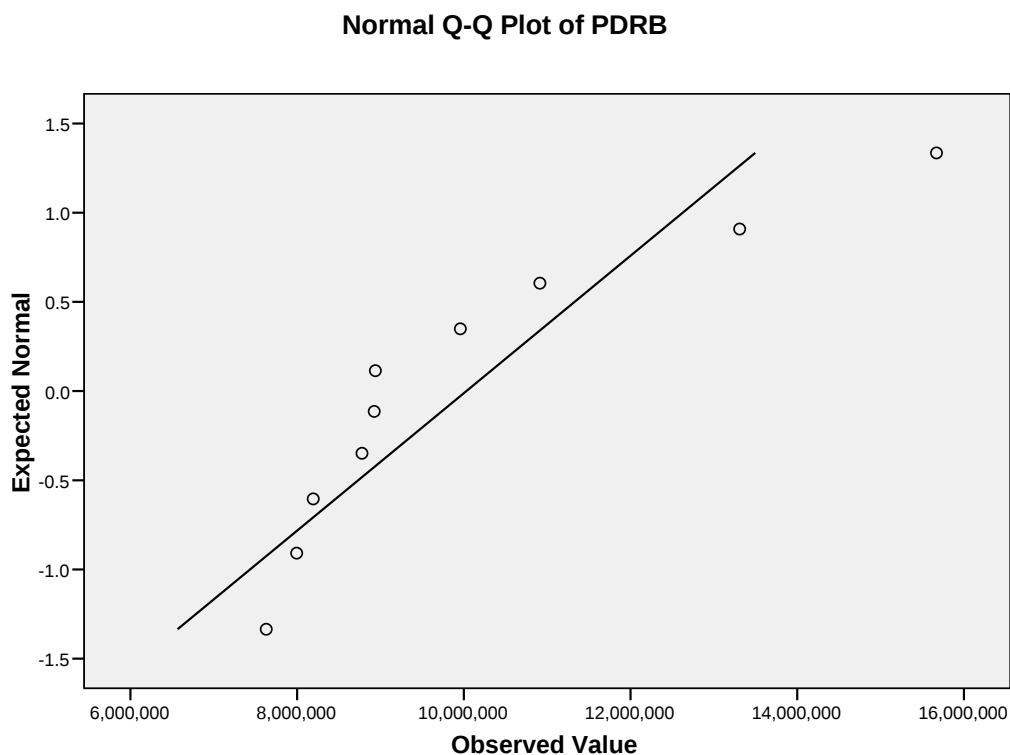
Analisis:

Output pada Gambar, merupakan output uji normalitas. Ada dua uji yang muncul, yaitu **Kolmogorov Smirnov Test** dan **Shapiro Wilk Test**. Adapun kriteria pengujiannya adalah:

- Jika Nilai Signifikansi pada kolmogorov Smirnov < 0.05 , data tidak menyebar normal.
- Jika nilai Signifikansi pada Kolmogorov Smirnov > 0.05 , maka data menyebar normal.

Demikian juga kriteria yang berlaku pada Saphiro Wilk test. Pada output yang diuji pada data **PDRB**, dapat dilihat bahwa nilai signifikansi pada kedua uji < 0.05 (0.048 dan 0.033). Sehingga dapat disimpulkan bahwa data **PDRB** tidak menyebar normal dan tidak dapat dilakukan analisis lebih lanjut dengan menggunakan statistika parametrik.

- Output selanjutnya yaitu seperti yang muncul pada dibawah ini.

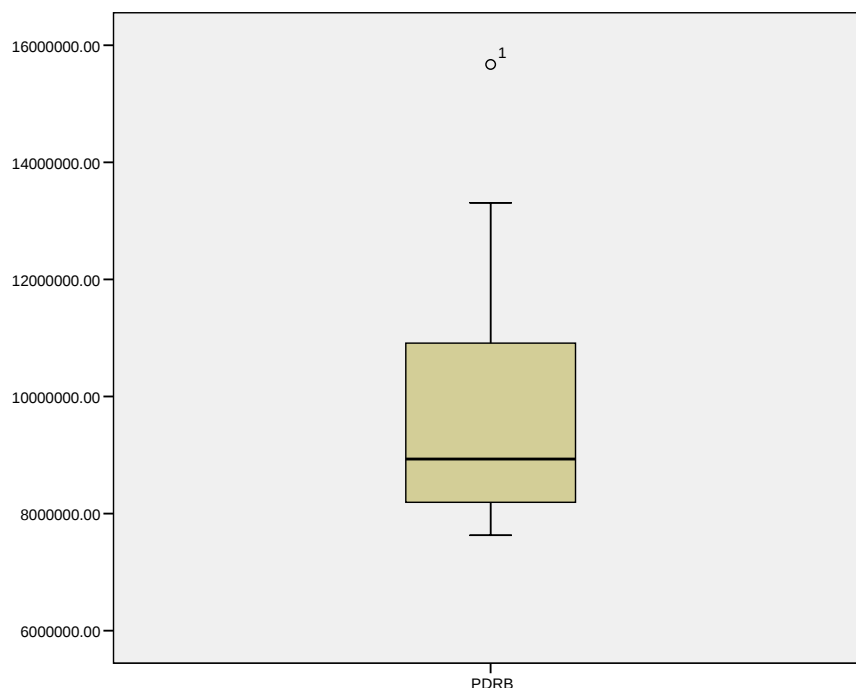


Analisis:

Normal Q-Q Plot dapat digunakan sebagai alat pengujian normalitas secara visual. Kriterianya adalah, jika titik-titik pengamatan berada di sekitar garis diagonal, maka dapat disimpulkan bahwa **data menyebar normal**. Seperti terlihat pada gambar, titik-titik pengamatan tidak berada di sekitar Garis Diagonal sehingga secara visual dapat dikatakan bahwa data **PDRB** tidak menyebar normal. Namun

pengujian secara visual ini harus tetap didukung dengan uji **Kolmogorov Smirnov** ataupun **Saphiro Wilk**.

7. Output terakhir yang muncul adalah **Box Plot**. Garis tengah horizontal Box Plot adalah letak Median, sedangkan dua garis lainnya adalah letak Quartil 1 dan Quartil 3. Titik yang berada di luar Box Plot merupakan **pengamatan yang berada jauh dari rata-rata** atau disebut dengan **Outlier**. Terkadang outlier menyebabkan hasil analisis menjadi bias karena keunikannya. Oleh karena itu, dalam berbagai penelitian, outlier disarankan untuk dibuang.

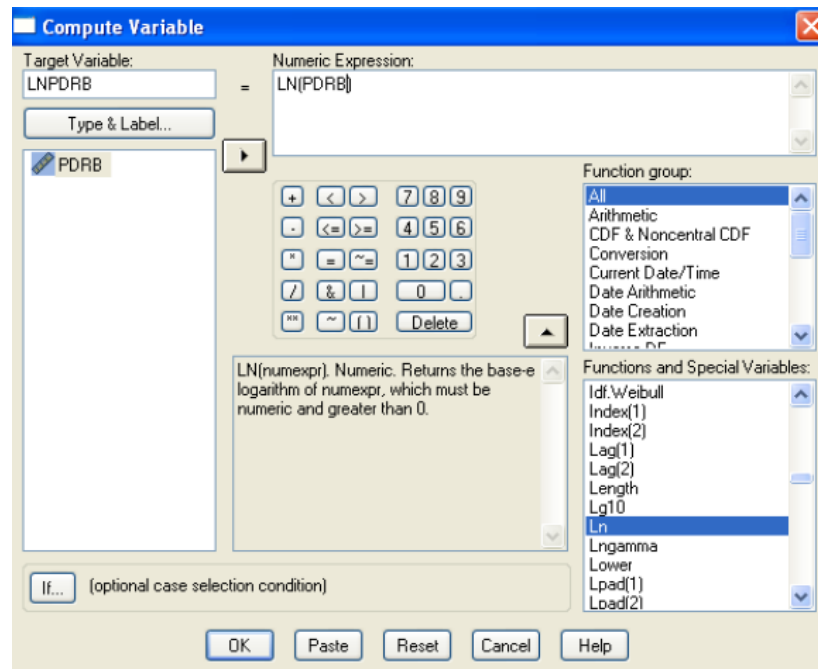


a. TRANSFORMASI DATA

Penelitian dapat dilanjutkan dengan menggunakan metode Statistika Parametrik jika diketahui data menyebar normal. Namun akan muncul pertanyaan, bagaimana jika setelah diuji ternyata data tidak menyebar normal?

Data yang tidak menyebar normal perlu ditransformasi terlebih dahulu. Langkah untuk mentransformasi data dalam SPSS adalah:

1. Buka data **PDRB** (bahwa data **PDRB** tidak menyebar normal).
2. Pilih menu **Transform → Compute Variable** sehingga muncul Dialog Box seperti gambar dibawah ini.



3. **Target Variable** merupakan kolom yang akan digunakan untuk data hasil transformasi. **Targer Variable** dapat diberi nama apapun. Untuk keseragaman, isi dengan nama **LNPD RB**.
4. Pilih **All** pada **Function Group**, kemudian pilih **Ln** pada **Functions and Special Variables** dengan cara *double click*. Selanjutnya masukkan variabel **PDRB** pada kotak **Numeric Expression** → **OK**.
5. Output yang dihasilkan adalah berupa kolom baru pada **Data View**, seperti dibawah ini, Kolom tersebut adalah data hasil transformasi yang akan dianalisis lebih lanjut.

	PDRB	LNPD RB
1	15670973	16.57
2	8939135	16.01
3	8924911	16.00
4	13309402	16.40
5	10912204	16.21
6	8777758	15.99
7	9957953	16.11
8	8193297	15.92
9	7994917	15.89
10	7629590	15.85
11		

Kolom Transformed Variable

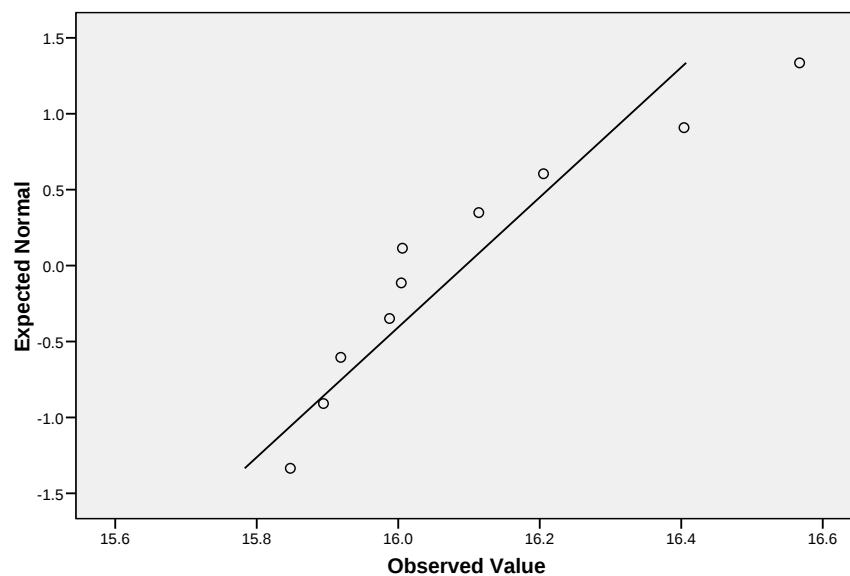
Tests of Normality

	Kolmogorov-Smirnov ^a			Shapiro-Wilk		
	Statistic	df	Sig.	Statistic	df	Sig.
LNPDRB	.248	10	.081	.880	10	.131

a. Lilliefors Significance Correction

Dapat dilihat bahwa nilai signifikansi pada kedua uji > 0.05 (0.081 dan 0.131). Sehingga dapat disimpulkan bahwa data **LNPDRB** menyebar normal dan tidak dapat dilakukan analisis lebih lanjut dengan menggunakan statistika parametrik.

Normal Q-Q Plot of LNPDRB

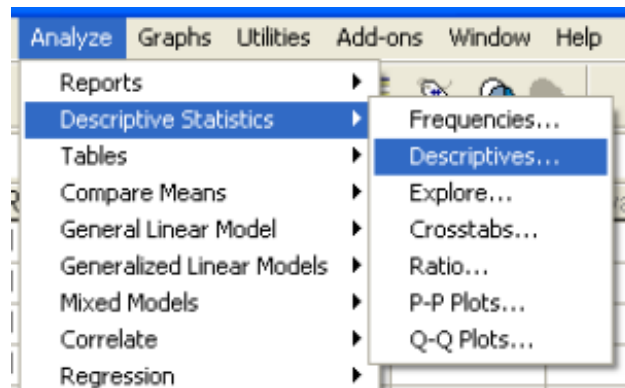


DETEKSI OUTLIER

Outlier adalah pengamatan yang memiliki simpangan yang cukup jauh dari rata-rata. Cara untuk mendeteksi outlier sangat tergantung pada tingkatan analisis data, apakah tergolong analisis data univariate, bivariate, atau multivariate. Pada bab ini akan dibahas deteksi outlier pada data univariate. Deteksi dari secara visual telah dibahas sebelumnya yaitu dengan menggunakan **Box Plot**. Cara lain adalah melalui nilai **z-score**.

Langkah-langkahnya adalah sebagai berikut:

1. Buka file yang berisi **PDRB** tersebut, Pilih **Analyze** ➔ **Descriptive Statistics** ➔ **Descriptives**



Descriptives memunculkan nilai Zscore

2. Masukkan variabel **PDRB** pada kolom **Variable(s)**, kemudian cek list **Save standardized values as variables** ➔ **OK**. Output akan muncul berupa kolom baru pada sheet **Data View**.

	PDRB	LNPDRB	ZPDRB
1	15670973	16.57	2.17324
2	8939135	16.01	-.42073
3	8924911	16.00	-.42621
4	13309402	16.40	1.26326
5	10912204	16.21	.33955
6	8777758	15.99	-.48292
7	9957953	16.11	-.02815
8	8193297	15.92	-.70813
9	7994917	15.89	-.78457
10	7629590	15.85	-.92534

3. Z kredit adalah nilai z-score dari masing-masing pengamatan. Kriteria penentuan outlier dipengaruhi oleh banyaknya sampel, yaitu :
 - Jika banyaknya sampel ≤ 80 , maka pengamatan dengan Zscore > 2.5 atau < -2.5 adalah outlier
 - Jika banyaknya sampel > 80 . Maka pengamatan dengan Zscore > 3 atau < -3 adalah outlier (Hair,dkk, 1998)



ANALISIS REGRESI

Analisis Regresi linier (*Linear Regression analysis*) adalah teknik statistika untuk membuat model dan menyelidiki pengaruh antara satu atau beberapa variabel bebas (*Independent Variables*) terhadap satu variabel respon (*dependent variable*). Ada dua macam analisis regresi linier:

1. Regresi Linier Sederhana: Analisis Regresi dengan satu *Independent variable* , dengan formulasi umum:

$$Y = a + b_1X_1 + e \quad (9.1)$$

2. Regresi Linier Berganda: Analisis regresi dengan dua atau lebih *Independent Variable* , dengan formulasi umum:

$$Y = a + b_1X_1 + b_2X_2 + \dots + b_nX_n + e \quad (9.2)$$

Dimana:

Y = Dependent variable

a = konstanta

b_1 = koefisien regresi X_1 , b_2 = koefisien regresi X_2 , dst.

e = Residual / Error

Fungsi persamaan regresi selain untuk memprediksi nilai *Dependent Variable* (Y), juga dapat digunakan untuk mengetahui arah dan besarnya pengaruh *Independent Variable* (X) terhadap *Dependent Variable* (Y).

Asumsi yang harus terpenuhi dalam analisis regresi (Gujarati, 2003) adalah:

1. Residual menyebar normal (asumsi normalitas)
2. Antar Residual saling bebas (Autokorelasi)
3. Kehomogenan ragam residual (Asumsi Heteroskedastisitas)
4. Antar Variabel independent tidak berkorelasi (multikolinearitas)

Asumsi-asumsi tersebut harus diuji untuk memastikan bahwa data yang digunakan telah memenuhi asumsi analisis regresi.

1. Input data **Keuntungan, Penjualan dan Biaya Promosi** dalam file SPSS. Definisikan variabel-variabel yang ada dalam sheet **Variable View**.

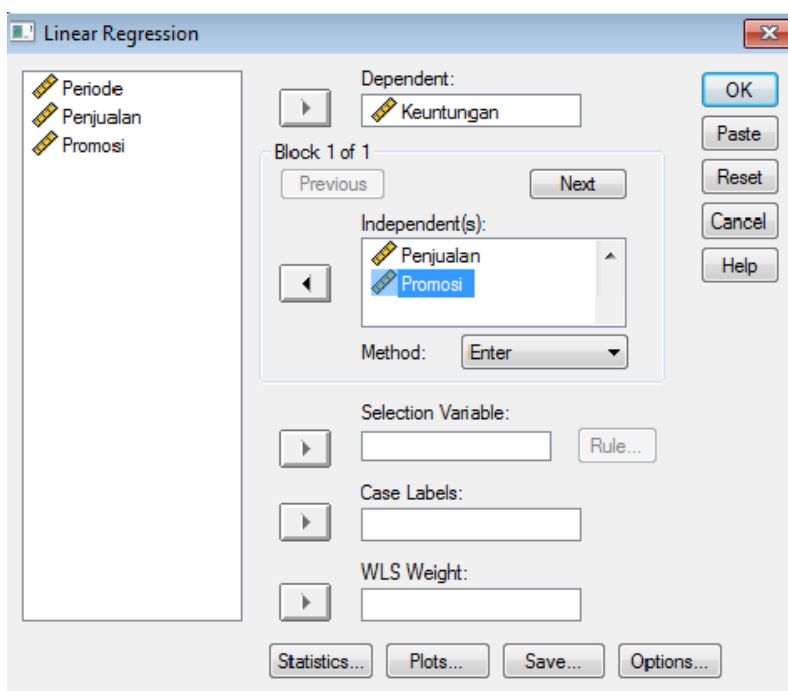
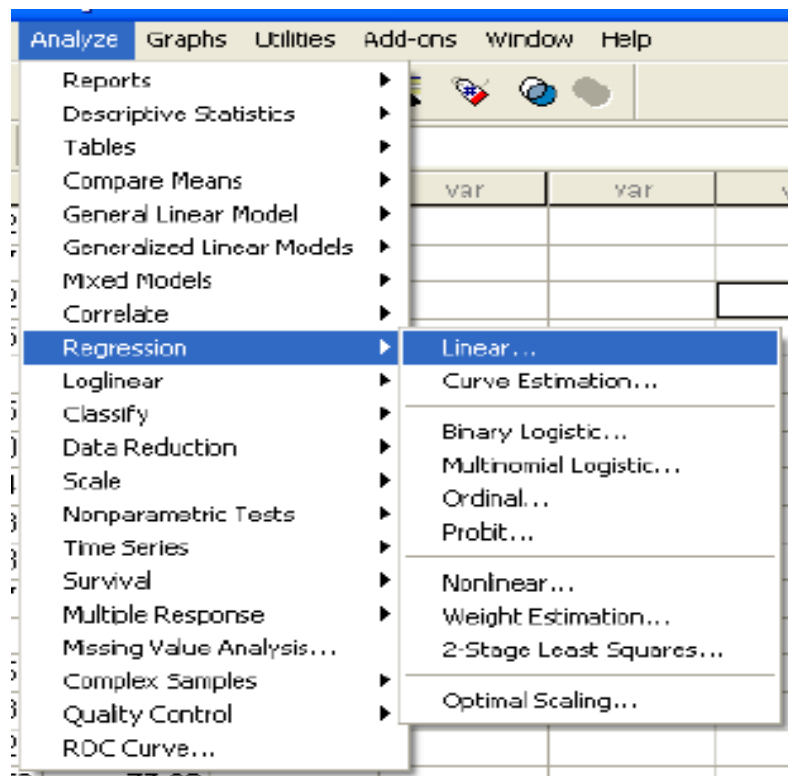
Periode	Keuntungan	Penjualan	Biaya Promosi
2012.01	100.000	1.000.000	55.000
2012.02	110.000	1.150.000	56.000
2012.03	125.000	1.200.000	60.000
2012.04	131.000	1.275.000	67.000
2012.05	138.000	1.400.000	70.000
2012.06	150.000	1.500.000	74.000
2012.07	155.000	1.600.000	80.000
2012.08	167.000	1.700.000	82.000
2012.09	180.000	1.800.000	93.000
2012.10	195.000	1.900.000	97.000
2012.11	200.000	2.000.000	100.000
2012.12	210.000	2.100.000	105.000
2013.01	225.000	2.200.000	110.000
2013.02	230.000	2.300.000	115.000
2013.03	240.000	2.400.000	120.000
2013.04	255.000	2.500.000	125.000
2013.05	264.000	2.600.000	130.000
2013.06	270.000	2.700.000	135.000
2013.07	280.000	2.800.000	140.000
2013.08	290.000	2.900.000	145.000
2013.09	300.000	3.000.000	150.000
2013.10	315.000	3.100.000	152.000
2013.11	320.000	3.150.000	160.000
2013.12	329.000	3.250.000	165.000
2014.01	335.000	3.400.000	170.000
2014.02	350.000	3.500.000	175.000
2014.03	362.000	3.600.000	179.000
2014.04	375.000	3.700.000	188.000
2014.05	380.000	3.800.000	190.000
2014.06	400.000	3.850.000	192.000
2014.07	405.000	3.950.000	200.000
2014.08	415.000	4.100.000	207.000
2014.09	425.000	4.300.000	211.000
2014.10	430.000	4.350.000	215.000
2014.11	440.000	4.500.000	219.000
2014.12	450.000	4.600.000	210.000

Sumber : Data Hipotesis

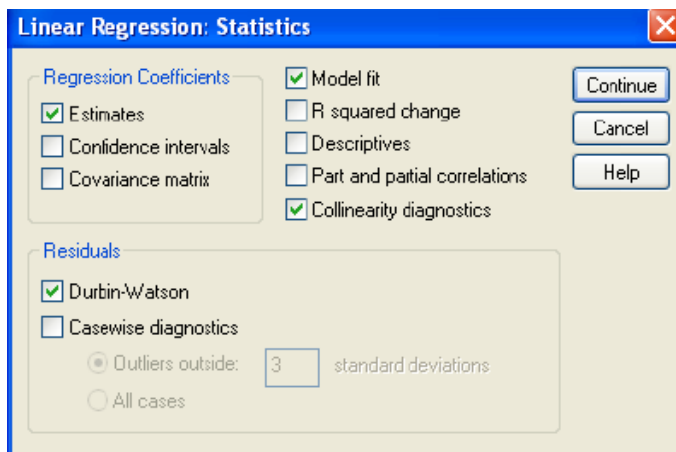
Masukan data diatas ke dalam program SPSS, sehingga akan seperti tampilan dibawah ini,

	Periode	Keuntunga	Penjualan	Promosi
1	2012,01	100000,0	1000000	55000,00
2	2012,02	110000,0	1150000	56000,00
3	2012,03	125000,0	1200000	60000,00
4	2012,04	131000,0	1275000	67000,00
5	2012,05	138000,0	1400000	70000,00
6	2012,06	150000,0	1500000	74000,00
7	2012,07	155000,0	1600000	80000,00
8	2012,08	167000,0	1700000	82000,00
9	2012,09	180000,0	1800000	93000,00
10	2012,10	195000,0	1900000	97000,00
11	2012,11	200000,0	2000000	100000,0
12	2012,12	210000,0	2100000	105000,0
13	2013,01	225000,0	2200000	110000,0
14	2013,02	230000,0	2300000	115000,0
15	2013,03	240000,0	2400000	120000,0
16	2013,04	255000,0	2500000	125000,0
17	2013,05	264000,0	2600000	130000,0
18	2013,06	270000,0	2700000	135000,0
19	2013,07	280000,0	2800000	140000,0
20	2013,08	290000,0	2900000	145000,0
21	2013,09	300000,0	3000000	150000,0
22	2013,10	315000,0	3100000	152000,0
23	2013,11	320000,0	3150000	160000,0
24	2013,12	329000,0	3250000	165000,0
25	2014,01	335000,0	3400000	170000,0
26	2014,02	350000,0	3500000	175000,0
27	2014,03	362000,0	3600000	179000,0
28	2014,04	375000,0	3700000	188000,0
29	2014,05	380000,0	3800000	190000,0
30	2014,06	400000,0	3850000	192000,0
31	2014,07	405000,0	3950000	200000,0

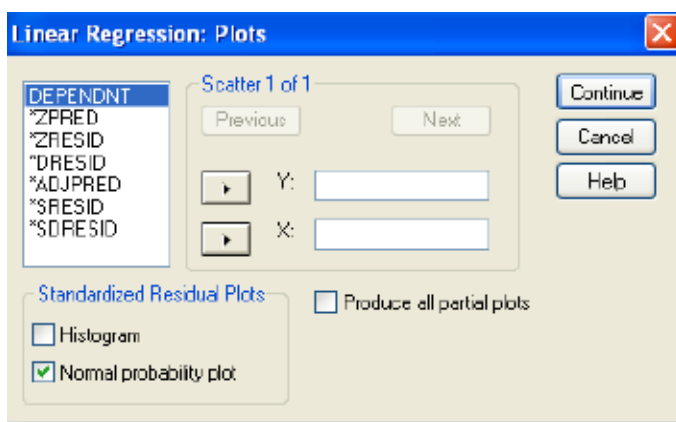
- Pilih Menu **Analyze** → **Regression** → **Linear** , sehingga muncul **Dialog Box** sesuai dibawah ini. Masukkan variabel **Kredit** pada kolom **Dependent Variable**, dan tiga variabel lain sebagai **Independent(s)**,



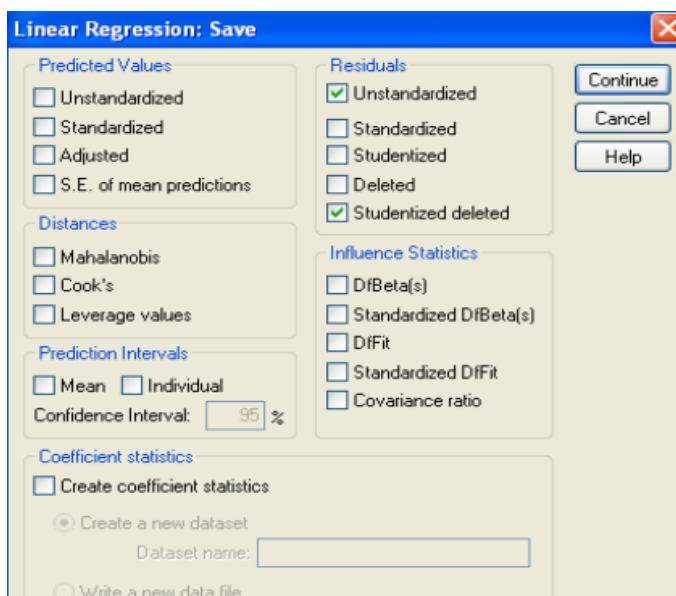
3. Pilih **Statistics**, cek list **Estimates**, **Collinearity Diagnostics**, dan **Durbin Watson** → **Continue**



4. Pilih **Plots**, cek List **Normal Probability Plot** → **Continue**,



5. Pilih **Save**, cek list **Unstandardized** dan **Studentized deleted Residuals**,



6. **Continue → OK,**
7. Langkah pertama yang harus dilakukan adalah membuang data outlier sehingga hasil output analisis yang dihasilkan tidak lagi terpengaruh oleh pengamatan yang menyimpang,

a. Uji Outlier

Perhatikan pada sheet **Data View** kita, maka kita akan temukan dua variabel baru, yaitu RES_1 (Residual) dan SDR (Studentized deleted Residual),

	Periode	Keuntungan	Penjualan	Promosi	RES_1	SDR_1
1	2012,01	100000,0	1000000	55000,00	-3501,56282	-,88997
2	2012,02	110000,0	1150000	56000,00	-3337,95445	-,84592
3	2012,03	125000,0	1200000	60000,00	5385,67299	1,37157
4	2012,04	131000,0	1275000	67000,00	1153,59360	,28444
5	2012,05	138000,0	1400000	70000,00	-1814,69395	-,44634
6	2012,06	150000,0	1500000	74000,00	902,64313	,22160
7	2012,07	155000,0	1600000	80000,00	-5015,06089	-1,24943
8	2012,08	167000,0	1700000	82000,00	-662,68271	-,16376
9	2012,09	180000,0	1800000	93000,00	-2667,98947	-,65139
10	2012,10	195000,0	1900000	97000,00	3049,34761	,74185
11	2012,11	200000,0	2000000	100000,0	-415,79476	-,10000
12	2012,12	210000,0	2100000	105000,0	-515,97823	-,12385
13	2013,01	225000,0	2200000	110000,0	4383,83830	1,06873
14	2013,02	230000,0	2300000	115000,0	-716,34517	-,17144
15	2013,03	240000,0	2400000	120000,0	-816,52864	-,19522
16	2013,04	255000,0	2500000	125000,0	4083,28789	,98971
17	2013,05	264000,0	2600000	130000,0	2983,10442	,71759
18	2013,06	270000,0	2700000	135000,0	-1117,07905	-,26683
19	2013,07	280000,0	2800000	140000,0	-1217,26252	-,29086
20	2013,08	290000,0	2900000	145000,0	-1317,44598	-,31502
21	2013,09	300000,0	3000000	150000,0	-1417,62945	-,33932
22	2013,10	315000,0	3100000	152000,0	5934,74872	1,47180
23	2013,11	320000,0	3150000	160000,0	1388,29397	,33730
24	2013,12	329000,0	3250000	165000,0	288,11050	,07005
25	2014,01	335000,0	3400000	170000,0	-6818,36333	-1,71280
26	2014,02	350000,0	3500000	175000,0	-1918,54680	-,46391
27	2014,03	362000,0	3600000	179000,0	798,79028	,19258
28	2014,04	375000,0	3700000	188000,0	428,52462	,10643
29	2014,05	380000,0	3800000	190000,0	-2219,09720	-,54217
30	2014,06	400000,0	3850000	192000,0	13139,57134	3,87060
31	2014,07	405000,0	3950000	200000,0	5586,82622	1,44044

Variabel Baru yang terbentuk

SDR adalah nilai-nilai yang digunakan untuk mendeteksi adanya outlier, Dalam deteksi outlier ini kita membutuhkan tabel distribusi t, Kriteria pengujiannya adalah jika nilai absolute $|SDR| > t_{n-k-1}^{\alpha/2}$, maka pengamatan tersebut merupakan outlier,

n = Jumlah Sampel, dan k = Jumlah variabel bebas

Nilai t pembanding adalah sebesar 2,056, Pada kolom SDR, terdapat 1 pengamatan yang memiliki nilai $|SDR| > 2,056$, yaitu pengamatan ke 17,

Berikut ini adalah outputnya,

Model Summary^b

Model	R	R Square	Adjusted R Square	Std. Error of the Estimate	Durbin-Watson
1	,999 ^a	,999	,998	186,51013	1,641

a. Predictors: (Constant), Promosi, Penjualan
b. Dependent Variable: Keuntungan

(Gambar 4,6 R Square)

Analisis:

b. R Square sebagai ukuran kecocokan model

Tabel **Variables Entered** menunjukkan variabel independent yang dimasukkan ke dalam model, Nilai **R Square** pada Tabel **Model Summary** adalah prosentase kecocokan model, **atau nilai yang menunjukkan seberapa besar variabel independent menjelaskan variabel dependent**, R^2 pada persamaan regresi rentan terhadap penambahan variabel independent, dimana semakin banyak variabel Independent yang terlibat, maka nilai R^2 akan semakin besar, Karena itulah digunakan R^2 *adjusted* pada analisis regresi linier Berganda, dan digunakan R^2 pada analisis regresi sederhana, Pada gambar output 4,6, terlihat nilai R Square adjusted sebesar 0,999, artinya variabel independent dapat menjelaskan variabel dependent sebesar 99,8%, sedangkan 0,2% dijelaskan oleh faktor lain yang tidak terdapat dalam model,

c. Uji F

Uji F dalam analisis regresi linier berganda bertujuan untuk mengetahui pengaruh variabel independent secara simultan, yang ditunjukkan oleh **dalam table ANOVA**,

ANOVA(b)

Model	Sum of Squares	df	Mean Square	F	Sig.
1 Regression	394212835607,795	2	197106417803,898	11245,958	,000(a)
Residual	578386614,427	33	17526867,104		
Total	394791222222,222	35			

a Predictors: (Constant), Promosi, Penjualan
b Dependent Variable: Keuntungan

Rumusan hipotesis yang digunakan adalah:

- H_0 Kedua variabel independent secara simultan tidak berpengaruh signifikan terhadap variabel Jumlah Kemiskinan,
 H_1 Kedua variabel independent secara simultan berpengaruh signifikan terhadap variabel Jumlah Kemiskinan,

Kriteria pengujiannya adalah:

Jika nilai signifikansi $> 0,05$ maka keputusannya adalah terima H_0 atau variabel independent secara simultan tidak berpengaruh signifikan terhadap variabel dependent.

Jika nilai signifikansi $< 0,05$ maka keputusannya adalah tolak H_0 atau variabel dependent secara simultan berpengaruh signifikan terhadap variabel dependent,

Berdasarkan kasus, Nilai **Sig**, yaitu sebesar 0,000, sehingga dapat disimpulkan bahwa Promosi dan penjualan secara simultan berpengaruh signifikan terhadap Besarnya Keuntungan.

d. Uji t

Uji t digunakan untuk mengetahui pengaruh masing-masing variabel independent secara parsial, ditunjukkan oleh Tabel **Coefficients** pada (Gambar 4,8),

Coefficients ^a								
		Unstandardized Coefficients		Standardized Coefficients	t	Sig.	Collinearity Statistics	
		B	Std. Error	Beta			Tolerance	VIF
1	(Constant)	-1587,875	2093,274		-,759	,453		
	Penjualan	,060	,009	,602	6,344	,000	,005	202,913
	Promosi	,818	,195	,398	4,191	,000	,005	202,913

a. Dependent Variable: Keuntungan

Rumusan hipotesis yang digunakan adalah:

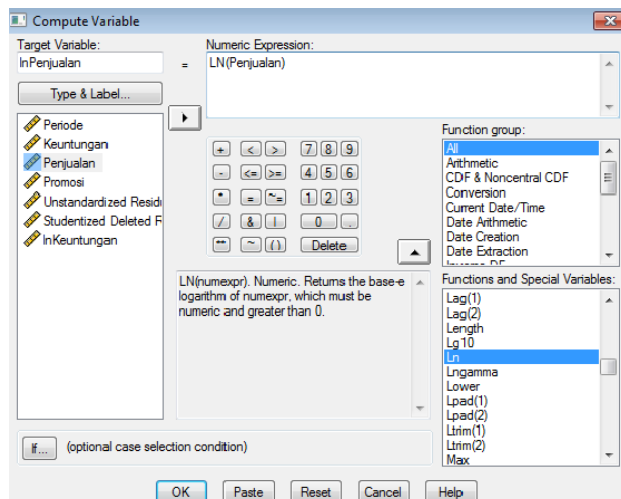
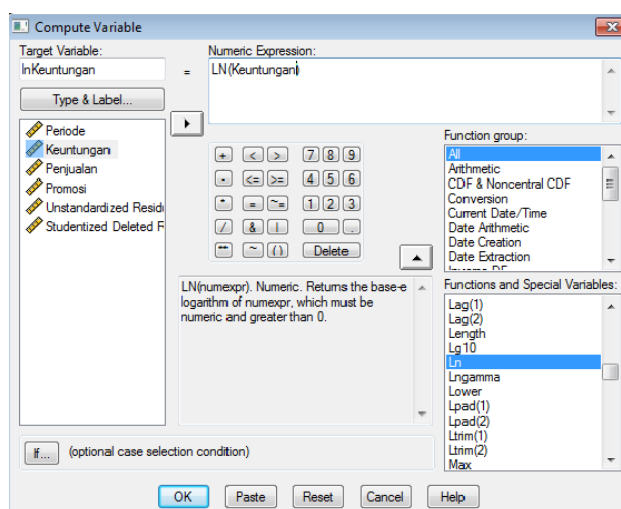
- H_0 : Penjualan tidak mempengaruhi besarnya Jumlah Keuntungan secara signifikan
 H_1 : Penjualan mempengaruhi besarnya Jumlah Keuntungan secara signifikan

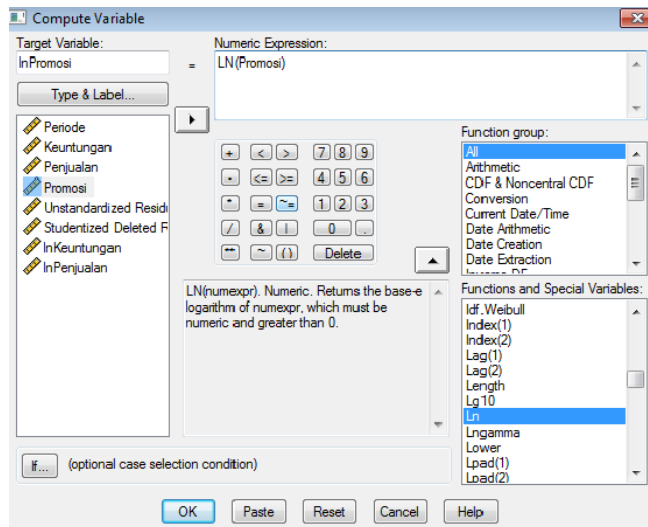
Hipotesis tersebut juga berlaku untuk variabel Inflasi, Perhatikan nilai **Unstandardized coefficients B** untuk masing-masing variabel, Variabel Penjualan mempengaruhi Jumlah Keuntungan yang disalurkan sebesar 0,06, Nilai ini positif artinya semakin besarnya Penjualan, maka semakin besar pula jumlah keuntungan, artinya jika penjualan naik sebesar 1.000 satuan maka keuntungan akan naik

sebesar 60 satuan. Demikian juga variabel Promosi berpengaruh positif terhadap jumlah Keuntungan sebesar 0,818, artinya jika promosi naik 1000 satuan maka keuntungan akan naik sebesar 818 satuan.

Signifikansi pengaruh variabel independent terhadap variabel dependent dapat dilihat dari nilai **Sig** pada kolom terakhir, Nilai signifikansi untuk variabel Penjualan yaitu sebesar 0,000, artinya variabel ini berpengaruh secara signifikan terhadap Jumlah Keuntungan, Hal ini berlaku juga untuk variabel promosi, dimana nilai signifikansinya $< 0,05$, sehingga kesimpulannya adalah ditolaknya H_0 atau dengan kata lain Penjualan dan Promosi mempunyai pengaruh signifikan terhadap Jumlah Keuntungan,

Dengan Model Ln





lnKeuntungan	lnPenjualan	lnPromosi
11,61	13,96	10,93
11,74	14,00	11,00
11,78	14,06	11,11
11,84	14,15	11,16
11,92	14,22	11,21
11,95	14,29	11,29
12,03	14,35	11,31
12,10	14,40	11,44
12,18	14,46	11,48
12,21	14,51	11,51
12,25	14,56	11,56
12,32	14,60	11,61
12,35	14,65	11,65
12,39	14,69	11,70
12,45	14,73	11,74
12,48	14,77	11,78
12,51	14,81	11,81
12,54	14,85	11,85
12,58	14,88	11,88
12,61	14,91	11,92
12,66	14,95	11,93
12,68	14,96	11,98
12,70	14,99	12,01
12,72	15,04	12,04
12,77	15,07	12,07
12,80	15,10	12,10
12,83	15,12	12,14
12,85	15,15	12,15
12,90	15,16	12,17
12,91	15,19	12,21
12,94	15,22	12,24

Model Summary^b

Model	R	R Square	Adjusted R Square	Std. Error of the Estimate	Durbin-Watson
1	,999 ^a	,999	,998	,01685	1,812

a. Predictors: (Constant), lnPromosi, lnPenjualan

b. Dependent Variable: lnKeuntungan

ANOVA^b

Model		Sum of Squares	df	Mean Square	F	Sig.
1	Regression	6,562	2	3,281	11560,184	,000 ^a
	Residual	,009	33	,000		
	Total	6,571	35			

a. Predictors: (Constant), lnPromosi, lnPenjualan

b. Dependent Variable: lnKeuntungan

Coefficients^a

Model		Unstandardized Coefficients		Standardized Coefficients	t	Sig.	Collinearity Statistics	
		B	Std. Error	Beta			Tolerance	VIF
1	(Constant)	-1,420	,314		-4,527	,000		
	lnPenjualan	,664	,111	,662	5,971	,000	,004	284,794
	lnPromosi	,347	,114	,337	3,043	,005	,004	284,794

a. Dependent Variable: lnKeuntungan

Analisis

Dari data diatas persamaan regresi dapat disusun sebagai berikut :

$$\text{lnKeuntungan} = b_0 + b_1 \text{lnPenjualan} + b_2 \text{lnPromosi} + e$$

Atau

$$\text{lnKeuntungan} = \text{antiln} (-1,420) + 0,664 \text{lnPenjualan} + 0,347 \text{lnPromosi} + e$$

Baik variable Penjualan maupun Promosi memiliki pengaruh terhadap Keuntungan. R Square 0,999 artinya variable Promosi dan Penjualan 99,9 persen dapat menjelaskan terhadap variable terikat (keuntungan) dan sisanya 0,1 persen dijelaskan oleh variable diluar model.



UJI ASUMSI KLASIK

Diketahui data keuntungan, Penjualan dan Biaya Promosi di suatu perusahaan periode Januari 2012 sampai desember 2014 sebagai berikut :

No	Periode	Keunt.	Penjualan	Biaya Promosi	No	Periode	Keunt.	Penjualan	Biaya Promosi
1	2012.01	100.000	1.000.000	55.000	19	2013.07	280.000	2.800.000	140.000
2	2012.02	110.000	1.150.000	56.000	20	2013.08	290.000	2.900.000	145.000
3	2012.03	125.000	1.200.000	60.000	21	2013.09	300.000	3.000.000	150.000
4	2012.04	131.000	1.275.000	67.000	22	2013.10	315.000	3.100.000	152.000
5	2012.05	138.000	1.400.000	70.000	23	2013.11	320.000	3.150.000	160.000
6	2012.06	150.000	1.500.000	74.000	24	2013.12	329.000	3.250.000	165.000
7	2012.07	155.000	1.600.000	80.000	25	2014.01	335.000	3.400.000	170.000
8	2012.08	167.000	1.700.000	82.000	26	2014.02	350.000	3.500.000	175.000
9	2012.09	180.000	1.800.000	93.000	27	2014.03	362.000	3.600.000	179.000
10	2012.10	195.000	1.900.000	97.000	28	2014.04	375.000	3.700.000	188.000
11	2012.11	200.000	2.000.000	100.000	29	2014.05	380.000	3.800.000	190.000
12	2012.12	210.000	2.100.000	105.000	30	2014.06	400.000	3.850.000	192.000
13	2013.01	225.000	2.200.000	110.000	31	2014.07	405.000	3.950.000	200.000
14	2013.02	230.000	2.300.000	115.000	32	2014.08	415.000	4.100.000	207.000
15	2013.03	240.000	2.400.000	120.000	33	2014.09	425.000	4.300.000	211.000
16	2013.04	255.000	2.500.000	125.000	34	2014.10	430.000	4.350.000	215.000
17	2013.05	264.000	2.600.000	130.000	35	2014.11	440.000	4.500.000	219.000
18	2013.06	270.000	2.700.000	135.000	36	2014.12	450.000	4.600.000	210.000

Sumber : data hipotesis

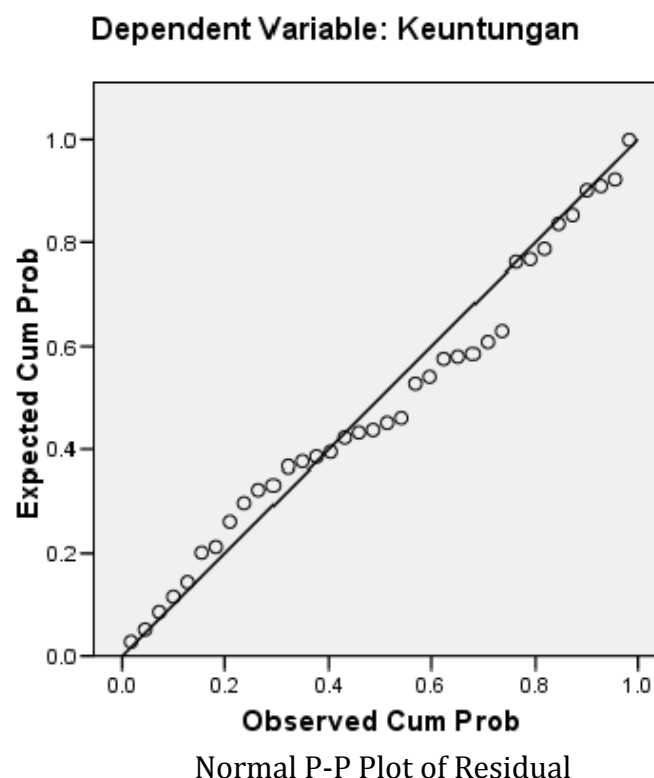
UJI ASUMSI KLASIK ANALISIS REGRESI

a. Uji Normalitas

Uji normalitas berguna untuk menentukan data yang telah dikumpulkan berdistribusi normal atau diambil dari populasi normal. Metode klasik dalam pengujian normalitas suatu data tidak begitu rumit. Berdasarkan pengalaman empiris beberapa pakar statistik, data yang banyaknya lebih dari 30 angka ($n > 30$), maka sudah dapat diasumsikan berdistribusi normal. Biasa dikatakan sebagai sampel besar.

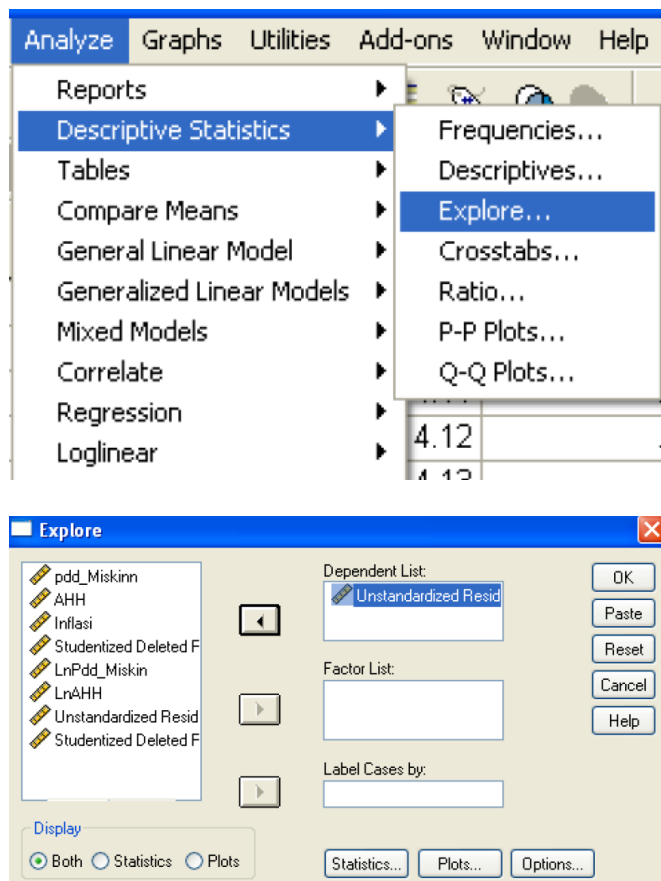
Namun untuk memberikan kepastian, data yang dimiliki berdistribusi normal atau tidak, sebaiknya digunakan uji statistik normalitas. Karena belum tentu data yang lebih dari 30 bisa dipastikan berdistribusi normal, demikian sebaliknya data yang banyaknya kurang dari 30 belum tentu tidak berdistribusi normal, untuk itu perlu suatu pembuktian. uji statistik normalitas yang dapat digunakan diantaranya **Chi-Square, Kolmogorov Smirnov, Lilliefors, Shapiro Wilk, Jarque Bera**.

Salah satu cara untuk melihat normalitas adalah secara visual yaitu melalui **Normal P-P Plot**, Ketentuannya adalah jika titik-titik masih berada di sekitar garis diagonal maka dapat dikatakan bahwa residual menyebar normal,



Namun pengujian secara visual ini cenderung kurang valid karena penilaian pengamat satu dengan yang lain relatif berbeda, sehingga dilakukan **Uji Kolmogorov Smirnov** dengan langkah-langkah:

1. Pilih Analyze → Descriptives → Explore, Setelah muncul Dialog Box seperti pada **Gambar 4,10**, masukkan variabel Unstandardized residual pada kolom **Dependent List**, Pilih **Plots** kemudian Cek list **Box Plot** dan **Normality plots with test** → OK



2. Output yang muncul adalah seperti pada gambar dibawah ini, Sesuai kriteria, dapat disimpulkan bahwa residual menyebar normal.

Tests of Normality						
	Kolmogorov-Smirnov ^a			Shapiro-Wilk		
	Statistic	df	Sig.	Statistic	df	Sig.
Unstandardized Residual	,116	36	,200*	,957	36	,170

*. This is a lower bound of the true significance.

a. Lilliefors Significance Correction

Test normality dapat dilihat dari nilai sig.

jika nilai sig lebih besar dari 5% maka dapat disimpulkan bahwa residual menyebar normal, dan jika nilai sig lebih kecil dari 5% maka dapat disimpulkan bahwa residual menyebar tidak normal.

Dari hasil test of normality diketahui nilai statistik 0,116 atau nilai sig 0,20 atau 20% lebih besar dari nilai α 5%, sehingga maka dapat disimpulkan bahwa residual menyebar normal

b. Uji Autokorelasi

Uji autokorelasi digunakan untuk mengetahui ada atau tidaknya penyimpangan asumsi klasik autokorelasi yaitu korelasi yang terjadi antara residual pada satu pengamatan dengan pengamatan lain pada model regresi. Prasyarat yang harus terpenuhi adalah tidak adanya autokorelasi dalam model regresi. Metode pengujian yang sering digunakan adalah dengan uji Durbin-Watson (uji DW) dengan ketentuan sebagai berikut:

1. Jika d lebih kecil dari dL atau lebih besar dari $(4-dL)$ maka hipotesis nol ditolak, yang berarti terdapat autokorelasi.
2. Jika d terletak antara dU dan $(4-dU)$, maka hipotesis nol diterima, yang berarti tidak ada autokorelasi.
3. Jika d terletak antara dL dan dU atau diantara $(4-dU)$ dan $(4-dL)$, maka tidak menghasilkan kesimpulan yang pasti.

Nilai dU dan dL dapat diperoleh dari tabel statistik Durbin Watson yang bergantung banyaknya observasi dan banyaknya variabel yang menjelaskan.

Sebagai contoh kasus kita mengambil contoh kasus pada uji normalitas pada pembahasan sebelumnya. Pada contoh kasus tersebut setelah dilakukan uji normalitas, multikolinearitas, dan heteroskedastisitas maka selanjutnya akan dilakukan pengujian autokorelasi.

Nilai Durbin Watson pada output dapat dilihat pada Gambar yaitu sebesar 1,641, Sedangkan nilai tabel pembandingan berdasarkan data keuntungan dengan melihat pada Tabel 4,3, nilai $d_{L,\alpha} = 1,153$, sedangkan nilai $d_{U,\alpha} = 1,376$, Nilai $d_{U,\alpha} < dw < 4 - d_{U,\alpha}$ sehingga dapat disimpulkan bahwa **residual tidak mengandung autokorelasi**.

Model Summary^b

Model	R	R Square	Adjusted R Square	Std. Error of the Estimate	Durbin-Watson
1	,999 ^a	,999	,998	4186,51013	1,641

a. Predictors: (Constant), Promosi, Penjualan

b. Dependent Variable: Keuntungan

Model Dengan Ln

Model Summary^b

Model	R	R Square	Adjusted R Square	Std. Error of the Estimate	Durbin-Watson
1	,999 ^a	,999	,998	,01685	1,812

a. Predictors: (Constant), lnPromosi, lnPenjualan

b. Dependent Variable: lnKeuntungan

Nilai Durbin Watson **dalam model ln** pada output dapat dilihat pada Gambar yaitu sebesar 1,812, Sedangkan nilai tabel pembandingan berdasarkan data keuntungan dengan melihat pada Tabel 4,3, nilai $d_{L,\alpha} = 1,153$, sedangkan nilai $d_{U,\alpha} = 1,376$, Nilai $d_{U,\alpha} < dw < 4 - d_{U,\alpha}$ sehingga dapat disimpulkan bahwa **residual tidak mengandung autokorelasi**.

c. Uji Multikolinearitas

Multikolinearitas atau *Kolinearitas Ganda (Multicollinearity)* adalah adanya hubungan linear antara peubah bebas X dalam Model Regresi Ganda. Jika hubungan linear antar peubah bebas X dalam Model Regresi Ganda adalah korelasi sempurna maka peubah-peubah tersebut berkolinearitas ganda sempurna (*perfect multicollinearity*). Sebagai ilustrasi, misalnya dalam menduga faktor-faktor yang memengaruhi konsumsi per tahun dari suatu rumah tangga, dengan model regresi ganda sebagai berikut :

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

dimana :

X_1 : pendapatan per tahun dari rumah tangga

X_2 : pendapatan per bulan dari rumah tangga

Peubah X_1 dan X_2 berkolinearitas sempurna karena $X_1 = 12X_2$. Jika kedua peubah ini dimasukkan ke dalam model regresi, akan timbul masalah Kolinearitas Sempurna, yang tidak mungkin diperoleh pendugaan koefisien parameter regresinya.

Jika tujuan pemodelan hanya untuk peramalan nilai Y (peubah respon) dan tidak mengkaji hubungan atau pengaruh antara peubah bebas (X) dengan peubah respon (Y) maka masalah multikolinearitas bukan masalah yang serius. Seperti jika menggunakan Model ARIMA dalam peramalan, karena korelasi antara dua parameter selalu tinggi, meskipun melibatkan data sampel dengan jumlah yang besar. Masalah multikolinearitas menjadi serius apabila digunakan unruk mengkaji hubungan antara peubah bebas (X) dengan peubah respon (Y) karena simpangan baku koefisiennya regresinya tidak signifikan sehingga sulit memisahkan pengaruh dari masing-masing peubah bebas

Pendeteksian multikolinearitas dapat dilihat melalui nilai *Variance Inflation Factors* (VIF) pada table dibawah ini (model tanpa ln dan Model dengan Ln), Kriteria pengujiannya yaitu apabila nilai $VIF < 10$ maka tidak terdapat mutikolinearitas diantara variabel independent, dan sebaliknya, Pada **tabel** ditunjukkan nilai VIF seluruhnya > 10 , sehingga **asumsi model tersebut mengandung multikolinieritas**.

Coefficients^a

Model		Unstandardized Coefficients		Standardized Coefficients	t	Sig.	Collinearity Statistics	
		B	Std. Error	Beta			Tolerance	VIF
1	(Constant)	-1587,875	2093,274		-,759	,453		
	Penjualan	,060	,009	,602	6,344	,000	,005	202,913
	Promosi	,818	,195	,398	4,191	,000	,005	202,913

a. Dependent Variable: Keuntungan

Model Dengan Ln

Coefficients^a

Model		Unstandardized Coefficients		Standardized Coefficients	t	Sig.	Collinearity Statistics	
		B	Std. Error	Beta			Tolerance	VIF
1	(Constant)	-1,420	,314		-4,527	,000		
	lnPenjualan	,664	,111	,662	5,971	,000	,004	284,794
	lnPromosi	,347	,114	,337	3,043	,005	,004	284,794

a. Dependent Variable: lnKeuntungan

Cara mengatasi multikolinearitas

Beberapa cara yang bisa digunakan dalam mengatasi masalah multikolinearitas dalam Model Regresi Ganda antara lain, Analisis komponen utama yaitu analisis dengan mereduksi peubah bebas (X) tanpa mengubah karakteristik peubah-peubah bebasnya, penggabungan data *cross section* dan data *time series* sehingga terbentuk data panel, metode regresi step wise, metode best subset, metode backward elimination, metode forward selection, mengeluarkan peubah variabel dengan korelasi tinggi walaupun dapat menimbulkan kesalahan spesifikasi, menambah jumlah data sampel, dan lain-lain.

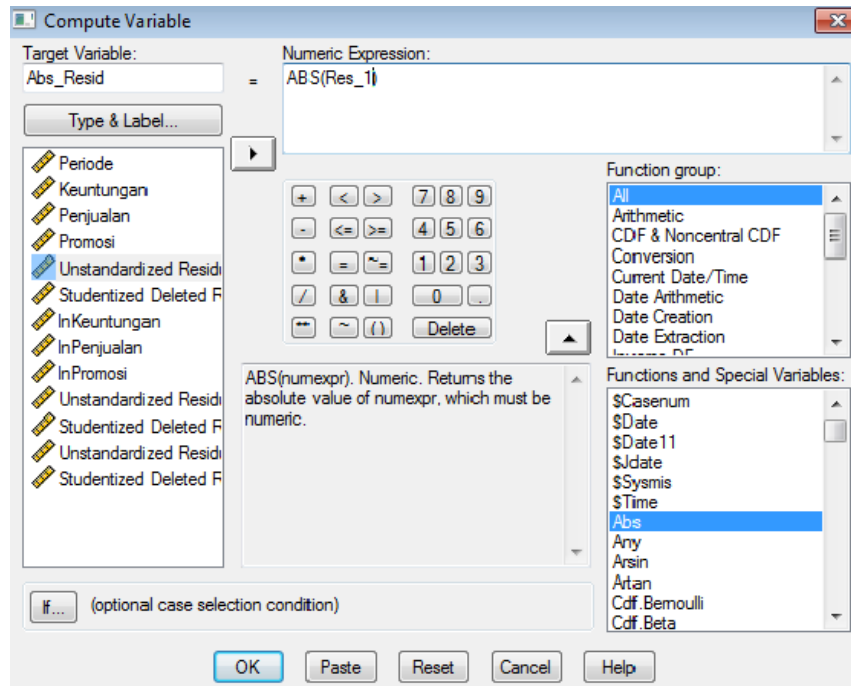
d. Uji Heteroskedastisitas

Heteroskedastisitas adalah adanya ketidaksemaan varian dari residual untuk semua pengamatan pada model regresi.

Mengapa dilakukan uji heteroskedastitas? jawabannya adalah untuk mengetahui adanya penyimpangan dari syarat-syarat **asumsi klasik** pada **model regresi**, di mana dalam model regresi harus dipenuhi syarat tidak adanya heteroskedastisitas.

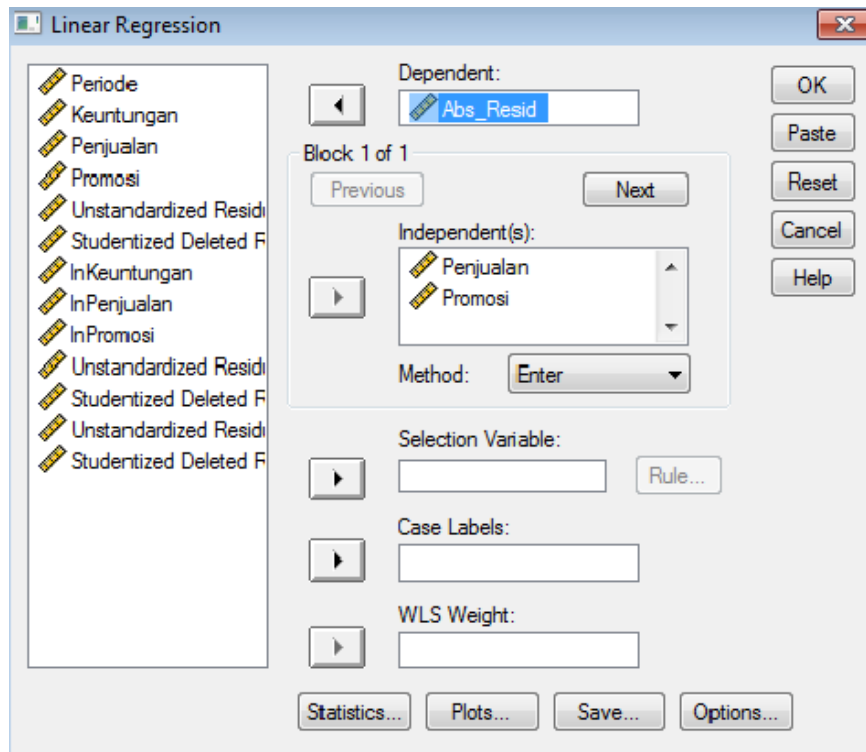
Uji heteroskedastisitas dilakukan dengan cara meregresikan nilai absolute residual dengan variabel – variabel independent dalam model, Langkah-langkahnya adalah:

1. Pilih **Transform → Compute Variable**



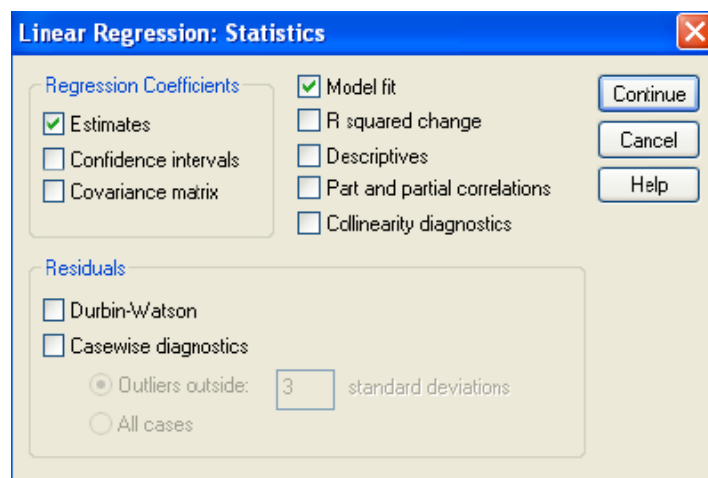
Compute Variable

2. Pilih **All** pada **Function Group** kemudian pilih **Abs** pada **Functions and Special Variables** dengan cara melakukan double klik, Selanjutnya ketik **Abs_Res** pada **Target Variable** dan masukkan **Unstandardized Residual_1** pada **Numeric Expression**, → **OK**
3. Outputnya adalah berupa variabel baru pada **Data View**,
4. Next, pilih **Analyze → Regression → Linear** → Masukkan **Abs_Res** sebagai dependent Variable Sedangkan variabel **Penjualan dan Promosi** sebagai variabel independent.



Linear Regression untuk Uji Glejser

5. Pilih **Estimates** dan **Model Fit** pada Menu **Statistics** → **Continue** → **OK**



Statistics Uji Glejser

6. Perhatikan output regresi antara Residual dengan Variabel-variabel independent lainnya seperti terlihat pada table koefisien dibawah ini, Output menunjukkan tidak adanya hubungan yang signifikan antara seluruh variabel independent terhadap nilai absolute residual, sehingga dapat disimpulkan bahwa **asumsi non-heteroskedastisitas terpenuhi**.

Coefficients^a

Model		Unstandardized Coefficients		Standardized Coefficients	t	Sig.
		B	Std. Error	Beta		
1	(Constant)	1215,233	1335,265		,910	,369
	Penjualan	,004	,006	1,494	,631	,532
	Promosi	-,064	,124	-1,212	-,512	,612

a. Dependent Variable: Abs_Resid

(Gambar 4,16 Output uji Glejser)

DAFTAR PUSTAKA

- Agus Tri Basuki dan Imamudin Yuliadi, *Elektronik Data Prosesing*, Penerbit Danisa Media, Yogyakarta 2014
- Agus Widarjono, *Ekonometrika Teori dan Aplikasi untuk Ekonomi dan Bisnis*, Edisi Kedua, Cetakan Kesatu, Penerbit Ekonisia Fakultas Ekonomi UII Yogyakarta 2007.
- Budiyuwono, Nugroho, *Pengantar Statistik Ekonomi & Perusahaan*, Jilid 2, Edisi Pertama, UPP AMP YKPN, Yogyakarta, 1996.
- Barrow, Mike. *Statistics of Economics: Accounting and Business Studies*. 3rd edition. Upper Saddle River, NJ: Prentice-Hall, 2001
- Catur Sugiyanto. 1994. *Ekonometrika Terapan*. BPFE, Yogyakarta
- Dajan, Anto. *Pengantar Metode Statistik*. Jakarta: Penerbit LP3ES, 1974
- Daniel, Wayne W. *Statistik Nonparametrik Terapan*. Terjemahan Alex Tri Kantjono W. Jakarta: PT Gramedia
- Gujarati, Damodar N. 1995. *Basic Econometrics*. Third Edition. Mc. Graw-Hill, Singapore.
- Hasan, Iqbal M, *Pokok-pokok Materi Statistik 2 (statistic deskriptif)*, Bumi Aksara, Jakarta, 1999.
- Sritua Arif. 1993. *Metodologi Penelitian Ekonomi*. BPFE, Yogyakarta.
- Singgih Santosa, *Berbagai Masalah Statistik dengan SPSS versi 11.5*, Cetakan ketiga, Penerbit PT Elex Media Komputindo Jakarta 2005.
- Johnston, J. and J. Dinardo (1997), *Econometric Methods*, McGraw-Hill
- Koutsoyiannis, A (1977). *Theory of Econometric An Introductory Exposition of Econometric Methods* 2nd Edition, Macmillan Publishers LTD.
- Maddala, G.S (1992). *Introduction to Econometric*, 2nd Edition, Mac-Millan Publishing Company, New York.
- Wonnacott, Thomas H. and Ronald J. Wonnacott. *Introductory Statistics for Business and Economics*. Third edition. New York: John Wiley & Sons, 1990



AGUS TRI BASUKI adalah Dosen Fakultas Ekonomi di Universitas Muhammadiyah Yogyakarta sejak tahun 1994. Mengajar Mata Kuliah Statistik, Ekonometrik, Matematika Ekonomi dan Pengantar Teori Ekonomi. S1 diselesaikan di Program Studi Ekonomi Pembangunan Universitas Gadjah Mada Yogyakarta tahun 1993, kemudian pada tahun 1997 melanjutkan Magister Sains di Pascasarjana Universitas Padjadjaran Bandung jurusan Ekonomi Pembangunan. Dan saat ini penulis sedang melanjutkan Program Doktor Ilmu Ekonomi di Universitas Sebelas Maret Surakarta.

Penulis selain mengajar di Universitas Muhammadiyah Yogyakarta juga mengajar diberbagai Universitas di Yogyakarta. Selain sebagai dosen, penulis juga menjadi konsultan di berbagai daerah di Indonesia.

Selain Buku **Penggunaan SPSS dalam Statistik** , penulis juga menyusun Buku **Statistik Untuk Ekonomi dan Bisnis, Analisis Regresi Untuk Ekonomi dan Bisnis, Elektronik Data Prosesing dan Pengantar Teori Ekonomi**.