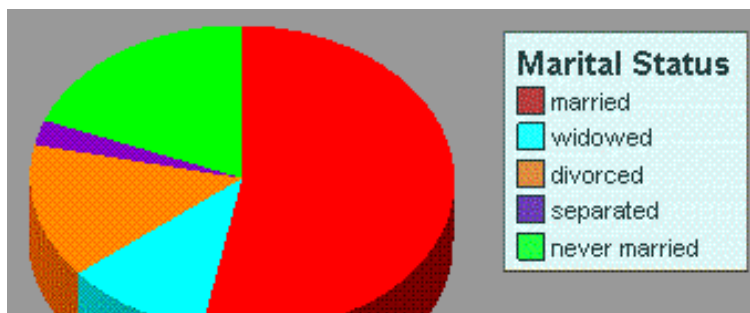
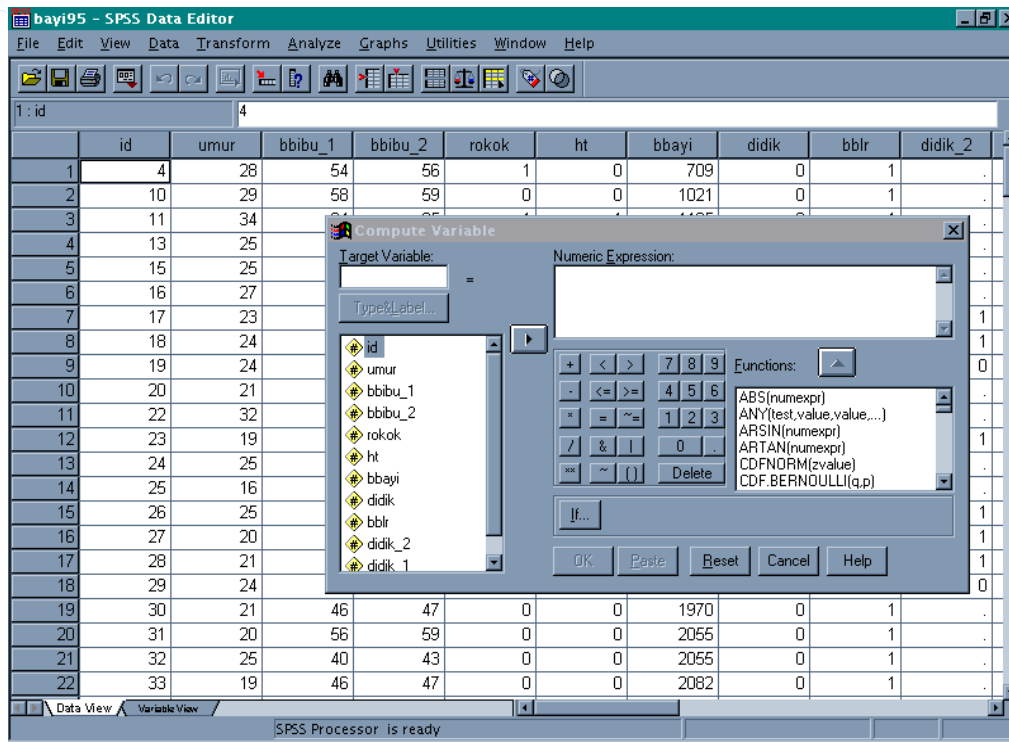


PENGOLAHAN dan ANALISA DATA-1

Menggunakan SPSS



Oleh: BESRAL
Departemen Biostatistika - Fakultas Kesehatan Masyarakat
Universitas Indonesia

KATA PENGANTAR

Pengolahan dan Analisa Data merupakan dua proses yang sangat menentukan dalam pengelolaan data menjadi suatu informasi. Kecepatan dalam pengolahan dan ketepatan dalam analisa akan sangat menentukan kualitas informasi dan penulisan laporan dalam suatu kegiatan monitoring dan evaluasi, baik menggunakan data rutin maupun menggunakan data survei.

Telah banyak buku panduan yang disusun untuk mengolah dan menganalisa data, namun hanya sedikit yang memberikan contoh-contoh secara nyata yang mudah dipahami oleh peserta pemula.

Buku panduan ini sengaja disusun secara sistematis, dengan memberikan contoh persoalan nyata dalam pengolahan dan analisa data. Pengguna buku ini harus dilengkapi dengan file-file data (BAYI95.SAV dan TNG.SAV dan Lebak-1.SAV dan Lebak-2.SAV) untuk dapat memperlihatkan contoh-contoh soal dan penyelesaiannya. Di setiap akhir Bab, diberikan contoh TABEL bagaimana cara penyajian data dan bagaimana menuliskan interpretasi dari hasil analisis data tersebut. Analisis yang dibahas dalam buku ini dibatasi hanya sampai pada tahap hubungan sederhana antara dua variabel saja (*bivariate*), bisa jadi kesimpulan yang didapat belum akurat, karena ukuran yang dihasilkan masih kasar (*crude analysis*). Untuk mendapatkan kesimpulan yang akurat dari hasil analisis data, terutama data survei atau data penelitian bukan eksperimen haruslah dilakukan analisis multivariat. Buku analisis multivariat untuk melihat pengaruh dari beberapa variabel sekaligus mengontrol variabel lainnya tersedia dalam versi berikutnya “Pengolahan dan Analisis Data-2”.

Semoga buku ini dapat dimanfaatkan oleh pengguna untuk membantu dalam pengolahan dan analisis data, baik untuk skripsi/thesis, maupun untuk monitoring/evaluasi program. Kritik dan saran kami terima dengan senang hati untuk kesempurnaan buku ini.

Depok, Agustus 2010

BESRAL

DAFTAR ISI

Kata Pengantar	1
Daftar Isi	2
1. Pengantar SPSS	4
1.1. MEMULAI SPSS	5
1.2. JENDELA SPSS	6
1.3. JENDELA SPSS OUTPUT	8
1.4. MEMASUKKAN (ENTRY) DATA	8
1.5. MENGEDIT DATA (DELETE & COPY)	12
1.6. MENYIMPAN (SAVE) DATA	14
1.7. MEMBUKA (OPEN) DATA SPSS	15
1.8. MEMBUKA (OPEN) DATA .DBF	15
2. Statistika Deskriptif	18
2.1. BUKU KODE	19
2.2. ANALYSIS DESKRIPTIF DATA KATEGORIK	19
2.3. PENYAJIAN DATA KATEGORIK	21
2.4. ANALYSIS DESKRIPTIF DATA NUMERIK	23
2.5. GRAFIK HISTOGRAM PADA DATA NUMERIK	27
2.6. UJI NORMALITAS DISTRIBUSI DATA NUMERIK	28
2.7. PENYAJIAN DATA NUMERIK	30
3. Transformasi Data	31
3.1. PENGERTIAN TRANSFORMASI DATA	32
3.2. ANALISA DESKRIPTIF	35
3.3. TRANSFORMASI DATA DG PERINTAH "RECODE"	38
3.4. TRANSFORMASI DATA DG PERINTAH "COMPUTE"	41
4. Merge File Data	43
4.1. PENGERTIAN MERGE	44
4.2. MERGER dengan ADD VARIABEL	45
4.3. MERGER dengan ADD CASES	47
4.4. MERGER antara INDIVIDU dengan RUMAHTANGGA	48
5. Uji Beda 2-Rata-rata (<i>t-test</i>)	52
5.1. Pengertian	52
5.2. Konsep Uji Beda Dua Rata-rata	52
5.3. Aplikasi Uji-t Dependen pada Data Berpasangan	53
5.4. Penyajian Hasil Uji-t Dependen pada Data Berpasangan	55
5.5. Aplikasi Uji-t pada Data Independen	55
5.6. Penyajian Hasil Uji-t Independen	57

6. Uji Beda $\geq$ 2-Rata-rata (ANOVA)	58
6.1. Pengertian	58
6.2. Konsep Uji ANOVA	59
6.3. Aplikasi Uji ANOVA	59
6.4. Penyajian Hasil Uji ANOVA	63
6.5. Transformasi Jika Varians Tidak Homogen	64
 7. Uji Beda Proporsi (χ^2) Chi-Square	 65
7.1. Pengertian	65
7.2. Konsep Uji Chi Square	66
7.3. Aplikasi Uji χ^2 pada Tabel Silang 2 x 2	67
7.4. Aplikasi Uji χ^2 pada Tabel Silang 2 x 3	70
7.5. Dummy Variabel	71
7.6. Regresi Logistik Sederhana	74
7.6. Penyajian Hasil Uji Beda proporsi	76
 7. Uji Korelasi & Regresi Linier	 77
8.1. Pendahuluan	77
8.2. Asumsi Normalitas	78
8.3. Aplikasi Uji Korelasi Pearson	78
8.4. Aplikasi Regresi Linier (Sederhana)	81
8.5. Penyajian dan Interpretasi Regresi Linier	84
8.6. Prediksi nilai Y	84

DAFTAR PUSTAKA

1 Pengantar SPSS

SPSS Windows merupakan perangkat lunak statistik multiguna yang bermanfaat untuk mengolah dan menganalisis data penelitian. SPSS menggunakan menu serta kotak dialog untuk memudahkan dalam memproses data. Sebagian besar perintah SPSS dapat dilakukan dengan mengarahkan dan mengklik mouse.

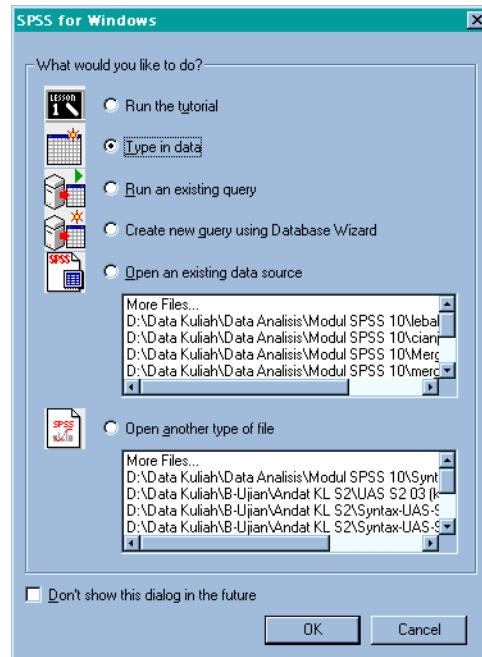
Setelah mempelajari BAB ini, anda akan mengetahui:

- 1. Membuka atau mengaktifkan program SPSS
- 2. Berbagai jenis jendela yang ada di program SPSS (**Data Editor, Output**, dll)
- 3. Membuat Variabel (*Name, Type, Lebar, Decimal, Label, Value, etc.*)
- 4. Memasukkan (*entry*) data pada SPSS
- 5. Menyimpan (*save*) file SPSS
- 6. Membuka (*open*) File SPSS
- 7. Membuka file dari program pengolah data lainnya seperti *dBASE*

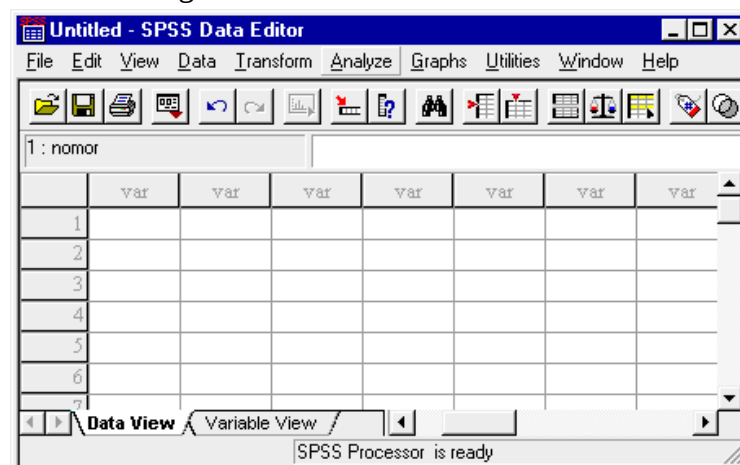
1.1. MEMULAI SPSS

Pertamkali anda harus memastikan bahwa komputer anda sudah diinstall program SPSS for Windows. Sama seperti program Windows lainnya, untuk mengaktifkan SPSS dimulai dari menu **Start**

1. Klik **Start** → **Program** → **SPSS for Windows** → **SPSS 10.0 for Windows**.
2. Pada menu SPSS tertentu (versi 10.x) akan muncul jendela sebagai berikut:



3. Silakan klik (.) **Type in data** kemudian tekan Enter atau klik OK.
4. Layar akan terbuka “**Untitled - SPSS Data Editor**” seperti pada gambar berikut: Selanjutnya disebut sebagai **Jendela Data Editor**. Karena belum ada data, maka tampilannya masih kosong.



5. Perhatikan di kiri bawah ada dua Jendela yaitu (1) **Data View** dan (2) **Variabel View**.

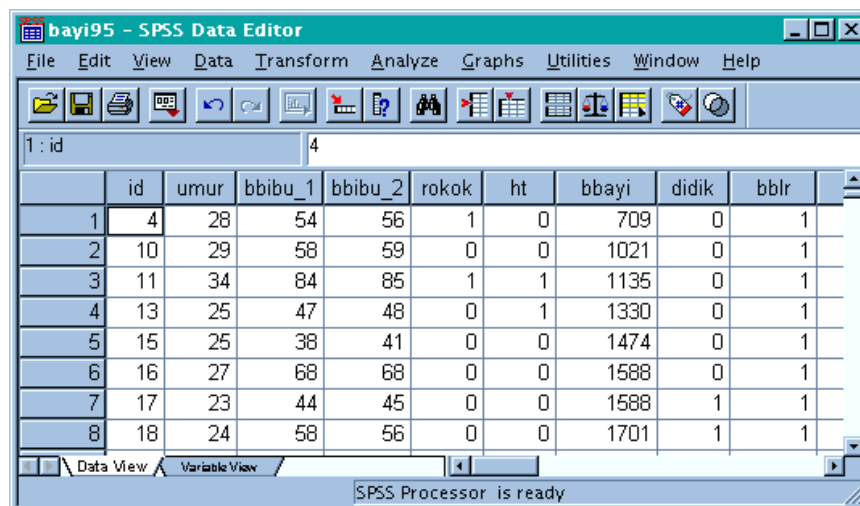
1.2. JENDELA SPSS

Setelah mengaktifkan SPSS akan muncul 2 jendela yaitu “SPSS Data Editor” dan “SPSS Output”.

1.2.1. JENDELA “SPSS DATA EDITOR”

Jendela SPSS Data Editor (selanjutnya disebut **jendela data**) mempunyai 2 tampilan yaitu (1) **Data View** dan (2) **Variabel View**. Data view akan menampilkan database dalam bentuk angka, sedangkan Variabel view menampilkan keterangan tentang variabel yang mencakup: Nama Variabel, Type, Label, Values, dll.

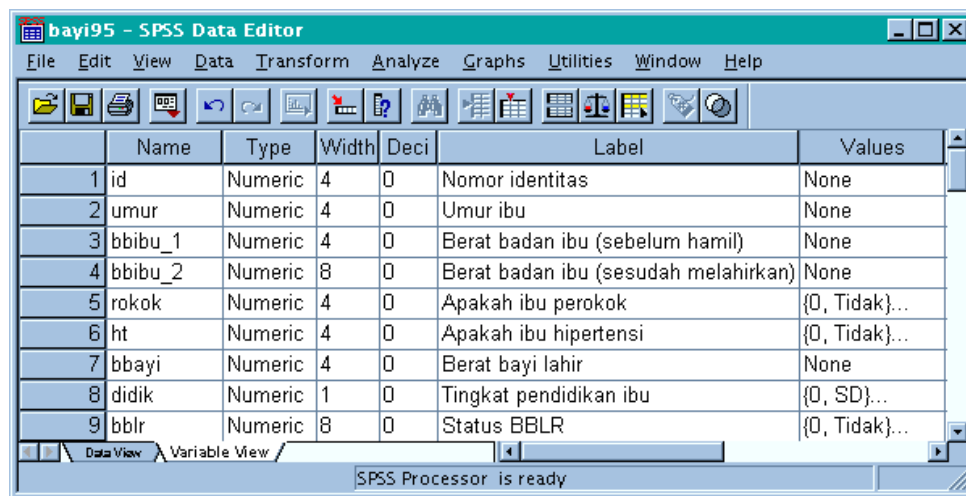
1.2.1.A DATA VIEW



	id	umur	bbibu_1	bbibu_2	rokok	ht	bbayi	didik	bblr
1	4	28	54	56	1	0	709	0	1
2	10	29	58	59	0	0	1021	0	1
3	11	34	84	85	1	1	1135	0	1
4	13	25	47	48	0	1	1330	0	1
5	15	25	38	41	0	0	1474	0	1
6	16	27	68	68	0	0	1588	0	1
7	17	23	44	45	0	0	1588	1	1
8	18	24	58	56	0	0	1701	1	1

Apabila sudah ada data dalam format SPSS (BAYI.SAV), anda bisa membuka data tersebut kemudian bentuk tampilannya pada jendela data atau Data view adalah seperti gambar di atas. (*Prosedur lengkap untuk membuka data BAYI.SAV dapat dilihat pada bagian 1.6*).

1.2.1.B. VARIABEL VIEW



Name atau nama variabel: Aturan pemberian nama variabel adalah 1) Wajib diawali dengan Huruf, dan 2) tidak boleh lebih dari 8 karakter, serta 3) tidak boleh ada spasi (spacebar). Misalnya, anda tidak bisa mengetik “Jenis Kelamin” atau “Je-kel” sebagai variabel, tetapi hanya bisa “Kelamin” saja.

Type atau jenis data: Jenis data yang akan dientry kedalam SPSS dibedakan hanya 2 saja, yaitu 1) Angka atau Numerik (angka: misalnya “18” tahun) dan 2) Huruf atau String (huruf: misalnya Amin, Laki-laki, Jalan Petasan)

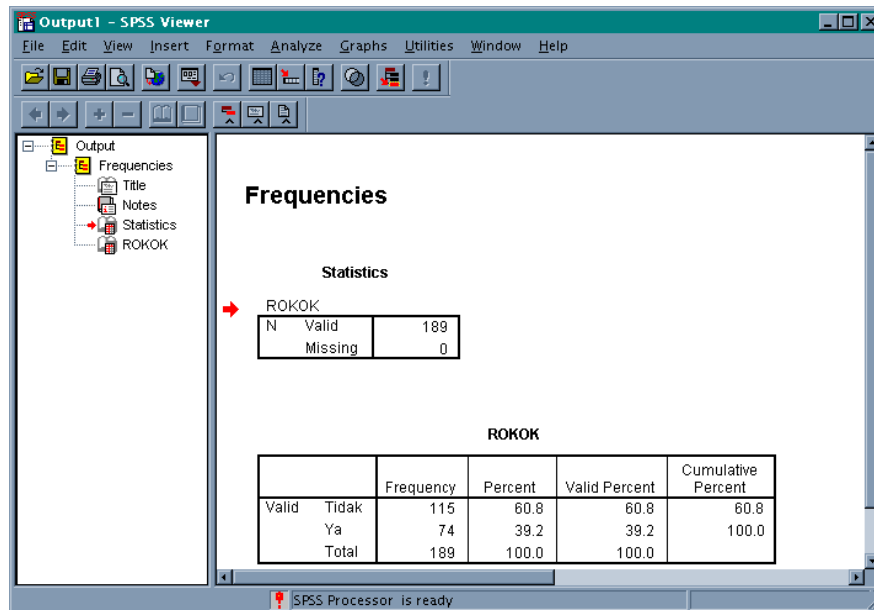
Label atau keterangan variabel: Karena nama variabel tidak boleh lebih dari 8 karakter, biasanya pemberian nama variabel menggunakan singkatan, supaya singkatan tersebut dapat dimengerti maka anda bisa memberi keterangan atau penjelasan terhadap variabel tersebut di kolom label. Misalnya pada variabel “Kelamin” anda bisa memberi label “Jenis Kelamin Anak Balita”, variabel “Food_exp” bisa diberi label dengan “Food expenditure per month” atau “Pengeluaran keluarga untuk makanan satu bulan”.

Values atau kode variabel: Jenis kelamin dapat anda masukkan dengan mengetik “Laki” atau “Perempuan”, tetapi hal ini tidak efisien (waktu dan tenaga hilang percuma). Sebaiknya anda beri kode 1=“Laki” dan 2=“Perempuan”, sehingga anda cukup memasukkan angka 1 atau 2. Supaya nantinya output SPSS yang muncul untuk Kelamin bukan angka 1 dan 2 tetapi yang muncul adalah Laki dan Perempuan, maka anda perlu mengisi Values.

1.3. JENDELA “SPSS OUTPUT”

Walaupun tidak muncul pada saat pertama kali menjalankan program SPSS, ada jendela lain yang terbuka tetapi belum aktif yaitu jendela **Output SPSS Viewer**. Jendela output viewer akan menampilkan hasil-hasil analysis statistik dan graphic yang anda buat. (Selanjutnya disebut **Jendela Output**).

Sebagai contoh pada gambar berikut ditampilkan Jendela Output SPSS Viewer hasil analysis deskriptif distribusi frekuensi dari PEROKOK:



Output SPSS Viewer

1.4. MEMASUKKAN (ENTRY) DATA

Apabila anda belum punya data SPSS (masih mulai dari awal untuk memasukkan data), maka jendela data yang muncul masih kosong. Untuk memulainya, anda dapat membuka jendela Variabel View terlebih dahulu dengan cara meng-klik-nya, selanjutnya mulailah membuat variabel yang dibutuhkan dengan cara mengetik nama variabel yang diinginkan.

Setelah proses pembuatan variabel selesai, selanjutnya buka jendela Data View dan masukkan datanya. Sebagai latihan gunakan contoh data berikut:

Contoh data untuk latihan memasukkan/entry data

<i>No</i>	<i>Nama</i>	<i>Kelamin</i>	<i>Umur</i> }	→ Variabel/field
1	Amin	Laki	28	} → Data/record/responden
2	Aminah	Perempuan	20	
3	Yoyo	Lelaki	36	
4	Yamin	Laki	30	
5	Yongki	Laki	32	
6	Yayang	Perempuan	24	
7	Yovi	Perempuan	22	
8	Yeny	Perempuan	26	
9	Yellow	Perempuan	25	
10	Yeti	Perempuan	21	

1.4.1 PEMBERIAN NAMA, TYPE, & LABEL VARIABEL

Untuk dapat memasukkan data di atas kedalam program SPSS, maka terlebih dahulu anda harus membuat mendefinisikan dan membuat VARIABEL atau FIELD pada jendela Data Editor → Variable View.

Bukalah jendela Data Editor, kemudian klik **Variabel View**, kemudian ketik nama variabel sbb:

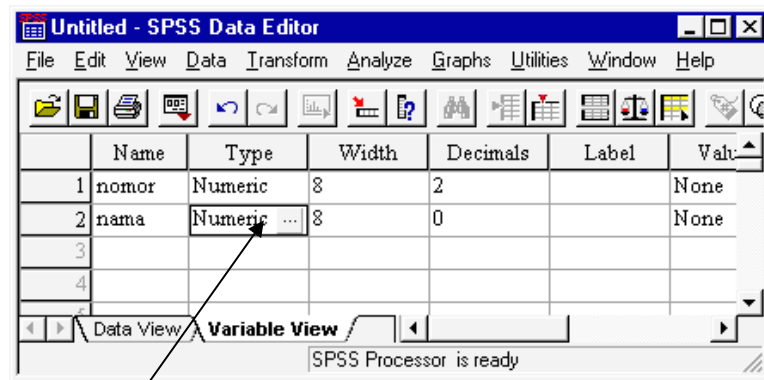
a. variabel NOMOR:

Pada kolom **Name** baris pertama, ketiklah “**nomor**” kemudian tekan enter. Biarkan **Type**-nya **Numerik** karena pada variabel **NOMOR** data yang ingin dimasukkan adalah berbentuk angka. Kemudian kolom **Label** ketik kalimat berikut “**Jenis Kelamin Responden**”.

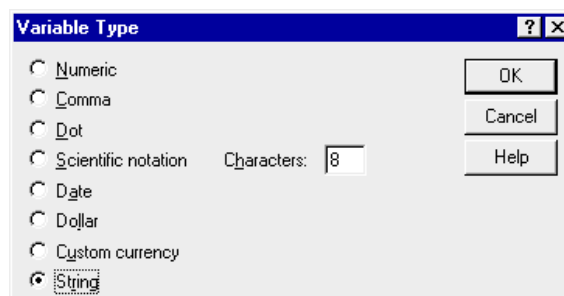
b. variabel NAMA:

Pada kolom **Name** baris kedua, ketiklah “**nama**” kemudian tekan enter. **Type**-nya ganti dengan **String** karena pada variabel **NAMA** data yang ingin dimasukkan adalah berbentuk huruf. Kemudian kolom **label** ketik kalimat berikut “**Nama Responden**”.

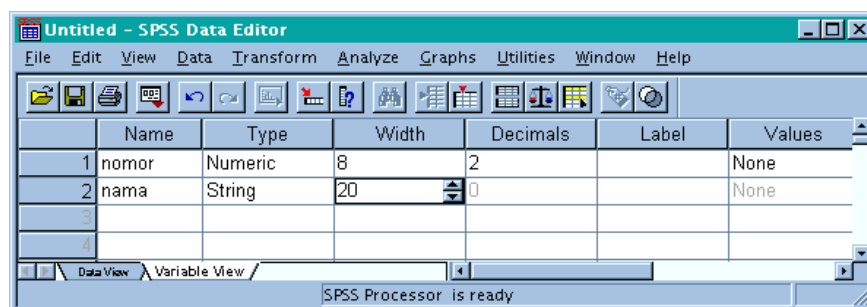
Cara mengganti type dari Numerik menjadi String adalah dengan mengklik bagian akhir dari “Numerik”, sehingga muncul menu **Variabel Type** sebagai berikut:



Klik di sini, untuk merubah Type Variabel, seperti gambar dibawah ini



Gantilah Numerik dengan mengklik **String**, kemudian klik **OK**, hasilnya sbb:



Karena nama responden membutuhkan ruang yang cukup luas, misalnya anda ingin mengetik nama responden sampai 20 karakter, maka silakan ganti **With** dari 8 menjadi 20, dengan cara klik angka 8 tersebut dan ganti dengan mengetik angka 20.

1.4.2 PEMBERIAN KODE VALUE LABELS

Penting untuk diingat pada **data kategorik** atau kualitatif (kelamin, pendidikan, pekerjaan, dll) data yang dimasukkan ke komputer (entry) biasanya untuk efisiensi maka data tersebut dirobah kedalam bentuk **kode angka** (1=laki, 2=Perempuan). Supaya pada saat analisis data tidak terjadi kebingungan, sebaiknya kode tersebut diberi label, dengan langkah sebagai berikut:

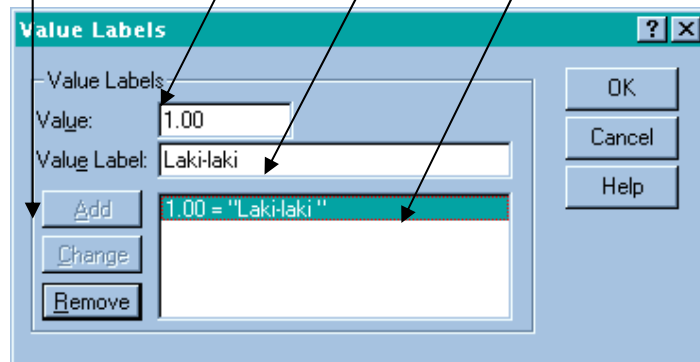
c. variabel KELAMIN:

Pada kolom name baris ketiga, ketiklah “**Kelamin**” kemudian tekan enter. Type-nya

biarkan numerik karena pada variabel **KELAMIN** data yang ingin dimasukkan adalah berbentuk angka 1 atau 2. Kemudian kolom label ketik “**Jenis Kelamin Responden**”.

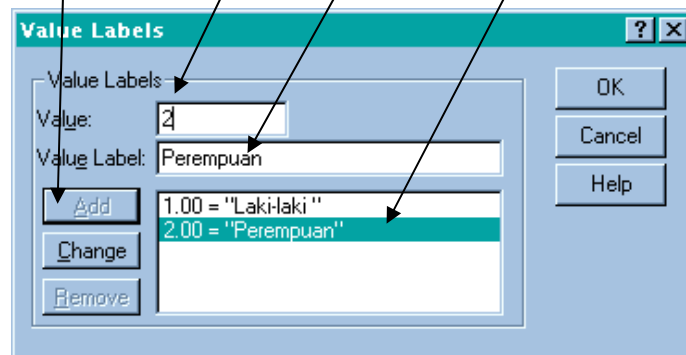
Untuk membuat value label bahwa kode 1 adalah “Laki-laki” dan kode 2 adalah “Perempuan”, maka klik kolom **Values** dan isi sebagai berikut:

1. Pada kotak **Value** isi dengan angka “1”
2. Pada kotak **Value Label** ketik “**Laki Laki**”
3. Kemudian klik **Add**. Sehingga muncul 1=“Laki-laki” pada kotak bawah.



Ulangi prosedur tersebut untuk kode 2=Perempuan,

1. Pada kotak **Value** isi dengan angka “2”
 2. Pada kotak **Value Label** ketik “**Perempuan**”
 3. Kemudian klik **Add**. Sehingga muncul 2=“Perempuan” pada kotak bawah.
- Setelah selesai klik **OK**.

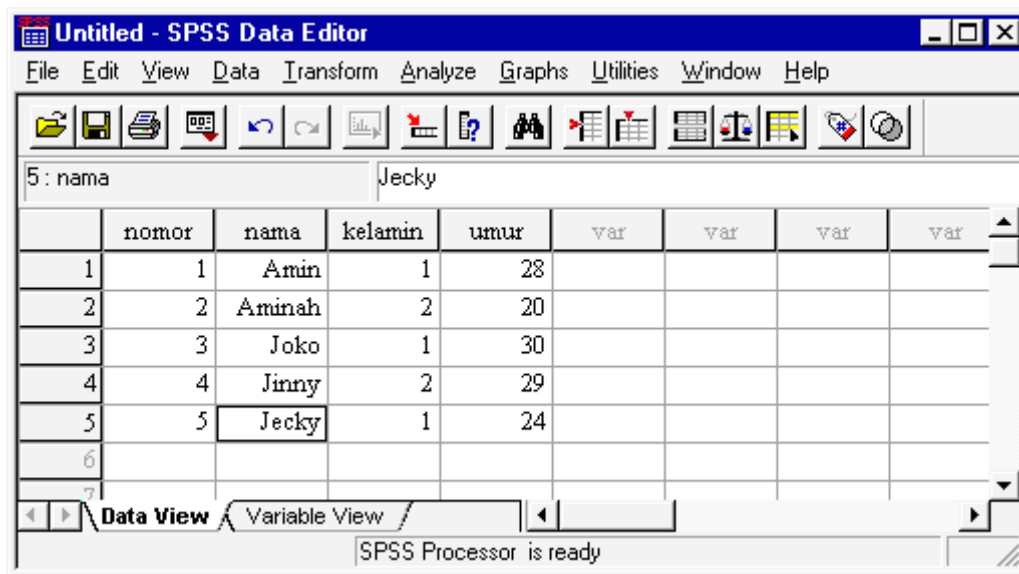


d. variabel UMUR:

Pada kolom **Name** baris keempat, ketiklah “**umur**” kemudian tekan enter. **Type**-nya biarkan **numerik**. Jika angka desimal tidak diperlukan, rubahlah **Decimals** pada kolom ke tiga, sehingga isinya menjadi angka 0 (nol).

1.4.3 MEMASUKKAN DATA

Bukalah **Data View** dengan cara mengkliknya, Kemudian ketik data berikut, seperti data contoh latihan entry data di halaman 5:



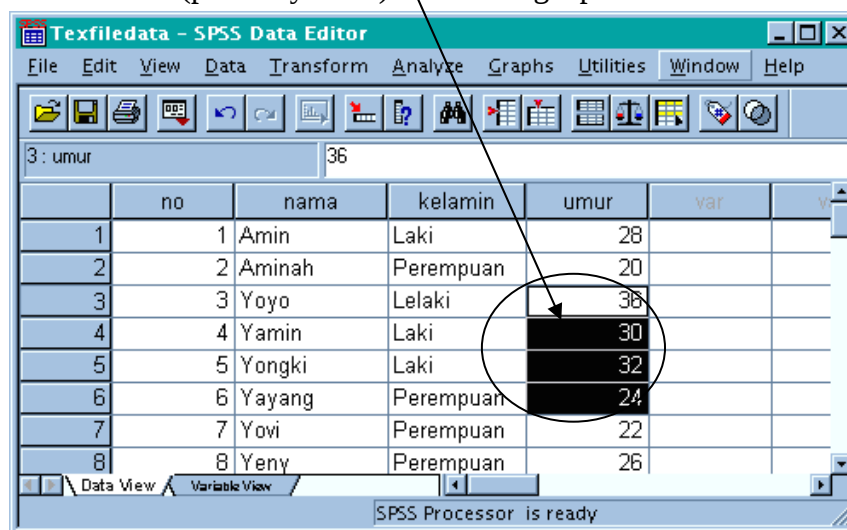
1.5. MENGEDIT DATA (DELETE & COPY)

Editing data biasanya dilakukan untuk menghapus (delete), menggandakan (copy), atau memindahkan (remove) data atau sekelompok data.

1.5.1 MENGHAPUS (DELETE) DATA PADA SEL TERTENTU

Misalnya, ada data yang salah ketik dan ingin dihapus atau diganti dengan data yang benar. Lakukan prosedur sbb:

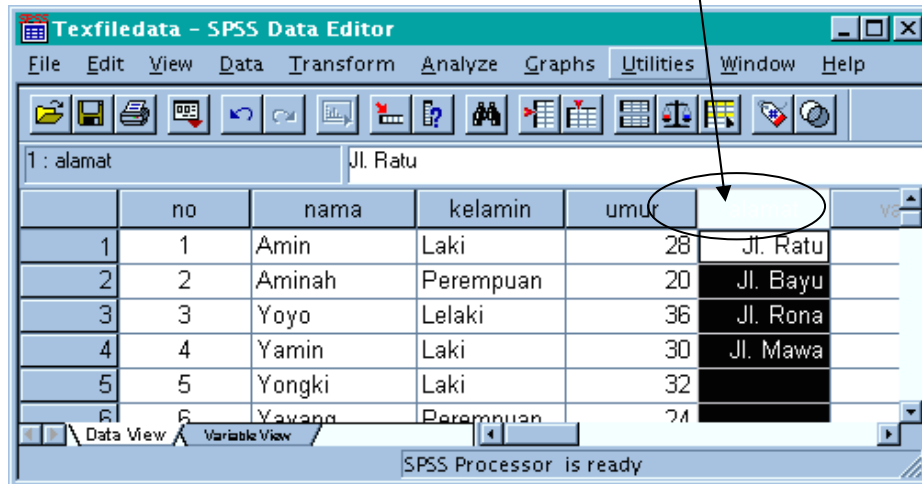
1. **Pilih** sel atau data yang akan dihapus dengan meng-klik (bisa dipilih sekelompok data sekaligus dengan cara **mem-blok** angka dari 36 sampai dengan 24)
2. Tekan tombol **Delete** (pada keyboard) untuk menghapus data tersebut.



1.5.2 MENGHAPUS (DELETE) DATA VARIABEL

Misalnya, ada variabel yang salah ketik dan ingin dihapus atau diganti dengan variabel lainnya. Lakukan prosedur sbb:

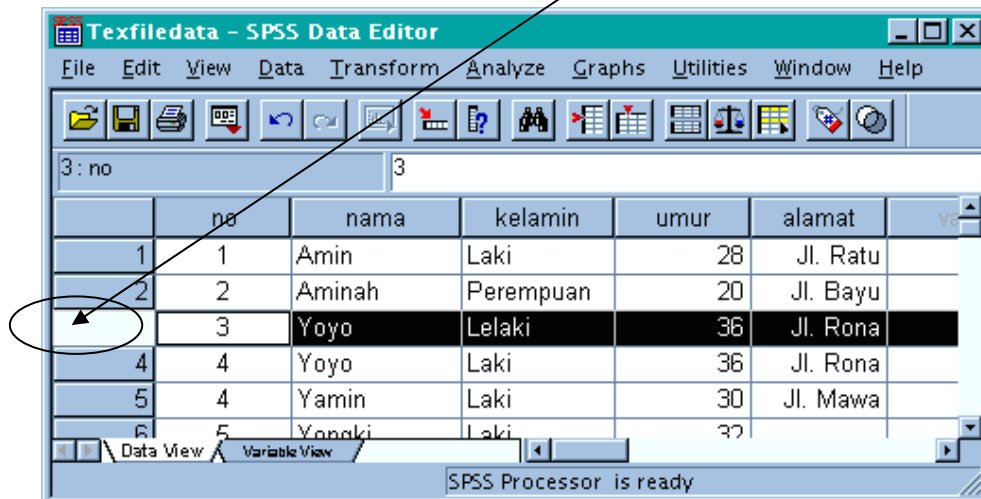
1. **Pilih** variabel yang akan dihapus (mis. alamat) dengan cara **meng-klik**
2. Tekan tombol **Delete** (pada keyboard) untuk menghapus variabel tersebut.



1.5.3 MENGHAPUS (DELETE) DATA RECORD

Misalnya, ada record yang salah ketik (diketik 2 kali) dan ingin dihapus atau diganti dengan variabel lainnya. Lakukan prosedur sbb:

1. **Pilih** record yang akan dihapus (mis. record nomor 3) dengan cara **meng-klik**
2. Tekan tombol **Delete** (pada keyboard) untuk menghapus variabel tersebut.

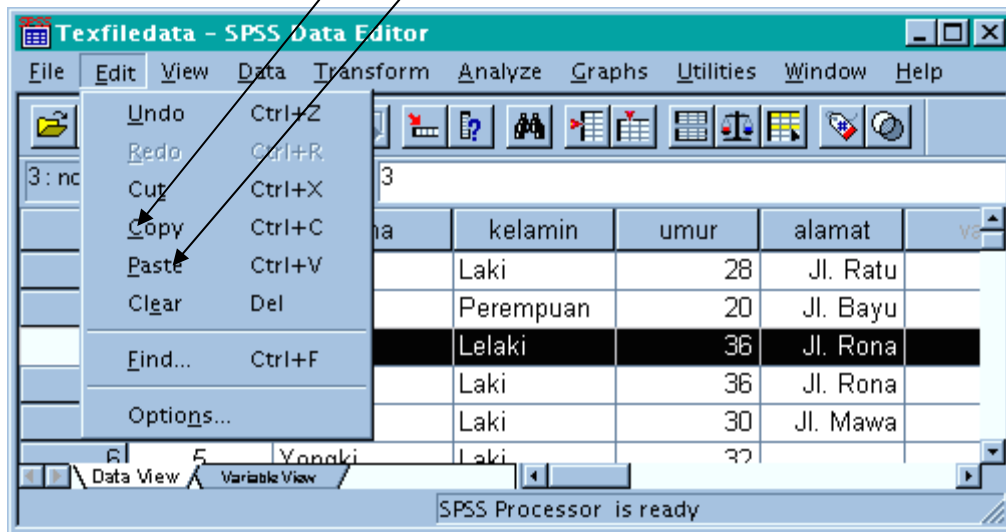


1.5.4 MENGGANDAKAN (COPY) DATA

Prosedur penggandaan (copy) data pada SPSS mirip dengan prosedur meng-copy pada umumnya dalam perintah komputer. Sebagai berikut:

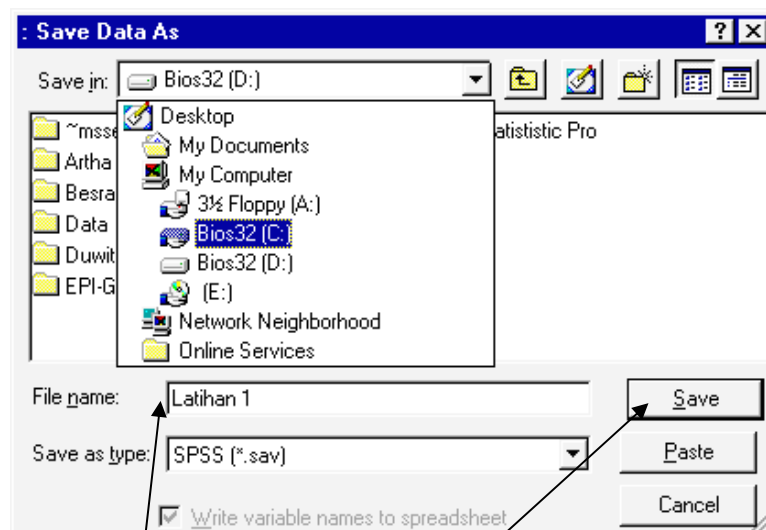
1. Dimulai dengan **memilih** data atau sel yang akan dicopy dengan cara **meng-klik** (pemilihan dapat dilakukan pada **sekelompok data, variabel, atau record**)

2. Kemudian pilih menu **Edit → Copy** (atau Ctrl + C, pada key board)
3. Kemudian letakkan kursor pada lokasi yang akan dicopykan
4. Kemudian pilih menu **Edit → Paste** (atau Ctrl + V, pada key board)



1.6. MENYIMPAN (SAVE) DATA

Pilihlah (kemudian klik) gambar disket yang ada di kiri atas atau Pilih **File → Save**. Atau **File → Save As..**



Jika anda baru menyimpan untuk pertamakali, maka akan muncul menu seperti gambar di atas (menu Save As..). Menu ini hanya muncul pertama kali saja, selanjutnya tidak muncul lagi, kecuali dengan perintah Save As.

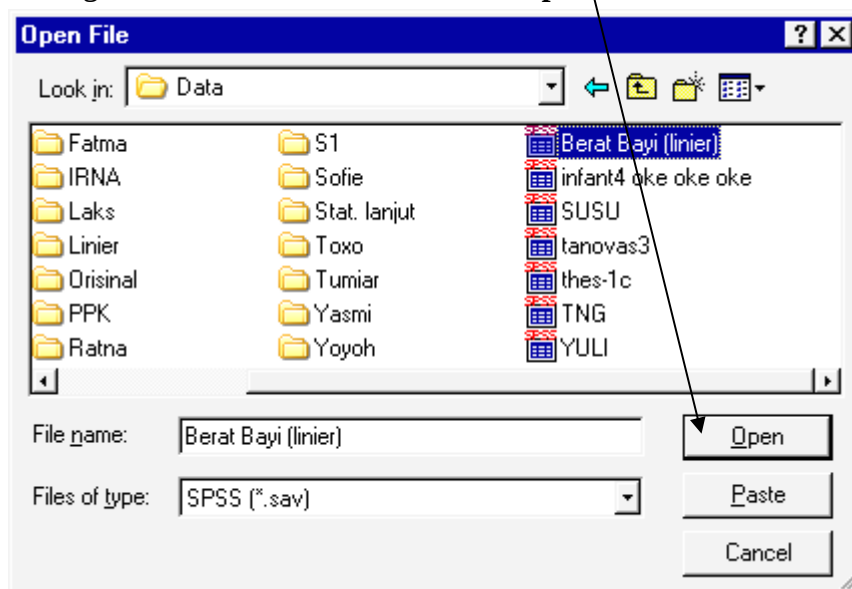
Isi kotak **File name** dengan "**Latihan 1**". Pilihlah **Save in** untuk menentukan apakah anda akan menyimpan di **Disket** (Floppy: A) atau di **Hardisk**:C. Jika anda pilih hardisk, jangan lupa untuk menentukan lokasi **Directory** mana tempat penyimpanan tersebut. Klik **save** untuk menjalankan proses penyimpanan.

Selesai proses Saving, perhatikan di kiri atas “Untitled – SPSS Data Editor” sudah berubah menjadi “Latihan 1 – SPSS Data Editor”

1.7. MEMBUKA (OPEN) DATA SPSS

Jika anda sudah mempunyai data dalam format SPSS yang disimpan di Disket atau di Hardisk, silakan buka dengan SPSS, sebagai berikut:

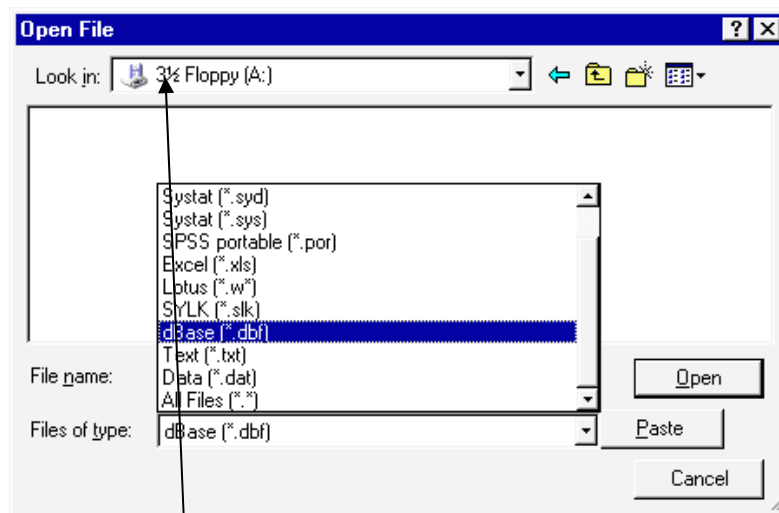
1. Pastikan anda berada di layar “SPSS Data Editor”, kemudian pilihlah menu **File** → **Open**
2. Pada **File of type**, pilihan standarnya adalah SPSS (*.sav), jika bukan ini yang muncul maka anda harus memilihnya terlebih dahulu
3. Pada **Look in**, pilihlah Drive yang sesuai (A:C:D) dan Directory tempat data tersimpan (mis. C:\Data\....)
4. Akan muncul daftar File yang ber-extensi.SAV, pilihlah file yang akan anda buka dengan mengklik file tersebut, kemudian klik **Open**



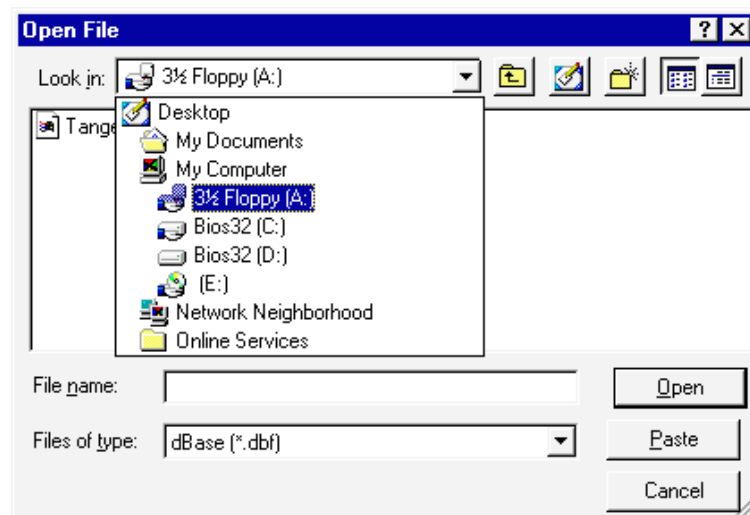
1.8. MEMBUKA (OPEN) DATA.DBF

SPSS punya kemampuan untuk membuka data dari Format lain seperti Dbase, Lotus, Excell, Foxpro, dll. Misalnya anda punya data Tangerang.DBF yang disimpan di Disket atau di Hardisk, silakan buka dengan SPSS, sebagai berikut:

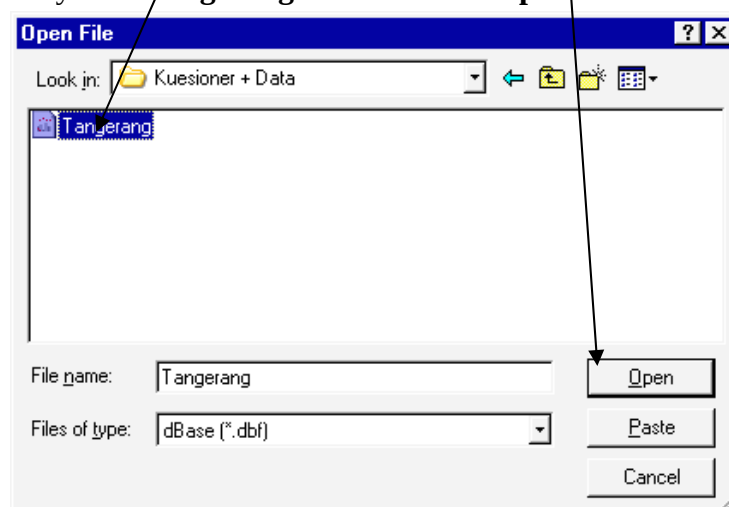
1. Pastikan anda berada di layar “SPSS Data Editor”, kemudian pilihlah menu **File** → **Open**
2. Pada **File of type**, pilihlah **dBase (*.dbf)**. (Selain dBASE anda bisa memilih program pengolah kata lainnya yang sesuai dengan keinginan)



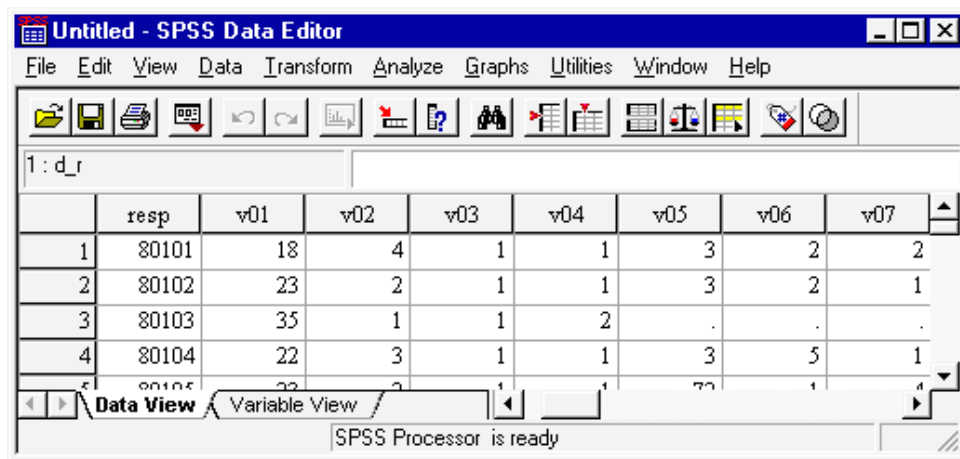
3. Pada **Look in**, pilihlah Floppy:A, jika data anda ada di Disket



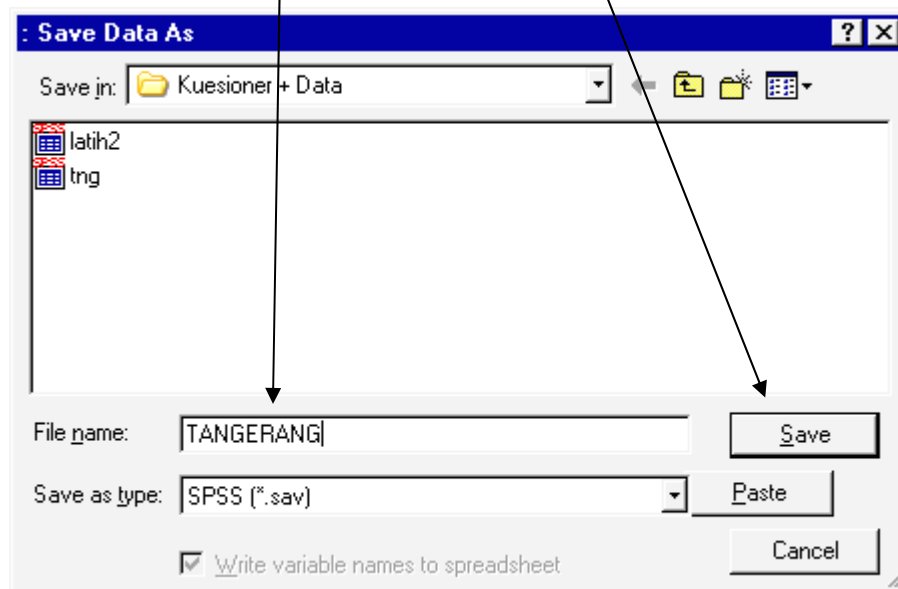
4. Secara otomatis akan muncul list file yang berekstensi DBF, klik file yang ingin dibuka, misalnya file **Tangerang** kemudian klik **Open**.



5. Maka data Tangerang.DBF akan muncul di **"Untitled – SPSS Data Editor"**. Laporan dari proses konversi data dari dBase tersebut akan dimunculkan di **"Output – SPSS Viewer"** dan Datanya sendiri akan muncul di Data View



6. Agar data tersebut tersimpan dalam bentuk file SPSS (*.SAV), maka anda harus menyimpannya dengan cara mengklik gambar disket di kiri atas atau pilih menu **File** → **Save**. Isi kotak File dengan nama yang anda inginkan, misalnya “DATA TNG” atau “TANGERANG”. Klik **Save** untuk menjalankan prosedur penyimpanan.



Setelah klik save, pastikan kiri atas layar monitor anda yang sebelumnya muncul “**Untitled – SPSS Data Editor**” telah berubah menjadi “**TANGERANG – SPSS Data Editor**”.

Pastikan anda menyimpan setiap saat data yang sudah diolah, agar jika sewaktu-waktu komputer mengalami kerusakan (mis. Listrik mati, komputer hang), maka anda tidak kehilangan data.

2 Statistik Deskriptif

Statistik deskriptif berupa frekuensi dan nilai pusat (*central tendency*).

Frekuensi biasanya dimunculkan dalam bentuk proporsi atau persentase untuk data atau variabel kategorik. Sedangkan nilai pusat berupa nilai tengah dan nilai sebaran (mean, median, SD, SE, dll) untuk data atau variabel numerik. Statistik deskriptif ini juga akan dilengkapi dengan grafik histogram untuk data numerik.

Setelah mempelajari BAB ini, anda akan mengetahui:

- 1. Buku Kode
- 2. Analisis Deskriptif Data Kategorik
- 3. Penyajian Data Kategorik
- 4. Analisis Deskriptif Data Numerik
- 5. Grafik Histogram
- 6. Uji Normalitas
- 7. Penyajian Data Numerik

2.1. BUKU KODE

Mulai Bab 2 kita akan membicarakan prosedur statistik deskriptif yang sering digunakan dalam melakukan analisis data. Untuk data latihan, kita akan menggunakan file BAYI95.SAV yang berisi variabel yang mempengaruhi berat bayi lahir. Agar kita bisa mengolah data tersebut, maka kita harus mengetahui keterangan dari variabel dan value-nya yang biasanya dimuat dalam buku kode. Buku kode untuk file tersebut adalah sbb:

Variabel	Keterangan
ID	Nomor identifikasi responden
UMUR	Umur ibu (tahun)
BBIBU_1	Berat badan ibu (kg) sebelum hamil (Pre-)
BBIBU_2	Berat badan ibu (kg) sesudah melahirkan (Post-)
ROKOK	Kebiasaan merokok dari ibu 0 = Tidak 1 = Ya
HT	Penyakit hipertensi pada ibu 0 = Tidak 1 = Ya
BBAYI	Berat bayi lahir (gram)
DIDIK	Pendidikan ibu 0 = Rendah 1 = Sedang 2 = Tinggi
BBLR	Status berat bayi lahir rendah 0 = Tidak 1 = Ya

Dalam melakukan analisis data, kita harus memahami terlebih dahulu konsep dari jenis data statistik yaitu data **Numerik** dan data **Kategorik**. Data numerik adalah data yang berbentuk angka (kombinasi dari 0,1,2...9), yang merupakan gambaran dari hasil mengukur atau menghitung. Sedangkan data kategorik merupakan data yang berbentuk pernyataan, kualitas, atau pengelompokan (misalnya: laki/perempuan, baik/buruk, setuju/tidak setuju, SD/SMP/SMU/PT, rendah/sedang/tinggi, dll).

Analisis data numerik akan berbeda dengan analisis data kategorik, termasuk cara penyajian dan cara interpretasinya. Data numerik biasanya ditampilkan dalam bentuk nilai tengah dan nilai sebaran (misalnya nilai rata-rata dan standar deviasi). Sedangkan data kategorik ditampilkan dalam bentuk persentase atau proporsi.

2.2. ANALYSIS DESKRIPTIF DATA KATEGORIK

Cara yang paling sering digunakan untuk menampilkan data katagorikal adalah

dengan menggunakan tabel distribusi frekuensi. Kita akan coba membuat tabel distribusi frekuensi pendidikan ibu dari file BAYI95.SAV.

1. Bukalah file BAYI95.SAV, sehingga data tampak di jendela **Data Editor** (*prosedur untuk membuka file dapat dilihat pada bagian 1.7*).
2. Prosedur untuk menampilkan distribusi frekuensi adalah sebagai berikut:

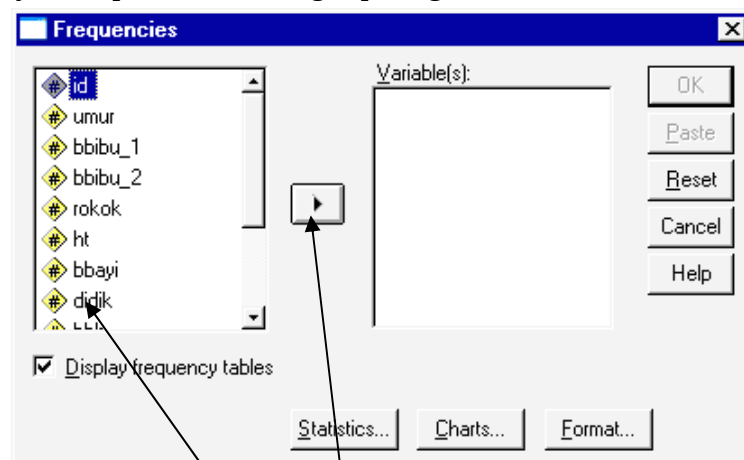
Dari menu utama, pilihlah:

Analyze

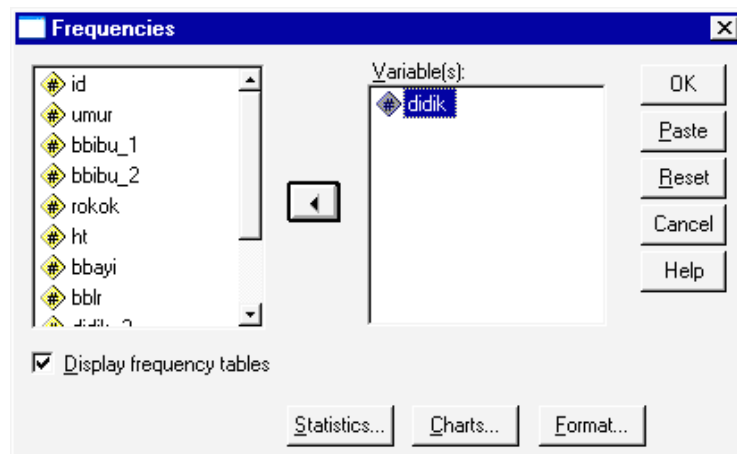
Descriptive Statistic <

Frequencies...

Pada layar tampak kotak dialog seperti gambar berikut:



3. Pada kotak dialog tersebut, klik pada variabel *DIDIK* yang terdapat pada kotak sebelah kiri. Kemudian klik tanda >, sehingga kotak dialog menjadi seperti gambar berikut:



4. Klik **OK** untuk menjalankan prosedur. Pada *jendela output* tampak hasil seperti berikut:

DIDIK

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	SD	47	24.9	24.9	24.9
	SMP	84	44.4	44.4	69.3
	SMA	58	30.7	30.7	100.0
	Total	189	100.0	100.0	

Pada kolom **Frequency** menunjukkan jumlah kasus dengan nilai yang sesuai. Jadi pada contoh di atas, ada 47 ibu yang berpendidikan SD dari 189 ibu yang ada. Proporsi dapat dilihat pada kolom **Percent**, pada contoh di atas, ada 24,9% ibu yang berpendidikan SD.

Kolom **Valid Percent** menampilkan proporsi jika missing cases tidak diikutsertakan sebagai penyebut. Pada contoh di atas, kolom **Percent** dan **Valid Percent** memberikan hasil yang sama karena pada data ini tidak ada missing cases. **Cumulative Percent** menjelaskan tentang persen kumulatif, jadi pada contoh di atas, ada 69,3% ibu yang berpendidikan SD dan SMP (24.9% + 44.4%).

2.3. PENYAJIAN DATA KATEGORIK

Penyajian data mempunyai prinsip efisiensi, artinya sajikan hanya informasi penting saja, jangan semua output komputer disajikan dalam laporan. Contoh penyajian data kategorik sbb:

Tabel 1. TINGKAT PENDIDIKAN RESPONDEN

	Frequency	Percent
SD	47	24.9
SMP	84	44.4
SMA	58	30.7
Total	189	100.0

Contoh Interpretasi:

“Distribusi frekuensi tingkat pendidikan responden dapat dilihat pada Tabel-1, terlihat bahwa sebagian besar responden adalah tamat SMP (44.4%), kemudian diikuti oleh tamat SMA sebanyak 30,7% dan sisanya hanya tamat SD (24,9%).”

LATIHAN ANALYSIS DATA KATEGORIK

Latihan:

1. Buatlah tabel distribusi frekuensi untuk variabel *HT*, *ROKOK*,
 - a. Sajikan
 - b. Interpretasikan

2. Buatlah distribusi frekuensi dari variabel *UMUR_KEL* dan *BBLR* setelah Anda melakukan pengelompokan ulang (lihat Bab 3: *Transformasi Data* untuk mengetahui prosedur pengelompokan ulang),
 - a. Sajikan
 - b. Interpretasikan

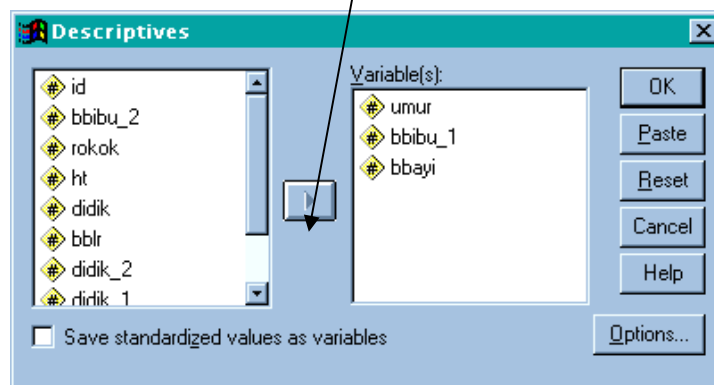
2.4. ANALYSIS DESKRIPTIF DATA NUMERIK

Pada data numerik atau kontinyu, peringkasan data dapat dilakukan dengan melaporkan ukuran tengah dan sebarannya. Ukuran tengah yang dapat digunakan adalah rata-rata, median dan modus. Sedangkan ukuran sebaran yang dapat digunakan adalah nilai minimum, maksimum, range, standar deviasi dan persentil. Dari ukuran-ukuran tersebut, yang paling sering digunakan adalah rata-rata dan standar deviasi. Sebagai contoh, kita akan coba mencari ukuran tengah dan sebaran dari *UMUR*, *BBIBU* dan *BBAYI*.

1. Bukalah file BAYI95.SAV (jika file ini belum dibuka), sehingga data tampak di jendela Data Editor. (prosedur untuk membuka file dapat dilihat pada bagian 1.7).

Perintah Descriptive..

2. Dari menu utama, pilihlah:
Analyze
Descriptive Statistic <
Descriptive ...
3. Pada kotak dialog tersebut, klik pada variabel *UMUR* yang terdapat pada kotak sebelah kiri. Tekan **Ctrl** (jangan dilepas), Klik variabel *BBIBU_1*, dan klik variabel *BBAYI*, lepaskan **Ctrl**.. Dengan cara ini kita memilih 3 variabel sekaligus. Kemudian klik tanda <, sehingga kotak dialog menjadi seperti gambar berikut:



4. Klik **OK** untuk menjalankan prosedur. Pada layar tampak hasil seperti berikut:

	N	Minimum	Maximum	Mean	Std. Deviation
UMUR	189	14	45	23.24	5.30
BBIBU_1	189	36	112	58.39	13.76
BBAYI	189	709	4990	2944.66	729.02
Valid N (listwise)	189				

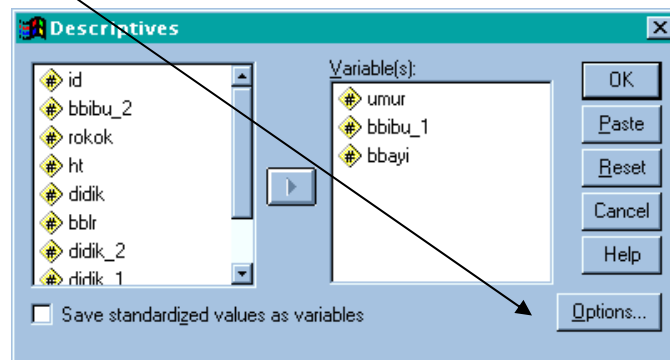
Nilai rata-rata dapat dilihat pada kolom **Mean**, sedangkan nilai standar deviasi dapat

dilihat pada **Std Deviation**. Pada contoh di atas, rata-rata umur ibu adalah 23,24 tahun dengan standar deviasi 5,30 tahun dan umur minimum 14 tahun serta umur maksimum 45 tahun.

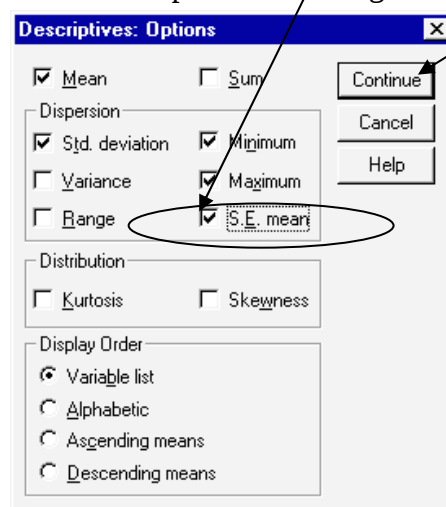
Dengan cara di atas, kita dapat memperoleh nilai rata-rata, minimum, maksimum serta standar deviasi. Tetapi kita tidak memperoleh nilai standar error, padahal nilai ini diperlukan untuk melakukan estimasi interval pada parameter populasi.

Perintah Option..

5. Jika Anda juga ingin agar SPSS menampilkan **standar error**, anda dapat memilih menu **Options**.



Misalnya anda menginginkan standar error maka klik **SE Mean**, kemudian klik **Continue** dan **OK** hasilnya pada Jendela Output adalah sebagai berikut:



Descriptive Statistics

	N	Minimum	Maximum	Mean		Std.
	Statistic	Statistic	Statistic	Statistic	Std. Error	Statistic
UMUR	189	14	45	23.24	.39	5.30
BBIBU_1	189	36	112	58.39	1.00	13.76
BBAYI	189	709	4990	2944.66	53.03	729.02
Valid N (listwise)	189					

Dari hasil tersebut kita dapat melakukan estimasi interval dari berat bayi. Kita dapat menghitung 95% *confidence interval* berat bayi, yaitu $2944,66 \pm 1,96 \times 53,03$

(mean $\pm$ SE mean). Jadi kita 95% yakin bahwa rata-rata berat bayi di populasi berada pada selang 2840,72 sampai 3048,60 gram.

Perintah Explore..

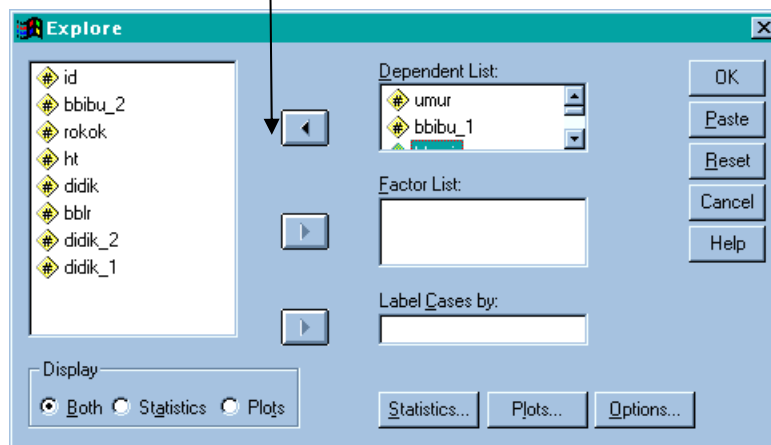
6. Cara yang lain untuk mengeluarkan nilai statistik deskriptif dari data numerik (nilai rata-rata/mean std. Dev) beserta 95% *confidence interval* adalah sebagai berikut:
Dari menu utama, pilihlah:

Analyze

Descriptive Statistic <

Explore...

7. Pada kotak dialog tersebut, klik pada variabel *UMUR* yang terdapat pada kotak sebelah kiri. Tekan **Ctrl** (jangan dilepas), Klik variabel *BBIBU_1*, dan klik variabel *BBAYI*, lepaskan **Ctrl**. Dengan cara ini kita memilih 3 variabel sekaligus. Kemudian klik tanda $\leftarrow$, sehingga ketiga variabel tersebut masuk ke kota **Dependent List** seperti gambar berikut:



8. Klik **OK** untuk menjalankan prosedur, sehingga hasilnya seperti gambar berikut:

Descriptives			Statistic	Std. Error
BBIBU_1	Mean		58.39	1.00
	95% Confidence Interval for Mean	Lower Bound	56.42	
		Upper Bound	60.37	
	5% Trimmed Mean		57.29	
	Median		54.00	
	Variance		189.463	
	Std. Deviation		13.76	
	Minimum		36	
	Maximum		112	
	Range		76	
	Interquartile Range		13.00	
	Skewness		1.395	.177
	Kurtosis		2.366	.352
BBAYI	Mean		2944.66	53.03
	95% Confidence Interval for Mean	Lower Bound	2840.05	
		Upper Bound	3049.26	
	5% Trimmed Mean		2957.83	
	Median		2977.00	
	Variance		531473.7	
	Std. Deviation		729.02	
	Minimum		709	
	Maximum		4990	
	Range		4281	
	Interquartile Range		1069.00	
	Skewness		-.210	.177
	Kurtosis		-.081	.352

Dari hasil tersebut kita mendapatkan estimasi titik dan estimasi interval dari variabel numerik yang diukur. Kita dapat melihat nilai rata-rata dan 95% *confidence interval* dari BIBU_1 yaitu 58,39 kg (56,42—60,37), artinya kita 95% yakin bahwa rata-rata berat ibu di populasi berada pada selang 56,42 sampai 60,37 kg. Untuk BBAYI yaitu 2944,66 gram (2840,05—3049,26), kita 95% yakin bahwa rata-rata berat bayi di populasi berada pada selang 2840,05 sampai 3049,26 gram. Nilai ini tidak jauh berbeda dengan nilai yang dihitung dari output yang didapat pada langkah no.5 sebelumnya.

2.5. GRAFIK HISTOGRAM PADA DATA NUMERIK

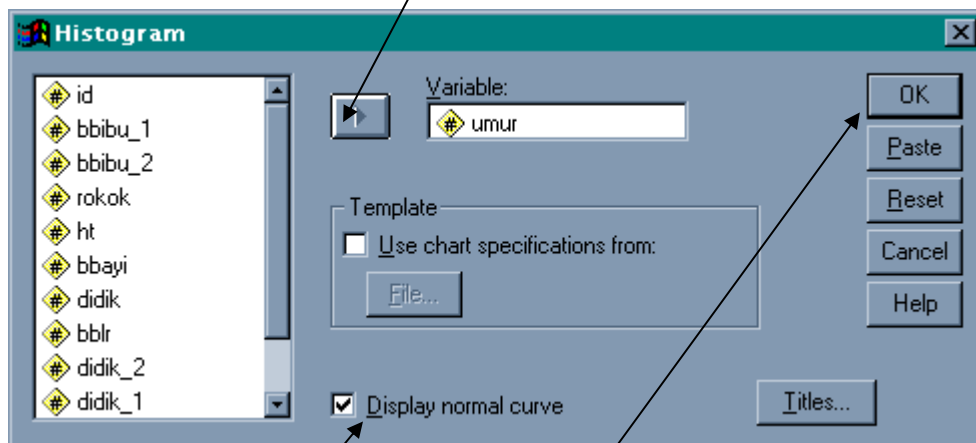
Analisis data Numerik akan lebih lengkap apabila dilengkapi dengan grafik. Salah satu Grafik yang cocok untuk data numerik adalah HISTOGRAM.

1. Dari menu utama, pilihlah:

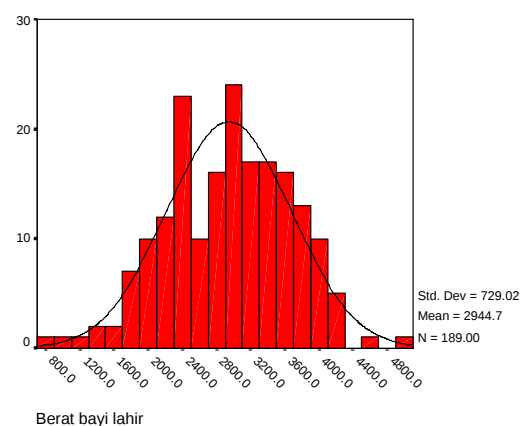
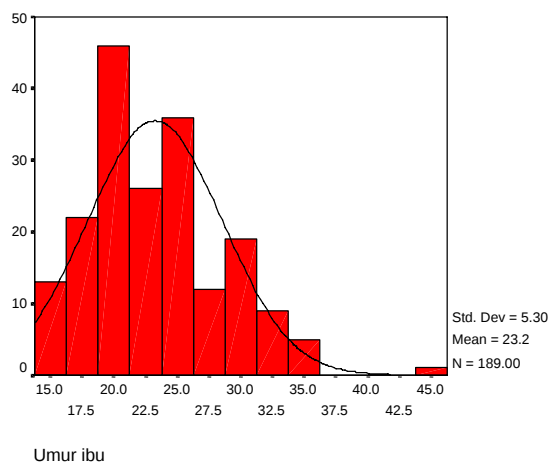
Graphs

Histogram...

2. Pada kotak dialog tersebut, klik pada variabel **UMUR** yang terdapat pada kotak sebelah kiri. Kemudian klik tanda <, sehingga kotak dialog seperti berikut:



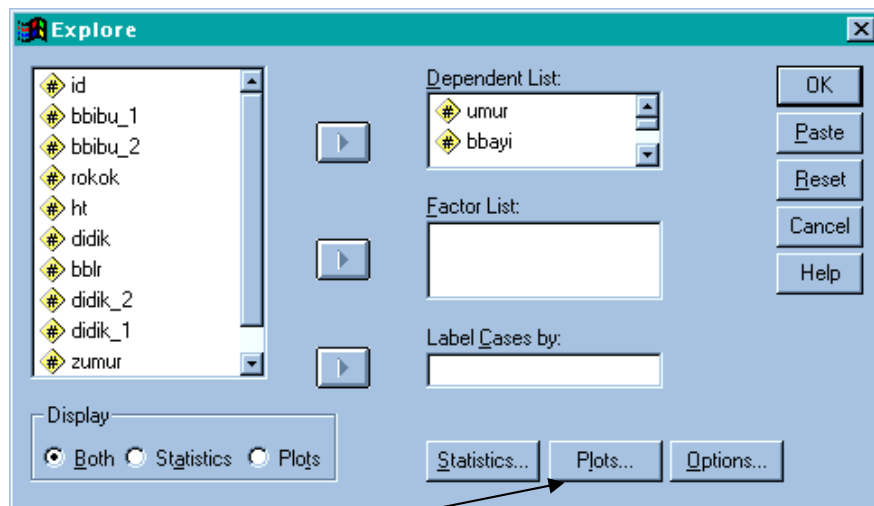
3. Klik **Display normal curve** (untuk menampilkan garis distribusi normal), Kemudian klik **OK** untuk menjalankan prosedur. Hasilnya sbb: (Lakukan prosedur yang sama untuk menampilkan grafik HISTOGRAM berat bayi **BBAYI**)



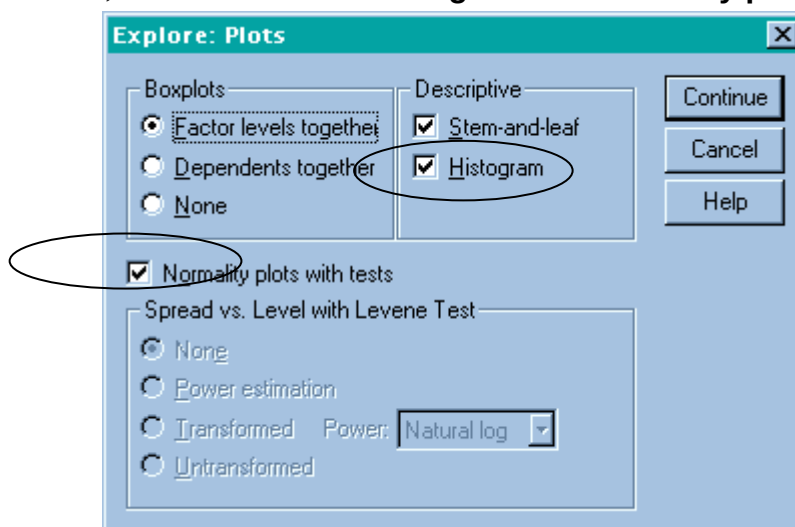
2.6. UJI NORMALITAS DISTRIBUSI DATA NUMERIK

Analisis data Numerik akan lebih lengkap apabila dilengkapi UJI NORMALITAS. Terutama jika akan dilakukan uji statistik parametrik terhadap variabel tersebut maka distribusi normal merupakan salah prasyarat yang harus dipenuhi. Uji normalitas dapat dilakukan melalui perintah **Explore..**

1. Dari menu utama, pilihlah:
Analyze
Descriptive Statistic <
Explore...
2. Pada kotak dialog tersebut, pilih variabel **UMUR** dan **BBAYI**, Kemudian klik tanda panah ke kanan >, untuk memasukkannya ke kotak **Dependent list**:



3. Klik **Plots...**, kemudian aktifkan **Histogram** dan **Normality plot with test**.



4. Klik **Continue** dan **OK** untuk menjalankan prosedur, hasilnya selain telah ditampilkan pada bagian 2.4 halaman 22 juga ada penambahan sbb:

Tests of Normality

	Kolmogorov-Smirnov ^a		
	Statistic	df	Sig.
Umur ibu	.095	189	.000
Berat bayi lahir	.043	189	.200*

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

a. Lilliefors Significance Correction

Hasil uji test normalitas

Dengan uji Kolmogorov-Smirnov, disimpulkan bahwa pada alpha 0.05 distribusi data umur ibu adalah tidak normal (nilai-p = 0.000) sedangkan distribusi data berat bayi adalah normal (nilai-p = 0.200).

Apabila diperhatikan grafik HISTOGRAM (*pada halaman 23*), maka terlihat bahwa data umur ibu memang tidak normal, tepatnya distribusi tersebut miring ke kanan (miring positif +). Kemiringan positif ini dapat juga dilihat dari nilai **Skewness**-nya yang bertanda positif (1.395)

Kesimpulan normal atau tidaknya suatu data didasarkan pada prinsip uji hipotesis yang berpatokan pada H_0 dan H_a . Dalam hal ini, H_0 berbunyi “Distribusi data **sama dengan distribusi normal**”, H_a berbunyi “Distribusi data **tidak sama dengan distribusi normal**”. Apabila nilai-p kurang dari alpha 0.05 (mis 0.000), maka H_0 ditolak dan disimpulkan “Distribusi data adalah tidak normal”. Sedangkan apabila nilai-p lebih dari atau sama dengan alpha 0.05 (mis. 0.222), maka H_0 gagal ditolak dan disimpulkan “Distribusi data adalah normal”.

2.7. PENYAJIAN DATA NUMERIK

Penyajian data mempunyai prinsip efisiensi, artinya sajikan hanya informasi penting saja, jangan semua output komputer disajikan dalam laporan. Contoh penyajian data numerik sbb:

Variabel	Jumlah	Min-Max	Mean	Median	SD	95% CI Mean
Umur ibu	189	14—45	23.24	23.0	5.30	22.48—24.0
Berat ibu	189	36—112	58.39	24.0	13.76	56.42—60.37
Berat bayi	189	709—4990	4990	2977.0	729.02	2840.05—3049.26

3 Transformasi Data

Transformasi data adalah suatu proses dalam merubah bentuk data. Misalnya merubah data numerik menjadi data kategorik atau merubah dari beberapa variabel yang sudah ada dibuat satu variabel komposit yang baru. Beberapa perintah SPSS yang sering digunakan adalah RECODE dan COMPUTE.

Setelah mempelajari BAB ini, anda akan mengetahui:

- 1. Pengertian Transformasi Data
- 2. Analisis Deskriptif
- 3. Transformasi Data dengan perintah RECODE
- 4. Transformasi Data dengan perintah COMPUTE
- 5. Interpretasi Hasil

3.1. PENGERTIAN TRANSFORMASI DATA

Transformasi data merupakan suatu proses untuk merubah bentuk data sehingga data siap untuk dianalisis. Banyak cara yang dapat dilakukan untuk merubah bentuk data namun yang paling sering digunakan antara lain adalah RECODE dan COMPUTE.

Perubahan bentuk data yang paling sederhana adalah pengkategorian data numerik menjadi data kategorik, misalnya *UMUR* dikelompokkan menjadi 3 kategori yaitu < 20 th, 20—30 th, dan >30 th. Atau dapat juga dilakukan pengelompokkan data kategorik menjadi beberapa kelompok yang lebih kecil, misalnya *DIDIK* dikelompokkan menjadi 2 kategori yaitu rendah (SD/SMP) dan tinggi (SMU/PT). Proses pengelompokan atau pengkategorian ulang tersebut lebih dikenal dengan istilah **RECODE**.

Perubahan bentuk data lainnya adalah penggunaan fungsi matematik dan algoritma. Misalnya penjumlahan skor pengetahuan, skor sikap, atau skor persepsi. Atau dapat juga dilakukan proses perkalian dan pembagian sekaligus, misalnya untuk menghitung Index Massa Tubuh ($IMT = BB/TB^2$). Atau dapat juga dilakukan pengelompokkan beberapa variabel sekaligus menggunakan fungsi algoritma, misalnya jika *TAHU*=1 dan *SIKAP*=1 dan *PRILAKU*=1 maka *KONSISTEN*=1 (jika ke-3 kondisi tersebut terpenuhi maka dikategorikan sebagai konsisten atau *KONSISTEN*=1, namun jika salah satu tidak terpenuhi maka dikategorikan tidak konsisten atau *KONSISTEN*=0). Proses penggunaan fungsi matematik dan algoritma tersebut lebih dikenal dengan istilah **COMPUTE**.

Berikut ini merupakan contoh transformasi data dari Survei Cepat Kesehatan Ibu dan Anak di 4 Kabupaten di Jawa Barat yaitu Tangerang, Cianjur, Lebak, dan Cirebon. Agar konsep transformasi data lebih mudah dipahami, maka langsung ditampilkan dalam bentuk contoh nyata dalam pengolahan data.

Sebagai contoh data kita gunakanlah file TNG.REC (hasil survei cepat di Kabupaten Tangerang yang telah dientry dengan program EPI INFO). Dengan menggunakan program EPI-INFO atau EPI Data lakukanlah Export data TNG.REC ke TNG.DBF. Kemudian buka file TANGERANG.DBF, dan Save ke TANGERANG.SAV. (Lihat Bab I: *Pendahuluan* untuk prosedur membuka file DBF dan menyimpan datanya).

LATIHAN MEMBUAT LABEL & VALUE:

Dengan program SPSS, buatlah **LABEL** untuk setiap **Variabel** dan **VALUE** untuk **Kode** tertentu yang diperlukan dari data TNG tersebut. Anda memerlukan BUKU KODE untuk dapat membuat LABEL dan VALUE (*Lihat Bab I: Pendahuluan untuk prosedur membuat label dan value*). Buku kode untuk membuat label tersebut ada di halaman berikutnya.

Buku kode**Survei Cepat Kesehatan Ibu di Kabupaten Tangerang, Lebak, Cianjur, Cirebon**

Nama variabel	No. Pertanyaan	Nilai	Label
Klaster	--	--	Nomor klaster
RESP	--	--	Nomor responden
V01	1	Kontinyu 15-45	Umur ibu (tahun)
V02	2	1 2 3 4 5 6	Pendidikan ibu Tidak sekolah Tidak tamat SD Tamat SD Tamat SLTP/ sederajat Tamat SLTA/ sederajat Akademi/ perguruan tinggi
V03	3	1 2 3 4 5 6 7 8	Pekerjaan utama ibu Tidak bekerja Buruh Pedagang Petani Jasa Pegawai swasta Pegawai negeri/ABRI Lain-lain
V04	4	1 2	Apakah ibu melakukan pemeriksaan kehamilan ? Ya Tidak
V05	5	Kontinyu	Berapa kali ibu periksa hamil ?
V06	6	1 2 3 4 5 6 7 8 9	Siapa yang menganjurkan ibu untuk periksa hamil ? Keinginan sendiri Keluarga Tetangga/teman Kader kesehatan Bidan Paraji Petugas puskesmas Dokter praktek swasta Lain-lain
V07	7	1 2 3 4 5 6 7 8 9	Tempat pemeriksaan kehamilan yg paling sering dikunjungi Posyandu Bidan praktek swasta Puskesmas Rumah sakit Pondok bersalin Dokter praktek swasta Rumah bersalin Paraji Lain-lain
V08	8	1 2 3 4 5 6 7 8	Alasan utama mengunjungi tempat pemeriksaan kehamilan tersebut Biaya murah Sabar/simpatik Teliti Jaraknya dekat Tradisi keluarga Aman/selamat Dianjurkan Lain-lain

Nama variabel	No. pertanyaan	Nilai	Label
---------------	----------------	-------	-------

V09a	9.a	1 2	Pada saat periksa hamil, apakah dilakukan penimbangan ? Ya Tidak
V09b	9.b	1 2	Pada saat periksa hamil, apakah dilakukan imunisasi TT ? Ya Tidak
V09c	9.c	1 2	Pada saat periksa hamil apakah diberikan pil Fe ? Ya Tidak
V09d	9.d	1 2	Pada saat periksa hamil apakah dilakukan pemeriksaan tinggi fundus? Ya Tidak
V09e	9.e	1 2	Pada saat periksa hamil, apakah dilakukan pemeriksaan tek. darah ? Ya Tidak
V10	10	1 2 3 4 5	Berapa pil Fe yg diminum selama hamil ? 1-30 pil 31-60 pil 61-90 pil > 90 pil Tidak pernah
V11	11	1 2 3 4 5 6	Siapa yang menolong ibu melahirkan pada kehamilan terakhir ? Tetangga/keluarga Dukun Kader Bidan Dokter Lain-lain
V12	12	1 2 3 4 5 6 7 8	Di mana ibu melahirkan ? Rumah sendiri/orang tua Rumah paraji Puskesmas Praktek bidan swasta Pondok bersalin Rumah sakit Rumah bersalin Lain-lain
V13	13	1 2	Apakah bayi ditimbang setelah lahir ? Ya Tidak
V14	14	Kontinyu	Berat bayi lahir (gram)
V15	15	1 2	Apakah ibu memperoleh nasehat perawatan nifas ? Ya Tidak

3.2. ANALISA DESKRIPTIF

Setelah semua variabel dibuat LABEL dan VALUE, jawablah pertanyaan di bawah ini, dan sajikan dalam bentuk tabel yang sesuai dan tuliskan interpretasinya.

PERTANYAAN:

1. Bagaimana distribusi pendidikan ibu di Kabupaten tsb ?
2. Bagaimana distribusi pekerjaan ibu di Kabupaten tsb ?
3. Berapa persen ibu yang melakukan pemeriksaan kehamilan ?
4. Dari ibu yang melakukan pemeriksaan kehamilan, berapa kali rata-rata mereka memeriksakan kehamilannya ?
5. Dari ibu yang melakukan pemeriksaan kehamilan, berapa persen yang melakukan pemeriksaan kehamilan 4 kali atau lebih ? *Buat variabel baru dg nama PERIKSA*
6. *Dari ibu yang melakukan pemeriksaan kehamilan, berapa persen yang dianjurkan oleh tenaga kesehatan (kader, bidan, puskesmas, dokter), berapa persen yang dianjurkan oleh non tenaga kesehatan (keluarga, tetangga, paraji, lain-lain) dan berapa persen karena keinginan sendiri ? Buat variabel baru dg nama ANJURAN*
7. Dari ibu yang periksa hamil, berapa persen ibu yang periksa hamil 4 kali atau lebih dan kualitasnya baik (ditimbang, diimunisasi TT, diberi pil Fe, diperiksa tinggi fundus dan diperiksa tekanan darah) dan dapat pil Fe > 90 pil ? Kombinasi variabel ini merupakan proksi dari kualitas K4. *Buat variabel baru dg nama K4*
8. *Berapa persen ibu yang pada saat melahirkan ditolong oleh tenaga kesehatan (Bidan/dokter)? Buat variabel baru dg nama PENOLONG*
9. *Dari bayi yang ditimbang, berapa rata-rata berat badan bayi lahir dan berapa standar deviasinya?*
10. *Dari bayi yang ditimbang, berapa persen yang BBLR ? (BBLR = Berat lahir kurang dari 2500 gram) Buat variabel baru dg nama BBLR*

Langkah-langkah untuk menjawab pertanyaan no. 1 s.d no. 5 dan 7 akan dipandu selangkah demi selangkah dalam uraian buku ini, sedangkan pertanyaan no.6, 8 s.d 10 harus anda kerjakan sendiri sebagai latihan.

Jawaban Pertanyaan no. 1 sampai no. 3

Pertanyaan no. 1 s.d no. 3 berkaitan dengan jenis **data kategorik**, sehingga analysis data disesuaikan dengan prosedur analysis data kategorik (*Lihat Bagian 2.1 untuk prosedur lengkapnya*) yaitu sebagai berikut:

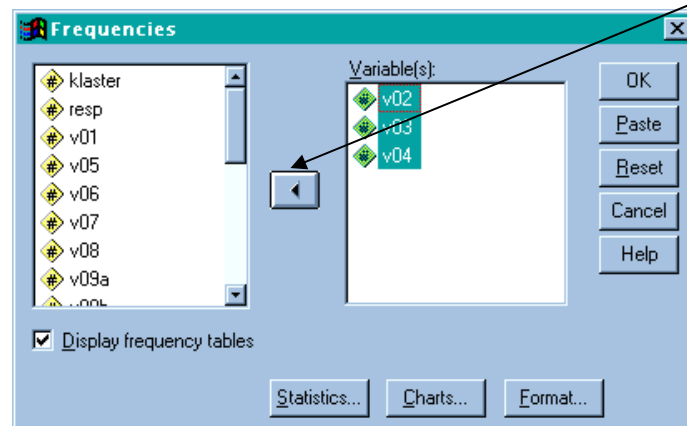
1. Bukalah file TANGERANG.SAV, sehingga data muncul di Data editor window.
2. Dari menu utama, pilihlah:

Analyze <

Descriptive Statistics <

Frequencies....

Pilih variabel **V02 V03 V04** dengan cara mengklik masing-masing variable tersebut, dan masukkan ke kotak Variable(s) di sebelah kanan dengan cara mengklik tanda < seperti berikut:



Klik **OK** untuk menjalankan prosedur. Pada layar Output tampak hasil seperti berikut:

Pendidikan Ibu

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	Tidak sekolah	42	14.1	14.1	14.1
	Tdk Tamat SD	98	32.9	32.9	47.0
	Tamat SD	87	29.2	29.2	76.2
	Tamat SMP	37	12.4	12.4	88.6
	Tamat SMU	33	11.1	11.1	99.7
	Tamat PT	1	.3	.3	100.0
	Total	298	100.0	100.0	

Pekerjaan Ibu

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	Tidak bekerja	274	91.9	91.9	91.9
	Buruh	3	1.0	1.0	93.0
	Pedagang	11	3.7	3.7	96.6
	Petani	1	.3	.3	97.0
	Jasa	1	.3	.3	97.3
	Pegawai swasta	5	1.7	1.7	99.0
	Pengawai negeri/ABRI	3	1.0	1.0	100.0
	Total	298	100.0	100.0	

Jawaban Pertanyaan no. 4

Pertanyaan no. 4 berkaitan dengan jenis **data numerik**, sehingga analysis data disesuaikan dengan prosedur analysis data numerik (*Lihat Bagian 2.4 untuk prosedur lengkapnya*) yaitu sebagai berikut:

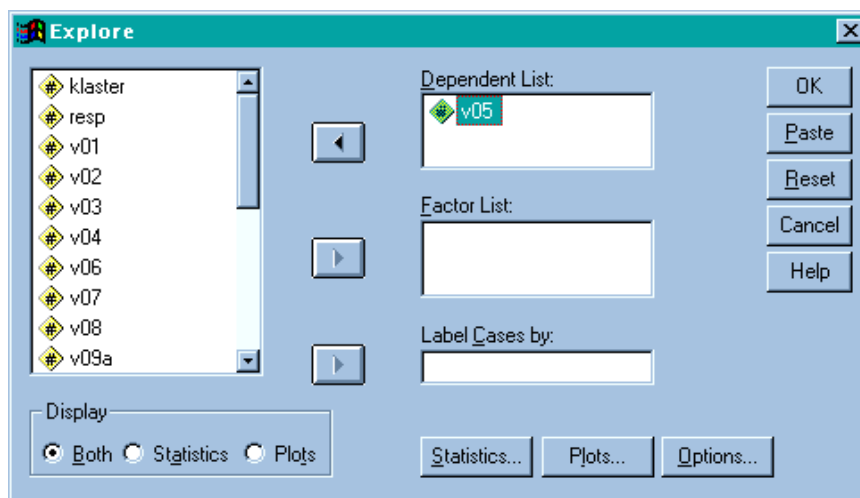
1. Dari menu utama, pilihlah:

Analyze <

Descriptive Statistics <

Explore

2. Pada kotak yang tersedia, pilih variabel **V05** dengan cara mengklik variable tersebut, dan masukkan ke kotak Variable(s) di sebelah kanan dengan cara mengklik tanda < seperti berikut:



3. Untuk menjalankan prosedur, klik OK sehingga outputnya sebagai berikut:

Descriptives

			Statistic	Std. Error
Berapa Kali Periksa Kehamilan	Mean		6.49	.42
	95% Confidence Interval for Mean	Lower Bound	5.67	
		Upper Bound	7.31	
	5% Trimmed Mean		5.79	
	Median		5.00	
	Variance		48.149	
	Std. Deviation		6.94	
	Minimum		1	
	Maximum		81	
	Range		80	
	Interquartile Range		5.00	
	Skewness		7.727	.147
	Kurtosis		76.446	.293

Catatan: Untuk penyajian dan Interpretasi dapat dilihat Bab 2: *Analysis Deskriptif*.

Nilai maksimum adalah 81, anda harus mempertanyakan apakah data ini benar atau tidak? Lakukan terlebih dahulu "Cleaning Data".

3.3. TRANSFORMASI DATA DG PERINTAH “RECODE”

Jawaban Pertanyaan no. 5

Pada pertanyaan no. 5, anda harus membuat kategori baru dari variabel *V05* menjadi variabel *PERIKSA*, dimana nilai 1--3 pada *V05* menjadi kode=1 pada *PERIKSA* dan nilai 4--Max pada *V05* menjadi kode=2 pada *PERIKSA*. Dapat ditulis ulang sebagai berikut:

1—3 → 1 = “Periksa kurang dari 4 kali”

4—max → 2 = “Periksa 4 kali atau lebih”

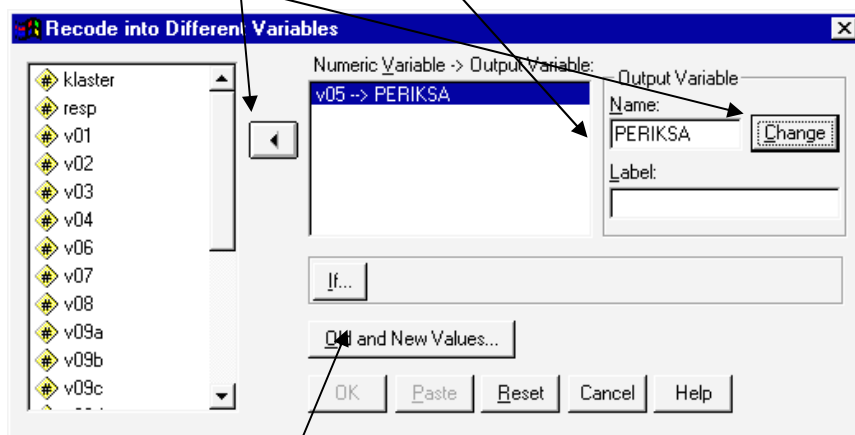
1. Dari menu utama, pilihlah:

Transform <

Recode <

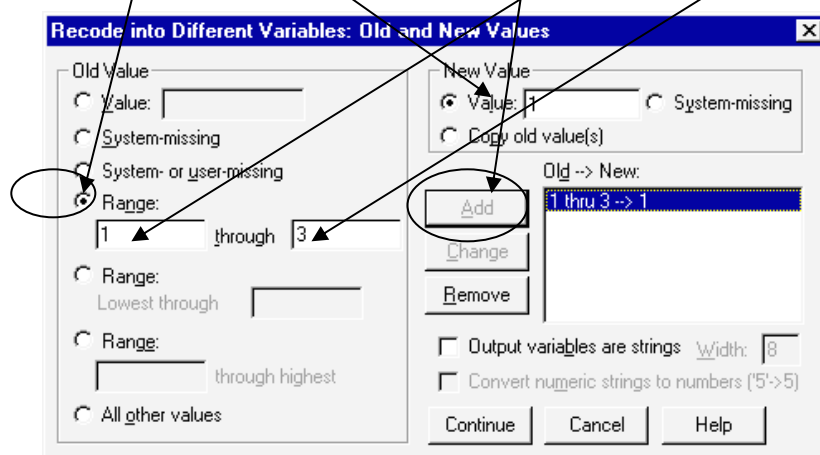
Into Different Variable....

2. Pilih variabel *V05* klik tanda < untuk memasukkannya ke kotak sebelah kanan
3. Isi Kotak **Name** dengan variabel baru **PERIKSA**
4. Klik **Change**, sehingga “*V05* → ...” berubah menjadi “*V05* → PERIKSA” seperti berikut:

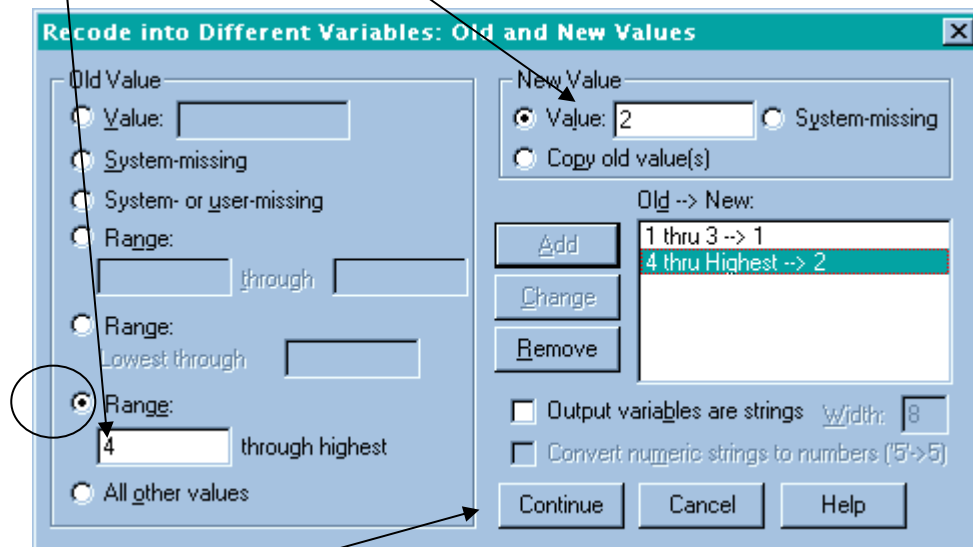


5. Klik **OLD AND NEW VALUES...**

6. Pada **OLD** Value, Pilih (•) Range through dan isi **1** through **3**. Kemudian pada **NEW** Value isi **1**, selanjutnya klik **ADD**. Hasilnya dapat dilihat pada gambar berikut:



7. Berikutnya, pada OLD Value, Pilih (.) Range through highest dan isi kotak 4 through highest.
Kemudian pada NEW Value isi 2, kemudian klik ADD



8. Klik **Continue** dan kemudian **OK** untuk menjalankan prosedur
Proses transformasi selesai, lihat pada jendela Data-View, kolom paling kanan

Pemberian LABEL dan VALUE..

9. Beri **Label** PERIKSA → Jumlah Kunjungan Periksa Hamil
10. Beri **Value** PERIKSA kode 1 → “Kurang 4 kali” kode 2 → “4 kali atau lebih”
11. Tampilkan distribusi frekuensi untuk variabel *PERIKSA* sebagai berikut:

Jumlah Kunjungan Periksa Hamil

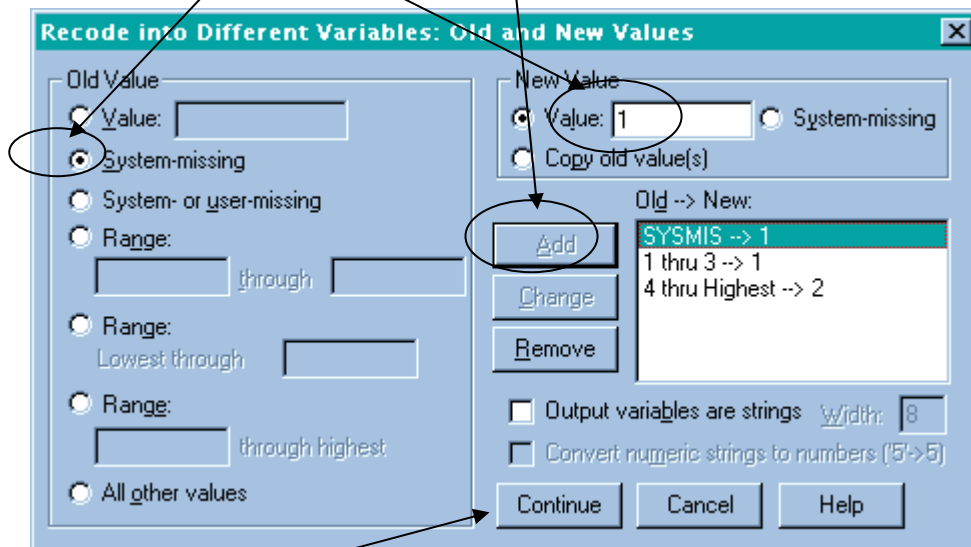
		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	Kurang dari 4 kali	76	25.5	27.6	27.6
	4 kali atau lebih	199	66.8	72.4	100.0
	Total	275	92.3	100.0	
Missing	System	23	7.7		
Total		298	100.0		

“Dari semua yang periksa hamil (275), sebanyak 199 (72.4%) memeriksakan kehamilannya 4 kali atau lebih, ada 23 responden yang missing (artinya tidak pernah periksa hamil)”.

Catatan tambahan:

Jika anda menginginkan data yang missing tersebut juga diberi kode= 1 (Periksa kurang dari 4 kali/tidak periksa hamil), maka setelah langkah nomor 7 tambahkan perintah berikut:

12. Pada OLD Value, Pilih **System missing**, kemudian pada NEW Value isi **1**, kemudian klik ADD, hasilnya sbb:



13. Klik **Continue** dan **OK** untuk menjalankan prosedur.

14. Keluarkan distribusi frekuensi dari variabel *PERIKSA*, hasilnya sbb:

Jumlah Kunjungan Periksa Hamil

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	Kurang dari 4 kali	99	33.2	33.2	33.2
	4 kali atau lebih	199	66.8	66.8	100.0
	Total	298	100.0	100.0	

Interpretasinya berbeda dengan output sebelumnya:

"Dari semua reponden (298), sebanyak 199 (66.8%) memeriksakan kehamilannya 4 kali atau lebih"

Catatan: Untuk penyajian dan Interpretasi lebih detail dapat dilihat Bab 2: *Analysis Deskriptif*.

3.4. TRANSFORMASI DATA DG PERINTAH “COMPUTE”

Pertanyaan no. 7

Dari ibu yang periksa hamil, berapa persen ibu yang periksa hamil 4 kali atau lebih dan kualitasnya baik (ditimbang, diimunisasi TT, diberi pil Fe, diperiksa tinggi fundus dan diperiksa tekanan darah) dan dapat pil Fe > 90 pil ?. Kombinasi variabel ini merupakan proksi dari kualitas K4. *Buat variabel baru dg nama K4*

Jawaban no.7

Untuk menjawab pertanyaan nomor 7 anda terlebih dahulu harus membuat variabel baru yang namanya K4. Jika $V05 \geq 4$ dan ($V09a=1$ dan $V09b=1$ dan $V09c=1$ dan $V09e=1$) dan $v10=4$ maka $K4=1$ (K4 berkualitas baik) selain itu $K4=0$ (K4 tidak berkualitas tidak)

1. Dari menu utama, pilihlah:

**Transform <
Compute <**

2. Isi Target Variabel dengan **K4**

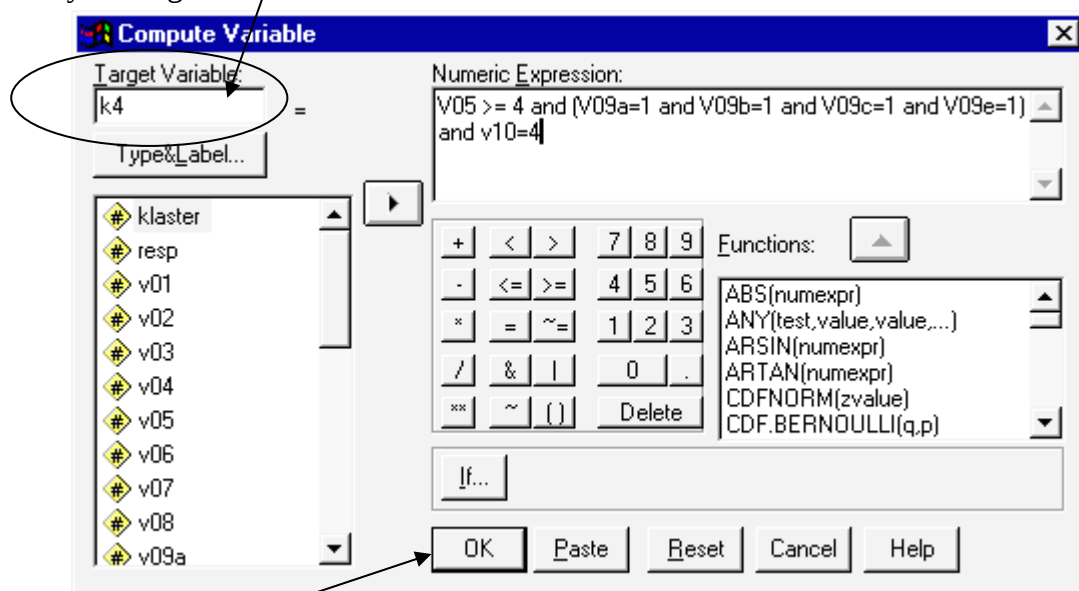
3. Isi Kotak Numeric Expression dengan persamaan berikut:

$V05 \geq 4$ and ($V09a=1$ and $V09b=1$ and $V09c=1$ and $V09e=1$) and $v10=4$

Pilih variabel yang sesuai di kotak kiri bawah, kemudian klik tanda > untuk memasukkannya ke kotak bagian kanan atas (Numeric Expression)

(Jangan biasakan mengetik nama variabel, cukup pakai klik dan pilih tanda >, untuk mengurangi kesalahan akibat pengetikan)

4. Hasilnya Sebagai berikut:



Klik OK untuk menjalankan prosedur

Kemudian keluarkan distribusi frekuensi dari **K4** (Analysis deskriptif data kategorik), sehingga muncul hasil seperti berikut:

K4

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	.00	221	74.2	82.5	82.5
	1.00	47	15.8	17.5	100.0
	Total	268	89.9	100.0	
Missing	System	30	10.1		
Total		298	100.0		

Buat **Label** untuk variabel K4="Pemeriksakan kehamilan dengan kualitas baik",
 Buat **VALUE** kode 0="Kualitas K4 tidak baik" dan kode 1="Kualitas K4 baik",
 Keluarkan kembali tabel frekuensinya sbb:

Ibu memeriksakan kehamilan dengan kualitas baik

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	Kualitas K4 tidak baik	221	74.2	82.5	82.5
	Kualitas K4 baik	47	15.8	17.5	100.0
	Total	268	89.9	100.0	
Missing	System	30	10.1		
Total		298	100.0		

Contoh interpretasi:

"Dari semua responden ibu hamil (298), sebanyak 47 (15.8%) memeriksakan kehamilan dengan kualitas K4 yang baik"

Hati-hati dengan interpretasi lain yang berbeda:

"Dari semua yang pernah periksa hamil (268), sebanyak 47 (17.5%) mendapatkan pemeriksaan kehamilan dengan kualitas K4 yang baik"

Interpretasi mana yang akan dipilih harus disesuaikan dengan tujuan dan substansi yang ingin diukur oleh peneliti

Perintah Compute tersebut dapat juga diketik pada SPSS Syntax sbb:

**COMPUTE K4 = V05 >= 4 and V09a=1 and V09b=1 and V09c=1 and V09d=1 and V09e=1 and v10=4.
 FREQ K4.**

***Pemberian Variabel Label**

VAR LAB K4 "Ibu memeriksakan kehamilan K4 dengan kualitas baik".

***Pemberian Value Label**

VALUE LAB K4 1 "Kualitasw Baik" 0 "Kualitas Kurang".

Tabel Frekuensi:

FREQ K4.

4 Merge File Data

Merger atau menggabung beberapa file data menjadi satu file biasanya dilakukan pada survei besar, dimana proses entri data dilakukan oleh lebih dari satu orang pada saat yang bersamaan atau file entri data sengaja dipilah-pilah sesuai dengan topik penelitiannya agar lebih mudah dalam proses entri-nya. Sebelum datanya bisa dianalisa, maka file-file data yang terpisah itu harus digabungkan terlebih dahulu.

Setelah mempelajari BAB ini, anda akan mengetahui:

- 1. Pengertian Merge
- 2. Merger dengan “ADD VARIABEL”
- 3. Merger dengan “ADD CASES”
- 4. Merger antara data INDIVIDU dengan data RUMAH TANGGA

4.1. PENGERTIAN MERGE

Merge merupakan suatu proses yang diperlukan untuk menggabungkan beberapa file data yang ingin dijadikan satu file data saja. Secara umum ada tiga jenis merger, yaitu 1) merger untuk menambah record/kasus/responden, 2) merger untuk menambah variabel, dan 3) merger untuk menggabungkan antara data individu dengan data rumah tangga.

1. MERGER dengan ADD CASES:

Jenis merger ini biasanya dilakukan pada satu penelitian dengan **jumlah variabel yang relatif sedikit** tetapi jumlah **record/kasus/responden relatif banyak** dan pada saat melakukan ENTRY data biasanya dilakukan oleh lebih dari 1 orang supaya cepat selesai.

Contoh:

Penelitian survei cepat dengan topik yang sama (Antenatal Care) dilakukan di Cianjur (300 responden) dan Lebak (300 responden), proses ENTRY data dilakukan oleh 2 orang. Data tersebut dapat dianalisis terpisah satu persatu, namun peneliti ingin juga melakukan analisis gabungan. Sebelum dilakukan analisis gabungan, file tersebut harus dimerge terlebih dahulu. Pada merge jenis ini dilakukan penambahan record/kasus/responden (ADD CASES). Hasil gabungan 2 file tersebut akan didapatkan 600 responden, dengan jumlah “variabel” yang sama karena surveinya sama.

2. MERGER dengan ADD VARIABEL

Jenis merger ini biasanya dilakukan pada satu penelitian dengan jumlah **variabel yang relatif banyak**, atau pada beberapa penelitian dengan **topik yang berbeda** dengan **responden yang sama** dan pada saat ENTRY data biasanya dilakukan oleh lebih dari 1 orang supaya cepat selesai. Atau ENTRY data antara satu topik dengan topik lainnya sengaja dipisah supaya databasenya tidak terlalu besar.

Contoh:

Penelitian Survei Sosial Ekonomi Nasional (Susenas) mempunyai banyak topik yang diteliti (ISPA, DIARE, MENYUSUI, KB, dll) pada responden/keluarga yang sama. Proses ENTRY data biasanya dilakukan per topik (satu topik satu file data) sehingga jika ingin mengolah data tersebut harus dilakukan merger terlebih dahulu.

Pada merge jenis ini dilakukan penambahan variabel (ADD VARIABEL). Proses merge ini memerlukan variabel “ID” yang sama.

Contoh lain:

Survei cepat di Lebak (300 responden), proses ENTRY data dilakukan oleh 2 orang. Petugas ENTRY-1 melakukan pemasukan data untuk variabel V01 sampai 09 sedangkan Petugas ENTRY-2 melakukan pemasukan data untuk variabel V10 sampai 15. Sebelum dianalisis file tersebut harus digabung terlebih dahulu. Proses penggabungan memerlukan “ID” yang sama.

4.2. MERGER dengan ADD VARIABEL

Pada uraian berikut akan dijelaskan penggabungan file LEBAK-1.SAV (variabel id v01—v09) dengan LEBAK-2.SAV (variable id v10—v15).

Persyaratan:

1. Harus ada variabel “ID” yang sama, artinya nomor identitas responden pada file-1 harus sama dengan nomor identitas responden pada file-2
2. Variabel “ID” atau nomor identitas tersebut tidak boleh ada nomor yang sama atau nomor kembar (*double*), artinya dalam satu file hanya boleh ada satu nomor identitas. Tidak boleh ada responden-A memiliki nomor ID “10012” tetapi responden-B juga memiliki nomor ID “10012”.

Dalam contoh ini variabel “ID” yang dipakai adalah “RESP_ID”.

- SORT DATA-1.SAV

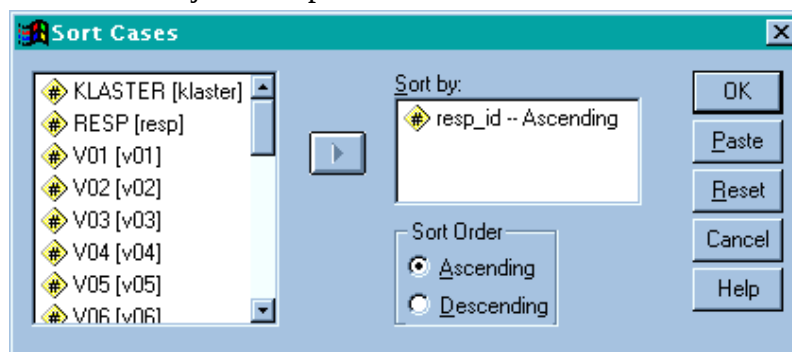
1. Bukalah file LEBAK-1.SAV, sehingga data tampak di Data editor window.
2. Lakukan SORT terhadap variabel “ID” (*RESP_ID*). Dari menu utama, pilihlah:

Data <

Sort Cases <

Pindahkan variabel **RESP_ID** ke kotak kanan dengan cara mengklik *RESP_ID* dan klik tanda <. Pastikan Sort Order yang dipilih adalah Sort by Ascending (Urutan dari nilai terkecil ke nilai terbesar).

Kemudian klik OK untuk menjalankan prosedur.



Pilih SAVE untuk menyimpan file data tersebut.

Tutup data LEBAK-1 tersebut dan Buka LEBAK-2

- SORT DATA-2.SAV

Bukalah file LEBAK-2.SAV, sehingga data tampak di Data editor window.

Lakukan SORT terhadap variabel “ID” (*RESP_ID*). Sesuai dengan prosedur no 2 sampai 4: Pilih SAVE untuk menyimpan file data tersebut.

- PROSES MERGE

Buka kembali file LEBAK-1.SAV

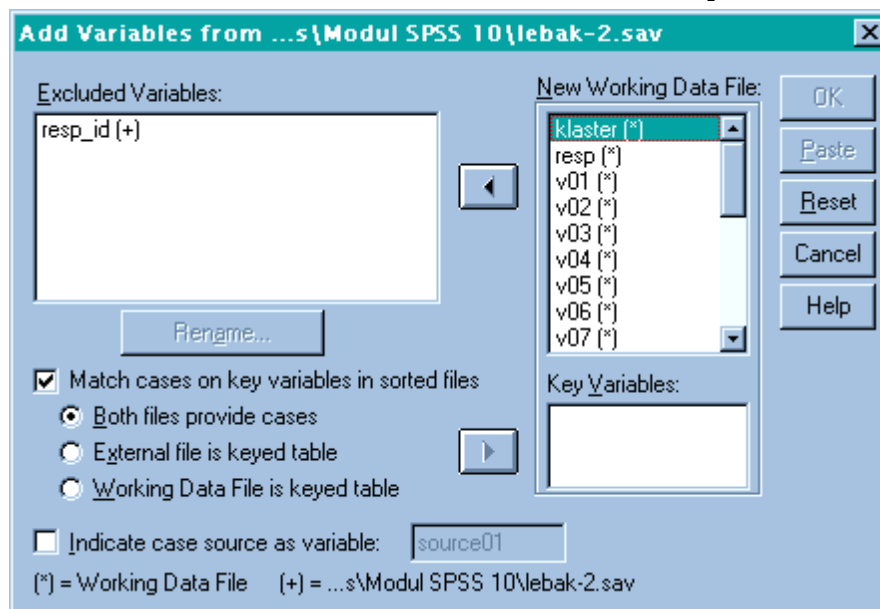
Mulai proses Merger dengan perintah:

Data <

Merge Files <

Add Variabel...

Pilih file LEBAK-2 kemudian klik Open



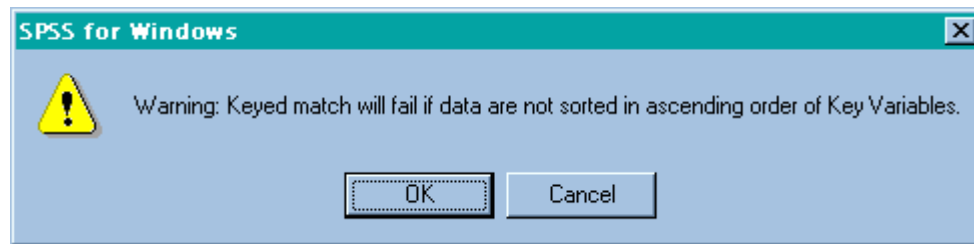
Pindahkan variabel *RESP_ID* dari kotak Excluded variabel ke kotak Key Variabel (di kanan bawah) dengan cara: klik *RESP_ID*, klik Match cases on key variables in sorted files, dan klik tanda panah ke kanan <.

Klik dan aktifkan Match cases on key variables in sorted files, dengan pilihan

- * Both file provide cases: kedua data digabung secara utuh (pilihan standar)
- * *External file is keyed table*: data-2 sebagai acuan untuk menggabung
- * *Working file is keyed table*: data-1 sebagai acuan untuk menggabung

Kemudian klik OK untuk menjalankan prosedur.

Selanjutnya akan muncul warning yang akan memberi tahu bahwa proses merger akan gagal jika Key variabel (*RESP_ID*) tidak di sort menurut ascending.



Pilih saja OK, karena kita sudah mensortnya, kemudian proses penggabungan akan berlangsung

Jika selesai, lihat Data View bagian paling kanan akan ditambahkan V10 sampai V15

Simpan file dengan nama LEBAK.SAV

4.3. MERGER dengan ADD CASES

Suatu survei yang dilakukan pada dua atau lebih tempat yang berbeda, untuk efisiensi maka proses ENTRY data dilakukan pada tempat yang berbeda pula. Pada saat ANALYSIS data, beberapa file data yang terpisah tadi perlu digabung terlebih dahulu agar analisis secara menyeluruh dapat dilakukan.

Persyaratan:

1. Kedua file harus mempunyai NAME variabel yang sama, artinya jika file-1 ada 15 variabel maka file-2 juga harus mempunyai 15 variabel yang sama. Kesamaan NAME ke-15 variabel harus mencakup juga kesamaan dalam TYPE dan WITH, serta DECIMAL
2. *“ID” atau nomor identitas responden tidak terlalu penting*

Uraian di bawah ini akan menjelaskan langkah-langkah proses penggabungan file LEBAK.SAV dan CIANJUR.SAV dengan prosedur sebagai berikut:

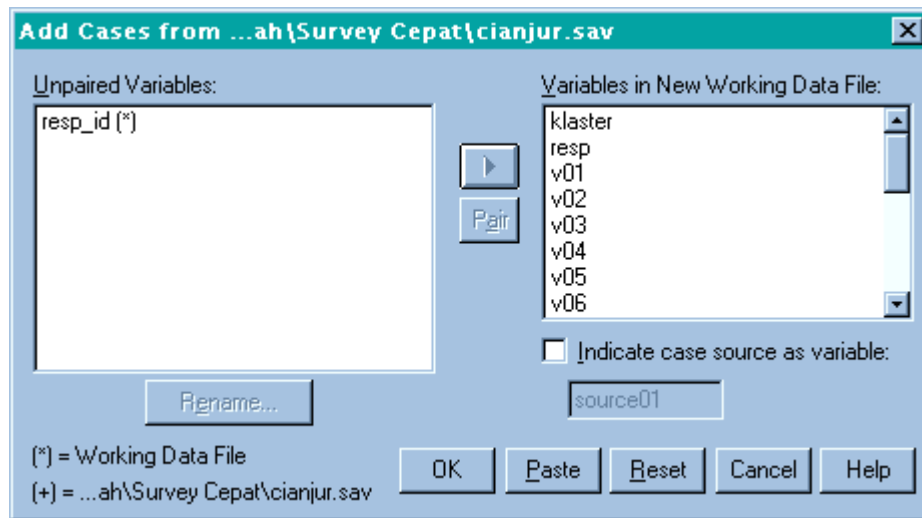
1. Bukalah file LEBAK.SAV, sehingga data tampak di Data editor window.
2. Lakukan merger dengan perintah berikut: Dari menu utama, pilihlah:

Data <

Merge files <

Add cases.....

Pilih file Cianjur kemudian klik Open.



3. Di menu tersebut terlihat bahwa SPSS memberi tahu bahwa ada variabel yang tidak sama atau tidak tersedia pada ke-2 file data yaitu RESP_ID, variabel yang tidak sama tersebut akan ditempatkan dalam kotak Unpaired variables. Abaikan saja variabel tersebut, sehingga nantinya tidak akan dimasukkan dalam data gabungan. Jika variabel tersebut merupakan variabel penting yang harus masuk, maka lakukan perubahan terhadap variabel itu terlebih dahulu, harus dicek apakah NAME, TYPE, WITH, atau DECIMAL-nya yang berbeda.
4. Pilih OK untuk menjalankan proses merger
5. Pastikan jumlah record/responden dan variabel sudah sesuai dengan keinginan, dalam hal ini (LEBAK+CIANJUR) respondennya adalah $300+300 = 600$ dan variabelnya V01 sampai dengan V15.
6. Jika selesai, Simpan file dengan nama baru, misalnya: LEBAK-CIANJUR.SAV

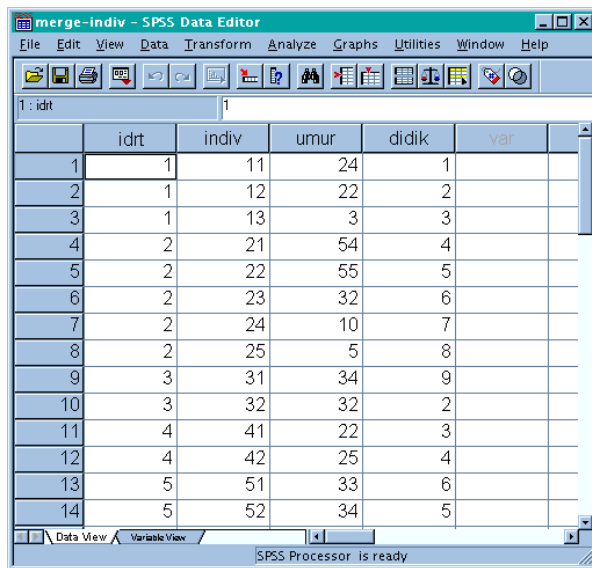
4.4. MERGER antara data INDIVIDU dengan data RUMAHTANGGA

Pada survei besar seperti SUSENAS database antar topik atau antara rumah tangga dengan individu dipisahkan demi efisiensi. Pada saat analysis, seringkali kita membutuhkan penggabungan antara data tersebut. Suatu hubungan file yang memiliki hirarkhi tersebut, seperti data rumah tangga dengan data individu dalam rumah tangga, ingin digabungkan untuk analysis lebih lanjut.

Prosedur penggabungannya adalah sama dengan cara Merge Add Variable. Namun, pada proses penggabungan, perlu diingat bahwa file data yang memiliki hirarkhi lebih tinggi (rumah tangga) harus dijadikan sebagai acuan (Match cases) pada variabel kunci.

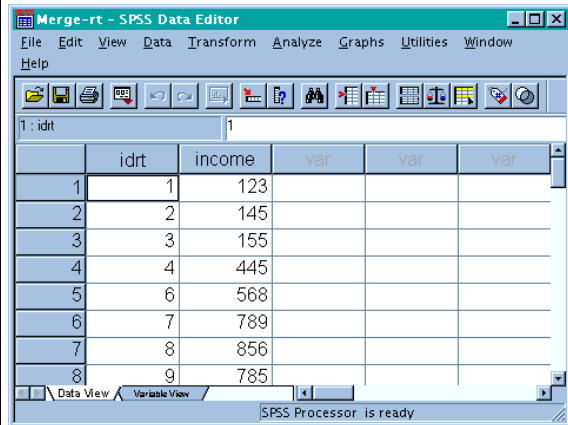
Contoh berikut ini akan menggabungkan file RT.SAV (berisi variabel *idrt* dan *income*) dengan file INDIV.SAV (berisi variabel *idrt*, *umur*, & *didik*). File yang aktif adalah

INDIV dan file yang akan digabungkan adalah RT. **Harus dipastikan tidak ada IDRT yang double pada file RT.SAV**



	idrt	indiv	umur	didik	var
1	1	11	24	1	
2	1	12	22	2	
3	1	13	3	3	
4	2	21	54	4	
5	2	22	55	5	
6	2	23	32	6	
7	2	24	10	7	
8	2	25	5	8	
9	3	31	34	9	
10	3	32	32	2	
11	4	41	22	3	
12	4	42	25	4	
13	5	51	33	6	
14	5	52	34	5	

File INDIV.SAV



	idrt	income	var	var	var
1	1	123			
2	2	145			
3	3	155			
4	4	445			
5	6	568			
6	7	789			
7	8	856			
8	9	785			

File RT.SAV

PERSIAPAN MERGE

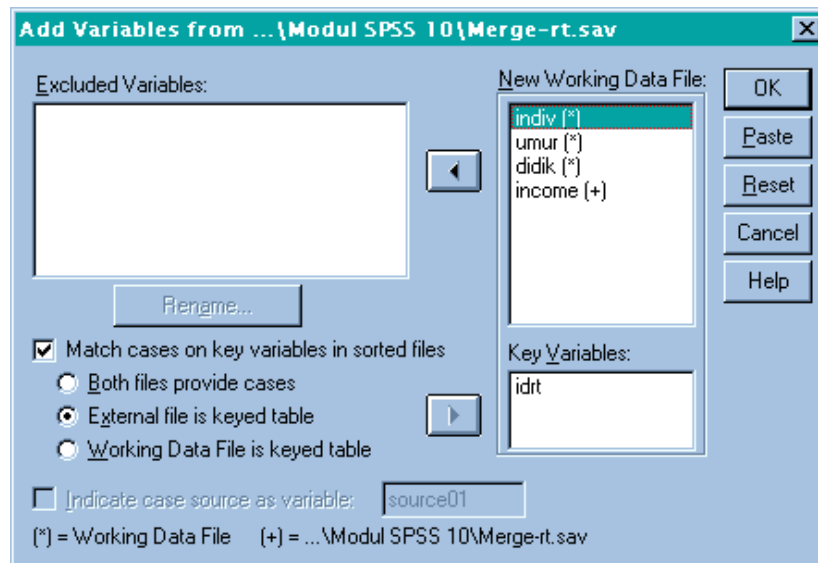
1. Bukalah file RT.SAV, sehingga data tampak di Data editor window.
2. Lakukan SORT terhadap variabel "ID" (IDRT). Dari menu utama, pilihlah:
Data <
Sort Cases <
3. Sort IDRT by Ascending, SAVE, dan CLOSE
4. Bukalah file INDIV.SAV, sehingga data tampak di Data editor window.
5. Lakukan SORT terhadap variabel "ID" (IDRT). Dari menu utama, pilihlah:
Data <
Sort Cases <
6. Sort IDRT by Ascending, dan SAVE

PROSEDUR MERGE

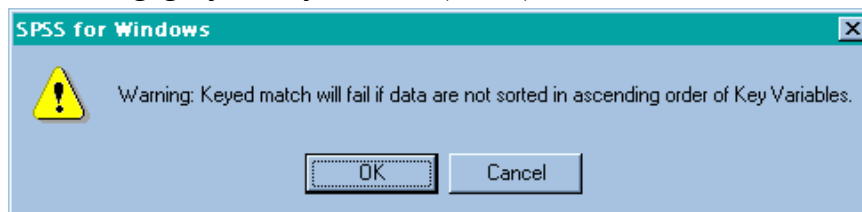
7. Pastikan data yang aktif adalah data INDIV.SAV
Mulai proses Merger dengan perintah:
Data <
Merge Files <
Add Variabel...
Pilih file RT.SAV kemudian klik Open
8. Pindahkan variabel IDRT dari kotak Excluded variabel ke kotak Key Variabel dengan cara: klik IDRT, aktifkan Match cases on key variables in sorted files, klik tanda panah ke kanan >.
9. Klik dan aktifkan Match cases on key variables in sorted files, dengan pilihan

* *External file is keyed table*: data-external (RT.SAV) sebagai acuan untuk digabungkan ke data yang sedang aktif (INDIV.SAV)

10. Kemudian klik OK untuk menjalankan prosedur.



11. Selanjutnya akan muncul warning yang akan memberi tahu bahwa proses merger akan gagal jika Key variabel (IDRT) tidak di sort menurut ascending.



12. Pilih saja OK, kemudian proses penggabungan akan berlangsung, hasilnya sbb:

	idrt	indiv	umur	didik	income	var
1	1	11	24	1	123	
2	1	12	22	2	123	
3	1	13	3	3	123	
4	2	21	54	4	145	
5	2	22	55	5	145	
6	2	23	32	6	145	
7	2	24	10	7	145	
8	2	25	5	8	145	
9	3	31	34	9	155	
10	3	32	32	2	155	
11	4	41	22	3	445	
12	4	42	25	4	445	
13	5	51	33	6	.	
14	5	52	34	5	.	
15						

Perhatikan bahwa rumah tangga nomor 5 tidak mempunyai data income karena memang data aslinya pada RT.SAV IDRT nomor 5 tidak ada, dari IDRT no 4 langsung no 6.

Selain itu, income RT nomor 6 sampai 10 tidak masuk dalam file gabungan, karena memang *IDRT* 6 –10 tidak tersedia pada file *INDIV.SAV*.

Jangan lupa untuk menyimpan file gabungan dengan nama yang berbeda dengan perintah *SAVE AS..*, tulis nama file *GAB-INDIV-RT.SAV*.

5 Uji Beda 2-Rata-rata (*t-test*)

5.1. Pengertian

Di bidang kesehatan sering kali kita harus membuat kesimpulan apakah suatu intervensi berhasil atau tidak. Untuk mengukur keberhasilan tersebut kita harus melakukan uji untuk melihat apakah parameter (rata-rata) dua populasi tersebut berbeda atau tidak. Misalnya, apakah ada perbedaan rata-rata tekanan darah populasi intervensi (kota) dengan populasi kontrol (desa). Atau, apakah ada perbedaan rata-rata berat badan antara sebelum dengan sesudah mengikuti program diet.

Sebelum kita melakukan uji statistik dua kelompok data, kita perlu perhatikan apakah dua kelompok data tersebut berasal dari dua kelompok yang independen atau berasal dari dua kelompok yang dependen/berpasangan. Dikatakan kedua kelompok data independen bila populasi kelompok yang satu tidak tergantung dari populasi kelompok kedua, misalnya membandingkan rata-rata tekanan darah sistolik orang desa dengan orang kota. Tekanan darah orang kota adalah independen (tidak tergantung) dengan orang desa. Dilain pihak, dua kelompok data dikatakan dependen/pasangan bila datanya saling mempunyai ketergantungan, misalnya data berat badan sebelum dan sesudah mengikuti program diet berasal dari orang yang sama (data sesudah dependen/tergantung dengan data sebelum).

5.2. Konsep Uji Beda Dua Rata-rata

Uji beda rata-rata dikenal juga dengan nama uji-t (*t-test*). Konsep dari uji beda rata-rata adalah membandingkan nilai rata-rata beserta selang kepercayaan tertentu (*confidence interval*) dari dua populasi. Prinsip pengujian dua rata-rata adalah melihat perbedaan variasi kedua kelompok data. Oleh karena itu dalam pengujian ini diperlukan informasi apakah varian kedua kelompok yang diuji sama atau tidak. Varian kedua kelompok data akan berpengaruh pada nilai standar error yang akhirnya akan membedakan rumus pengujiannya.

Dalam menggunakan uji-t ada beberapa syarat yang harus dipenuhi. Syarat/asumsi utama yang harus dipenuhi dalam menggunakan uji-t adalah data harus berdistribusi normal. Jika data tidak berdistribusi normal, maka harus dilakukan transformasi data terlebih dahulu untuk menormalkan distribusinya. Jika transformasi yang dilakukan tidak mampu

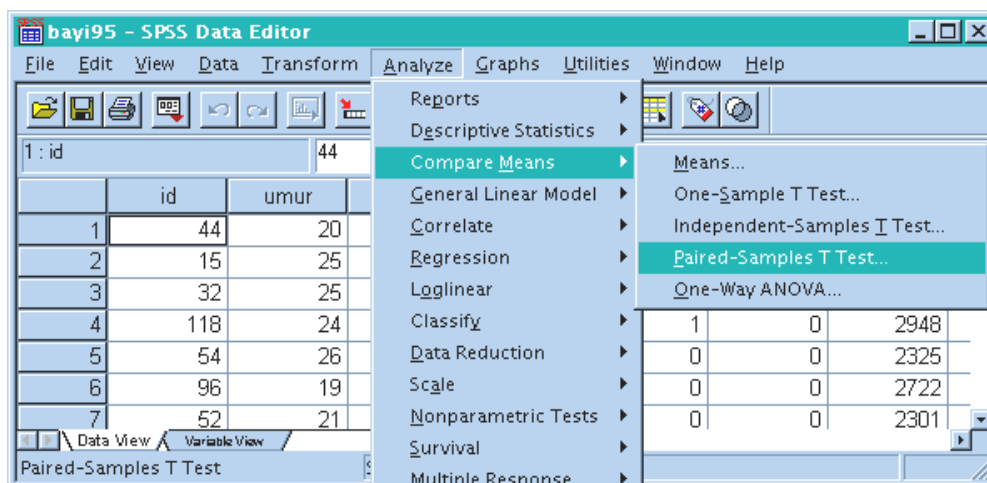
menormalkan distribusi data tersebut, maka uji-t tidak valid untuk dipakai, sehingga disarankan untuk melakukan uji non-parametrik seperti Wilcoxon (data berpasangan) atau Mann-Whitney U (data independen).

Berdasarkan karakteristik datanya maka uji beda dua rata-rata dibagi dalam dua kelompok, yaitu: uji beda rata-rata independen dan uji beda rata-rata berpasangan.

5.3. Aplikasi Uji-t Dependen pada Data Berpasangan

Uji-t untuk data berpasangan berarti setiap subjek diukur dua kali. Misalnya sebelum dan sesudah dilakukannya suatu intervensi atau pengukuran yang dilakukan terhadap pasangan orang kembar. Dalam contoh ini akan membandingkan data sebelum dengan sesudah intervensi. Dalam BAYI95.SAV sudah ada data berpasangan yaitu pengukuran berat badan ibu yang dilakukan sebelum hamil. Sebelum merencanakan kehamilan, subjek melakukan penyesuaian diet (mengikuti program makanan tambahan) selama 2 bulan. Pengukuran berat badan yang pertama (BBIBU_1) dilakukan sebelum kegiatan penyesuaian diet dilakukan, dan pengukuran berat badan yang kedua (BBIBU_2) dilakukan setelah dua bulan menjalani penyesuaian diet.

Kita akan melakukan uji hipotesis untuk menilai apakah ada perbedaan berat badan ibu antara sebelum dengan sesudah mengikuti program diet, langkah-langkahnya sebagai berikut.



1. Bukalah file BAYI95.SAV, sehingga data tampak di Data editor window.

2. Dari menu utama, pilihlah: (pada SPSS 10.0)
 Analyze <
 Compare Mean <
 Paired-Sample T-test....
3. Pilih variabel *BBIBU_1* dan *BBIBU_2* dengan cara mengklik masing-masing variable tersebut.
4. Kemudian klik tanda < untuk memasukkannya ke dalam kotak Paired-Variables.
5. Pada menu “**Options**” pilihlah derajat kepercayaan yang diinginkan, misalnya 95%.
 Kemudian pilih **Continue**.
6. Klik OK untuk menjalankan prosedur. Pada layar Output tampak hasil seperti berikut:

Paired Samples Statistics

		Mean	N	Std. Deviation	Std. Error Mean
Pair 1	BBIBU_1	58.39	189	13.76	1.00
	BBIBU_2	60.12	189	13.72	1.00

Dari 189 subjek yang diamati terlihat bahwa rata-rata (mean) berat badan dari ibu sebelum intervensi (*BBIBU_1*) adalah 58.39, dan rata-rata berat badan sesudah intervensi (*BBIBU_2*) adalah 60.12. Jika kita ingin mengeneralisasikan pada populasi, apakah di populasi ada perbedaan yang signifikan antara sebelum dan sesudah intervensi. Uji ‘t’ yang dilakukan terlihat pada tabel berikut:

Paired Samples Test		Paired Differences							
		Mean	Std. Dev	Std. Error Mean	95% CI of the Difference		t	df	Sig. (2-tailed)
					Lower	Upper			
Pair 1	BBIBU_1 - BBIBU_2	-1.730	1.773	0.129	-1.985	-1.476	-13.413	188.000	0.000

Dari hasil uji-t berpasangan tersebut terlihat bahwa rata-rata perbedaan antara *BBIBU_1* dengan *BBIBU_2* adalah sebesar -1.73. Tanda minus (-) berarti berat sesudah lebih besar daripada berat sebelum intervensi. Artinya ada peningkatan berat badan sesudah intervensi dengan rata-rata peningkatan tersebut adalah 1.73 kg.

Hasil perhitungan nilai “t” adalah sebesar 13.41 dengan p-value 0.000 (uji 2-arah). Hal ini berarti kita menolak H_0 dan menyimpulkan bahwa **pada populasi** (dari mana sampel tersebut diambil) **secara statistik ada perbedaan yang bermakna antara rata-rata berat badan sebelum dengan sudah intervensi**.

5.4. Penyajian Hasil Uji-t Dependen pada Data Berpasangan

Tabel Distribusi nilai rata-rata berat ibu antara sebelum dengan sesudah hamil

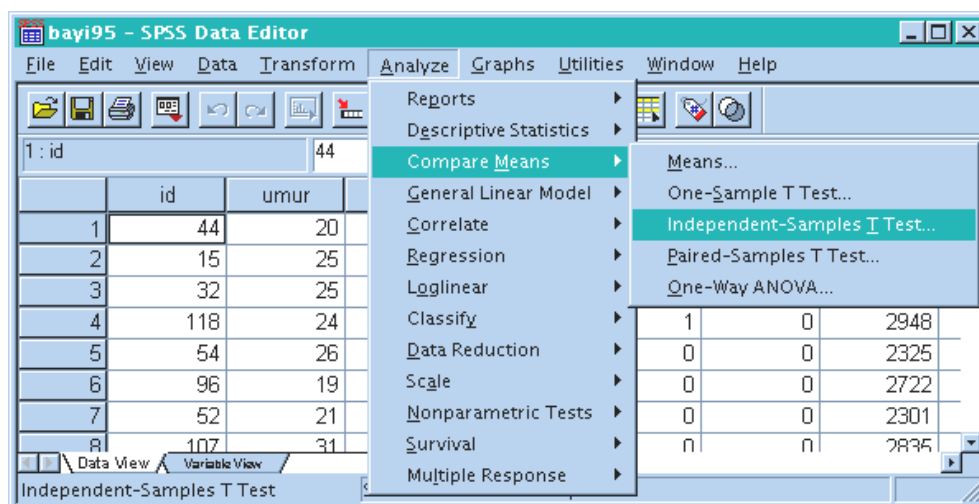
Variabel	n	Mean	SD	p-value
Berat Badan ibu				
- Sebelum hamil	189	58.39	13.76	0.000
- Sesudah hamil	189	60.12	13.72	

Dari 189 subjek yang diamati terlihat bahwa rata-rata (mean) berat badan dari ibu sebelum hamil adalah 58.39 dan rata-rata berat badan sesudah hamil adalah 60.12. secara statistik ada perbedaan yang bermakna antara rata-rata berat badan sebelum dengan sudah proses kehamilan.

5.5. Aplikasi Uji-t pada Data Independen

Uji-t untuk data independen dilakukan terhadap dua kelompok data yang tidak saling berkaitan antara satu dengan lainnya. Misalnya membandingkan kelompok intervensi dengan kelompok kontrol atau kelompok ibu-ibu perokok dengan ibu-ibu bukan perokok adalah dua kelompok yang tidak saling berkaitan.

Pada analisis ini kita akan melihat apakah ada perbedaan berat bayi yang lahir dari ibu perokok dengan bayi yang lahir dari ibu bukan perokok. Kita akan melakukan uji hipotesis apakah ada perbedaan rata-rata berat bayi yang lahir dari ibu bukan perokok dengan rata-rata berat bayi yang lahir dari ibu perokok, dengan langkah-langkah sebagai berikut.



1. Bukalah file BAYI95.SAV, sehingga data tampak di Data editor window.
2. Dari menu utama, pilihlah:

Analize <

Compare Mean <

Independent-Samples T-test....

7. Pilih variabel *BBAYI* dengan cara mengklik variable tersebut.
8. Kemudian klik tanda < untuk memasukkannya ke dalam kotak **Test variable(s)**.
9. Pilih variabel *ROKOK* dan masukkan ke dalam kotak **Grouping variable**.
10. Kemudian klik menu **Define group**, dan isi angka 0 (nol) -kode untuk bukan perokok- pada Group-1 dan isi angka 1 (satu) -kode untuk perokok- pada Group-2. Kemudian pilih **Continue**. (Kodenya bisa saja 1 dengan 2 tergantung data yang dipakai)
11. Pada menu "**Options**" pilihlah derajat kepercayaan yang diinginkan, misalnya 95%. Kemudian pilih **Continue**.
12. Klik **OK** untuk menjalankan prosedur. Pada layar Output tampak hasil seperti berikut:

Group Statistics

ROKOK		N	Mean	Std. Deviation	Std. Error Mean
BBAYI	Tidak	115	3054.96	752.41	70.16
	Ya	74	2773.24	660.08	76.73

Hasil tersebut memperlihatkan bahwa ada 115 ibu yang tidak perokok dan mereka mempunyai rata-rata berat bayi sebesar 3054.96 gram. Sedangkan 74 ibu yang perokok melahirkan bayi yang lebih rendah beratnya daripada kelompok sebelumnya yakni dengan rata-rata 2773.24 gram.

Independent Samples Test

		Levene's Test for Equality of Variances		t-test for Equality of Means				
		F	Sig.	t	df	Sig. (2-tailed)	Mean Difference	Std. Error Difference
BBAYI	Equal variances assumed	1.508	.221	2.634	187	.009	281.71	106.97
	Equal variances not assumed			2.709	170.0	.007	281.71	103.97

Uji-t independen menyajikan dua buah uji statistik. Pertama adalah uji Levene's untuk melihat apakah ada perbedaan varians antara kedua kelompok atau tidak. Kedua adalah uji-t untuk melihat apakah ada perbedaan rata-rata kedua kelompok atau tidak.

Jika p-value (Sig.) dari uji Levene's besar dari nilai α (0.05), hal ini berarti varians kedua kelompok adalah sama, maka signifikansi uji-t yang dibaca adalah pada baris pertama (*Equal variances assumed*). Tetapi jika p-value dari uji Levene's kecil atau sama dengan nilai α

(0.05), hal ini berarti bahwa varians kedua kelompok adalah tidak sama, maka signifikansi uji-t yang dibaca adalah pada baris kedua (*Equal variances not assumed*).

Pada contoh diatas signifikansi uji Levene's adalah **0.221**, berarti varians kedua kelompok adalah sama, maka hasil uji-t pada baris pertama memperlihatkan p-value (sig.) adalah **0.009** untuk uji 2-sisi. (Jika uji yang kita lakukan adalah uji 1-sisi maka nilai p-value harus dikalikan 2 sehingga menjadi 0.018). Dapat kita simpulkan bahwa **secara statistik rata-rata berat bayi yang lahir dari populasi ibu yang tidak perokok lebih tinggi dari populasi ibu perokok.**

5.6. Penyajian Hasil Uji-t Independen

Tabel Distribusi nilai rata-rata berat bayi yang dilahirkan oleh ibu perokok dibandingkan dengan bukan ibu perokok

Variabel	n	Mean	SD	T (t-test)	p-value
Ibu Perokok					
- Tidak	115	3054,96	752,41	2,634	0.009
- Ya	74	2773,24	660,08		

Hasil tersebut memperlihatkan bahwa ada 115 ibu yang tidak perokok dan mereka mempunyai rata-rata berat bayi sebesar 3054.96 gram. Sedangkan 74 ibu yang perokok melahirkan bayi dengan berat yang lebih rendah yakni rata-rata 2773.24 gram. Dari hasil uji statistik dapat kita simpulkan bahwa terdapat perbedaan yang bermakna antara berat bayi dari populasi ibu perokok dibandingkan dengan ibu bukan perokok (nilai-p = 0,009). atau Secara statistik rata-rata berat bayi yang lahir dari populasi ibu yang tidak perokok lebih tinggi dari populasi ibu perokok (p-value = 0,018).

6 Uji Beda ≥ 2 -Rata-rata (ANOVA)

6.1. Pengertian

Jika untuk menguji perbedaan rata-rata antara 2 kelompok independen digunakan Uji-t, maka untuk melakukan uji terhadap perbedaan rata-rata antara 3 kelompok independen atau lebih, kita tidak boleh menggunakan uji t berulang-ulang. Misalnya kita ingin mengetahui apakah ada perbedaan rata-rata hasil antara 3 kelompok intervensi, apakah ada perbedaan rata-rata berat badan bayi lahir menurut tingkat pendidikan ibu (rendah, menengah, & tinggi). Dalam menganalisis data seperti ini (lebih dari dua kelompok) sangat tidak dianjurkan menggunakan uji-t. Ada dua kelemahan jika menggunakan uji-t yaitu pertama: kita harus melakukan pengujian berulang kali sesuai kombinasi yang mungkin, kedua: bila melakukan uji-t berulang-ulang akan meningkatkan (inflasi) nilai α , inflasi nilai α sebesar $= 1 - (1-\alpha)^n$, artinya akan meningkatkan peluang mendapatkan hasil yang keliru.

Untuk mengatasi masalah tersebut maka uji statistik yang dianjurkan (uji yang tepat) dalam menganalisis beda lebih dari dua mean kelompok independen adalah Uji ANOVA atau uji-F. Analisis varian (ANOVA) mempunyai dua jenis yaitu analisis varian satu faktor (one way anova) dan analisis varian dua faktor (two ways anova). Pada bab ini hanya akan dibahas analisis varian satu faktor.

Beberapa asumsi yang harus dipenuhi pada uji Anova adalah:

1. Sampel berasal dari kelompok yang independen
2. Varian antar kelompok harus homogen
3. Data masing-masing kelompok berdistribusi normal

Asumsi pertama harus dipenuhi pada saat pengambilan sampel yang dilakukan secara random terhadap beberapa (≥ 2) kelompok yang independen, yang mana nilai pada satu kelompok tidak tergantung pada nilai di kelompok lain. Sedangkan pemenuhan terhadap asumsi kedua dan ketiga dapat dicek jika data telah dimasukkan ke komputer, jika asumsi ini tidak terpenuhi dapat dilakukan transformasi terhadap data. Apabila proses transformasi tidak juga dapat memenuhi asumsi ini maka uji Anova tidak valid untuk dilakukan, sehingga harus menggunakan uji non-parametrik misalnya Kruskal Wallis.

6.2. Konsep Uji ANOVA

Apabila kita ingin membandingkan efek 3 jenis obat terhadap penurunan kadar kolesterol serum darah tikus atau membandingkan rata-rata berat bayi lahir dari ibu yang perokok berat, perokok ringan/pasif, dan bukan perokok. Dengan menggunakan uji Anova yang pada prinsipnya uji Anova adalah melakukan analisis variabilitas data menjadi dua sumber variasi yaitu variasi didalam kelompok (within) dan variasi antar kelompok (between). Bila variasi within dan between sama (nilai perbandingan kedua varian mendekati angka satu) maka berarti tidak ada perbedaan efek dari intervensi yang dilakukan, dengan kata lain nilai mean yang dibandingkan tidak ada perbedaan. Sebaliknya bila variasi antar kelompok lebih besar dari variasi didalam kelompok, artinya intervensi tersebut memberikan efek yang berbeda, dengan kata lain nilai mean yang dibandingkan menunjukkan adanya perbedaan.

$$F = \frac{MS_b^2}{MS_w^2} = \frac{\text{Varian_between}/(k-1)}{\text{Varian_within}/(n-k)}$$

df = k-1 → untuk pembilang

n - k → untuk penyebut

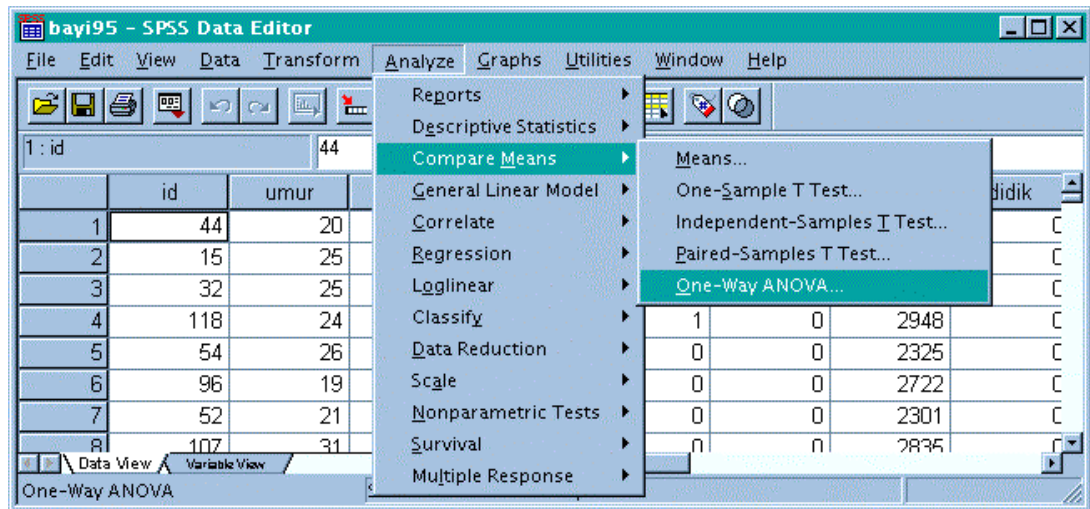
6.3. Aplikasi Uji ANOVA

Uji-Anova digunakan untuk melihat perbedaan rata-rata dari dua atau lebih kelompok independen (data yang tidak saling berkaitan antara satu dengan lainnya). Misalnya membandingkan pengaruh dari 3 jenis intervensi atau membandingkan rata-rata berat bayi dari kelompok ibu-ibu perokok berat dengan perokok ringan atau bukan perokok.

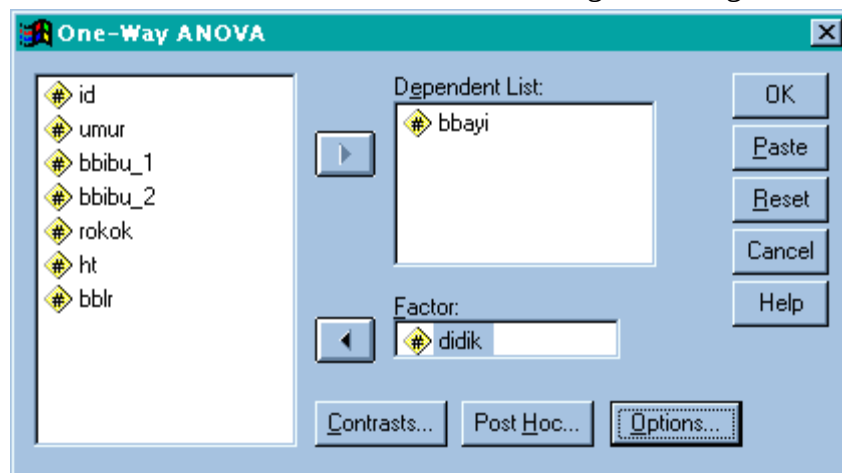
Pada contoh analisis ini kita akan melihat apakah ada perbedaan berat bayi yang lahir dari ibu yang berpendidikan SD, ibu yang berpendidikan SMP, dengan ibu yang berpendidikan SMA. Kita akan melakukan uji hipotesis apakah ada perbedaan rata-rata berat bayi yang lahir dari ibu dari jenis pendidikan yang berbeda, dengan langkah-langkah sebagai berikut.

1. Bukalah file BAYI95.SAV sampai tampak pada Data editor window.
2. Dari menu utama, pilihlah
Analyze <
Compare Means <

One-way ANOVA ...



3. Klik variabel *BBAYI* dan klik tanda < untuk memasukkannya ke kotak **Dependent List**.
4. Klik variabel *DIDIK* dan klik tanda < untuk memasukkannya ke kotak **Factor**.
5. Pada menu **Options..** klik Descriptive and Homogeneity of variances.
6. Klik **Continue** dan Anda akan kembali ke kotak dialog awal dengan isian lengkap



7. Klik **OK** untuk menjalankan prosedur. Hasilnya tampak di output seperti berikut:

Descriptives

BBAYI								
	N	Mean	Std. Deviation	Std. Error	95% Confidence Interval for Mean		Minimum	Maximum
					Lower Bound	Upper Bound		
SD	47	2400.43	695.90	101.51	2196.10	2604.75	709	3940
SMP	84	2915.17	555.33	60.59	2794.65	3035.68	1588	4153
SMA	58	3428.38	655.32	86.05	3256.07	3600.69	1729	4990
Total	189	2944.66	729.02	53.03	2840.05	3049.26	709	4990

Pada hasil di atas terlihat bahwa rata-rata berat bayi pada ibu dengan pendidikan SD adalah 2400.43 gram, pada ibu dengan pendidikan SMP adalah 2915.17 gram, dan pada ibu berpendidikan SMA adalah 3428.38 gram. Standar deviasi, nilai minimum-maximum, dan interval 95% tingkat kepercayaan juga diperlihatkan.

Test of Homogeneity of Variances

	Levene Statistic	df1	df2	Sig.
BBAYI	1.300	2	186	.275

Salah satu asumsi dari uji Anova adalah varians masing-masing kelompok harus sama. Untuk itu dilakukan uji homogenitas varians yang hasilnya memperlihatkan bahwa p-value (sig.) lebih besar dari nilai $\alpha=0.05$, berarti varians antar kelompok adalah sama. Jika varians tidak sama, uji Anova tidak valid untuk dipakai. Catatan: dalam hal ini kita tidak melakukan uji normalitas.

ANOVA

BBAYI

	Sum of Squares	df	Mean Square	F	Sig.
Between Groups	27565146	2	13782572.92	35.432	.000
Within Groups	72351907	186	388988.746		
Total	99917053	188			

Pada hasil di atas diperoleh nilai ANOVA $F = 35.43$ dengan $p\text{-value}=0.000$ (dalam keadaan ini boleh juga ditulis $p < 0.001$). Hipotesis nol pada uji ANOVA adalah tidak ada perbedaan rata-rata berat bayi antara kelompok ibu dengan pendidikan SD, SMP, dan SMA. Sedangkan hipotesis alternatifnya adalah salah satu nilai rata-rata berat bayi berbeda dengan lainnya. Dengan menggunakan $\alpha=0.05$, dari hasil di atas kita menolak hipotesis nol. Sehingga kita menyimpulkan ada perbedaan berat badan bayi dari ke tiga kelompok ibu tersebut (setidaknya salah satu nilai mean berbeda dengan lainnya). Namun, kita belum tahu kelompok mana yang berbeda antara satu dengan yang lainnya.

Dengan uji ANOVA saja kita belum tahu kelompok mana yang berbeda, apakah antara pendidikan SD dengan SMP, SD dengan SMA, atau SMP dengan SMA. Untuk menjawab pertanyaan ini kita harus melakukan uji banding ganda. Untuk melakukan uji banding ganda, kita harus klik menu **Post Hoc...** pada kotak dialog ANOVA.

Silakan kembali ke langkah awal (Analyze, Compare means, One-way-Anova) Pada kotak dialog tersebut ada banyak pilihan uji komparasi ganda. Anda harus membuka buku statistik untuk memahami kelebihan dan kekurangan masing-masing uji. Pada contoh kali ini, kita akan menggunakan uji Tukey honestly significant different (Tukey HSD), suatu uji yang sering digunakan. Kemudian klik **Tukey** untuk meminta agar uji tersebut dilakukan oleh

komputer. Hasil output SPSS adalah sama dengan hasil uji ANOVA sebelumnya dan ditambah dengan tampilan berikut:

Post Hoc Tests

Multiple Comparisons

Dependent Variable: BBAYI

Tukey HSD

(I) DIDIK	(J) DIDIK	Mean Difference (I-J)	Std. Error	Sig.	95% Confidence Interval	
					Lower Bound	Upper Bound
SD	SMP	-514.74*	113.61	.000	-781.01	-248.47
	SMA	-1027.95*	122.41	.000	-1314.84	-741.07
SMP	SD	514.74*	113.61	.000	248.47	781.01
	SMA	-513.21*	106.48	.000	-762.76	-263.66
SMA	SD	1027.95*	122.41	.000	741.07	1314.84
	SMP	513.21*	106.48	.000	263.66	762.76

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

Pada hasil di atas terlihat perbedaan yang “bermakna” pada $\alpha=0.05$ yang ditandai dengan tanda *. Pada baris pertama (SD) dapat dilihat perbandingan antara berat bayi dari ibu berpendidikan SD dengan berat bayi dari ibu berpendidikan SMP atau SMA. Begitu juga dengan baris ke-2, terlihat perbandingan antara berat bayi dari ibu berpendidikan SMP dengan berat bayi dari ibu berpendidikan SD dan SMA.

Dari hasil di atas dapat disimpulkan bahwa ada perbedaan rata-rata berat bayi antara ibu berpendidikan SD dengan ibu berpendidikan SMP, antara ibu berpendidikan SD dengan ibu berpendidikan SMP, dan ibu berpendidikan SMP dengan ibu berpendidikan SMA.

6.4. Penyajian Hasil Uji ANOVA

Tabel Distribusi nilai rata-rata berat bayi menurut status pendidikan ibu

Variabel	N	Mean	SD	F (Anova)	p-value
Pendidikan ibu					
- SD	47	2400,43	695,9	35,4	0,000
- SMP	84	2915,17	555,3		
- SMA	58	3428,38	655,3		

Pada hasil di atas terlihat bahwa rata-rata berat bayi meningkat sesuai dengan peningkatan status pendidikan ibu. Ibu dengan pendidikan SD rata-ratanya adalah 2400.43 gram, ibu dengan pendidikan SMP adalah 2915.17 gram, dan ibu berpendidikan SMA adalah 3428.38 gram. Hasil uji Anova memperlihatkan bahwa ada perbedaan yang bermakna pada rata-rata berat bayi menurut tingkat pendidikan ibu (nilai-p 0.000).

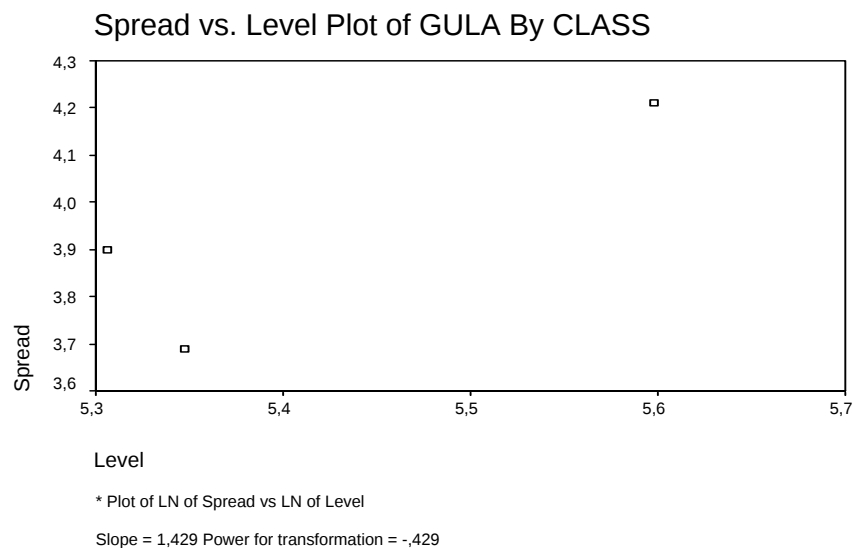
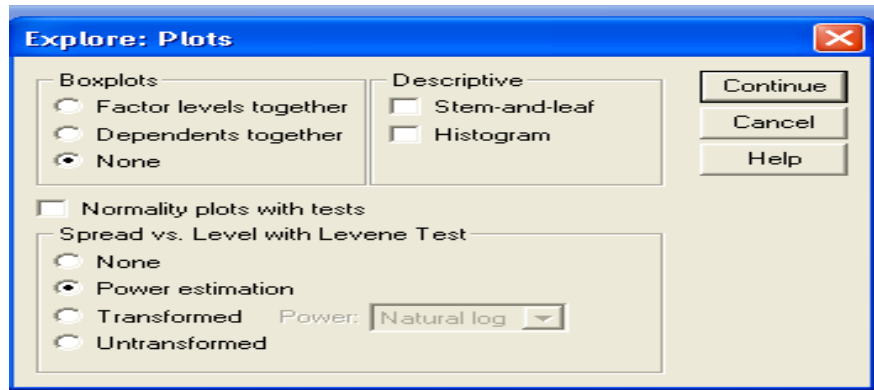
Tabel Signifikansi perbedaan rata-rata berat bayi menurut Pendidikan ibu
(Hasil Uji-Tukey)

Pendidikan	p-value	Simpulan
- SD vs SMP	0,000	Berbeda bermakna
- SD vs SMU	0,000	Berbeda bermakna
- SMP vs SMU	0,000	Berbeda bermakna

Analisis lebih lanjut memperlihatkan bahwa ada perbedaan rata-rata berat bayi antara ibu berpendidikan SD dengan SMP, SD dengan SMU, dan SMP dengan SMU.

6.5. Transformasi Jika Varians Tidak Homogen

Jika ingin mendapatkan hasil uji Anova yang valid, maka varians antar kelompok harus homogen. Namun asumsi ini terkadang tidak bisa dipenuhi oleh data kita, maka kita dapat melakukan transformasi agar varians antar kelompok menjadi homogen dengan cara sbb:



Nilai slope dan nilai power adalah panduan untuk menentukan jenis transformasi. Berikut ini ditampilkan tabel transformasi yang dianjurkan untuk menghomogenkan varians berdasarkan nilai slope dan power.

Tabel 1 : Panduan Mencari Bentuk Transformasi Terbaik Dengan
Memperhitungkan Factor Slope Dan Power.

Slope	Power	Bentuk transformasi
-1	2	Square (kuadrat)
0	1	Tidak perlu transformasi
0,5	0,5	Square root (akar)
1	0	Logaritma
1,5	-0,5	1/ square root
2	-1	Reciprocal (1/n)

7 Uji Beda Proporsi (χ^2 :Chi-square)

7.1. Pengertian

Dalam penerapan praktis, kita ingin menguji apakah ada hubungan antara dua variabel kategorik. Atau kita ingin menguji apakah ada perbedaan proporsi pada populasi. Jika perbedaan proporsi itu eksist dapat kita katakan bahwa adanya keterkaitan atau hubungan antara dua variabel kategorik tersebut.

Misalnya kita ingin menguji apakah proporsi hipertensi pada populasi perokok lebih tinggi dari proporsi hipertensi pada populasi bukan perokok. Pengamatan dilakukan terhadap kebiasaan merokok dan pengukuran dilakukan terhadap tekanan darahnya (yang setelah diukur dikategorikan menjadi normotensi dan hipertensi). Apabila pengamatan diatas disusun didalam suatu tabel, maka tabel tersebut dinamakan tabel kontingensi (tabel silang). Dari data tersebut dapat dilakukan uji statistik untuk melihat ada tidaknya asosiasi antara dua sifat/variabel tadi (kebiasaan merokok dan hipertensi)

Uji statistik untuk melihat hubungan antara dua variabel yang dikategorikan sering digunakan uji “**chi-square**” (χ^2). Secara spesifik uji chi square dapat digunakan untuk menentukan/menguji:

- 1) Ada tidaknya hubungan/asosiasi antara 2 variabel (*test of independency*)
- 2) Apakah suatu kelompok homogen dengan sub kelompok lain (*test of homogeneity*)
- 3) Apakah ada kesesuaian antara pengamatan dengan parameter tertentu yang dispesifikasikan (*Goodness of fit*).

Secara umum tidak ada asumsi yang harus dipenuhi untuk uji χ^2 , karena distribusi χ^2 ini termasuk free-distribution. Hanya saja, jumlah pengamatan tidak boleh terlalu sedikit, frekuensi harapan (*expected frequency*) tidak boleh kurang dari satu dan frekuensi harapan yang kurang dari lima tidak boleh lebih dari 20%. Jika asumsi ini tidak terpenuhi maka harus dilakukan pengelompokan ulang sampai hanya menjadi dua kelompok saja (tabel 2 x 2), Pada tabel 2 x 2 gunakan Fisher Exact test yang merupakan nilai-p sebenarnya, yang secara otomatis sudah ada di output SPSS.

7.2. Konsep Uji Chi Square

Dasar dari uji kaid kuadrat adalah membandingkan frekuensi yang diamati dengan frekuensi yang diharapkan. Misalnya sebuah uang logam dilambungkan seratus kali, kemudian diamati permukaan uang yang muncul yaitu A (Angka) sebanyak 55 kali dan B (Gambar) sebanyak 45 kali. Kalau uang logam tersebut seimbang tentu permukaan A dan B diharapkan muncul sama banyak yaitu 50 kali. Hal ini berarti tidak ada perbedaan antara frekuensi yang diamati (Observed = O) adalah 55 kali dengan frekuensi yang diharapkan (Expected= E) yakni 50 kali. Jadi tidak ada perbedaan antara pengamatan dengan yang diharapkan (O - E), seandainya terjadi perbedaan, apakah perbedaan itu cukup berarti (bermakna) atau hanya karena faktor kebetulan saja. Hasil percobaan melambungkan mata uang tadi disajikan seperti tabel dibawah ini:

Tabel: 1 Hasil pelambungan 100 kali sebuah mata uang logam

	(1)	(2)	(3)	(4)	(5)
	O(observed)	E (expected)	O - E	(O - E) ²	$\frac{(O - E)^2}{E}$
A (Angka)	45	50	-5	25	0.5
B (Huruf)	55	50	5	25	0.5
Total	100	100	0	200	$\chi^2 = 1.0$

Perhitungan nilai χ^2 dilakukan dengan rumus berikut, dan dari tabel tersebut diatas dapat dilihat bahwa nilai χ^2 adalah 1.0.

$$\chi^2 = \sum \frac{(O - E)^2}{E}$$

E

Pertanyaan berikutnya ialah apakah nilai χ^2 yang telah dihitung = 1.0 memiliki kemungkinan besar untuk terjadi atau hanya terjadi secara kebetulan (merupakan peristiwa yang jarang terjadi), misalnya kemungkinannya kecil dari 5%. Untuk menjawab pertanyaan ini, perlu diketahui distribusi kuantitas χ^2 yakni distribusi probabilitas untuk statistik. Para ahli statistik telah membuktikan, bahwa distribusi ini mempunyai kemencengan positif, dengan menghitung luas daerah diluar harga 1.0 pada distribusi χ^2 , dapat ditentukan nilai-p (p-value) serta keputusan untuk menerima atau menolak hipotesis dengan membandingkannya dengan nilai α . Ternyata setelah dihitung nilai-p adalah 0.15 artinya terjadi kesalahan 15% jika kita menyatakan “55 berbeda dengan 50”, sehingga kita lebih memilih untuk menyatakan “tidak adanya perbedaan antara 55 dengan 50”.

7.3. Aplikasi Uji χ^2 pada Tabel Silang 2 x 2

Cara yang paling sering digunakan untuk menampilkan hubungan antara dua data katagori adalah dengan menggunakan tabel silang (*crosstabs*). Dalam contoh ini, kita akan menguji apakah ada hubungan antara merokok dengan BBLR. Agar memudahkan dalam penyajian datanya, kita akan membuat tabel silang antara merokok dan BBLR dari file BAYI95.SAV.

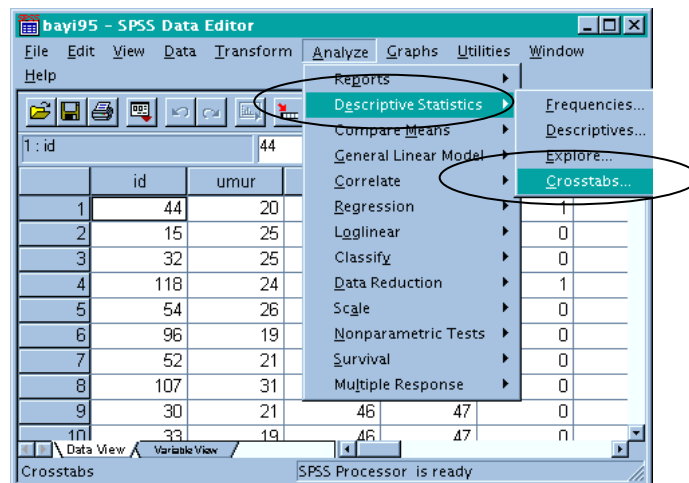
1. Bukalah file BAYI95.SAV, sehingga data tampak di Data editor window.
2. Dari menu utama, pilihlah: (pada SPSS 10.0)

Analyze <

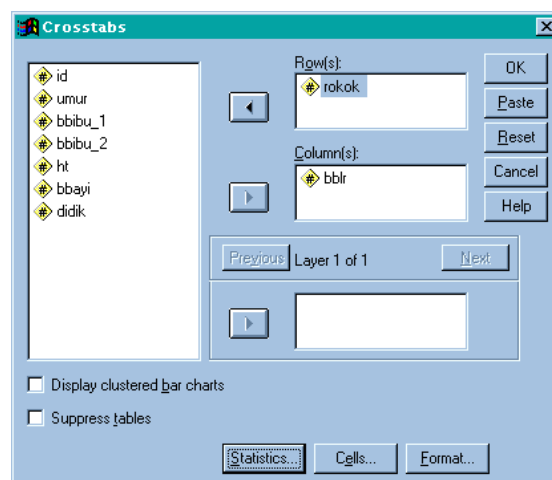
Descriptif statistic <

Crosstabs...

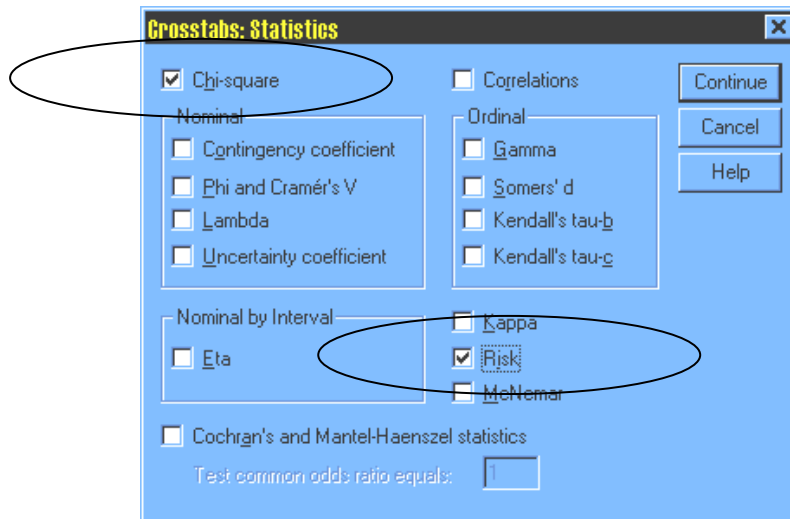
Seperti gambar berikut:



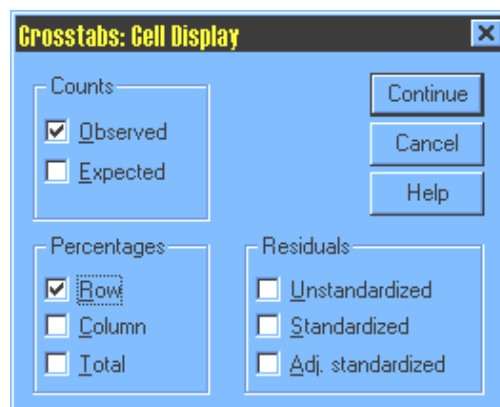
13. Pilih variabel *ROKOK*, kemudian klik tanda > untuk memasukkannya ke kotak **Row(s)**
14. Pilih variabel *BBLR*, kemudian klik tanda > untuk memasukkannya ke kotak **Colom(s)**.



15. Pada menu “**Statistics**” pilih **Chi-Square** dan **Risk** dengan mengklik kotak disampingnya hingga muncul tanda “√”. Jika anda klik sekali lagi, maka tanda “√” akan hilang. Kemudian Klik Continue.



16. Klik menu “**Cells**”, kemudian aktifkan **Observed** pada menu Count dan aktifkan **Rows** pada menu Percentages hingga muncul tanda “√”. Kemudian Klik Continue.



17. Klik **OK** untuk menjalankan prosedur. Pada layar tampak hasil seperti berikut:

ROKOK * BBLR Crosstabulation

			BBLR		Total
			Tidak	Ya	
ROKOK	Tidak	Count	86	29	115
		% within ROKOK	74.8%	25.2%	100.0%
	Ya	Count	44	30	74
		% within ROKOK	59.5%	40.5%	100.0%
Total	Count		130	59	189
	% within ROKOK		68.8%	31.2%	100.0%

Dari tabel silang tersebut terlihat bahwa dari 74 ibu-ibu perokok, ada 30 orang (40.5%) melahirkan bayi dengan BBLR. Dari 115 ibu-ibu yang bukan perokok, hanya ada 29 orang (25.2%) yang melahirkan bayi BBLR. Artinya **proporsi BBLR pada ibu perokok lebih besar dari proporsi BBLR pada ibu yang bukan perokok**.

Walaupun secara proporsional terlihat ada hubungan antara merokok dan BBLR yang terlihat dari proporsi bayi BBLR lebih besar pada ibu perokok dari pada ibu tidak perokok, namun untuk menguji apakah hubungan tersebut bermakna secara statistik, maka kita harus melakukan uji chi-square dengan melihat hasil output berikut:

Chi-Square Tests

	Value	df	Asymp. Sig. (2-sided)	Exact Sig. (2-sided)	Exact Sig. (1-sided)
Pearson Chi-Square	4.924 ^b	1	.026		
Continuity Correction ^a	4.236	1	.040		
Likelihood Ratio	4.867	1	.027		
Fisher's Exact Test				.036	.020
Linear-by-Linear Association	4.898	1	.027		
N of Valid Cases	189				

a. Computed only for a 2x2 table

b. 0 cells (.0%) have expected count less than 5. The minimum expected count is 23.10.

Output SPSS menampilkan semua nilai chi-square dari berbagai macam uji, seperti Pearson Chi-square, Continuity Correction, atau Fisher's Exact Test. Masing-masing uji tersebut dilengkapi dengan p-value untuk test 2-sisi.

Untuk memilih nilai χ^2 atau p-value yang paling sesuai, kita harus berpedoman pada asumsi-asumsi yang terkait dengan uji χ^2 . Antara lain:

1. Pada tabel lebih dari 2x2 (misalnya 3x2 atau 3x3), apabila nilai frekuensi harapan (*expected*) yang kurang dari 5 tidak lebih dari 20%, maka nilai χ^2 atau p-value dari *Pearson Chi-square* atau *Likelihood Ratio* dapat kita laporkan. *Catt: Jika nilai expected yang kurang dari 5 lebih dari 20% atau ada nilai expected yang kurang dari 1.0 (karena ada sell yang kosong), maka hasil uji chi-square tidak valid, harus dilakukan pengelompokan ulang terlebih dahulu.*
2. Untuk tabel 2 x 2, nilai χ^2 atau p-value dari *Continuity Correction* dapat kita laporkan. Tetapi jika nilai frekuensi harapan kurang dari 5, maka nilai p-value dari *Fisher's Exact Test* yang harus kita laporkan.

Nilai-p *Fisher's Exact Test* merupakan p-value yang cukup valid, sehingga dapat juga kita laporkan meskipun frekuensi harapan tidak ada yang kurang dari 5. Dalam hal ini, kita pakai nilai tersebut **dengan p-value = 0.036. Artinya hubungan antara merokok dengan BBLR secara statistik cukup bermakna dan bukanlah terjadi secara kebetulan belaka.**

Dari tabel Risk Estimate terlihat bahwa **OR=2.022**. Hal ini berarti bahwa **ibu yang perokok mempunyai kecenderungan (risiko) sebesar 2 kali lebih besar untuk melahirkan bayi dengan BBLR dibandingkan dengan ibu yang bukan perokok**.

Risk Estimate

	Value	95% Confidence Interval	
		Lower	Upper
Odds Ratio for ROKOK (0 / 1)	2.022	1.081	3.783
For cohort BBLR = 0	1.258	1.013	1.561
For cohort BBLR = 1	.622	.409	.945
N of Valid Cases	189		

Untuk estimasi resiko (OR atau RR), nilai perhitungannya dari tabel silang hanya akan keluar jika tabel silang yang dibuat adalah tabel 2 x 2. Jika tabel silang yang dibuat lebih dari tabel 2 x 2 (misalnya 2x3, 3x3), maka nilai estimasi resiko tidak akan keluar, karena SPSS tidak bisa menghitungnya. Untuk menghitung nilai OR pada tabel 2x3 atau 3x3 kita dapat memilih salah satu dari 3 alternatif berikut yaitu 1) menghitung secara manual dari tabel silang tersebut, 2) membuat dummy variabel kemudian dilakukan crosstab, atau 3) melalui regresi logistik sederhana.

7.4. Aplikasi Uji χ^2 pada Tabel Silang 2 x 3

Pada contoh ini, kita akan menguji apakah ada perbedaan proporsi BBLR pada populasi dengan tingkat pendidikan yang berbeda-beda (SD, SMP, dan SMA), kita akan membuat tabel silang antara DIDIK dan BBLR dari file BAYI95.SAV. Dengan langkah yang sama seperti pada tabel 2x2 kita lakukan prosedur untuk Crosstabs. Pilih variabel **ROKOK**, kemudian klik tanda < untuk memasukkannya ke kotak **Row(s)**. Pilih variabel **BBLR**, kemudian klik tanda < untuk memasukkannya ke kotak **Colom(s)**. Pada menu “**Statistics**” aktifkan **Chi-Square**. Pada menu “**Cells**” aktifkan **Observed** dan aktifkan **Rows**. Pilih continue dan klik OK untuk menjalankan analisis. Hasilnya sebagai berikut:

DIDIK * BBLR Crosstabulation

			BBLR		Total
			Tidak	Ya	
DIDIK	SD	Count	18	29	47
		% within DIDIK	38.3%	61.7%	100.0%
	SMP	Count	61	23	84
		% within DIDIK	72.6%	27.4%	100.0%
	SMA	Count	51	7	58
		% within DIDIK	87.9%	12.1%	100.0%
Total	Count	130	59	189	
	% within DIDIK	68.8%	31.2%	100.0%	

Tabel silang tersebut memperlihatkan bahwa dari 47 ibu-ibu berpendidikan SD, ada 29 orang (61.7%) melahirkan bayi dengan BBLR. Dari 84 ibu-ibu yang berpendidikan SMP, ada 23 orang

(27.4%) yang melahirkan bayi BBLR. Dari 58 ibu-ibu yang berpendidikan SMA, ada 7 orang (12.1%) yang melahirkan bayi BBLR Artinya **semakin rendah tingkat pendidikan ibu akan semakin besar proporsi BBRL**. Walaupun secara proporsional terlihat ada hubungan antara PENDIDIKAN dengan BBLR yang mana ibu berpendidikan rendah cenderung melahirkan bayi BBLR, namun untuk menguji apakah hubungan tersebut bermakna secara statistik, maka kita lakukan uji chi-square dengan melihat hasil output sebagai berikut:

Chi-Square Tests

	Value	df	Asymp. Sig. (2-sided)
Pearson Chi-Square	30.822 ^a	2	.000
Likelihood Ratio	30.774	2	.000
Linear-by-Linear Association	28.715	1	.000
N of Valid Cases	189		

a. 0 cells (.0%) have expected count less than 5. The minimum expected count is 14.67.

Dalam tabel tersebut terlihat bahwa nilai χ^2 baik Pearson maupun Likelihood Ratio memperlihatkan hasil yang sama yaitu 30.8 dengan **p-value = 0.000**. Artinya secara statistik ada hubungan yang bermakna antara pendidikan ibu dengan BBLR dan kejadian tersebut sangat kecil kemungkinannya untuk terjadi secara kebetulan.

7.5. Dummy Variabel

Output SPSS tidak bisa menampilkan nilai OR, karena nilai OR hanya bisa dihitung pada tabel 2 x 2, padahal tabel untuk pendidikan dengan BBLR adalah tabel 3 x 2. Untuk bisa mendapatkan nilai OR dan CI-nya pada tabel 3 x 2 ada dua cara yang dapat dilakukan yaitu 1) harus dibuat dummy variabel tabel terlebih dahulu kemudian baru dilakukan Crosstabs atau 2) lakukan analisis regresi logistik sederhana.

Untuk membuat dummy variabel dari pendidikan (SD, SMP, & SMA), pertama-tama harus ditetapkan kelompok mana yang akan dijadikan sebagai pembanding, kelompok pembanding akan diberi kode = 0 (nol).

Dalam hal ini sebagai pembanding kita tetapkan SMA sehingga SMA diberi kode 0 pada variabel dummy. Dari DIDIK (0=SD, 1=SMP, 2=SMA) dibuat 2-variabel dummy dari menu Transformasi data dengan perintah RECODE.

DIDIK_1 (0=SMA, 1=SD)

DIDIK_2 (0=SMA, 1=SMP)

Selanjutnya lakukan crosstabs dari 2 variabel dummy itu dengan BBLR, hasilnya sbb:

DIDIK_1 dengan BBLR

Crosstab

			BBLR		Total
			Tidak	Ya	
DIDIK_1	SMA	Count	51	7	58
		% within DIDIK_1	87.9%	12.1%	100.0%
	SD	Count	18	29	47
		% within DIDIK_1	38.3%	61.7%	100.0%
Total		Count	69	36	105
		% within DIDIK_1	65.7%	34.3%	100.0%

Proporsi BBLR lebih tinggi pada ibu dengan pendidikan SD (61.7%) dibandingkan dengan ibu pendidikan SMA (12.1%). Hasil ini sama dengan tabel 3 x 2 sebelumnya.

Chi-Square Tests

	Value	df	Asymp. Sig. (2-sided)	Exact Sig. (2-sided)	Exact Sig. (1-sided)
Pearson Chi-Square	28.386 ^b	1	.000		
Continuity Correction ^a	26.226	1	.000		
Likelihood Ratio	29.732	1	.000		
Fisher's Exact Test				.000	.000
Linear-by-Linear Association	28.116	1	.000		
N of Valid Cases	105				

a. Computed only for a 2x2 table

b. 0 cells (.0%) have expected count less than 5. The minimum expected count is 16.11.

Nilai-p dari χ^2 dan Fisher Exact memperlihatkan hasil yang sama dan bermakna secara statistik ($p=0.000$).

Dari Nilai OR dapat disimpulkan bahwa ibu yang berpendidikan SD mempunyai kecenderungan untuk melahirkan bayi BBLR sebesar 11.7 kali lebih besar dibandingkan dengan ibu yang berpendidikan SMA.

Risk Estimate

	Value	95% Confidence Interval	
		Lower	Upper
Odds Ratio for DIDIK_1 (SMA / SD)	11.738	4.384	31.429
For cohort BBLR = Tidak	2.296	1.578	3.341
For cohort BBLR = Ya	.196	.094	.406
N of Valid Cases	105		

DIDIK_2 dengan BBLR**Crosstab**

			BBLR		Total
			Tidak	Ya	
DIDIK_2	SMA	Count	51	7	58
		% within DIDIK_2	87.9%	12.1%	100.0%
	SMP	Count	61	23	84
		% within DIDIK_2	72.6%	27.4%	100.0%
Total		Count	112	30	142
		% within DIDIK_2	78.9%	21.1%	100.0%

Proporsi BBLR lebih tinggi pada ibu yang berpendidikan SD (27.4%) dibandingkan dengan ibu yang berpendidikan SMA (12.1%), dan hubungan ini bermakna secara statistik ($p = 0.036$)

Chi-Square Tests

	Value	df	Asymp. Sig. (2-sided)	Exact Sig. (2-sided)	Exact Sig. (1-sided)
Pearson Chi-Square	4.827 ^b	1	.028		
Continuity Correction ^a	3.952	1	.047		
Likelihood Ratio	5.099	1	.024		
Fisher's Exact Test				.036	.022
Linear-by-Linear Association	4.793	1	.029		
N of Valid Cases	142				

a. Computed only for a 2x2 table

b. 0 cells (.0%) have expected count less than 5. The minimum expected count is 12.25.

Dari Nilai OR dapat disimpulkan bahwa ibu yang berpendidikan SMP mempunyai kecenderungan untuk melahirkan bayi BBLR sebesar 2.7 kali lebih besar dibandingkan dengan ibu yang berpendidikan SMA.

Risk Estimate

	Value	95% Confidence Interval	
		Lower	Upper
Odds Ratio for DIDIK_2 (SMA / SMP)	2.747	1.090	6.922
For cohort BBLR = Tidak	1.211	1.029	1.424
For cohort BBLR = Ya	.441	.203	.959
N of Valid Cases	142		

7.6. Regresi Logistik Sederhana

Seperti telah dijelaskan sebelumnya bahwa Crosstabs pada tabel 2 x 2 tidak bisa menampilkan nilai OR, misalnya pendidikan dengan BBLR yang merupakan tabel 3 x 2. Untuk bisa mendapatkan nilai OR dan CI-nya pada tabel 3 x 2 ada dua cara yang dapat dilakukan yaitu 1) harus dibuat dummy variabel tabel terlebih dahulu kemudian baru dilakukan Crosstabs atau 2) lakukan analisis regresi logistik sederhana. Langkah-langkah dengan dummy variabel telah diperlihatkan pada penjelasan di atas, sedangkan langkah-langkah pada regresi logistic sederhana akan diuraikan berikut ini.

Pada contoh ini, kita akan membandingkan risiko kejadian BBLR pada populasi dengan tingkat pendidikan yang berbeda-beda (variabel didik dengan kode sbb: 0=SD, 1=SMP, dan 2=SMA). Sebagai kelompok pembanding kita tetapkan SMA. Lakukan perintah analisis dengan SPSS sebagai berikut:

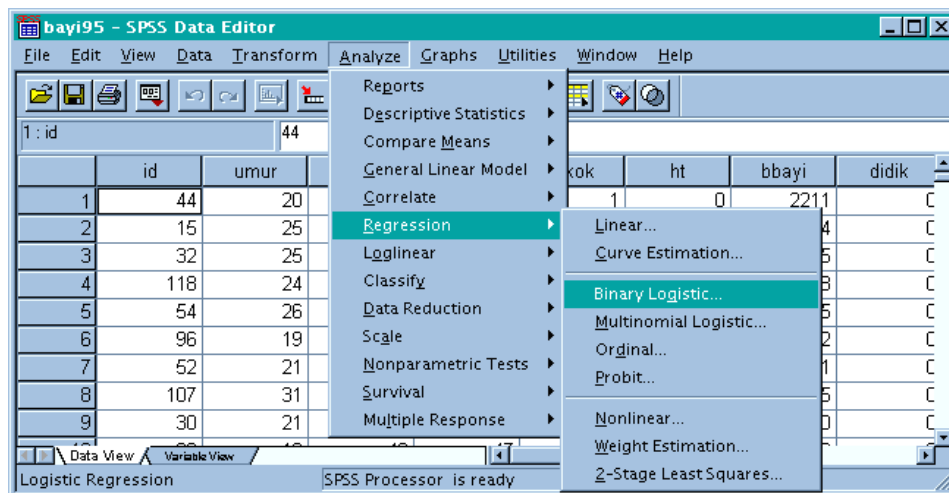
1. Bukalah file BAYI95.SAV, sehingga data tampak di Data editor window.
2. Dari menu utama, pilihlah:

Analyze <

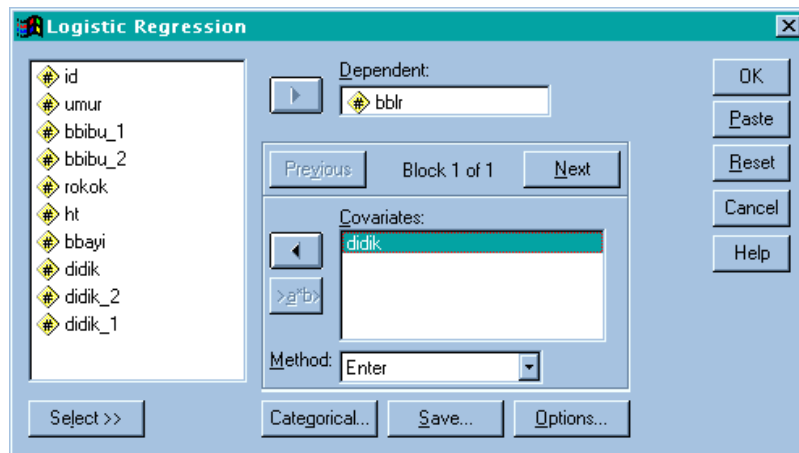
DescriptRegression <

Binary Logistic...

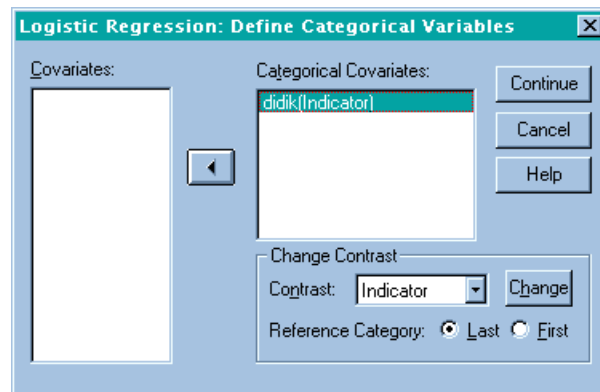
Seperti gambar berikut:



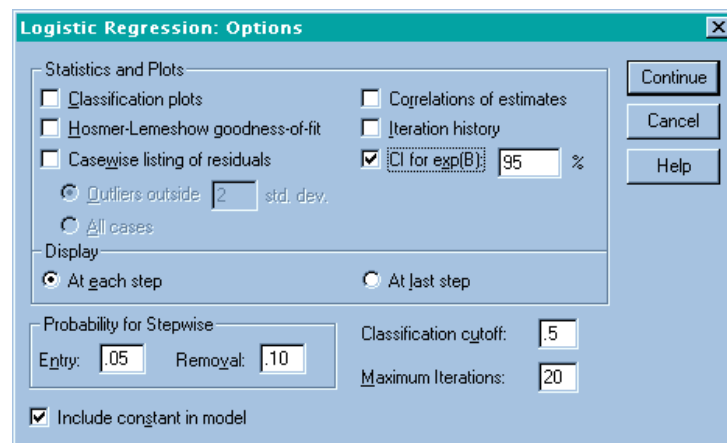
3. Pilih variabel *BBLR*, kemudian masukkan ke kotak **Dependent**
4. Pilih variabel *DIDIK*, kemudian masukkan ke kotak **Covariates**



5. Pada menu “**Categorical**” pilih variabel *DIDIK* dan klik tanda < untuk memasukkannya ke kotak Categorical Covariates
6. Pastikan Reference Kategori adalah Last (artinya kelompok pembanding adalah kode tertinggi, dalam hal ini kode 2=SMA).



7. Klik **Continue** jika sudah selesai, SPSS akan kembali ke menu utama.
8. Klik **Option** kemudian aktifkan CI for exp(B) seperti gambar berikut.



9. Kemudian klik **Continue** jika sudah selesai, SPSS akan kembali ke menu utama.
10. Klik **OK** untuk menjalankan prosedur. Pada layar output akan tampak hasil regresi logistic.

Pada output ini, kita hanya mengambil bagian yang paling akhir saja, yang berkaitan dengan perbandingan risiko BBLR pada berbagai tingkat pendidikan (OR atau Exp(B)) seperti berikut:

Variables in the Equation

		B	S.E.	Wald	df	Sig.	Exp(B)	95.0% C.I. for EXP(B)	
								Lower	Upper
Step 1	DIDIK			26.820	2	.000			
	DIDIK(1)	2.463	.502	24.022	1	.000	11.738	4.384	31.428
	DIDIK(2)	1.011	.472	4.593	1	.032	2.747	1.090	6.922
	Constant	-1.986	.403	24.275	1	.000	.137		

a. Variable(s) entered on step 1: DIDIK.

Dari Nilai OR atau (Exp(B)) dapat disimpulkan bahwa ibu yang berpendidikan SD -*DIDIK(1)*- mempunyai kecenderungan untuk melahirkan bayi BBLR sebesar **11.7** kali lebih besar dibandingkan dengan ibu yang berpendidikan SMA (p-value=0.000). Sedangkan ibu yang berpendidikan SMP -*DIDIK(2)*- mempunyai kecenderungan untuk melahirkan bayi BBLR sebesar **2.7** kali lebih besar dibandingkan dengan ibu yang berpendidikan SMA (p-value=0.032).

7.7. Penyajian Hasil Uji Beda Proporsi (Chi-Square)

Tabel Hubungan BBLR dengan Pendidikan ibu dan Status Rokok

Variabel	BBLR		Total n=189	OR (95%CI)	p-value
	Tidak n (%)	Ya n (%)			
Rokok					
- Tidak	86 (74,8)	29 (25,2)	115	2,0 (1,1—3,8)	0,020
- Ya	44 (59,5)	30 (40,5)	74		
Pendidikan ibu					
- SD	18 (38,3)	29 (61,7)	47	11,7 (4,3—31,4)	0,000
- SMP	61 (72,6)	23 (27,4)	84	2,7 (1,1—6,9)	0,032
- SMA	51 (87,9)	7 (12,1)	58	1,0	

Hubungan antara Pendidikan dengan BBLR terlihat bahwa semakin rendah tingkat pendidikan ibu akan semakin besar kemungkinan untuk melahirkan bayi BBLR. Dari 47 ibu-ibu berpendidikan SD, sebanyak 61.7% melahirkan bayi dengan BBLR. Dari 84 ibu-ibu yang berpendidikan SMP, sebanyak 27.4% melahirkan bayi BBLR. Dari 58 ibu-ibu yang berpendidikan SMA, sebanyak 12.1% yang melahirkan bayi BBLR

Dari Nilai OR dapat disimpulkan bahwa ibu yang berpendidikan SD mempunyai kecenderungan untuk melahirkan bayi BBLR sebesar **11.7** kali lebih besar dibandingkan dengan ibu yang berpendidikan SMA (p-value=0.000). Sedangkan ibu yang berpendidikan SMP mempunyai kecenderungan untuk melahirkan bayi BBLR sebesar **2.7** kali lebih besar dibandingkan dengan ibu yang berpendidikan SMA (p-value=0.032).

8 Uji Korelasi & Regresi Linier

8.1. Pendahuluan

Dalam penerapan praktis, kita ingin menguji apakah ada hubungan atau korelasi antara dua variabel numerik, jika ada seperti apa persamaan garis regresi liniernya. Misalnya kita ingin menguji apakah ada hubungan antara berat ibu sebelum hamil (x) dengan berat bayi yang dilahirkannya (y).

Uji statistik untuk melihat hubungan antara dua variabel numerik adalah uji **“uji korelasi”**. Koefisien korelasi ini dikembangkan oleh Pearson, sehingga dikenal dengan nama *Pearson Coeficient Correlation* dengan lambang **“r”** kecil atau **“R”** kapital. Nilai **“r”** berkisar antara 0.0 yang berarti tidak ada korelasi, sampai dengan 1.0 yang berarti adanya korelasi yang sempurna. Semakin kecil nilai **“r”** semakin lemah korelasi, sebaliknya semakin besar nilai **“r”** semakin kuat korelasi.

Selain itu, **“r”** juga mempunyai nilai negatif (-) atau minus yang menandakan adanya hubungan terbalik antara x dengan y . Artinya, semakin tinggi nilai x maka semakin rendah nilai y , misalnya korelasi antara umur dengan kemampuan daya ingat pada kelompok usia lanjut.

Jika korelasi yang ada bermakna secara statistik, kita bisa menganalisis lebih lanjut untuk memprediksi atau memperkirakan berapa nilai (y) jika nilai (x) diketahui. Prediksi tersebut dapat dilakukan jika kita mempunyai persamaan garis lurus yang biasanya disebut dengan istilah **“regresi linier”** dengan persamaan matematis $y = a + bx$. Besaran nilai **“b”** menggambarkan besarnya perubahan (peningkatan/penurunan) pada nilai y untuk setiap kenaikan nilai x sebesar satu satuan.

8.2. Asumsi Normalitas pada Uji Korelasi Pearson

Dasar dari uji korelasi Pearson adalah statistik Parametrik, yang berasumsi data mempunyai distribusi normal. Dalam hal ini variabel y harus berdistribusi normal. Apabila asumsi ini tidak terpenuhi, dapat dilakukan transformasi terlebih dahulu misalnya dengan LOG, AKAR, atau KUADRAT. Jika pada proses transformasi tidak berhasil membuat distribusi data menjadi normal, maka pilihan statistik non-parametrik lebih dianjurkan, yakni uji korelasi Spearman.

8.3. Aplikasi Uji Korelasi Pearson

Dalam contoh ini, kita akan menguji apakah ada korelasi antara berat badan ibu sebelum hamil dengan berat badan bayi yang akan dilahirkannya kelak. Kita akan menggunakan variabel *bbibu_1* dan *bbayi* dari file BAYI95.SAV.

8.3.1. Uji Normalitas

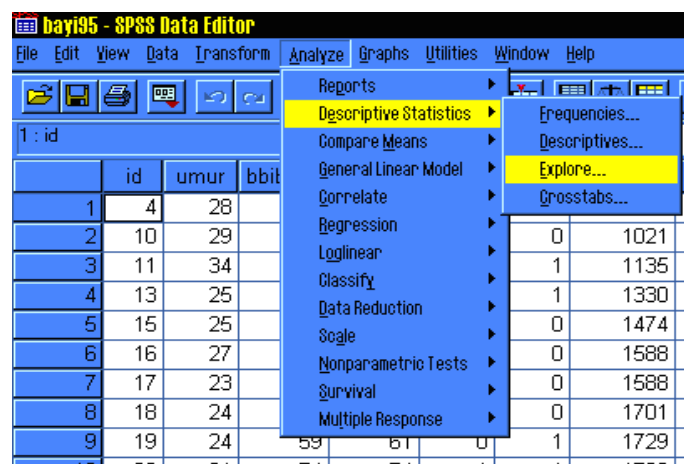
1. Bukalah file BAYI95.SAV, sehingga data tampak di Data editor window.
2. Dari menu utama, pilihlah: (pada SPSS 10.0)

Analyze <

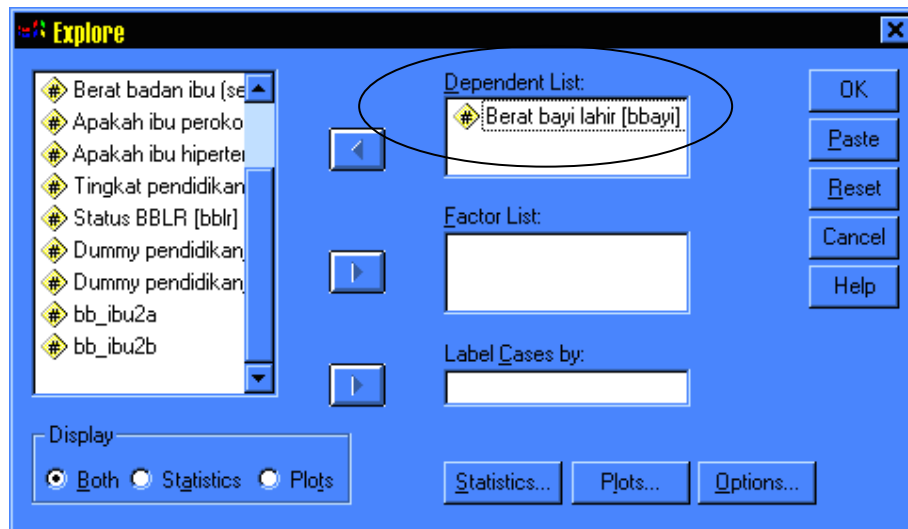
Descriptif statistic <

Explore...

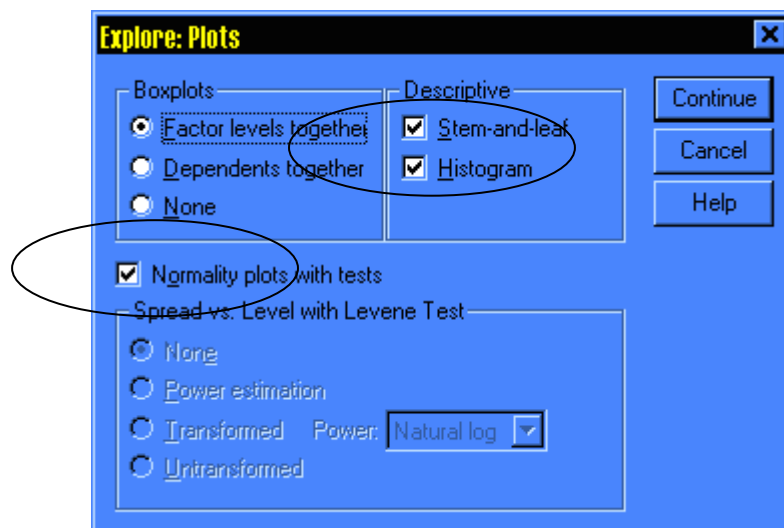
Seperti gambar berikut:



18. Pilih variabel *Berat Bayi Lahir (bbayi)*, kemudian klik tanda > untuk memasukkannya ke kotak **Dependent List**.



19. Pilih Plots..., kemudian aktifkan Histogram dan Normality plots with tests. Kemudian klik Continue.



20. Klik **OK** untuk menjalankan prosedur. Pada layar tampak hasil seperti berikut:

Descriptives				
			Statistic	Std. Error
Berat bayi lahir	Mean		2944.66	53.03
	95% Confidence Interval for Mean	Lower Bound	2840.05	
		Upper Bound	3049.26	
	5% Trimmed Mean		2957.83	
	Median		2977.00	
	Variance		531473.7	
	Std. Deviation		729.02	
	Minimum		709	
	Maximum		4990	
	Range		4281	
	Interquartile Range		1069.00	
	Skewness		-.210	.177
	Kurtosis		-.081	.352

Tests of Normality

	Kolmogorov-Smirnov ^a		
	Statistic	df	Sig.
Berat bayi lahir	.043	189	.200*

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

a. Lilliefors Significance Correction

Hasil uji test normalitas

Dengan uji Kolmogorov-Smirnov, disimpulkan bahwa distribusi data berat bayi adalah normal (nilai-p = 0.200).

8.3.2. Uji Korelasi

Setelah dilakukan uji normalitas, kita akan menguji apakah ada korelasi antara berat badan ibu sebelum hamil (*bbibu_1*) dengan berat badan bayi (*bbayi*) yang akan dilahirkannya kelak dengan prosedur sbb:

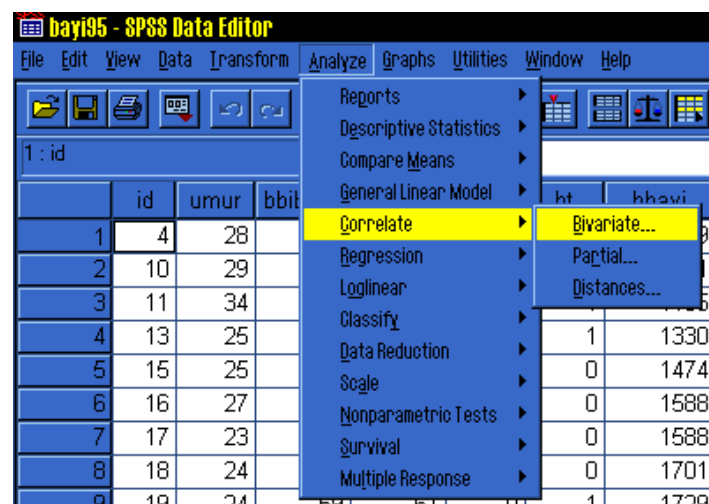
1. Bukalah file BAYI95.SAV, sehingga data tampak di Data editor window.
2. Dari menu utama, pilihlah:

Analyze <

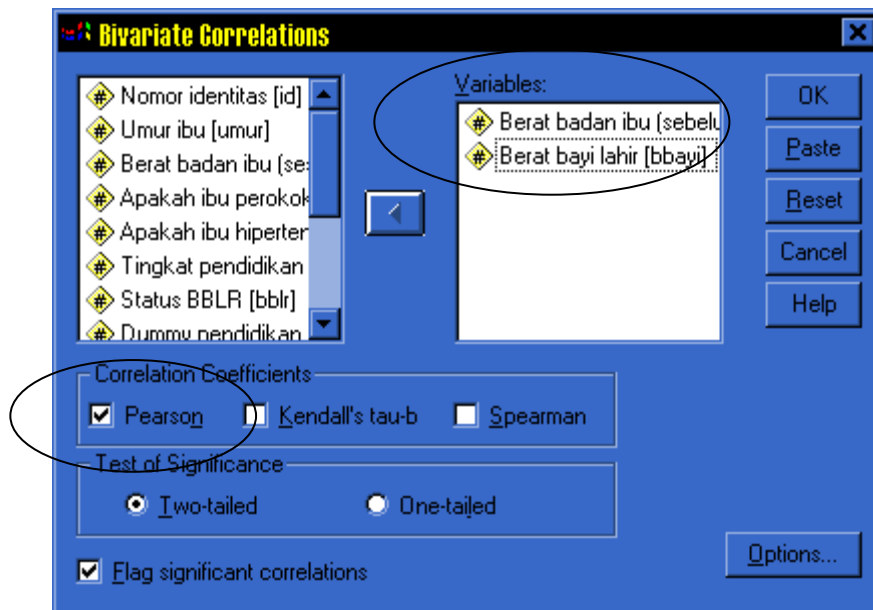
Correlate <

Bivariate...

Seperti gambar berikut:



11. Pilih variabel *bbibu_1* dan *bbayi*, kemudian masukkan ke kotak **Variables**
12. Pada Correlation Coefficient, atifkan Pearson, kemudian OK, dan hasilnya dapat dilihat sbb:



Correlations

		Berat badan ibu (sebelum hamil)	Berat bayi lahir
Berat badan ibu (sebelum hamil)	Pearson Correlation	1.000	.186*
	Sig. (2-tailed)	.	.011
	N	189	189
Berat bayi lahir	Pearson Correlation	.186*	1.000
	Sig. (2-tailed)	.011	.
	N	189	189

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

Hasil diatas memperlihatkan bahwa koefisien korelasi Pearson antara berat badan ibu sebelum hamil dengan berat bayi lahir adalah 0.186, korelasi itu bermakna secara statistik dengan nilai-p 0.011.

8.4. Aplikasi Regresi Linier (Sederhana)

Setelah dilakukan uji korelasi, kita menyimpulkan korelasi tersebut bermakna secara statistik. Selanjutnya kita akan membuat persamaan garis lurus untuk menggambarkan secara lebih rinci korelasi antara *bbibu* dengan *bbayi* serta dapat digunakan untuk memprediksi berat bayi jika berat ibunya diketahui. Analisa statistik yang kita gunakan adalah regresi linier, dalam

hal ini regresi linier sederhana, dengan prosedur sbb:

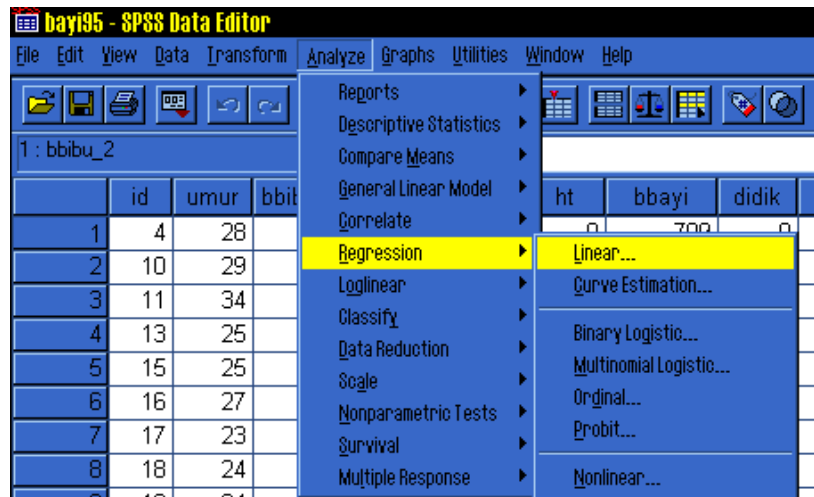
1. Bukalah file BAYI95.SAV, sehingga data tampak di Data editor window.
2. Dari menu utama, pilihlah:

Analyze <

Regressions <

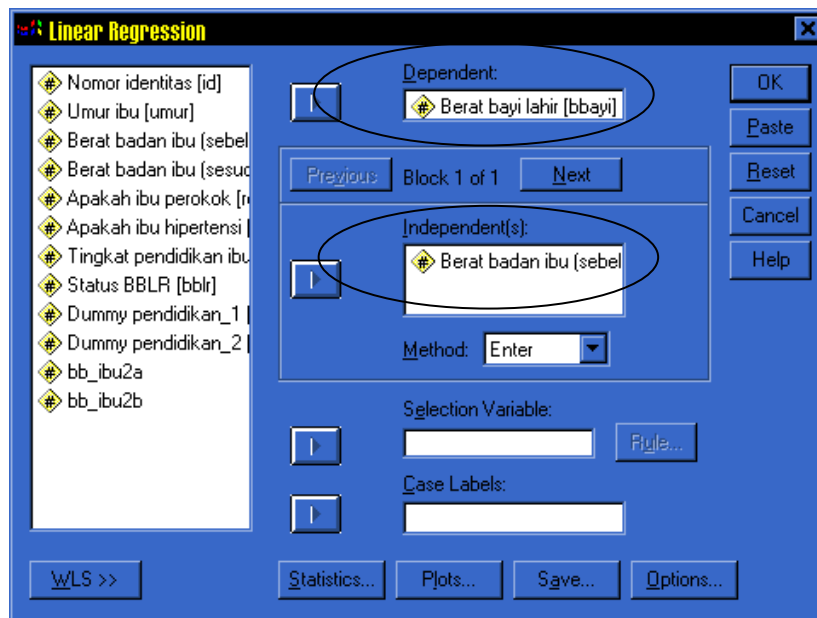
Linier...

Seperti gambar berikut:



13. Klik variabel *bbibu_1*, kemudian masukkan ke kotak **Dependent**

14. Klik variabel *bbayi*, kemudian masukkan ke kotak **Independent(s)**



15. Kemudian klik OK, dan hasilnya sbb:

Model Summary

Model	R	R Square	Adjusted R Square	Std. Error of the Estimate
1	.186 ^a	.034	.029	718.26

a. Predictors: (Constant), Berat badan ibu (sebelum hamil)

Nilai R yang ditampilkan merupakan nilai koefisien korelasi Pearson yang hasilnya sama dengan analisa Korelasi – Bivariat yang dikerjakan sebelumnya yaitu 0.186. R-square merupakan nilai r yang dikuadratkan, yang artinya besarnya variasi pada variabel *bbayi* yang dapat dijelaskan oleh variabel *bbibu_1* (atau oleh persamaan garis regresi yang kita peroleh) adalah 3,4%.

ANOVA^b

Model		Sum of Squares	df	Mean Square	F	Sig.
1	Regression	3444549	1	3444549.214	6.677	.011 ^a
	Residual	96472503	187	515895.740		
	Total	99917053	188			

a. Predictors: (Constant), Berat badan ibu (sebelum hamil)

b. Dependent Variable: Berat bayi lahir

Nilai signifikansi dari ANOVA yang ditampilkan merupakan gambaran apakah model persamaan garis yang kita peroleh sudah bermakna secara statistik. Dengan nilai-p 0.011 bila dibandingkan dengan alpha 0.05 kita simpulkan bahwa persamaan garis yang kita peroleh secara statistik memang bermakna.

Coefficients^a

Model		Unstandardized Coefficients		Standardized Coefficients	t	Sig.
		B	Std. Error	Beta		
1	(Constant)	2370.440	228.282		10.384	.000
	Berat badan ibu (sebelum hamil)	9.834	3.806	.186	2.584	.011

a. Dependent Variable: Berat bayi lahir

Nilai koefisien B yang ditampilkan merupakan gambaran untuk membuat model persamaan garis $y = a + bx$. Nilai B untuk variabel Constant (atau a) adalah 2370.44 dengan nilai-p 0.000, sedangkan nilai B untuk variabel berat badan ibu (atau b) adalah 9.834 dengan nilai-p 0.011. Persamaan garis lurus yang kita dapat adalah:

$$\text{Berat bayi lahir} = 2370.44 + 9.834 (\text{berat ibu})$$

8.5. Penyajian dan Interpretasi Korelasi & Regresi Linier

Setelah dilakukan uji korelasi dan Regresi Linier, kita harus memilih nilai-nilai tertentu untuk disajikan dalam suatu laporan singkat yang dapat dimengerti dengan baik oleh pembacanya, sebagai berikut:

Tabel 1. Analisis Korelasi dan Regresi Linier Berat Ibu sebelum hamil dengan Berat bayi lahir

Variabel	R	R ²	Persamaan garis	Nilai-p
1. Berat ibu sebelum hamil	0.186	0.034	Berat bayi lahir = 2370.44 + 9.834 (berat ibu)	0.011
2. ..				

Hubungan antara berat ibu sebelum hamil dengan berat bayi lahir menunjukkan korelasi yang positif dengan kekuatan/keeratan hubungan yang rendah ($R=0.186$). Artinya semakin tinggi berat ibu sebelum hamil maka semakin tinggi berat bayi yang akan dilahirkannya, setiap kenaikan satu kilogram berat ibu akan dapat meningkatkan 9.384 gram berat bayi.

Namun, variabel berat ibu hanya dapat menjelaskan 3,4% variasi pada variabel berat bayi atau variabel berat ibu kurang dapat menjelaskan variabel berat bayi. Walaupun hubungan ini bermakna secara statistik (nilai- 0.011).

8.6. Memprediksi nilai Y

Dari persamaan garis regresi linier yang didapatkan, kita bisa memperkirakan atau memprediksi nilai y, bila nilai x kita ketahui. Misalnya, jika diketahui berat badan ibu sebelum hamil adalah 70 kg, maka perkiraan berat bayi yang akan dilahirkannya adalah sebagai berikut:

$$\begin{aligned} \text{Berat bayi lahir} &= 2370.44 + 9.834 (\text{berat ibu}) \\ &= 2370.44 + 9.834 (70) \\ &= 2370.44 + 688.38 \end{aligned}$$

= 3058.82 gram

Daftar Pustaka