

BUKU AJAR STATISTIKA



Ig. Dodiet Aditya Setyawan, SKM., MPH.,
Ade Devriany, SKM., M.Kes.
Nuril Huda
Nina Rahmadiliyani, S.Kep., MPH
Ros Endah Happy Patriyani, S.Kp., Ns., M.Kep



BUKU AJAR **STATISTIKA**

Ig. Dodiet Aditya Setyawan, SKM., MPH.,

Ade Devriany, SKM., M.Kes.

Nuril Huda

Nina Rahmadiliyani, S.Kep., MPH

Ros Endah Happy Patriyani, S.Kp., Ns., M.Kep

Endang Caturini Sulustyowati, SKep, Ns., MKep.

BUKU AJAR STATISTIKA

Indramayu © 2021, Penerbit Adab

Penulis :

Ig. Dodiet Aditya Setyawan, SKM., MPH.,

Ade Devriany, SKM., M.Kes.

Nuril Huda

Nina Rahmadiliyani, S.Kep., MPH

Ros Endah Happy Patriyani, S.Kp., Ns., M.Kep

Endang Caturini Sulustyowati, S.Kep, Ns., M.Kep.

Editor : Muhamad Seto

Perancang Sampul : Nurul Musyafak

Layouter : Fitri Yanti

Diterbitkan oleh **Penerbit Adab**

CV. Adanu Abimata

Anggota IKAPI : 354/JBA/2020

Jln. Kristal blok F6 Pabean Udik Indramayu Jawa Barat

Kode Pos 45219 Telp : 081221151025

Surel : penerbitadab@gmail.com

Web : <https://penerbitadab.id>

Buku Ajar | Non Fiksi | R/D

vi + 142 hlm. ; 20,5 x 29 cm

No ISBN : 978-623-.....

Cetakan Pertama, November 2021



Hak Cipta dilindungi undang-undang.

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

All right reserved

KATA PENGANTAR

Segala Puji dan Syukur kami panjatkan selalu kepada Allah SWT, dan Hidayah yang sudah diberikan sehingga kami bisa menyelesaikan buku panduan yang berjudul "Buku Ajar Statistika " dengan tepat waktu. Tujuan dari penulisan buku ini tidak lain adalah untuk membantu para mahasiswa dalam memahami mengenai Statistika contoh materi distribusi frekuensi, ukuran pemusatan, dispersi, probabilitas, populasi dan sampel , Teknik pemilihan analisis statistika, Analisis statistik parametrik, analisis statistic non parametrik.

Buku ini juga akan memberikan informasi secara lengkap mengenai Statistika dari berbagai penulis atau peneliti yang namanya sudah terkenal dimana-mana.

Kami sadar bahwa penulisan buku ini bukan merupakan buah hasil kerja keras kami sendiri. Ada banyak pihak yang sudah berjasa dalam membantu kami di dalam menyelesaikan buku ini, seperti pembuatan cover, editing dan lain-lain. Maka dari itu, kami mengucapkan banyak terimakasih kepada semua pihak yang telah membantu memberikan wawasan dan bimbingan kepada kami sebelum maupun ketika menulis buku ajar Statistika ini.

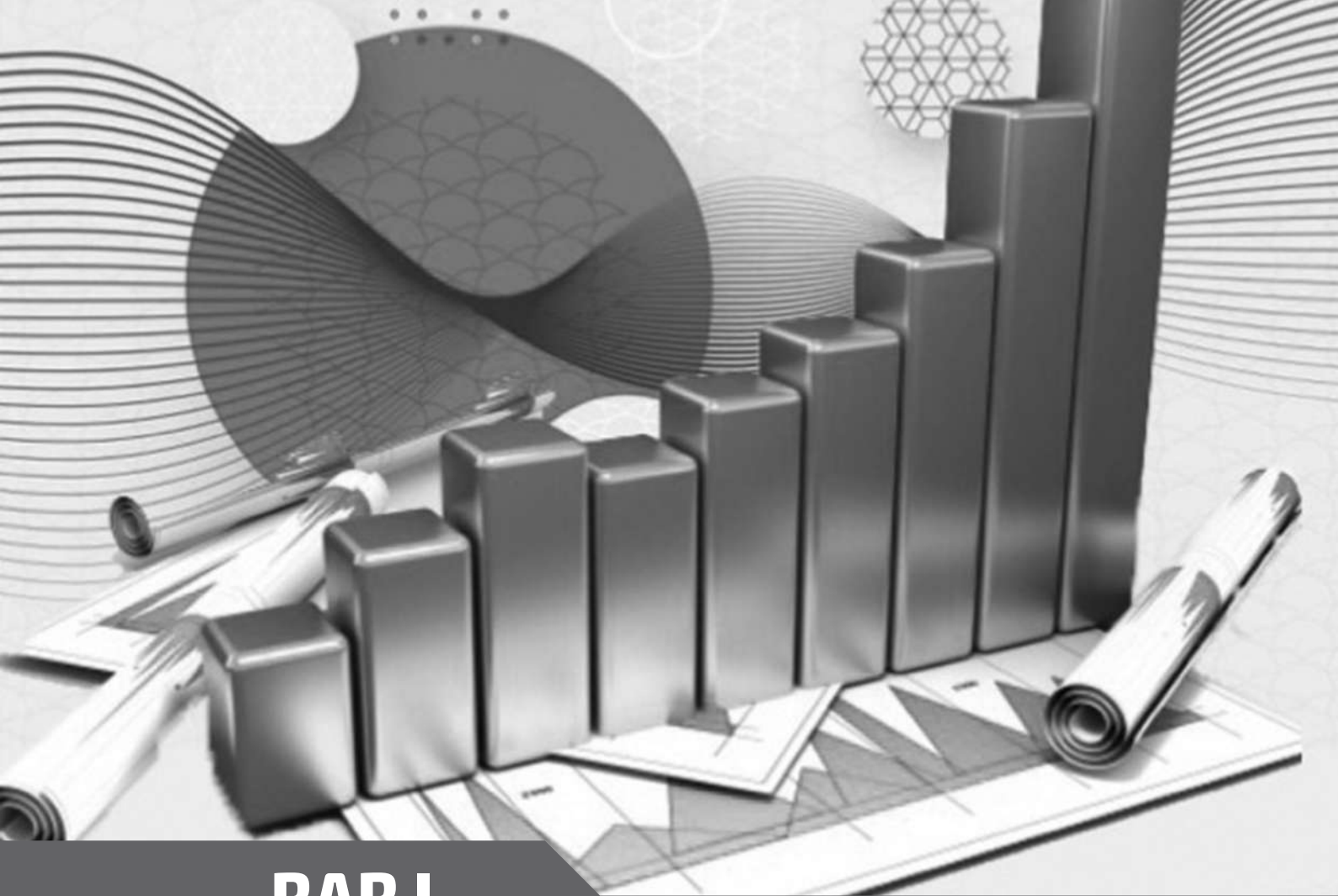
Kami juga sadar bahwa buku yang kami buat masih tidak belum bisa dikatakan sempurna. Maka dari itu, kami meminta dukungan dan masukan dari para pembaca, agar kedepannya kami bisa lebih baik lagi di dalam menulis sebuah buku.

Penulis

DAFTAR ISI

KATA PENGANTAR	iii
DAFTAR ISI	v
BAB I DISTRIBUSI FREKUENSI	1
A. Tujuan Pembelajaran.....	2
B. Materi	2
C. Rangkuman	8
D. Tugas.....	9
E. Referensi	9
BAB II UKURAN PEMUSATAN.....	11
A. Tujuan Pembelajaran.....	12
B. Materi	12
C. Rangkuman	23
D. Tugas.....	23
E. Referensi	24
BAB III DISPERSI DATA.....	25
A. Tujuan Pembelajaran.....	26
B. Materi	26
C. Rangkuman	33
D. Tugas.....	34
E. Referensi	34
BAB IV PROBABILITAS.....	35
A. Tujuan pembelajaran :	36
B. Materi	36
C. Rangkuman	45
D. Tugas.....	45
E. Referensi	46

BAB V	POPULASI DAN SAMPEL	47
A.	Tujuan Pembelajaran.....	48
B.	MATERI	48
C.	Rangkuman	57
D.	Tugas.....	58
E.	Referensi	60
BAB VI	TEKNIK PEMILIHAN STATISTIK.....	61
A.	Tujuan Pembelajaran.....	62
B.	Materi	62
C.	TUGAS	71
D.	KESIMPULAN	71
E.	DAFTAR PUSTAKA	71
BAB VII	UJI STATISTIK PARAMETRIK	73
A.	Tujuan Pembelajaran:	74
B.	Materi	74
C.	Rangkuman	101
D.	Tugas.....	101
E.	Referensi	102
BAB VIII	ANALISIS STATISTIK NON PARAMETRIK.....	103
A.	Tujuan pembelajaran :	104
B.	Materi	104
C.	Rangkuman	136
D.	Tugas.....	137
C.	Referensi	138
RIWAYAT HIDUP	PENULIS	139



BAB I

DISTRIBUSI FREKUENSI

Dodiet Aditya Setyawan, SKM., MPH.

A. Tujuan Pembelajaran

Setelah mempelajari materi ini mahasiswa diharapkan mampu:

1. Memahami pengertian Distribusi Frekuensi
2. Memahami pedoman umum membuat Tabel Distribusi Frekuensi.
3. Memahami Teknik Penyusunan Tabel Distribusi Frekuensi
4. Memahami macam-macam bentuk Tabel Distribusi Frekuensi
5. Membuat Tabel Distribusi Frekuensi dengan benar.

B. Materi

1. Pengertian

Distribusi Frekuensi merupakan teknik penyusunan data dalam bentuk kelompok mulai dari yang terkecil sampai yang terbesar berdasarkan kelas-kelas interval dan kategori tertentu (Hasibuan,dkk.2009). Manfaat dari penyajian data dalam bentuk Distribusi Frekuensi ini adalah untuk menyederhanakan teknik penyajian data sehingga menjadi lebih mudah untuk dibaca dan dipahami sebagai bahan informasi. Tabel Distribusi Frekuensi disusun apabila jumlah data yang akan disajikan cukup banyak, sehingga bila data tersebut disajikan dengan menggunakan tabel biasa menjadi tidak efektif dan efisien serta kurang komunikatif. Disamping itu tabel distribusi frekuensi juga dapat dibuat untuk bahan dalam melakukan uji normalitas data. Beberapa hal yang perlu diperhatikan dalam Tabel Distribusi frekuensi adalah:

- a. Tabel distribusi frekuensi mempunyai beberapa kelas yang jumlahnya disesuaikan dengan banyaknya data.
- b. Setiap kelas pada tabel distribusi frekuensi mempunyai kelas interval. Setiap kelas interval mempunyai panjang kelas yaitu interval kelas bawah sampai dengan interval kelas atas. Jadi panjang kelas interval merupakan jarak antara nilai batas bawah dengan batas atas pada setiap kelas.
- c. Setiap kelas interval mempunyai jumlah atau yang sering disebut frekuensi.

2. Pedoman Umum Membuat Tabel Distribusi Frekuensi

Membuat tabel distribusi frekuensi diawali dengan menentukan kelas interval dari sejumlah data yang sudah dikumpulkan atau di tabulasi. Berikut ini adalah bagian-bagian yang harus dibuat terlebih dahulu dalam sebuah tabel distribusi frekuensi.

- a. Kelas Interval/Jumlah Kelas Interval (*Class*)

Kelas interval merupakan kelompok-kelompok nilai atau variabel. Jumlah kelas menunjukkan jumlah kelompok nilai/variabel dari data yang diobservasi. Dalam menentukan Jumlah Kelas Interval terdapat 3 pedoman sebagai berikut:

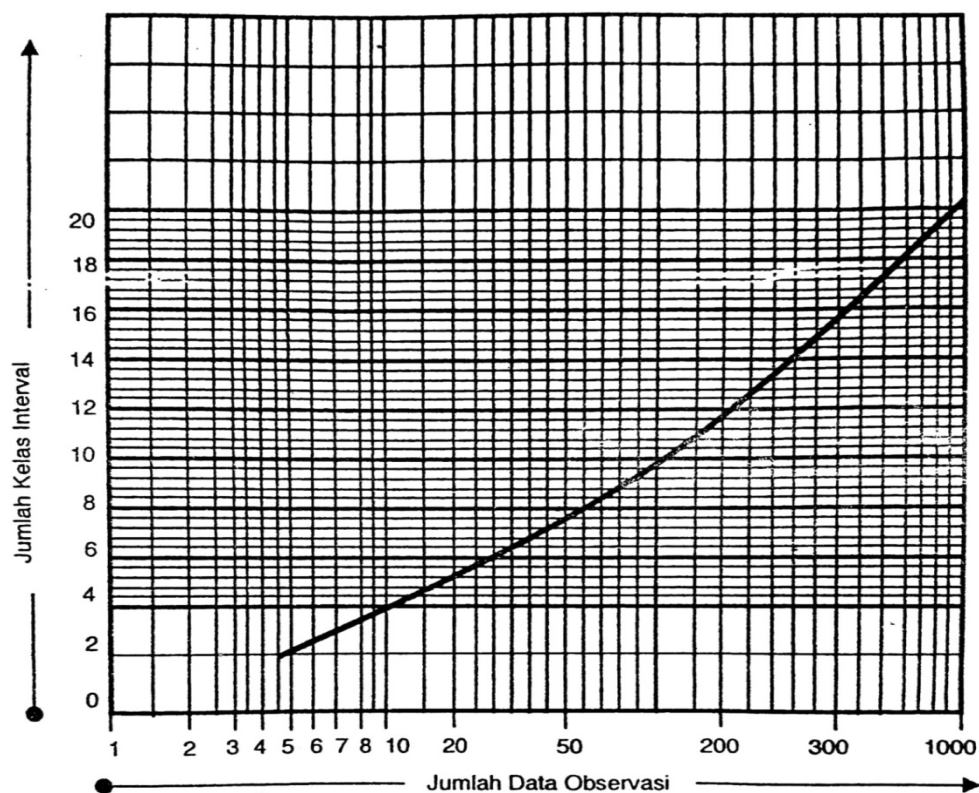
- 1) Ditentukan berdasarkan Pengalaman

Pada umumnya jumlah kelas interval yang dipergunakan dalam penyusunan Tabel Distribusi Frekuensi berkisar antara 6-15 kelas. Makin banyak data, maka makin banyak pula jumlah kelas intervalnya, tetapi jumlah yang paling banyak atau maksimal adalah 15 kelas interval dalam satu tabel distribusi frekuensi, sehingga tabel distribusi frekuensi tidak terlalu panjang.

2) Ditentukan dengan Membaca Grafik 'Jumlah Interval Kelas'

Dengan menggunakan Grafik yang menunjukkan hubungan antara banyaknya data (n) dengan jumlah kelas interval yang diperlukan, maka penentuan jumlah kelas interval akan lebih cepat. Dimana dalam grafik tersebut, Garis Vertikal menunjukkan Jumlah Kelas Interval dan Garis Horizontal menunjukkan Jumlah Data Observasi. Misalnya, bila jumlah data yang diobservasi 50, maka berdasarkan Tabel, Jumlah Kelas Intervanya kurang lebih 8.

Selanjutnya apabila jumlah data yang diobservasi sebanyak 200, maka jumlah kelas intervalnya kurang lebih 12, dan seterusnya. Contoh untuk menentukan jumlah kelas interval dengan cara tersebut dapat dilihat pada gambar grafik berikut ini:



Gambar 1.1 Grafik Untuk Menentukan Jumlah Kelas Interval

(Sugiyono, 2015)

3) Ditentukan dengan Rumus Sturges

Jumlah Interval Kelas Interval juga dapat dihitung dengan menggunakan rumus Sturges sebagai berikut:

$$K = 1 + 3,3 \text{ Log. } N$$

Dimana :

K = Jumlah Kelas Interval

n = Jumlah Data Observasi

Log = Logaritma

Contoh:

Misalnya Jumlah Data yang diobservasi sebanyak 150, maka jumlah Kelas Intervalnya dapat ditentukan dengan cara sebagai berikut:

$$K = 1 + 3,3 \cdot \log 150$$

$$K = 1 + 3,3 \cdot 2,17$$

$$K = 1 + 7,161$$

$$K = 8,161 \rightarrow \text{Dibulatkan menjadi } 8$$

Jadi dari sejumlah 150 data yang dikumpulkan dan diobservasi, maka jumlah kelas interval yang diperlukan adalah 8 kelas.

Berikut ini adalah gambaran tentang Kelas Interval (Jumlah Kelas Interval) dalam Tabel Distribusi Frekuensi:

	Nilai UAS Statistik	Frekuensi (f)
Jumlah Kelas Interval	50 – 55	3
	56 – 61	7
	62 – 67	2
	68 – 73	8
	74 – 79	4
	80 – 85	6
	Jumlah (n)	30

b. Batas Kelas (*Class Limits*)

Batas Kelas merupakan nilai-nilai yang membatasi antara kelas yang satu dengan kelas berikutnya. Batas Kelas terdiri atas 2 macam, yaitu:

1) Batas Kelas Bawah (*Lower Class Limits*)

Batas Kelas Bawah adalah nilai atau angka yang terdapat pada bagian sebelah kiri dari setiap kelas.

2) Batas Kelas Atas (*Upper Class Limits*)

Batas Kelas Atas adalah nilai atau angka yang berada pada bagian sebelah kanan dari setiap kelas.

Gambaran tentang yang dimaksud Batas Kelas Atas dan Batas Kelas Bawah pada Tabel Distribusi Frekuensi adalah sebagai berikut:

Nilai UAS Statistik	Frekuensi (f)
50 – 55	3
56 – 61	7
62 – 67	2
68 – 73	8
74 – 79	4
80 – 85	6
Jumlah (n)	30

c. Rentang Data (*Range*)

Rentang Data (*Range*) adalah selisih antara data tertinggi dengan data terendah (Data terbesar dikurangi Data terkecil). Pada Contoh Tabel di atas, dimana data terendah adalah 50 dan data tertinggi adalah 85. Sehingga untuk menentukan Rentang Data (*Range*) adalah $85 - 50 = 35$. Jadi rentang data pada tabel tersebut adalah 35.

d. Panjang Interval Kelas

Panjang Interval Kelas atau disebut juga Panjang Kelas atau Interval Size merupakan jarak antara tepi kelas atas dengan tepi kelas bawah. Dapat dihitung dengan cara membagi Rentang Data dengan Jumlah Kelas. Pada contoh Tabel di atas, sudah diketahui bahwa rentang data adalah 35 dan jumlah kelas adalah 6, maka Panjang Interval kelasnya dapat dihitung dengan cara Rentang dibagi Jumlah kelas, sehingga didapatkan $35/6 = 5,8$ (Dibulatkan = 6). Jadi Panjang Kelas interval adalah 6 pada setiap kelas.

e. Frekuensi Kelas (*Class Frequency*)

Frekuensi kelas merupakan banyaknya jumlah data yang terdapat pada kelas tertentu. Misalnya pada contoh tabel di atas, Frekuensi pada kelas interval 50-55 adalah 3; frekuensi pada kelas interval 56-61 adalah 7, frekuensi pada kelas interval 80-85 adalah 6, dan seterusnya.

3. Teknik Penyusunan Tabel Distribusi Frekuensi

Langkah-langkah untuk membuat sebuah Tabel Distribusi Frekuensi secara sistematis dapat dilakukan sebagai berikut:

- Mengurutkan semua data yang diobservasi mulai dari yang terkecil sampai yang terbesar.
- Menghitung Rentang/Range (R), yaitu Data terbesar dikurangi dengan Data terkecil.
- Menentukan jumlah kelas, dengan menggunakan rumus Sturges:

$$K = 1 + 3,3 \cdot \log n$$
- Menghitung Panjang Kelas atau Interval, dengan rumus:
 Panjang Kelas (P) = Rentang (R) : Jumlah Kelas
- Membuat tabel distribusi frekuensi sementara yang terdiri atas kolom Interval Kelas, Tally, dan Frekuensi.

Kelas Interval	Tally	Frekuensi (f)

- Menghitung jumlah Frekuensi dengan Tally atau melidi dalam Kolom Tally sesuai dengan banyaknya data.

Kelas Interval	Tally	Frekuensi (f)
50 – 55	III	3
56 - 61	HHH II	7
		dan seterusnya

- g. Setelah jumlah keseluruhan Frekuensi ditemukan, kemudian kolom Tally dihilangkan dalam Penyajian Data sehingga terbentuk Tabel Distribusi Frekuensi yang dimaksud.

Kelas Interval	Frekuensi (f)
50- 55	3
56 – 61	7
	dan seterusnya

4. Macam-Macam Tabel Distribusi Frekuensi

Secara Umum, Tabel Distribusi Frekuensi dapat dikategorikan menjadi beberapa macam, yaitu:

- a. Tabel Distribusi Frekuensi Data Tunggal

Distribusi Frekuensi Data Tunggal yaitu jenis tabel distribusi frekuensi yang menyajikan frekuensi dari data tunggal yang berdiri sendiri atau tidak dikelompokkan.

Contoh:

Tabel Distribusi Frekuensi Nilai UAS Statistik Semester II adalah sebagai berikut:

Nilai		Frekuensi (f)
Data Tunggal	50	5
	60	10
	70	15
	80	10
	90	5
	100	5
Jumlah (n)		50

- b. Tabel Distribusi Frekuensi Data Kelompok

Distribusi Frekuensi Data Kelompok merupakan tabel distribusi frekuensi yang menyajikan frekuensi dari data yang dikelompokkan.

Contoh:

Tabel Distribusi Frekuensi Umur Mahasiswa Poltekkes Surakarta

Nilai		Frekuensi (f)
Data Dikelompokkan	22 – 27	40
	28 – 33	20
	34 – 39	5
	dst	Dst
Jumlah (n)		

c. Tabel Distribusi Frekuensi Kumulatif

Distribusi Frekuensi Kumulatif merupakan tabel statistik yang menyajikan frekuensi dari data yang dihitung dengan ditambah-tambahkan baik dari atas ke bawah maupun dari bawah ke atas. Dibedakan menjadi 2, yaitu:

- 1) Frekuensi Kumulatif Atas atau $fk_{(a)}$ yaitu: Frekuensi yang angka-angkanya ditambahkan dari Bawah ke Atas.
- 2) Frekuensi Kumulatif Bawah atau $fk_{(b)}$ yaitu: Frekuensi yang angka-angkanya ditambahkan dari Atas ke Bawah

Contoh-1:

Tabel Distribusi Frekuensi Kumulatif Nilai Statistik

NILAI	f	fk(b)	fk(a)
44	2	40 (=n)	2
69	15	38	17
76	14	23	31
79	6	9	37
84	3	3	40 (=n)
(n) = 40			

Dari contoh tabel di atas, cara mendapatkan $fk(b)$ dapat dijelaskan sebagai berikut:

- $fk(b) = 40$, didapatkan dari penjumlahan semua frekuensi (f): $2+15+14+6+3 = 40$
- $fk(b) = 38$, didapatkan dari penjumlahan frekuensi (f) ke-2 - ke-5: $15+14+6+3 = 38$
- $fk(b) = 23$, didapatkan dari penjumlahan frekuensi (f) ke-3 - ke-5: $14+6+3 = 23$
- $fk(b) = 9$, didapatkan dari penjumlahan frekuensi (f) ke-4 - ke-5: $6+3 = 9$
- $fk(b) = 3$, didapatkan dari penjumlahan frekuensi (f) ke-5: yaitu $= 3$.

Sedangkan cara untuk mendapatkan $fk(a)$ dapat dijelaskan sebagai berikut:

- $fk(a) = 40$, didapatkan dari penjumlahan semua frekuensi (f): $3+6+14+15+2 = 40$
- $fk(a) = 37$, didapatkan dari penjumlahan frekuensi (f) ke-4- ke-1: $6+14+15+2 = 37$
- $fk(a) = 31$, didapatkan dari penjumlahan frekuensi (f) ke-3- ke-1: $14+15+2 = 31$
- $fk(a) = 17$, didapatkan dari penjumlahan frekuensi (f) ke-2- ke-1: $15+2 = 17$
- $fk(a) = 2$, didapatkan dari penjumlahan frekuensi (f) ke-1: yaitu $= 2$

Selanjutnya sebagai latihan, lakukan identifikasi terhadap contoh 2 berikut ini, apakah penentuan $fk(b)$ dan $fk(a)$ pada contoh tabel di bawah ini sudah benar atau salah?

Contoh-2:

Tabel Distribusi Frekuensi Kumulatif Umur Mahasiswa

NILAI	f	fk(b)	fk(a)
22-27	15	60	15
28-33	29	45	44
34-39	16	16	60
(n) = 60			

d. Tabel Distribusi Frekuensi Relatif (Tabel Persentase)

Tabel Distribusi Frekuensi Relatif adalah jenis tabel statistik yang di dalamnya menyajikan frekuensi dalam bentuk angka persentase (p). Nilai Persentase dihitung dengan menggunakan rumus sebagai berikut:

$$P (\%) =$$

Contoh:

Tabel Distribusi Frekuensi Relatif Umur Mahasiswa

Umur	Frekuensi (f)	Persentase (%)
22-27	15	25
28-33	29	48
34-39	16	27
(n) = 60		100

Berdasarkan contoh tabel diatas, maka cara mendapatkan persentase tersebut dapat dijelaskan sebagai berikut:

- Pesentase 25 (25%) didapatkan dengan cara $(15/60) \times 100 = 25\%$
- Pesentase 48 (48%) didapatkan dengan cara $(29/60) \times 100 = 48\%$
- Pesentase 27 (27%) didapatkan dengan cara $(16/60) \times 100 = 27\%$

C. Rangkuman

Distribusi Frekuensi adalah teknik penyusunan data dalam bentuk kelompok mulai dari yang terkecil sampai yang terbesar berdasarkan kelas-kelas interval dan kategori tertentu. Penggunaan Tabel Distribusi Frekuensi ditujukan untuk menyederhanakan teknik penyajian data sehingga menjadi lebih mudah untuk dibaca dan dipahami sebagai bahan informasi.

Membuat Tabel Distribusi Frekuensi diawali dengan menentukan kelas interval dari sejumlah data yang sudah dikumpulkan atau di tabulasi, selanjutnya berturut-turut harus menentukan batas kelas, rentang data (Range), panjang interval kelas dan frekuensi kelas.

Kelas Interval merupakan kelompok-kelompok nilai atau variabel. Jumlah kelas menunjukkan jumlah kelompok nilai/variabel dari data yang diobservasi.

Batas Kelas merupakan nilai-nilai yang membatasi antara kelas yang satu dengan kelas berikutnya.

Rentang Data (Range) adalah selisih antara data tertinggi dengan data terendah (Data terbesar dikurangi Data terkecil).

Panjang Interval Kelas atau disebut juga *Panjang Kelas* atau *Interval Size* merupakan jarak antara tepi kelas atas dengan tepi kelas bawah.

Frekuensi Kelas merupakan banyaknya jumlah data yang terdapat pada kelas tertentu.

Pada dasarnya Tabel Distribusi Frekuensi dapat dikategorikan menjadi 4 macam, yaitu Tabel Distribusi Frekuensi Data Tunggal, Tabel Distribusi Frekuensi Data Kelompok, Tabel Distribusi Frekuensi Kumulatif dan Tabel Distribusi Frekuensi Relatif (Persentase).

D. Tugas

Untuk meningkatkan pemahaman tentang Distribusi Frekuensi, silahkan mengerjakan latihan berikut ini:

Diketahui Nilai hasil Ujian Statistika pada sejumlah 50 mahasiswa diperoleh data nilai dari masing-masing mahasiswa sebagai berikut:

NO	NILAI	NO	NILAI	NO	NILAI	NO	NILAI	NO	NILAI
1	60	11	66	21	77	31	71	41	75
2	71	12	66	22	80	32	72	42	75
3	63	13	67	23	80	33	72	43	75
4	70	14	67	24	80	34	72	44	75
5	80	15	67	25	80	35	72	45	75
6	70	16	68	26	73	36	83	46	75
7	81	17	76	27	73	37	84	47	75
8	81	18	76	28	74	38	84	48	75
9	74	19	77	29	74	39	84	49	78
10	74	20	77	30	74	40	84	50	78

Berdasarkan data nilai yang diperoleh tersebut, maka buatlah Tabel Distribusi Frekuensi sesuai dengan langkah-langkah pembuatan Tabel Distribusi Frekuensi secara berurutan/ sistematis !

E. Referensi

Amin.I., Aswin.A., Fajar.I., Isnaeni, Iwan.S., Pudjirahaju.A., Sunindya.R.. 2009. *Statistika untuk Praktisi Kesehatan*. Yogyakarta. Graha Ilmu

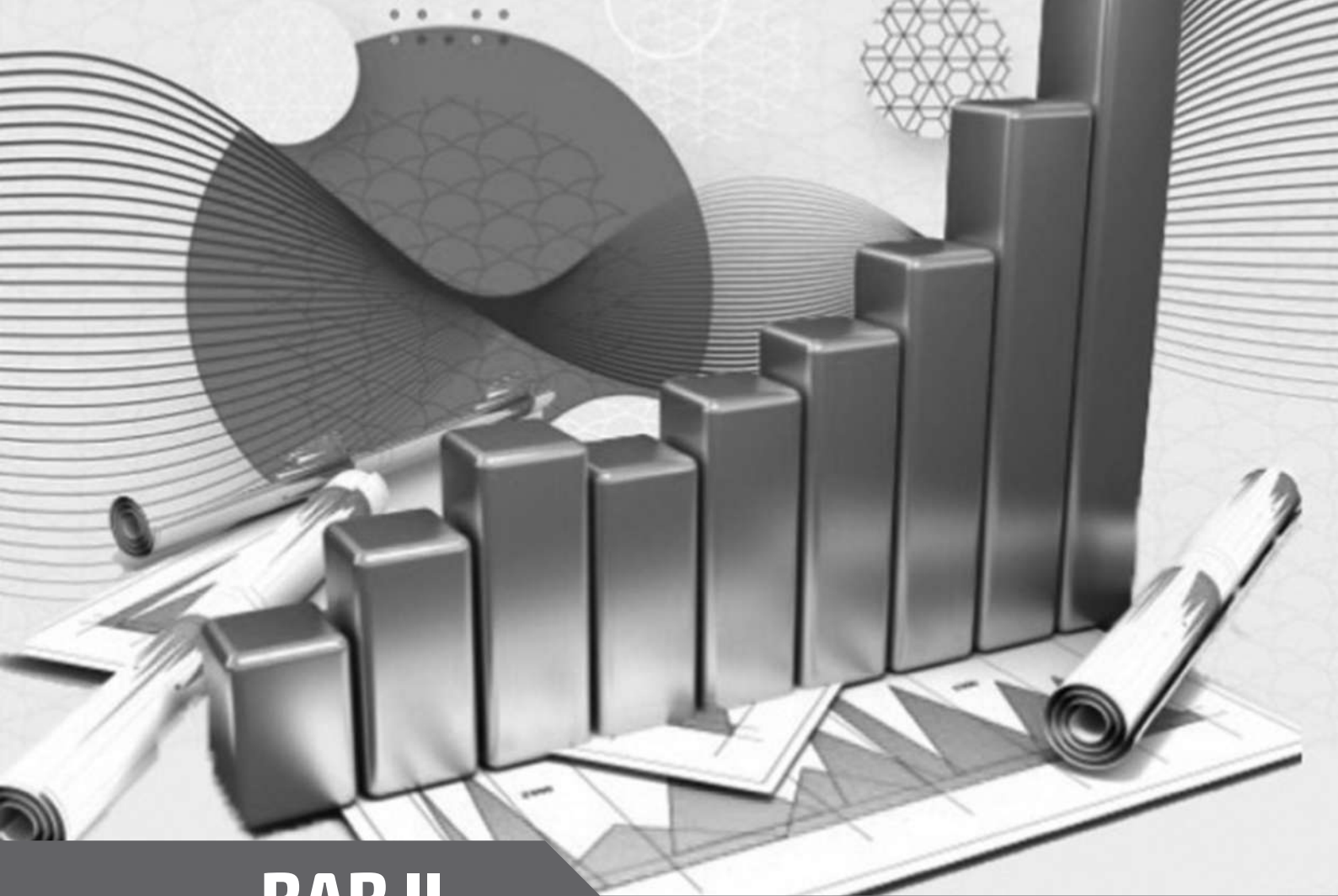
Hadi.S. 2002. *Statistik*. Jilid 2. Yogyakarta. Andi Offset.

Hasibuan.A.A.,Supardi, Syah.D. 2009. *Pengantar Statistik Pendidikan*. Jakarta. Gaung Persada Press.

Machfoedz, I., 2016. *Bio Statistika*. Edisi Revisi. Yogyakarta. Fitramaya.

Riduwan.2010. *Dasar-dasar Statistika*. Bandung. Alfabeta.

- Riwidikdo, H., 2013. *Statistik Kesehatan: Belajar mudah teknik analisis data dalam Penelitian Kesehatan (Plus Aplikasi Software SPSS)*. Yogyakarta. Mitra Cendikia Press.
- Riwidikdo, H., 2012. *Statistik Kesehatan*. Yogyakarta. Nuha Medika
- Sarwono, J. 2006. *Analisis Data Penelitian Menggunakan SPSS*. Yogyakarta. Andi Offset.
- Santjaka, A. 2011. *Statistik Untuk Penelitian Kesehatan: Deskriptif, Inferensial, Parametrik dan Non Parametrik*. Yogyakarta. Nuha Medika.
- Siswandari. 2009. *Statistika (Komputer Based)*. Surakarta. LPP UNS dan UNS Press
- Stang. 2018. *Cara Praktis Penentuan Uji Statistik dalam Penelitian Kesehatan dan Kedokteran. Edisi 2*. Jakarta. Mitra Wacana Media.
- Sugiyono. 2015. *Statistika untuk Penelitian*. Bandung. Alfabeta



BAB II

UKURAN PEMUSATAN

Dodiet Aditya Setyawan, SKM., MPH.

A. Tujuan Pembelajaran

Setelah mempelajari materi ini mahasiswa diharapkan mampu:

1. Memahami pengertian ukuran pemusatan
2. Memahami macam-macam bentuk ukuran pemusatan
3. Melakukan penghitungan Mean, Median, Modus pada data tunggal dan data kelompok.
4. Mengaplikasikan teknik penentuan Mean, Median, Modus berbasis komputer dengan menggunakan aplikasi SPSS.

B. Materi

1. Pengertian

Ukuran pemusatan merupakan salah satu bentuk analisis statistik Deskriptif untuk Data atau Variabel yang berskala Numerik (Interval atau Rasio). Ukuran pemusatan yang juga sering dikenal dengan Ukuran Nilai Pusat atau *Measures of Central Tendency* atau yang biasa disebut juga sebagai Ukuran Rata-Rata adalah suatu nilai tunggal yang merepresentasikan gambaran secara umum tentang keadaan dari nilai tersebut yang terdapat dalam suatu data. Ukuran pemusatan atau nilai pusat dapat mewakili data secara keseluruhan dan merupakan rata-rata (*average*), karena nilai rata-ratanya dihitung dari keseluruhan nilai yang terdapat dalam data tersebut. Nilai rata-rata ini sering disebut juga dengan Tendensi Pusat. Artinya jika nilai data-data yang ada diurutkan besarnya kemudian dimasukkan nilai rata-rata ke dalamnya, maka nilai rata-rata tersebut memiliki Kecendrungan (*Tendensi*) terletak di *tengah-tengah* atau *pada pusat* di antara data-data yang ada.

2. Macam-Macam Ukuran Nilai Tengah

Beberapa macam Ukuran Pemusatan atau Tendensi Sentral antara lain:

- a. *Rata-Rata Hitung atau MEAN atau Arithmetic Mean ($\bar{x}$)*
Rata-rata hitung atau MEAN merupakan nilai rata-rata yang dilambangkan dengan $\bar{x}$ (baca X Bar). Nilai rata-rata (Mean) merupakan nilai rata-rata dari keseluruhan data yang ada dan merupakan ukuran yang paling stabil untuk data kontinyu atau data numerik.
- b. *Rata-Rata Pertengahan atau MEDIAN atau MEDIUM (Me) atau (Md).*
Median/ Medium merupakan nilai tengah atau nilai yang terletak di tengah-tengah dari keseluruhan data setelah diurutkan dari yang terkecil hingga terbesar atau sebaliknya. Median membagi Data dalam 2 bagian yang sama.
- c. *Modus atau Mode (Mo)*
Modus merupakan nilai yang paling sering muncul dalam suatu kelompok atau sekumpulan data.

3. Penghitungan Ukuran Pemusatan.

a. MEAN

Untuk menghitung Nilai Rata-rata atau Mean dapat dilakukan dengan cara menjumlah keseluruhan angka atau data yang ada dibagi banyaknya angka atau

data tersebut. Untuk menentukan atau menghitung Mean dibedakan menjadi 2 macam, yaitu:

1) Mean pada Data Tunggal

Penghitungan nilai Mean pada Data Tunggal ini dapat dibedakan lagi menjadi 2 macam, yaitu Data Tunggal dengan Frekuensi =1 (Satu) dan Data Tunggal dengan Frekuensi >1 (Lebih dari Satu).

(a) Mean pada Data Tunggal dengan Frekuensi = 1

Untuk menghitung nilai Rata-rata atau MEAN pada Data Tunggal dengan Frekuensi yang sama dengan 1 (=1) dapat digunakan rumus:

$$\text{Rumus : } \bar{X} = \frac{\sum x}{N}$$

Dimana:

$\sum X$: Jumlah Nilai

N : Banyaknya Frekuensi

: Mean (Rata-Rata)

Contoh:

Misalnya terdapat data tentang Nilai Mata Kuliah Statistika pada 5 mahasiswa didapatkan tabel distribusi frekuensi sebagai berikut:

Nilai Mahasiswa (X)	Frekuensi (f)
60	1
70	1
80	1
90	1
100	1
$\sum X = 400$	$N = 5$

Berdasarkan data pada contoh tabel tersebut, maka Nilai Rata-Rata atau MEAN dari data tersebut adalah: $400/5 = 80$.

(b) Mean pada Data Tunggal dengan Frekuensi >1

Untuk menghitung nilai Rata-rata atau MEAN pada Data Tunggal dengan Frekuensi yang Lebih dari 1 (>1) dapat digunakan rumus sebagai berikut:

$$\text{Rumus : } \bar{X} = \frac{\sum fx}{\sum f}$$

Dimana:

$\sum f$: Jumlah Total Frekuensi

$\sum fx$: Jml. Frekuensi X Nilai ($f \times X$)

Contoh:

Misalnya didapatkan data distribusi tentang Nilai Statistika pada 70 mahasiswa sebagai berikut:

Nilai (X)	Frekuensi (f)	F(X)
5	1	5
6	3	18
7	10	70
8	30	240
9	20	180
10	6	60
	$\Sigma f = 70$	$\Sigma fx = 573$

Berdasarkan contoh tabel di atas, maka nilai rata-rata atau Mean dari data tersebut adalah: $(573/70) = 8,2$.

2) Mean pada Data Kelompok

Untuk dapat menghitung Nilai Rata-rata atau Mean pada Data Kelompok (Interval) harus terlebih dahulu menghitung atau mengetahui Nilai atau Titik Tengah atau MIDPOINT dari Data Kelompok tersebut. Adapun untuk menghitung atau menentukan MIDPOINT dari Data Kelompok pada Tiap-tiap Kelompok adalah $\frac{1}{2}$ (Batas Kelas Atas + Batas Kelas Bawah). Selanjutnya untuk menghitung nilai Rata-rata atau Mean pada Data Kelompok dapat digunakan rumus:

$$\text{Rumus : } \bar{X} = \frac{\Sigma fx}{\Sigma f}$$

$$\text{MIDPOINT (X)} = \frac{(\text{Nilai Batas Atas} + \text{Nilai Batas Bawah})}{2}$$

Keterangan:

Σf : Jumlah Total Frekuensi

X : Nilai MIDPOINT tiap-tiap Interval/ Kelompok

Σfx : Jml. Perkalian MIDPOINT dgn Frekuensi

$\bar{x}$: Mean

Sebagai contoh, misalnya didapatkan distribusi data tentang hasil ujian semester mata kuliah statistika pada 100 mahasiswa sebagai berikut:

Interval Nilai	Frekuensi (f)	Mid Point (X)	fX
34-38	1	36	36
39-43	4	41	164
44-48	5	46	230
49-53	7	51	357
54-58	10	56	560
59-63	25	61	1525
64-68	20	66	1320
69-73	10	71	710
74-78	9	76	684
79-83	9	81	729
	$\Sigma f = 100$		$\Sigma fX = 6315$

Berdasarkan contoh data pada tabel di atas, maka nilai rata-rata atau Mean dari data tersebut adalah: $(\sum fX/\sum f) = 6315/100 = 63,15$. Jadi nilai rata-rata dari data tersebut adalah 63,15.

b. MEDIAN

Median atau Medium merupakan nilai tengah atau nilai yang terletak di tengah-tengah dari keseluruhan data setelah diurutkan dari yang terkecil hingga terbesar atau sebaliknya. Median merupakan angka yang membagi Data dalam 2 bagian yang sama besarnya. Terdapat 2 macam penghitungan Median, yaitu:

1) MEDIAN pada Data Tunggal

Menentukan Median pada Data Tunggal juga dapat dibedakan menjadi 2 macam lagi, yaitu Median Data Tunggal dengan Jumlah Data Ganjil dan Median Data Tunggal dengan Jumlah Data Genap.

a) Median Data Tunggal dengan Jumlah Data Ganjil

Untuk menentukan Median pada Data Tunggal dengan Jumlah Data Ganjil dapat dilakukan dengan rumus berikut:

$$Me = X_{\frac{n+1}{2}} \text{ (baris ke)} \quad N = \text{Jml Data} = \sum f$$

Contoh:

X	F
30	1
40	1
50	1
60	1
70	1
80	1
90	1
Total	$\sum f = 7$

Berdasarkan Data tersebut, maka MEDIAN dari Data Tunggal dengan Jumlah Data Ganjil di atas dapat ditentukan dengan cara sebagai berikut:

$$\begin{aligned} X_{(n+1)/2} &= (7+1)/2 \\ &= 8/2 \\ &= 4. \end{aligned}$$

Jadi MEDIAN dari Data tersebut terletak pada BARIS ke-4, yaitu **60** ($Me = 60$).

b) Median Data Tunggal dengan Jumlah Data Genap

Untuk menentukan Median pada Data Tunggal dengan Jumlah Data Genap dapat dilakukan dengan rumus berikut:

$$Me = \frac{X_{\frac{n}{2}} + X_{\frac{n+2}{2}}}{2}$$

Keterangan:

Me : Median

X : Baris ke-

$n = \sum f$: Jumlah Data

Contoh:

X	F
30	1
40	1
50	1
60	1
70	1
80	1
90	1
100	1
Total	$\sum f = 8$

Berdasarkan Data tersebut, maka MEDIAN dari data di atas dapat ditentukan dengan cara sebagai berikut:

$$\begin{aligned}
 & (X_{(n/2)} + X_{(n+2)/2})/2 \\
 &= (X_{8/2} + X_{(8+2)/2})/2 \\
 &= (X_4 + X_5)/2 \\
 &= (60 + 70)/2 \\
 &= 130/2 \\
 &= 65
 \end{aligned}$$

Jadi MEDIAN dari Data tersebut adalah **65** ($Me = 65$).

2) MEDIAN pada Data Kelompok

Untuk menentukan Median pada Data Kelompok dapat dilakukan dengan rumus berikut:

$$Me = b + p \left(\frac{\frac{1}{2}n - F}{f} \right)$$

Keterangan:

B : Batas Bawah Kelas Median (Dimana Median akan terletak)

p : Panjang Kelas Median

n : Ukuran sampel (Banyaknya Data)

f : Frekuensi Kelas Median (Diambil Berdasarkan Frekuensi Terbanyak)

F : Jumlah Semua Frekuensi dengan Tanda Kelas Lebih Kecil dari Kelas Median.

Langkah-langkah yang harus dilakukan untuk menentukan MEDIAN pada Data Kelompok adalah:

- Menentukan FREKUENSI Kelas Median (f) Dengan Cara: Melihat Jumlah FREKUENSI yang Terbanyak/ Terbesar.
- Menentukan Kelompok atau Kelas atau Interval dimana Median akan terletak

- c) Menentukan Panjang Kelas Median (p) *Yaitu Menghitung Jarak antara Batas Bawah sampai Batas Atas Kelas Median.*
- d) Menentukan BATAS BAWAH Kelas Median (b) *Karena ini adalah Data Kelompok, maka cara menentukan BATAS BAWAH Kelas Median adalah dengan Menjumlah Batas Atas Kelas SEBELUM Kelas Median dengan Batas Bawah Kelas Median dibagi Dua.*
- e) Menentukan Jumlah semua Frekuensi dari Kelas atau Interval atau Kelompok yang LEBIH KECIL dari Kelas Median (F) *Yaitu dengan menjumlahkan Semua Frekuensi (f) pada Kelompok atau Kelas atau Interval yang LEBIH KECIL atau SEBELUM Kelas Median.*
- f) Mengidentifikasi Jumlah Data atau Banyaknya Data atau Sampel (n).

Contoh:

Didapatkan distribusi data tentang Nilai Statistika pada 30 Mahasiswa sebagai berikut:

Interval Nilai Ujian (X)	Frekuensi (f)
39-44	4
45-50	4
51-56	7
57-62	6
63-68	4
69-74	5
Total	$\Sigma f = 30$

Untuk dapat menentukan Median pada data tersebut, maka harus dilakukan langkah-langkah sebagai berikut:

- a) Menentukan FREKUENSI Kelas Median (f) = 7
- b) Menentukan Kelompok atau Kelas atau Interval dimana Median akan terletak = Interval 51-56
- c) Menentukan Panjang Kelas Median (p) = 6
- d) Menentukan BATAS BAWAH Kelas Median (b) = $(50+51)/2 = 50,5$
- e) Menentukan Jumlah semua Frekuensi dari Kelas atau Interval atau Kelompok yang LEBIH KECIL dari Kelas Median (F) = $(4+4) = 8$
- f) Mengidentifikasi Jumlah Data atau Banyaknya Data atau Sampel (n) = 30
- g) Selanjutnya angka-angka tersebut dimasukkan ke dalam rumus:

$$Me = b + p \left(\frac{\frac{1}{2}n - F}{f} \right)$$

$$\begin{aligned}
 Me &= 50,5 + 6 \{(1/2 (30)-8)\}/7 \\
 &= 50,5 + 6 \{(7)/7\} \\
 &= 50,5 + 6 \\
 &= \mathbf{56,5}.
 \end{aligned}$$

Jadi MEDIAN dari Data tersebut adalah **56,5** ($Me = 56,5$).

c. MODUS

Modus merupakan nilai yang paling sering muncul dalam suatu kelompok atau sekumpulan data. Modus juga dibedakan menjadi 2 macam, yaitu Modus pada Data Tunggal dan Modus pada Data Kelompok.

1) Menentukan MODUS pada Data Tunggal

MODUS pada Data Tunggal dapat ditentukan berdasarkan Data yang Jumlah Frekuensinya Paling Banyak.

Contoh:

Nilai	F
50	3
60	7
70	10
80	8
90	2
Total	$\sum f = 30$

Berdasarkan data tersebut, maka Modusnya adalah 70 (Berdasarkan data yang mempunyai Frekuensi terbanyak, yaitu 10 adalah Nilai 70). Jadi $Mo = 70$.

2) Menentukan MODUS pada Data Kelompok

MODUS pada Data Kelompok dapat ditentukan dengan cara menggunakan Rumus berikut ini:

$$Mo = b + p \left(\frac{b_1}{b_1 + b_2} \right)$$

Keterangan:

"FREKUENSI KELAS MODUS DITENTUKAN BERDASARKAN JUMLAH FREKUENSI DATA YANG TERBANYAK"

b : Batas Bawah Kelas MODUS (Yaitu: Kelas dimana MODUS akan terletak)

p : Panjang Kelas MODUS.

b_1 : Frekuensi Kelas Modus DIKURANGI Frekuensi Kelas Interval yang mempunyai Tanda Kelas Lebih Kecil (SEBELUM) Kelas Modus.

b_2 : Frekuensi Kelas Modus DIKURANGI Frekuensi Kelas Interval yang mempunyai Tanda Kelas Lebih Besar (SESUDAH) Kelas Modus.

Contoh:

Interval Nilai Ujian (X)	Frekuensi (f)
39-44	4
45-50	4
51-56	7
57-62	6
63-68	4
69-74	5
Total	$\sum f = 30$

Untuk dapat menentukan Modus pada data tersebut, maka harus dilakukan langkah-langkah sebagai berikut:

- Menentukan FREKUENSI Kelas Modus (f) = 7
- Menentukan Kelompok atau Kelas atau Interval dimana Modus akan terletak = Interval 51-56
- Menentukan Panjang Kelas Modus (p) = 6
- Menentukan BATAS BAWAH Kelas Modus (b) = $(50+51)/2 = 50,5$
- Menentukan b_1 , yaitu Frekuensi Kelas Modus DIKURANGI Frekuensi Kelas Interval yang mempunyai Tanda Kelas Lebih Kecil (SEBELUM) Kelas Modus = $7-4 = 3$
- Menentukan b_2 , yaitu : Frekuensi Kelas Modus DIKURANGI Frekuensi Kelas Interval yang mempunyai Tanda Kelas Lebih Besar (SESUDAH) Kelas Modus = $7-6 = 1$
- Selanjutnya angka-angka tersebut dimasukkan ke dalam rumus:

$$Mo = b + p \left(\frac{b_1}{b_1 + b_2} \right)$$

$$\begin{aligned} Mo &= 50,5 + 6 (3/4) \\ &= 50,50 + 4,5 \\ &= \mathbf{55} \end{aligned}$$

Jadi Modus atau **Mo** dari data tersebut di atas adalah **55**

4. Teknik Menentukan Ukuran Pemusatan dengan Aplikasi SPSS

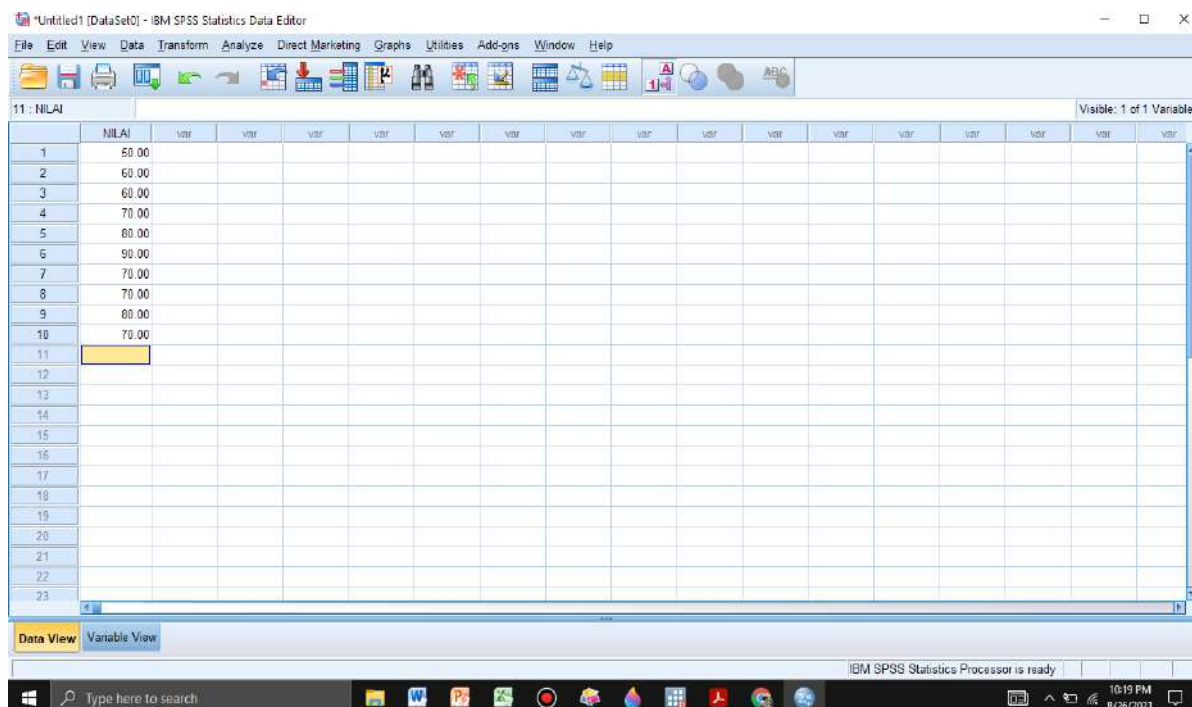
Menentukan Nilai Mean, Median dan Modus dapat dilakukan secara cepat dan mudah dengan menggunakan aplikasi SPSS. Karena ukuran pemusatan merupakan bagian dari statistik deskriptif, maka untuk menentukan Mean, Median dan Modus pada aplikasi SPSS menggunakan analisis deskriptif. Berikut dicontohkan cara menggunakan analisis tersebut beserta langkah-langkah praktis dari aplikasi SPSS.

Sebagai contoh, kita akan menentukan Mean, Median dan Modus pada data sebagai berikut:

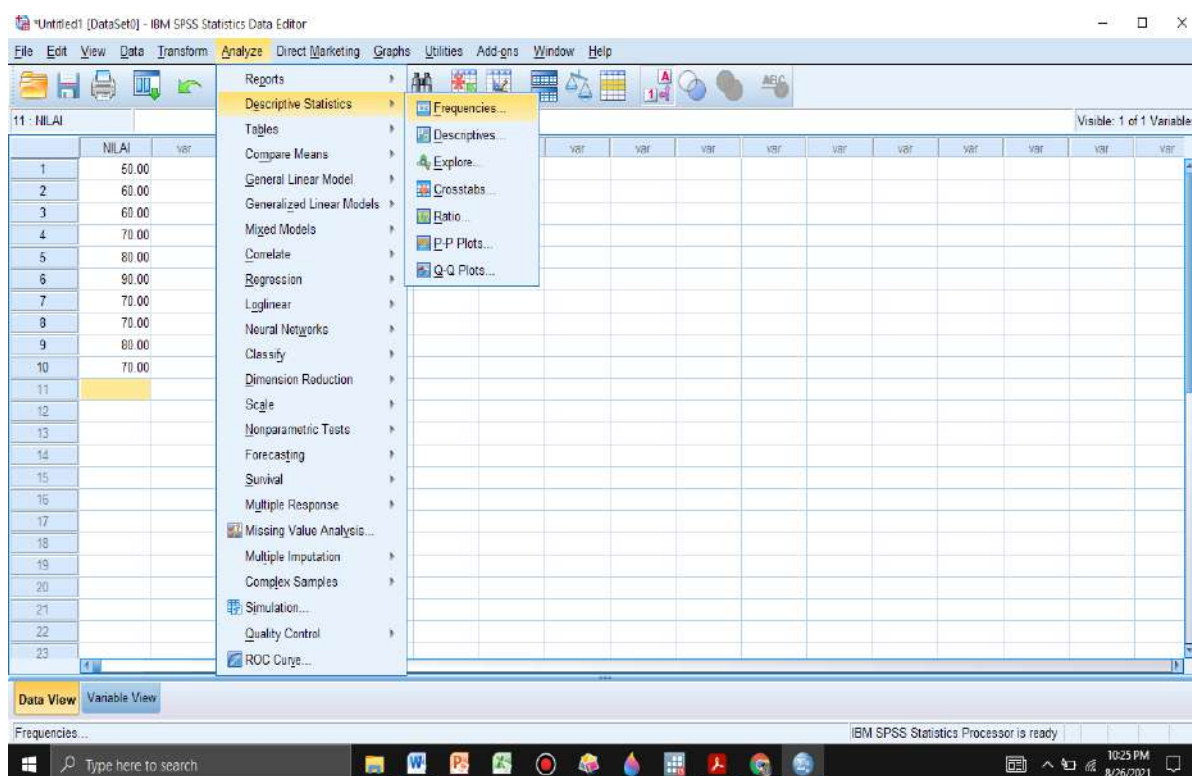
NO	NILAI
1	50
2	60
3	60
4	70
5	80
6	90
7	70
8	70
9	80
10	70

Langkah-langkah untuk menentukan ukuran pemusatan dengan SPSS dari data tersebut adalah sebagai berikut:

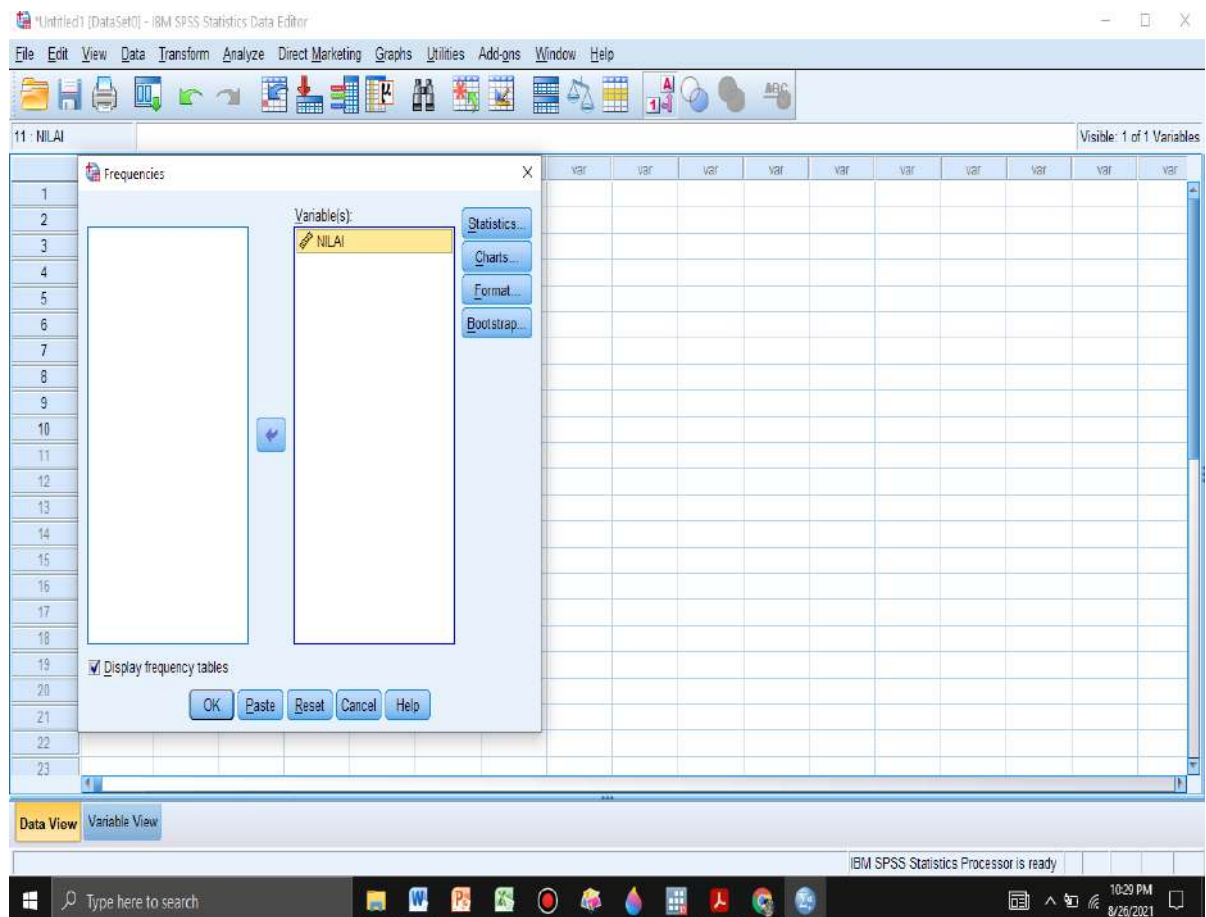
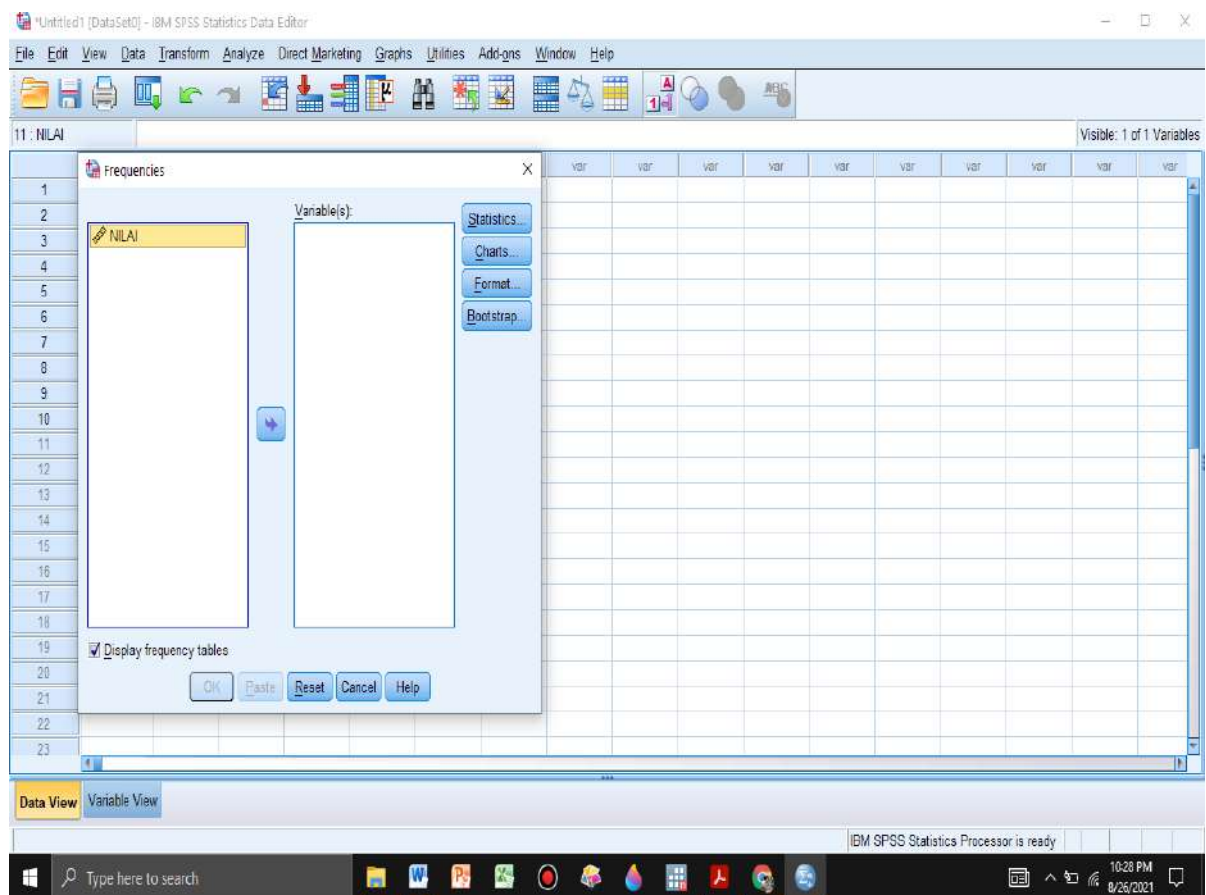
1. Entry data dari hasil tabulasi data di atas pada Data View yang terdapat pada SPSS. Perhatikan gambar berikut ini:



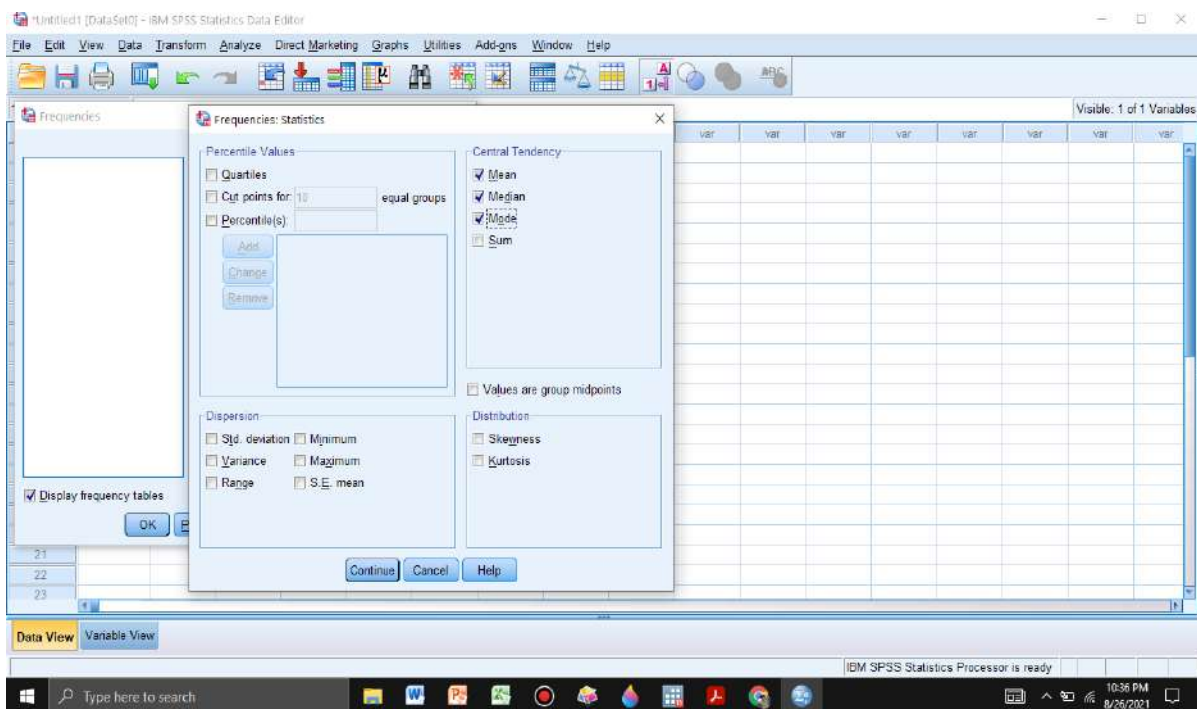
2. Setelah selesai entry data, kemudian klik pada menu ANALYZE, pilih atau klik DESCRIPTIVE STATISTICS, selanjutnya klik FREQUENCIES. Perhatikan tampilan gambar berikut ini:



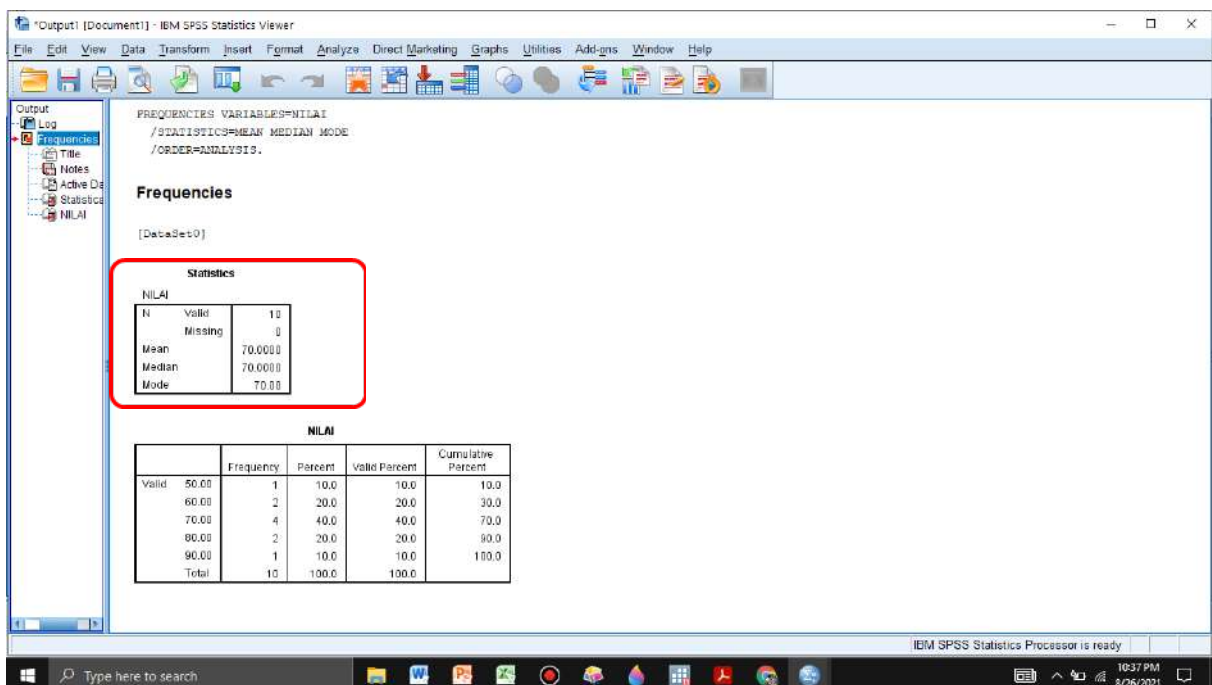
3. Selanjutnya masukkan data nilai ke dalam kotak Variable(s) sebagaimana tampilan gambar berikut:



4. Kemudian klik **Statistics...** untuk menentukan jenis analisis deskriptif yang akan dilakukan. Kemudian pilih dengan memberi tanda pada Mean, Median dan Modus. Seperti yang terlihat pada gambar berikut:



5. Selanjutnya klik CONTINUE dan klik OK, kemudian lihat hasil atau Out Putnya seperti berikut:



Berdasarkan hasil tersebut dapat diketahui nilai Mean, Median dan Modus dari data yang kita analisis.

C. Rangkuman

Ukuran Pemusatan atau Tendensi Sentral adalah suatu nilai tunggal yang merepresentasikan gambaran secara umum tentang keadaan dari nilai tersebut yang terdapat dalam suatu data. Ukuran pemusatan atau nilai pusat dapat mewakili data secara keseluruhan dan merupakan rata-rata (*average*), karena nilai rata-ratanya dihitung dari keseluruhan nilai yang terdapat dalam data tersebut.

Ukuran Pemusatan atau Tendensi Sentral terdiri atas Mean, Median dan Modus, dimana masing-masing mempunyai sifat dan kekhususan sendiri, sehingga penggunaannya harus disesuaikan dengan tujuan analisa data. Penghitungan atau penentuan Mean, Median dan Modus dapat dilakukan baik pada Data Tunggal maupun Data Kelompok, dengan menggunakan rumus dan ketentuan yang berbeda satu dengan yang lainnya.

D. Tugas

Untuk meningkatkan pemahaman tentang materi Ukuran Pemusatan ini, silahkan mengerjakan latihan berikut ini:

1. Ukuran Pemusatan pada Data Tunggal

Berdasarkan catatan pendaftaran pada sebuah klinik selama tahun 2019 didapatkan jumlah pengunjung setiap bulan adalah sebagai berikut:

NO	BULAN	JUMLAH PENGUNJUNG
1	Januari	35
2	Februari	39
3	Maret	50
4	April	42
5	Mei	41
6	Juni	40
7	Juli	35
8	Agustus	39
9	September	35
10	Oktober	45
11	November	39
12	Desember	40

Berdasarkan tabulasi data tersebut, maka:

- Tentukan Mean jumlah pengunjung klinik pada tiap bulan!
- Tentukan Median jumlah pengunjung klinik pada tiap bulan!
- Tentukan Modus jumlah pengunjung klinik pada tiap bulan!

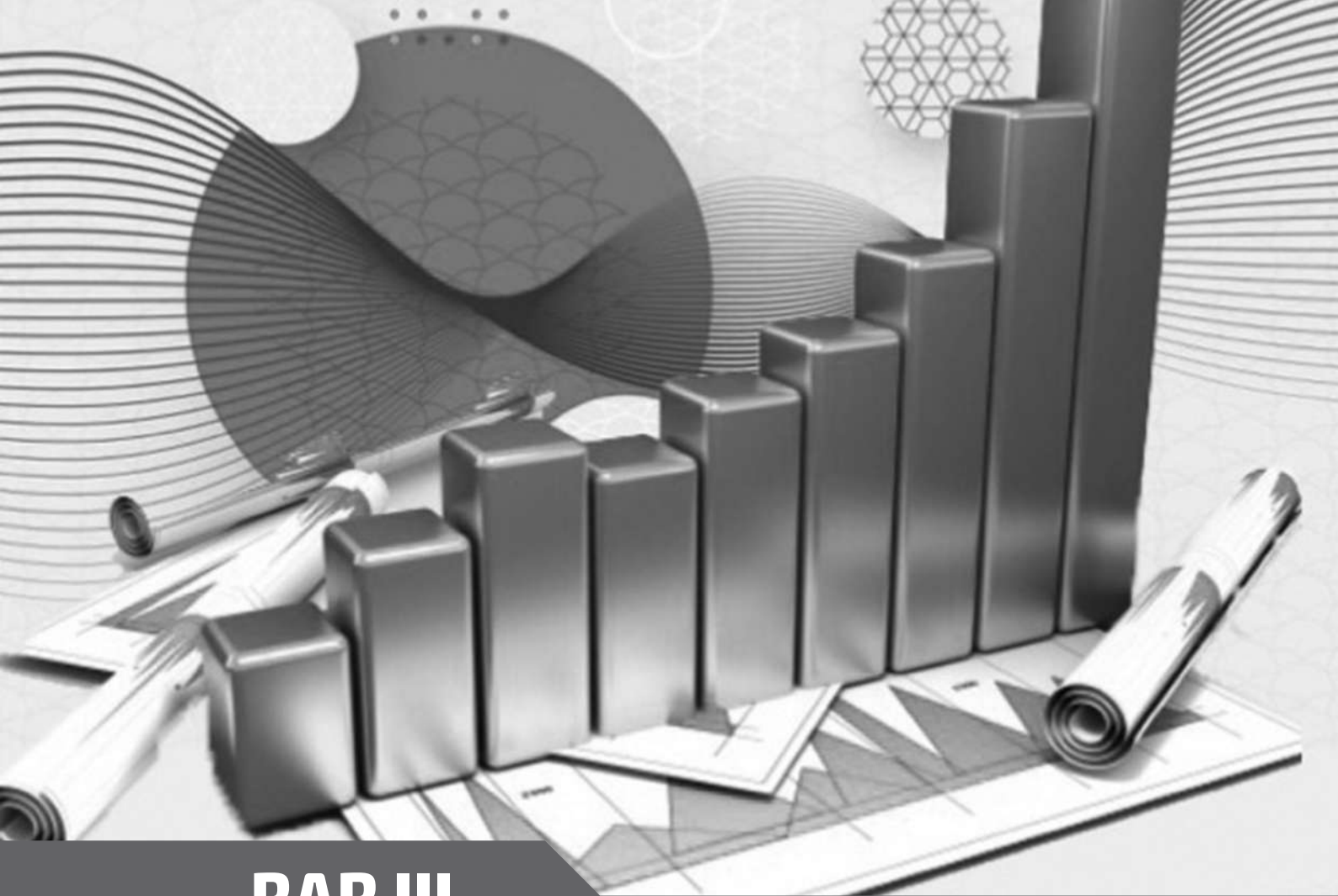
2. Ukuran Pemusatan pada Data Kelompok

Tentukan Mean, Median dan Modus dari data tentang distribusi umur responden adalah sebagai berikut:

NO	UMUR (Tahun)	FREKUENSI
1	15 - 19	2
2	21 – 24	30
3	25 – 29	40
4	30 – 34	42
5	35 – 39	25
6	40 - 44	11
Jumlah		150

E. Referensi

- Alhusin, S. 2003. *Aplikasi Statistik Praktis dengan SPSS for Windows*. Yogyakarta. Graha Ilmu
- Amin.I., Aswin.A., Fajar.I., Isnaeni, Iwan.S., Pudjirahaju.A., Sunindya.R.. 2009. *Statistika untuk Praktisi Kesehatan*. Yogyakarta. Graha Ilmu
- Hadi.S. 2002. *Statistik*. Jilid 2. Yogyakarta. Andi Offset.
- Hasibuan.A.A.,Supardi, Syah.D. 2009. *Pengantar Statistik Pendidikan*. Jakarta. Gaung Persada Press.
- Machfoedz, I., 2016. *Bio Statistika*. Edisi Revisi. Yogyakarta. Fitramaya.
- Riduwan.2010. *Dasar-dasar Statistika*. Bandung. Alfabeta.
- Riwidikdo, H., 2013. *Statistik Kesehatan: Belajar mudah teknik analisis data dalam Penelitian Kesehatan (Plus Aplikasi Software SPSS)*. Yogyakarta. Mitra Cendikia Press.
- Riwidikdo,H., 2012. *Statistik Kesehatan*. Yogyakarta. Nuha Medika
- Sarwono, J. 2006. *Analisis Data Penelitian Menggunakan SPSS*. Yogyakarta. Andi Offset.
- Santjaka, A. 2011. *Statistik Untuk Penelitian Kesehatan: Deskriptif, Inferensial, Parametrik dan Non Parametrik*. Yogyakarta. Nuha Medika.
- Siswandari. 2009. *Statistika (Komputer Based)*. Surakarta. LPP UNS dan UNS Press
- Stang. 2018. *Cara Praktis Penentuan Uji Statistik dalam Penelitian Kesehatan dan Kedokteran. Edisi 2*. Jakarta. Mitra Wacana Media.
- Sugiyono, 2000. *Statistik untuk Penelitian*, Bandung: Alfabeta.
- Sugiyono. 2015. *Statistika untuk Penelitian*. Bandung. Alfabeta
- Trihendradi, C. 2010. *Step by Step SPSS 18: Analisis Data Statistik*. Yogyakarta. Andi Offset.



BAB III

DISPERSI DATA

Dodiet Aditya Setyawan, SKM., MPH.

A. Tujuan Pembelajaran

Setelah mempelajari materi ini mahasiswa diharapkan mampu:

1. Memahami pengertian dispersi data.
2. Memahami macam-macam dispersi data
3. Melakukan penghitungan Range, Deviasi, Mean deviasi dan Standard Deviasi.
4. Mengaplikasikan teknik penentuan Standard Deviasi berbasis komputer dengan menggunakan aplikasi SPSS.

B. Materi

1. Pengertian

Dispersi merupakan ukuran variasi dan sering juga disebut sebagai ukuran penyimpangan atau ketidaksesuaian. Dengan demikian Dispersi merupakan sebuah ukuran yang dapat menunjukkan besar kecilnya penyimpangan data terhadap nilai rata-rata atau Mean-nya. Dispersi dapat menunjukkan persebaran data terhadap nilai sentralnya atau berpencarnya data. Dengan melakukan pengukuran terhadap dispersi, maka kita akan dapat mengetahui bagaimana persebaran data yang sesungguhnya.

Pengukuran dispersi dapat digunakan untuk menilai heterogenitas distribusi data yang sebenarnya. Apabila hasil pengukuran dispersi suatu data itu besar, maka menunjukkan bahwa data tersebut terlalu heterogen distribusinya, artinya bahwa sebaran data sangat bervariasi dari yang sangat tinggi sampai dengan yang sangat rendah. Sebaliknya apabila nilai dispersi kecil, maka menunjukkan persebaran data tersebut semakin homogen. Artinya bahwa data-data dengan nilai dispersi yang kecil menunjukkan bahwa nilai data-data tersebut hampir sama.

Nilai Dispersi juga menunjukkan bahwa nilai data mempunyai rata-rata yang sama, tetapi bervariasi dan luas penyebarannya berbeda. Oleh karena itu ukuran dispersi juga dapat dipergunakan untuk mencari luas penyebaran data, variasi data dan stabilitas data.

2. Macam-Macam Pengukuran Dispersi

Macam-macam pengukuran dispersi yang akan dibahas pada buku ini meliputi:

a. Rentang atau Range (R)

Range merupakan ukuran variasi yang paling sederhana, dimana data yang sudah diurutkan mulai dari yang terkecil sampai yang terbesar selanjutnya dihitung selisih antara data terbesar dengan data terkecil. Karena penghitungan Range atau Rentang Data hanya melibatkan data terbesar dan data terkecil saja dan mengabaikan keberadaan data-data yang lain, maka dikatakan bahwa Range merupakan pengukuran dispersi yang masih kasar. Jadi Range (R) adalah selisih dari nilai terbesar dengan nilai terkecil dari suatu rangkaian data. Untuk menghitung Range atau Rentang Data (R) Jadi untuk menghitung Range (R) adalah Data Terbesar dikurangi Data Terkecil.

Contoh:

Hasil nilai ujian pada mata kuliah Statistika dari 15 mahasiswa adalah sebagai berikut:

Mahasiswa	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15
Nilai	90	80	85	70	75	80	80	85	70	60	75	60	90	80	75

Berdasarkan tabulasi data tersebut, maka Range dari data tersebut dapat dihitung dengan cara sebagai berikut:

Data terbesar = 90

Data Terkecil = 60.

Dengan demikian nilai Range adalah: Nilai Terbesar – Nilai Terkecil

= 90-60

= 30.

Jadi Range dari data tersebut adalah 30.

b. Deviasi (D)

Deviasi merupakan ukuran penyimpangan yang menunjukkan selisih antara hasil pengukuran dengan nilai mean atau rata-ratanya. Sehingga dapat dirumuskan:

$$D_i = X_i - \bar{X}$$

Keterangan:

D_i = Deviasi Data Ke-i;

X_i = Nilai Data Ke-i;

$\bar{X}$ = Mean.

Contoh:

Hasil pengukuran kadar gula darah pada 5 orang adalah sebagai berikut:

Orang ke-	1	2	3	4	5
Kadar Gula Darah	10	200	210	180	250

Berdasarkan data tersebut, maka deviasi dari masing-masing data hasil pengukuran dapat dilakukan dengan cara:

- 1) Menentukan Nilai Rata-rata (Mean) dari data tersebut:

$$\text{MEAN} = (170+200+210+180+250)/5$$

$$= 1.010/5$$

$$= 202.$$

Dengan demikian Mean = 202.

- 2) Menghitung Deviasi dari masing-masing data:

Setelah nilai Rata-rata atau Mean diketahui, maka Deviasi dari masing-masing data dapat ditentukan sebagai berikut:

- 1) Data Kadar Gula Darah 170 (X_1), mempunyai nilai Deviasi (D_1) sebesar: [Nilai Data Ke-1 (X_1) – Mean]
 $= 170 - 202$
 $= -32$
- 2) Data Kadar Gula Darah 200 (X_2), mempunyai nilai Deviasi (D_2) sebesar: [Nilai Data Ke-2 (X_2) – Mean]
 $= 200 - 202$
 $= -2$
- 3) Data Kadar Gula Darah 210 (X_3), mempunyai nilai Deviasi (D_3) sebesar: [Nilai Data Ke-3 (X_3) – Mean]
 $= 210 - 202$
 $= 8$
- 4) Data Kadar Gula Darah 180 (X_4), mempunyai nilai Deviasi (D_4) sebesar: [Nilai Data Ke-4 (X_4) – Mean]
 $= 180 - 202$
 $= -22$
- 5) Data Kadar Gula Darah 250 (X_5), mempunyai nilai Deviasi (D_5) sebesar: [Nilai Data Ke-5 (X_5) – Mean]
 $= 250 - 202$
 $= 48$

Hasil penghitungan Deviasi dari data tersebut menunjukkan bahwa pada dasarnya Mean itu membagi data menjadi 2 (dua) kelompok, yaitu data yang kurang dari Mean dan data yang melebihi Mean. Data yang kurang dari Mean ditunjukkan dengan nilai yang bertanda Negatif, sedangkan data yang melebihi nilai Mean, ditunjukkan dengan nilai yang bertanda Positif.

Jadi nilai Deviasi mempunyai makna:

- 1) Tanda Positif: menunjukkan posisi dari sebaran data berada di atas Nilai Rata-rata atau Mean.
- 2) Tanda Negatif: menunjukkan posisi dari sebaran data berada di bawah Nilai Rata-rata atau Mean.
- 3) Jumlah dari Deviasi pasti sama dengan 0 (Nol).

c. Mean Deviasi (Deviasi Rata-Rata)

Mean Deviasi atau dikenal juga sebagai Deviasi Rata-Rata pada prinsipnya sama dengan konsep Mean atau Rata-rata, yaitu Jumlah seluruh data dibagi dengan banyaknya data. Dengan konsep Mean seperti ini, maka MEAN DEVIASI adalah jumlah seluruh deviasi dibagi dengan banyaknya deviasi. Dalam penghitungan Mean Deviasi mempunyai makna:

- 1) Jika Nilai Mean Deviasi Positif, menunjukkan bahwa nilai data tersebut lebih besar dari nilai Mean atau Rata-rata.
- 2) Jika Nilai Mean Deviasi Negatif, menunjukkan bahwa nilai data tersebut lebih kecil dari nilai Mean atau Rata-rata.

- 3) Semakin Kecil nilai Mean Deviasi, maka semakin baik karena mendekati data realitasnya.

Adapun rumus untuk penghitungan Mean Deviasi adalah:

$$MD = \frac{\sum_{i=1}^n Di}{n}$$

Dimana: $\sum_{i=1}^n Di$ adalah: Jumlah Semua Nilai Deviasi, **n** adalah Jumlah Data

Contoh:

Hasil pengukuran kadar gula darah pada 5 orang adalah 170, 150, 200, 180, 250, didapatkan nilai Deviasinya sebagai berikut:

Gula Darah	170	150	200	180	250
Deviasi	-20	-40	10	-10	60

Berdasarkan data tersebut, maka Mean Deviasi dari masing-masing data hasil pengukuran dapat dilakukan dengan cara: **Jumlah dari semua Nilai Deviasi dibagi dengan Jumlah Data.**

Jumlah dari semua deviasi:

$$= (20) + (40) + 10 + (10) + 60$$

$$= 140.$$

Dengan Jumlah $n = 5$, maka Mean Deviasi data tersebut:

$$= 140/5$$

$$= 28.$$

Sehingga dapat disimpulkan bahwa: Rata-Rata Penyimpangan masing-masing data terhadap Nilai Mean sebesar 28.

d. Standard Deviation

Standard Deviation (SD) atau Deviasi Standar atau dikenal juga dengan istilah Simpangan Baku merupakan pengukuran yang digunakan untuk mengetahui seberapa jauh persebaran data dari nilai rata-ratanya (Mean). Deviasi Standar pada Sampel dilambangkan dengan **s**. kegunaan dari Standar Deviasi ini yang sesungguhnya adalah untuk menentukan apakah persebaran data itu mendekati homogen atau mendekati heterogen. Karena data yang mempunyai rata-rata yang baik, belum tentu semua data tersebut adalah baik. Oleh karena itu, Standar Deviasi yang kecil menunjukkan data tersebut semakin homogen dan Standar Deviasi yang besar menunjukkan data tersebut semakin heterogen. Adapun rumus dari Standard Deviation (SD) adalah:

- a. Deviasi Standar untuk Sampel Besar ($n > 30$)

$$SD = \sqrt{\frac{\sum (X_i - \bar{X})^2}{n}}$$

Dimana:

X_i = Nilai Data Ke-i dari Sampel

$\bar{X}$ = Rata-Rata Sampel

n = Jumlah Sampel

- b. Deviasi Standar untuk Sampel Kecil ($n \leq 30$)

$$SD = \sqrt{\frac{\sum (X_i - \bar{X})^2}{n-1}}$$

Dimana:

X_i = Nilai Data Ke-i dari Sampel

$\bar{X}$ = Rata-Rata Sampel

n = Jumlah Sampel

Contoh:

Hasil dari ujian statistik pada 15 mahasiswa didapatkan nilai sebagai berikut:

MHS	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15
Nilai	70	60	80	75	65	80	70	75	60	65	70	80	80	75	70

Maka untuk menentukan berapa SD dari data tersebut adalah:

- 1) Hitung Mean data tersebut terlebih dahulu:
 $(70+60+80+75+65+80+70+75+60+65+70+80+80+75+70)/15$
 $= 1075/15$
 $= 72$
- 2) Buatlah Tabel penolong untuk menghitung nilai deviasi dan deviasi kuadrat dari masing-masing data seperti berikut:

NO	NILAI	$(X_i - \bar{X}) = D$	$(X_i - \bar{X})^2 = D^2$
1	70	-2	4
2	60	-12	144
3	80	8	64
4	75	3	9
5	65	-7	49
6	80	8	64
7	70	-2	4
8	75	3	9
9	60	-12	144

NO	NILAI	$(Xi - \bar{X}) = D$	$(Xi - \bar{X})^2 = D^2$
10	65	-7	49
11	70	-2	4
12	80	8	64
13	80	8	64
14	75	3	9
15	70	-2	4
Jumlah			685

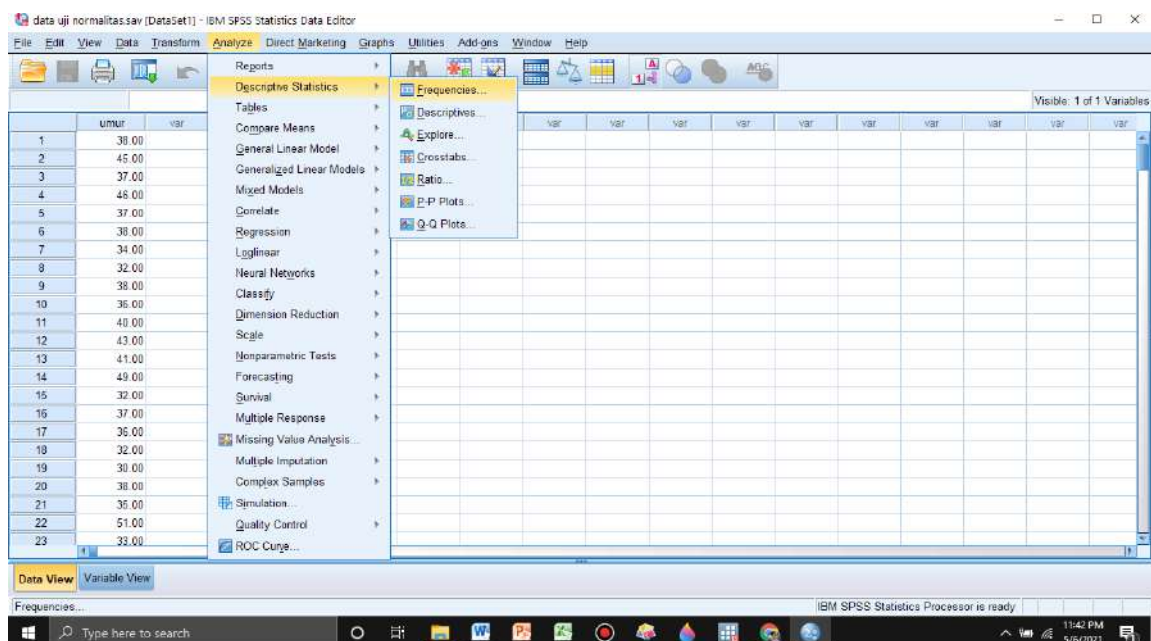
- 3) Hitunglah SD dengan menggunakan rumus: $\sqrt{\frac{685}{14}} = \sqrt{48,9} = 6,9$.

Jadi berdasarkan hasil penghitungan tersebut, menunjukkan bahwa Standar Deviasi untuk contoh data tersebut adalah 6,9. Standar deviasi juga dapat dihitung dengan menggunakan computer, yaitu dengan bantuan aplikasi SPSS, sehingga menjadi lebih mudah dan cepat.

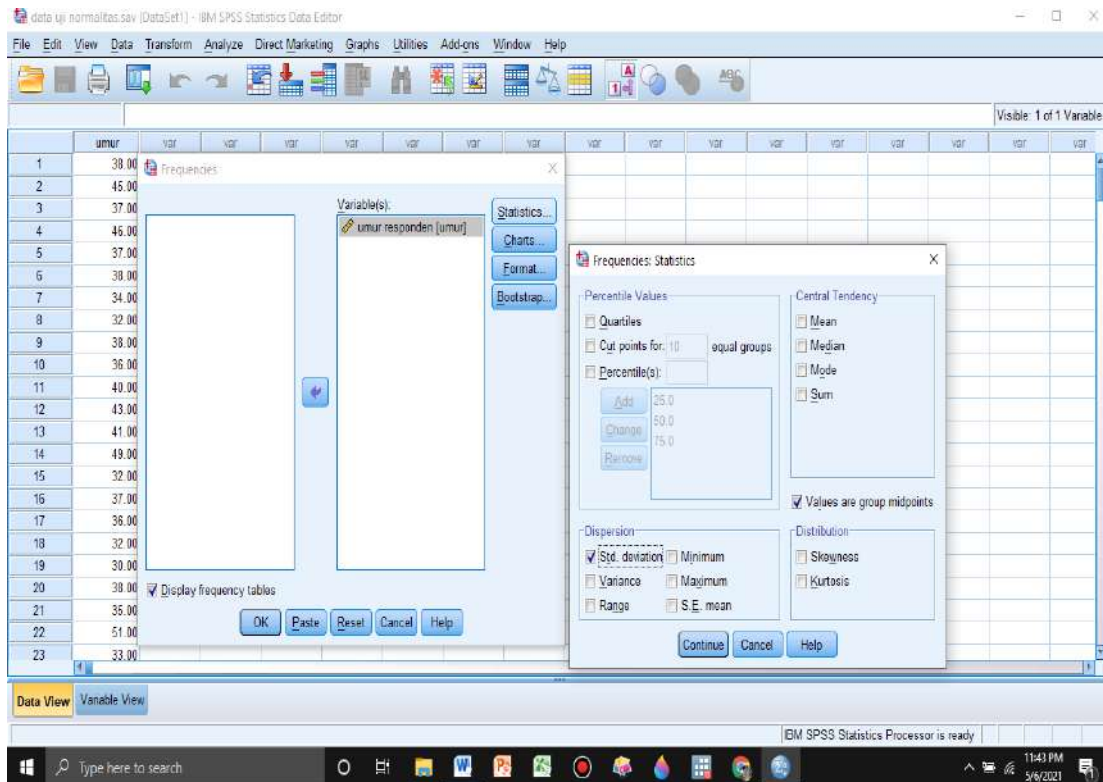
3. Cara Menentukan Standard Deviation Menggunakan SPSS

Analisis deskriptif pada pengukuran Standar Deviasi (SD) secara praktis dapat dilakukan dengan mudah menggunakan aplikasi SPSS. Adapun langkah-langkah untuk melakukan pengukuran Standar Deviasi (SD) pada program SPSS adalah sebagai berikut:

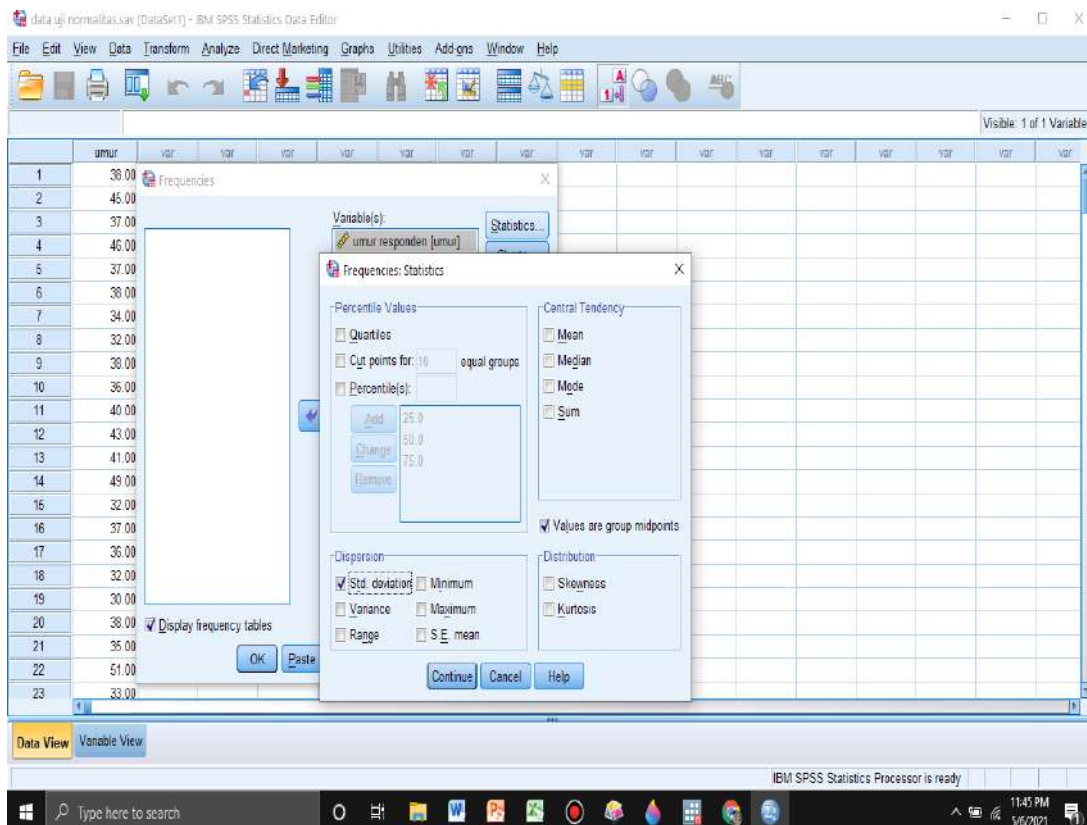
- Masukkan semua variabel (nama-nama variabel) pada Sheet Variabel View.
- Selanjutnya masukkan semua data dari masing-masing variabel kedalam Sheet Data View.
- Setelah entry data selesai, selanjutnya klik ANALYZE kemudian pilih DESKRIPTIF STATISTICS dan klik FREQUENCIES. Perhatikan langkah-langkah tersebut sebagaimana tampilan gambar seperti berikut:



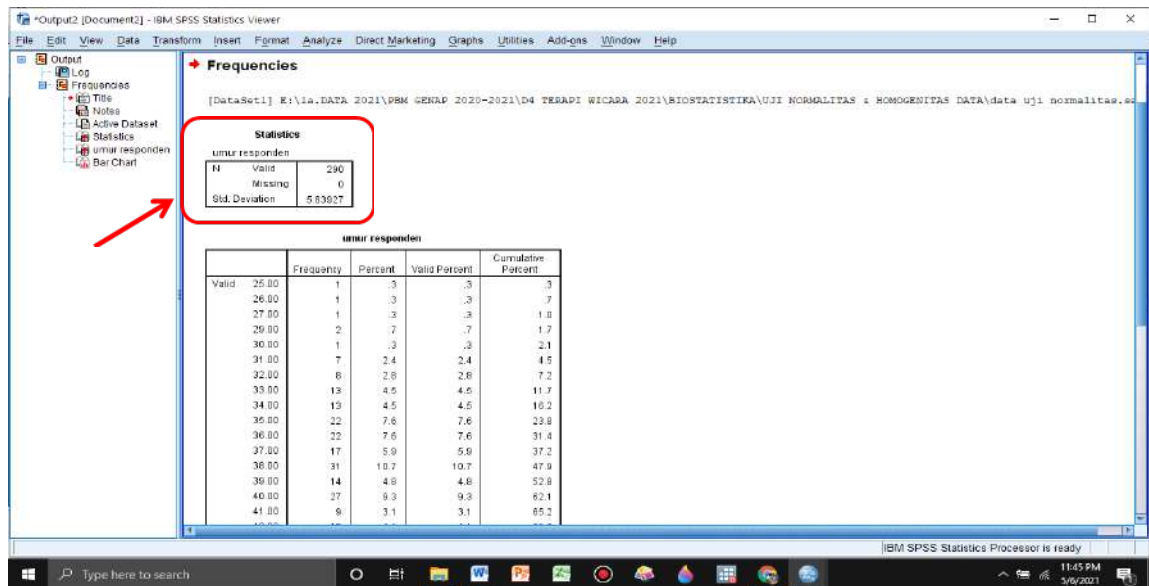
- d. Masukkan nama-nama variabel yang akan dianalisis ke dalam Kotak VARIABLE(S) dan selanjutnya klik STATISTICS untuk melakukan analisis Deskriptif. Perhatikan tampilan gambar berikut ini:



- e. Setelah klik Statistics, selanjutnya beri tanda centang dengan cara klik pada Std. Deviation. Untuk lebih jelasnya lihat tampilan gambar berikut ini:



- f. Selanjutnya klik CONTINUE dan Kemudian langkah terakhir adalah klik OK, maka akan muncul hasil atau Out put seperti gambar berikut ini:



C. Rangkuman

Dispersi merupakan sebuah ukuran yang dapat menunjukkan besar kecilnya penyimpangan data terhadap nilai rata-rata atau Meannya. Dispersi dapat menunjukkan persebaran data terhadap nilai sentralnya atau berpencarnya data. Dengan melakukan pengukuran terhadap dispersi, maka kita akan dapat mengetahui bagaimana persebaran data yang sesungguhnya.

Pengukuran dispersi dapat digunakan untuk menilai heterogenitas distribusi data yang sebenarnya. Apabila hasil pengukuran dispersi suatu data itu besar, maka menunjukkan bahwa data tersebut terlalu heterogen distribusinya, artinya bahwa sebaran data sangat bervariasi dari yang sangat tinggi sampai dengan yang sangat rendah. Sebaliknya apabila nilai dispersi kecil, maka menunjukkan persebaran data tersebut semakin homogeny. Artinya bahwa data-data dengan nilai dispersi yang kecil menunjukkan bahwa nilai data-data tersebut hampir sama.

Pengukuran dispersi dapat berupa Range, Deviasi, Mean Deviasi dan Standar Deviasi. Range merupakan ukuran variasi yang paling sederhana, dimana data yang sudah diurutkan mulai dari yang terkecil sampai yang terbesar selanjutnya dihitung selisish antara data terbesar dengan data terkecil. Deviasi merupakan ukuran penyimpangan yang menunjukkan selisih antara hasil pengukuran dengan nilai mean atau rata-ratanya. Mean Deviasi atau dikenal juga sebagai Deviasi Rata-Rata pada prinsipnya sama dengan konsep Mean atau Rat-rata, yaitu Jumlah seluruh data dibagi dengan banyaknya data. Standard Deviation (SD) atau Deviasi Standar atau dikenal juga dengan istilah Simpangan Baku merupakan pengukuran yang digunakan untuk mengetahui seberapa jauh persebaran data dari nilai rata-ratanya (Mean).

D. Tugas

Untuk meningkatkan pemahaman tentang materi Dispersi Data tersebut, silahkan mengerjakan latihan berikut ini:

1. Hitunglah nilai Range, Deviasi, Mean Deviasi dan Standard Deviasi dari data tentang Berat Badan Responden berikut ini:

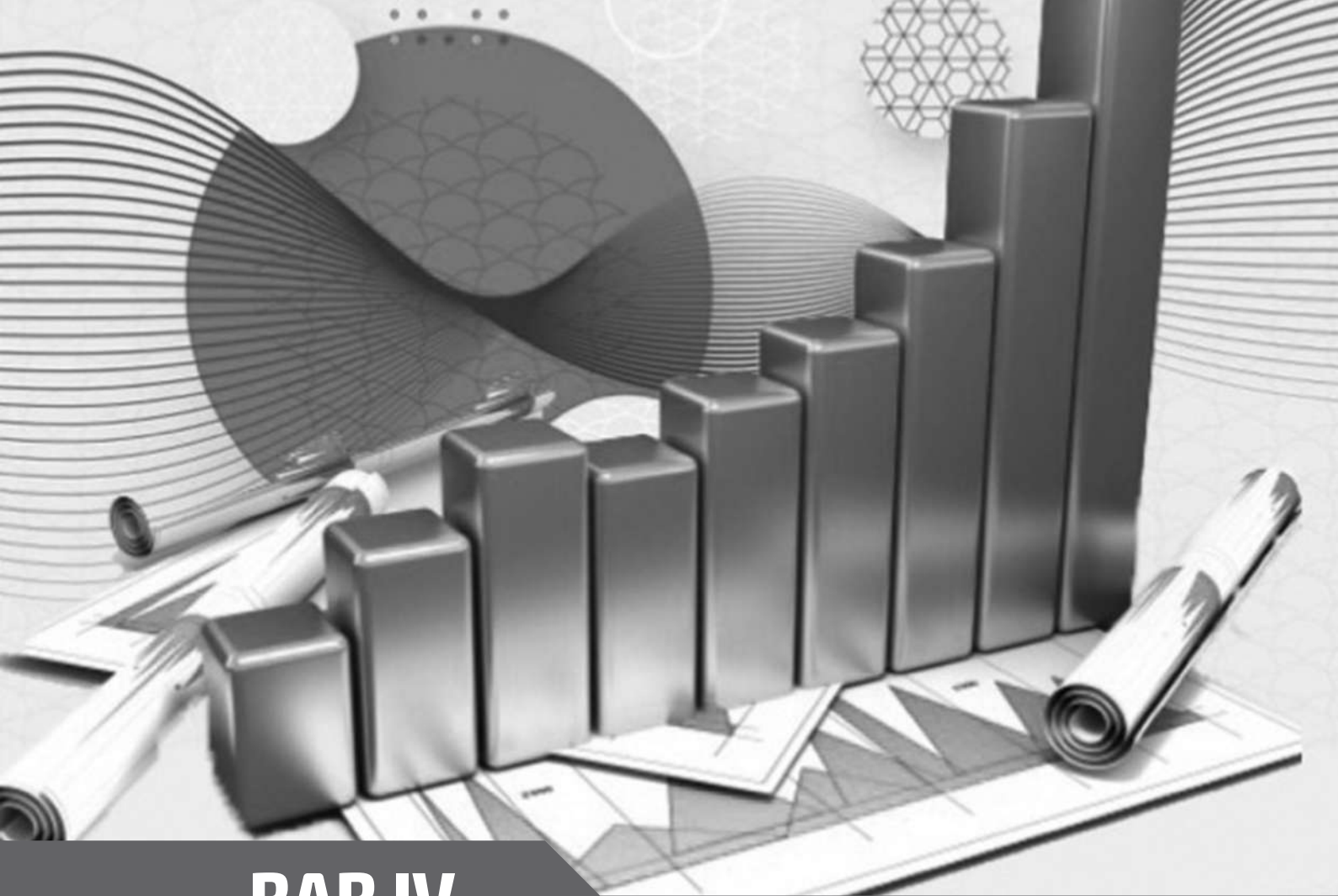
Responden	1	2	3	4	5	6	7	8	9	10
Berat Badan	60	55	70	68	72	63	49	48	52	56

2. Hitunglah nilai Range, Deviasi, Mean Deviasi dan Standard Deviasi dari data tentang Tinggi Badan Responden berikut ini:

Responden	1	2	3	4	5	6	7	8	9	10	11	12
TB	160	165	150	169	145	170	181	169	165	169	148	169

E. Referensi

- Alhusin, S. 2003. *Aplikasi Statistik Praktis dengan SPSS for Windows*. Yogyakarta. Graha Ilmu
- Amin.I., Aswin.A., Fajar.I., Isnaeni, Iwan.S., Pudjirahaju.A., Sunindya.R.. 2009. *Statistika untuk Praktisi Kesehatan*. Yogyakarta. Graha Ilmu
- Dahlan.M.S. 2014. *Statistik untuk Kedokteran dan Kesehatan: Deskriptif, Bivariat dan Multivariat Dielngkapi Aplikasi dengan Menggunakan SPSS*, Edisi 6. Jakarta. Epidemiologi Indonesia.
- Dahlan.M.S. 2012. *Statistik untuk Kedokteran dan Kesehatan*. Jakarta. Salemba Medika.
- Hadi.S. 2002. *Statistik*. Jilid 2. Yogyakarta. Andi Offset.
- Hasibuan.A.A.,Supardi, Syah.D. 2009. *Pengantar Statistik Pendidikan*. Jakarta. Gaung Persada Press.
- Riduwan.2010. *Dasar-dasar Statistika*. Bandung. Alfabeta.
- Riwidikdo, H., 2013. *Statistik Kesehatan: Belajar mudah teknik analisis data dalam Penelitian Kesehatan (Plus Aplikasi Software SPSS)*. Yogyakarta. Mitra Cendikia Press.
- Riwidikdo,H., 2012. *Statistik Kesehatan*. Yogyakarta. Nuha Medika
- Sarwono, J. 2006. *Analisis Data Penelitian Menggunakan SPSS*. Yogyakarta. Andi Offset.
- Santjaka, A. 2011. *Statistik Untuk Penelitian Kesehatan: Deskriptif, Inferensial, Parametrik dan Non Parametrik*. Yogyakarta. Nuha Medika.
- Siswandari. 2009. *Statistika (Komputer Based)*. Surakarta. LPP UNS dan UNS Press
- Stang. 2018. *Cara Praktis Penentuan Uji Statistik dalam Penelitian Kesehatan dan Kedokteran. Edisi 2*. Jakarta. Mitra Wacana Media.
- Sugiyono. 2015. *Statistika untuk Penelitian*. Bandung. Alfabeta
- Trihendradi, C. 2010. *Step by Step SPSS 18: Analisis Data Statistik*. Yogyakarta. Andi Offset.



BAB IV

PROBABILITAS

Ros Endah Happy Patriyani,SKp, Ns., MKep.

A. Tujuan pembelajaran :

Setelah menyelesaikan materi tentang probabilitas, mahasiswa mampu memahami tentang:

1. Konsep dasar probabilitas
2. Manfaat probabilitas
3. Perumusan klasik probabilitas
4. Aturan Probabilitas
5. Ruang sampel dan kejadian
6. Probabilitas kejadian majemuk
7. Dua kejadian saling lepas
8. Dua kejadian saling bebas
9. Probabilitas bersyarat
10. Probabilitas gabungan
11. Probabilitas marginal

B. Materi

1. Konsep dasar probabilitas

- a. Pengertian

Ada beberapa pengertian dari probabilitas.

- 1) Probabilitas merupakan besarnya kesempatan (kemungkinan/peluang) suatu peristiwa akan terjadi.
- 2) Probabilitas" ialah suatu nilai yang digunakan untuk mengukur tingkat terjadinya suatu kejadian acak.
- 3) Probabilitas dapat diartikan sebagai derajat/ tingkat kepastian atau keyakinan dari munculnya hasil percobaan statistik.
- 4) Probabilitas diartikan sebagai suatu ukuran seberapa besar suatu kejadian akan terjadi terhadap semua kejadian yang ada.
- 5) Probabilitas didefinisikan sebagai peluang atau kemungkinan atau besarnya kesempatan terjadinya suatu peristiwa.
- 6) Probabilitas atau Peluang adalah suatu ukuran tentang kemungkinan suatu peristiwa (*event*) akan terjadi di masa mendatang. Probabilitas dapat juga diartikan sebagai harga angka yang menunjukkan seberapa besar kemungkinan suatu peristiwa terjadi, di antara keseluruhan peristiwa yang mungkin terjadi.
- 7) Probabilitas sebagai suatu ukuran tentang kemungkinan suatu peristiwa (*event*) akan terjadi dimasa mendatang. Besarnya kesempatan dapat ditulis dalam bentuk bilangan desimal, pecahan atau persen.

- b. Kata kunci probabilitas

Probabilitas dilambangkan dengan P. Dalam mempelajari probabilitas, ada tiga kata kunci atau tiga hal penting yang harus diketahui, yaitu: eksperimen/percobaan, hasil kejadian atau peristiwa.

- 1) Percobaan adalah pengamatan terhadap beberapa aktivitas atau proses yang memungkinkan timbulnya paling sedikit 2 peristiwa tanpa memperhatikan peristiwa mana yang akan terjadi.
- 2) Hasil adalah suatu hasil dari sebuah percobaan.

- 3) Peristiwa adalah kumpulan dari satu atau lebih hasil yang terjadi pada sebuah percobaan atau kegiatan.

Contoh probabilitas pelemparan koin:

- 1) Dari eksperimen pelemparan sebuah koin.
- 2) Hasil outcome dari pelemparan sebuah koin tersebut adalah "Muka" atau "Belakang."
- 3) Kumpulan dari beberapa hasil tersebut dikenal sebagai kejadian *event*.

Suatu probabilitas dinyatakan antara 0 sampai 1 atau dalam presentase. Probabilitas 0 menunjukkan peristiwa yang tidak mungkin terjadi, sedangkan probabilitas 1 menunjukkan peristiwa yang pasti terjadi. Semakin dekat nilai probabilitas ke nilai 0, semakin kecil kemungkinan suatu kejadian akan terjadi. Sebaliknya semakin dekat nilai probabilitas ke nilai 1 semakin besar peluang suatu kejadian akan terjadi.

Probabilitas yang rendah menunjukkan kecilnya kemungkinan suatu peristiwa akan terjadi. Kita dapat menentukan probabilitas terjadinya hujan, munculnya muka 1 pada percobaan pelemparan dadu, probabilitas munculnya kartu AS pada penarikan kartu dari sekelompok kartu Bridge dan seterusnya.

2. Manfaat probabilitas

- a. Dalam kehidupan sehari-hari probabilitas dapat membantu kita dalam mengambil suatu keputusan, serta meramalkan kejadian yang mungkin terjadi.
- b. Dengan teori probabilitas dapat ditarik kesimpulan secara tepat atas hipotesis yang terkait tentang karakteristik populasi.
- c. Jika ditinjau dalam penelitian, probabilitas membantu peneliti dalam pengambilan keputusan yang lebih tepat. Pengambilan keputusan yang lebih tepat dimaksudkan tidak ada keputusan yang sudah pasti karena kehidupan mendatang tidak ada yang pasti kita ketahui dari sekarang, karena informasi yang didapat tidaklah sempurna.

3. Perumusan klasik probabilitas

Perumusan konsep dasar probabilitas dilakukan dengan tiga cara, yaitu perumusan klasik, cara frekuensi relatif dan pendekatan subjektif.

a. Perumusan Klasik

Bila kejadian E terjadi dalam m cara dari seluruh n cara yang mungkin terjadi dan masing-masing n cara itu mempunyai kesempatan atau kemungkinan yang sama untuk muncul, probabilitas kejadian E yang ditulis $P(E)$ dirumuskan sebagai berikut :

$$\text{Rumus } P(E) = m / n$$

Keterangan

$P(E)$ = Probabilitas *Event*

m = Jumlah kejadian yang diinginkan (peristiwa)

n = Keseluruhan kejadian yang mungkin terjadi

Contoh:

Sebuah dadu dilemparkan. Muka dadu ada 6. Semua muka dadu mempunyai kesempatan yang sama untuk muncul. Salah satu muka yang akan muncul dari muka-muka dadu itu ($m=1$) adalah muka dadu 1, muka dadu 2, muka dadu 3, muka dadu 4, muka dadu 5 atau muka dadu 6. Maka probabilitas kejadian E adalah:

$$P(E) = P(1) = P(2) = P(3) = P(4) = P(5) = P(6) = m/n = 1/6$$

b. Cara Frekuensi Relatif

Perumusan konsep probabilitas dengan cara klasik mempunyai kelemahan karena menuntut syarat semua hasil mempunyai kesempatan yang sama untuk muncul. Pengertian ini mengaburkan adanya probabilitas yang sama. Sehubungan dengan itu dikembangkan konsep probabilitas berdasarkan statistik, yaitu dengan pendekatan empiris. Probabilitas empiris dari suatu kejadian dirumuskan dengan memakai frekuensi relatif dari terjadinya suatu kejadian dengan syarat banyaknya pengamatan atau banyaknya sampel n adalah sangat besar. Bila n bertambah besar sampai tak terhingga ($n \rightarrow \infty$), probabilitas kejadian E sama dengan nilai limit dari frekuensi relatif kejadian E tersebut. Dengan demikian, jika kejadian E berlangsung sebanyak f kali dari keseluruhan pengamatan sebanyak n , dimana n mendekati tak berhingga, probabilitas kejadian E dirumuskan sebagai berikut:

$$\text{Rumus } P(E) = \lim f / n$$

Keterangan

$P(E)$ = probabilitas kejadian E dengan nilai limit frekuensi relatif

m = seluruh cara yang mungkin terjadi

n = berapa x cara yang mungkin terjadi

Walaupun mudah dan berguna dalam praktik, secara matematis perumusan konsep probabilitas dengan frekuensi relatif ini juga mempunyai kelemahan karena suatu nilai limit yang benar-benar mungkin sebenarnya tidak ada. Oleh karena itu, konsep probabilitas modern dikembangkan dengan memakai pendekatan aksiomatis, yaitu suatu kebenaran yang diterima secara apa adanya tanpa memerlukan bukti matematis, dimana konsep probabilitas tidak didefinisikan, seperti konsep titik dan konsep garis yang tidak didefinisikan dalam ilmu geometri (Boediono, 2006).

Contoh:

Dari 50 mahasiswa yang mengikuti ujian statistika, distribusi frekuensi nilai mahasiswa adalah seperti tabel berikut

Nilai (X)	50	55	60	70	75	85
Frekuensi (f)	5	10	15	5	10	5

Maka probabilitas kejadian E mahasiswa memperoleh nilai tersebut adalah

$$P(E) = P(50) = 5/50$$

$$P(E) = P(55) = 10/50$$

$$P(E) = P(60) = 15/50, \text{ dst}$$

c. Pendekatan Subjektif

Pendekatan subjektif yang digunakan untuk menentukan probabilitas suatu peristiwa didasarkan pada selera dan keyakinan individu seseorang. Misalnya, saya ingin menentukan bahwa besok probabilitas naiknya harga dolar Amerika adalah 0.65 atau 65%. Atas dasar apa saya menentukan probabilitas naiknya harga dolar itu 65%? Pengetahuan ini hanya didasarkan pada pengetahuan, pengalaman, dan keahlian yang dimiliki. Dengan demikian, probabilitas suatu peristiwa yang ditentukan dengan pendekatan subjektif menyebabkan penentuan probabilitas suatu peristiwa antara orang yang satu dengan yang lain dapat berbeda. Hal ini disebabkan oleh tingkat pengetahuan, penguasaan informasi, naluri dan faktor-faktor lain yang berkaitan dengan peristiwa itu.

4. Aturan Probabilitas

a. Definisi dan notasi

Sebelum mendiskusikan tentang aturan probabilitas, ada beberapa definisi yang perlu diketahui sebagai berikut.

- 1) Dua kejadian dikatakan *mutually exclusive* atau saling asing (*disjoint*) jika tidak terjadi pada waktu yang sama.
- 2) Probabilitas kejadian A terjadi, dengan syarat kejadian B telah terjadi, dikatakan sebagai probabilitas bersyarat (*conditional probability*). Probabilitas bersyarat di atas dinotasikan dengan simbol $\Pr(A|B)$.
- 3) Pelengkap atau complement dari suatu kejadian A adalah tidak terjadinya kejadian tersebut, disimbolkan dengan A' . Peluangnya dinotasikan dengan $\Pr(A')$.
- 4) Probabilitas kejadian A dan B keduanya terjadi adalah probabilitas A irisan B (A intersection B) dinotasikan dengan $\Pr(A \cap B)$. Jika A dan B saling asing maka $\Pr(A \cap B) = 0$.
- 5) Probabilitas kejadian A atau B terjadi adalah probabilitas A gabungan B (A union B) dinotasikan dengan $\Pr(A \cup B)$.
- 6) Jika terjadinya peristiwa A mengubah probabilitas kejadian B maka kejadian A dan B dikatakan dependent. Sebaliknya, apabila terjadinya peristiwa A tidak mempengaruhi probabilitas peristiwa B maka A dan B dikatakan independent.

Aturan probabilitas

- 1) Aturan pengurangan

Probabilitas kejadian A akan terjadi sama dengan 1 dikurangi probabilitas kejadian A tidak akan terjadi.

$$\Pr(A) = 1 - \Pr(A')$$

Misal sebagai contoh probabilitas Anda lulus universitas sama dengan 0.80. Berdasarkan aturan pengurangan, probabilitas Anda tidak akan lulus universitas adalah $1.00 - 0.80$ atau 0.20.

2) Aturan Perkalian.

Aturan perkalian digunakan pada situasi kita ingin mengetahui probabilitas irisan dua kejadian, yaitu kita ingin mengetahui kejadian A dan B terjadi kedua-duanya. Probabilitas kejadian A dan B terjadi bersamaan sama dengan probabilitas kejadian A terjadi dikalikan probabilitas kejadian B bersyarat A.

$$\Pr (A \cap B) = \Pr (A) \Pr (B | A)$$

3) Aturan penambahan

Aturan penambahan digunakan pada situasi di mana kita punya dua kejadian dan kita ingin tahu salah satu atau keduanya terjadi. Probabilitas kejadian A dan atau kejadian B terjadi sama dengan probabilitas kejadian A terjadi ditambah probabilitas kejadian B terjadi dikurangi probabilitas kedua kejadian A dan B terjadi bersama-sama.

$$\Pr (A \cup B) = \Pr (A) + \Pr (B) - \Pr (A) \Pr (B | A)$$

b. **Teorema Bayes atau Aturan Bayes**

Teorema Bayes (juga dikenal dengan aturan Bayes) sangat berguna sebagai alat untuk menghitung probabilitas bersyarat. *Teorema Bayes* dapat dinyatakan sebagai berikut: *Teorema Bayes*. Diberikan $A_1, A_2, \dots, A_n$ masing-masing adalah himpunan dari kejadian yang saling asing dan bersama-sama membentuk ruang sampel S . Misalkan, B merupakan sebarang kejadian dari ruang sampel yang sama, $\Pr (B) > 0$. Selanjutnya,

$$\Pr (A_k|B) = \frac{\Pr (A_k \cap B)}{\Pr (A_1 \cap B) + \Pr (A_2 \cap B) + \dots \Pr (A_n \cap B)}$$

Catatan:

$$\Pr (A_k|B) = \frac{\Pr (A_k) \Pr (B | A_k)}{\Pr (A_1) \Pr (B | A_1) + \dots + \Pr (A_n) \Pr (B | A_n)}$$

Kapan memakai *Teorema Bayes*?

Tantangan dalam mengaplikasikan *teorema Bayes* melibatkan pengenalan tipe-tipe masalah yang menjamin penggunaannya. *Teorema Bayes* dapat dipertimbangkan ketika beberapa syarat berikut ini terpenuhi.

- Ruang sampel dapat dipartisi menjadi himpunan kejadian yang saling asing $\{A_1, A_2, \dots, A_n\}$.
- Dalam ruang sampel, terdapat kejadian B , $\Pr (B) > 0$.
- Tujuan analisisnya adalah menghitung probabilitas bersyarat dalam bentuk $\Pr (A_k | B)$
- Anda tahu minimal satu dari dua himpunan probabilitas berikut.
 - $\Pr (A_k \cap B)$ untuk setiap A_k .
 - $\Pr (A_k)$ dan $\Pr (B | A_k)$ untuk setiap A_k .

5. Ruang sampel dan kejadian

Pada pelemparan sebuah uang logam, ada dua hasil yang mungkin muncul, yaitu muka (m) atau belakang (b). Dua hasil yang mungkin muncul ini dapat dihimpun menjadi $S = \{m,b\}$. dengan demikian dapat dikatakan bahwa kumpulan himpunan dari semua hasil yang mungkin muncul atau terjadi pada suatu percobaan statistik disebut ruang sampel, yang dilambangkan dengan himpunan S , sedangkan anggota-anggota dari S disebut titik sampel.

Perhatikan bahwa pada pelemparan sebuah uang logam tersebut $S = \{m,b\}$ dan $A = \{m\}$, sehingga $A \subset S$, A merupakan himpunan bagian dari S . berdasarkan kejadian A dan ruang sampel S tersebut, perumusan konsep probabilitas didefinisikan sebagai berikut. Bila kejadian A berlangsung dalam m cara pada ruang sampel S yang terjadi dalam n cara, probabilitas kejadian A adalah:

$$\text{Rumus } P(A) = n(A) / n(S) = m / n$$

Keterangan

Dimana $n(A)$ = banyaknya anggota A

$n(S)$ = banyaknya anggota S

Perhatikan bahwa definisi probabilitas tersebut tidak menuntut syarat bahwa semua titik sampel mempunyai kesempatan yang sama untuk muncul. Definisi probabilitas kejadian ini terlepas dari definisi probabilitas yang dirumuskan secara klasik maupun memakai frekuensi relative. Dengan menggunakan rumus 3, kita dapat menentukan probabilitas dari sembarang kejadian A yang didefinisikan pada S .

Contoh:

Pada pelemparan dua buah uang logam:

- a. Tentukanlah ruang sampel S

Hasil-hasil yang mungkin muncul adalah sebagai berikut:

Uang Logam 1	Uang Logam 2	
	M	B
M	(M,M)	(M,B)
B	(B,M)	(B,B)

Jadi ruang sampel S adalah = $\{(m,m), (m,b), (b,m), (b,b)\}$

Titik sampel (m,m) menyatakan munculnya sisi muka dari uang logam pertama dan kedua, titik sampel (m,b) menyatakan munculnya muka dari uang logam pertama dan belakang dari uang logam kedua, begitu seterusnya.

- b. Bila A menyatakan kejadian munculnya sisi yang sama dari dua uang logam tersebut, tentukanlah probabilitas kejadian A

A adalah kejadian munculnya sisi-sisi yang sama dari dua uang logam, maka $A = \{(m,m), (b,b)\}$. Dengan demikian, $n(A) = 2$ dan $n(S) = 4$, sehingga probabilitas kejadian A adalah

$$P(A) = n(A) / n(S) = 2/4 = 1/2$$

6. Probabilitas kejadian majemuk

Dengan mengingat kembali pengetahuan mengenai teori himpunan bahwa bila A dan B dua himpunan dalam himpunan semesta S, gabungan dari A dan B adalah himpunan baru yang anggotanya terdiri atas anggota A atau anggota B, atau anggota keduanya ditulis $A \cup B = \{x \in A \text{ atau } x \in B\}$.

Banyaknya anggota himpunan $A \cup B$ adalah $n(A \cup B) = n(A) + n(B) - n(A \cap B)$

Sejalan dengan himpunan gabungan tersebut, karena ada keterkaitan antara teori himpunan dengan teori probabilitas, kita dapat merumuskan kejadian gabungan A dan B, yaitu kejadian $A \cup B$ pada ruang sampel S. Bila A dan B kejadian sembarang pada ruang sampel S, gabungan kejadian A dan B yang ditulis $A \cup B$ adalah kumpulan semua titik sampel yang ada pada A atau B atau pada kedua-duanya. Kejadian $A \cup B$ disebut kejadian majemuk. Demikian halnya, kejadian $A \cap B$ yaitu kumpulan titik sampel yang ada pada A dan B, juga disebut kejadian majemuk. Probabilitas kejadian $A \cup B$ dirumuskan sebagai berikut

$$\text{Rumus } P(A \cup B) = P(A) + P(B) - P(A \cap B)$$

Penjelasan lahirnya rumus tersebut adalah sebagai berikut
Kita telah tahu bahwa
 $n(A \cup B) = n(A) + n(B) - n(A \cap B)$

Bila dua ruas persamaan dibagi dengan $n(S)$, diperoleh
 $n(A \cup B) / n(S) = n(A) / n(S) + n(B) / n(S) - n(A \cap B) / n(S)$

sehingga

$$P(A \cup B) = P(A) + P(B) - P(A \cap B)$$

Contoh

Peluang seorang mahasiswa lulus Statistik adalah $1/3$ dan peluang ia lulus Kewirausahaan adalah $5/9$. Bila peluang lulus sekurang-kurangnya satu mata kuliah di atas adalah $4/5$, berapa peluang ia lulus kedua mata kuliah itu?

Jawab

Misalkan A = kejadian lulus Statistik

B = kejadian lulus Kewirausahaan

$$P(A) = 1/3$$

$$P(B) = 5/9$$

$$P(A \cap B) = 4/5$$

$$\begin{aligned}
 P(A \cup B) &= P(A) + P(B) - P(A \cap B) \\
 &= 1/3 + 5/9 - 4/5 \\
 &= 4/45
 \end{aligned}$$

7. Dua kejadian saling lepas

Dalam menentukan probabilitas dengan aturan matematis penjumlahan dan pengurangan perlu diketahui sifat dua atau lebih peristiwa. Sifat dua atau lebih peristiwa tersebut adalah saling meniadakan (*mutually exclusive*) dan tidak saling meniadakan (*non-mutually exclusive*). Bila A dan B dua kejadian sembarang pada S dan berlaku $A \cap B = \emptyset$, A dan B dikatakan dua kejadian saling lepas atau saling bertentangan, atau saling terpisah (*mutually exclusive*). Hal ini menunjukkan bahwa peristiwa A dan peristiwa B dua kejadian saling lepas, $P(A \cap B) = P(\emptyset) = 0$, sehingga probabilitas kejadian A u B dirumuskan sebagai berikut

$$\text{Rumus } P(A \cup B) = P(A) + P(B)$$

Contoh

Bila A dan B dua kejadian saling lepas, dengan $P(A) = 0.35$ dan $P(B) = 0.25$, tentukanlah $P(A \cup B)$

Jawab

Karena A dan B saling lepas, berlaku :

$$\begin{aligned}
 P(A \cup B) &= P(A) + P(B) \\
 &= 0.35 + 0.25 = 0.60
 \end{aligned}$$

8. Dua kejadian saling bebas

Sifat dua atau lebih peristiwa dari suatu percobaan dapat independen dan dapat pula dependen. Dua atau lebih peristiwa dikatakan independen jika terjadinya suatu peristiwa tidak mempengaruhi terjadinya peristiwa yang lain. Sebaliknya, dua atau lebih peristiwa dikatakan bersifat dependen jika terjadinya suatu peristiwa akan mempengaruhi terjadinya peristiwa yang lain. Dapat dikatakan bahwa dua kejadian A dan B dalam ruang sampel S dikatakan saling bebas jika kejadian A tidak mempengaruhi kejadian B dan sebaliknya, kejadian B tidak mempengaruhi kejadian A (Wibisono, 2007). Jika A dan B merupakan dua kejadian saling bebas, berlaku rumus berikut

$$\text{Rumus } P(A \cap B) = P(A) \cdot P(B)$$

Contoh :

Jika diketahui dua kejadian A dan B saling bebas dengan $P(A) = 0.35$ dan $P(B) = 0.45$, berlaku

Jawab

$$\begin{aligned}
 P(A \cap B) &= P(A) \cdot P(B) \\
 &= 0.35 \cdot 0.45 = 0.1575 = 0.16
 \end{aligned}$$

9. Probabilitas bersyarat

Probabilitas bersyarat peristiwa tidak saling bebas adalah probabilitas terjadinya suatu peristiwa dengan syarat peristiwa lain harus terjadi dan peristiwa-peristiwa tersebut saling mempengaruhi. Jika peristiwa B bersyarat terhadap A, probabilitas terjadinya peristiwa tersebut adalah $P(B|A)$ dibaca probabilitas terjadinya B dengan syarat peristiwa A terjadi.

Contoh :

Sebuah kotak berisikan 11 bola dengan rincian : 5 buah bola putih bertanda + 1 buah bola putih bertanda – 3 buah bola kuning bertanda + 2 buah bola kuning bertanda – Seseorang mengambil sebuah bola kuning dari kotak –

Berapa probabilitas bola itu bertanda +?

Jawab:

Misalkan : A = bola kuning B+ = bola bertanda positif B- = bola bertanda negatif. $P(A) = \frac{5}{11}$ $P(B+|A) = \frac{3}{5}$

10. Probabilitas gabungan

Probabilitas gabungan peristiwa tidak saling bebas adalah probabilitas terjadinya dua atau lebih peristiwa secara berurutan bersamaan dan peristiwa-peristiwa itu saling mempengaruhi. Jika dua peristiwa A dan B gabungan, probabilitas terjadinya peristiwa tersebut adalah $P(A \cap B) = P(A) \times P(B|A)$ Jika tiga buah peristiwa A, B, dan C gabungan, probabilitas terjadinya adalah $P(A \cap B \cap C) = P(A) \times P(B|A) \times P(C|A \cap B)$

Contoh :

Dari satu set kartu bridge berturut-turut diambil kartu itu sebanyak 2 kali secara acak. Hitunglah probabilitasnya kartu king A pada pengambilan pertama dan as B pada pengambilan kedua, jika kartu pada pengambilan pertama tidak dikembalikan.

Jawab:

A = pengambilan pertama keluar kartu king. $P(A) = \frac{4}{52}$ B = pengambilan kedua keluar kartu as $P(B|A) = \frac{1}{51}$ $P(A \cap B) = P(A) \times P(B|A) = \frac{4}{52} \times \frac{1}{51} = 0,006$

11. Probabilitas marjinal

Probabilitas marjinal peristiwa tidak saling bebas adalah probabilitas terjadinya suatu peristiwa yang tidak memiliki hubungan dengan terjadinya peristiwa lain dan peristiwa tersebut saling mempengaruhi. Jika dua peristiwa A adalah marjinal, probabilitas terjadinya peristiwa A tersebut adalah $P(A) = \sum_{i=1}^n P(A \cap B_i) = \sum_{i=1}^n P(A_i) \times P(B_i|A)$, $i = 1, 2, 3, \dots$

Contoh :

Sebuah kotak berisikan 11 bola dengan rincian : 5 buah bola putih bertanda + 1 buah bola putih bertanda – 3 buah bola kuning bertanda + 2 buah bola kuning bertanda – Tentukan probabilitas memperoleh sebuah bola putih

Jawab:

Misalkan: A = bola putih B+ = bola bertanda positif B- = bola bertanda negatif PB+A = 511 PB-A = 111 PA = PB+A + PB-A = 511 + 111 = 611

C. Rangkuman

1. Probabilitas dari suatu kejadian merupakan suatu ukuran dari kemungkinan kejadian tersebut akan terjadi. Probabilitas kejadian A biasa dinotasikan dengan $Pr A()$.
2. Ukuran probabilitas harus memenuhi syarat-syarat sebagai berikut:
 - a. Probabilitas dari suatu kejadian bernilai antara 0 sampai 1.
 - b. Jumlahan dari probabilitas-probabilitas dari semua kejadian dalam suatu ruang sampel sama dengan 1.
 - c. Probabilitas dari suatu kejadian A merupakan jumlahan dari probabilitas-probabilitas dari semua sampel di A.
3. Nilai probabilitas dari suatu kejadian dapat diperoleh dengan menggunakan beberapa definisi berikut definisi klasik, definisi frekuensi relatif atau empiris, dan definisi subjektif.
4. Aturan probabilitas ada 3, yaitu: Aturan pengurangan, Aturan Perkalian, Aturan penambahan
5. Dalam menentukan probabilitas dengan aturan matematis penjumlahan dan pengurangan perlu diketahui sifat dua atau lebih peristiwa.
6. Sifat dua atau lebih peristiwa dari suatu percobaan dapat independen dan dapat pula dependen.
7. Probabilitas bersyarat peristiwa tidak saling bebas adalah probabilitas terjadinya suatu peristiwa dengan syarat peristiwa lain harus terjadi dan peristiwa-peristiwa tersebut saling mempengaruhi
8. Probabilitas gabungan peristiwa tidak saling bebas adalah probabilitas terjadinya dua atau lebih peristiwa secara berurutan bersamaan dan peristiwa-peristiwa itu saling mempengaruhi.
9. Probabilitas marginal peristiwa tidak saling bebas adalah probabilitas terjadinya suatu peristiwa yang tidak memiliki hubungan dengan terjadinya peristiwa lain dan peristiwa tersebut saling mempengaruhi.

D. Tugas

Untuk memperdalam pemahaman Anda mengenai materi di atas, kerjakanlah latihan berikut!

1. Diketahui data populasi penduduk 40% laki-laki dan 60% perempuan. Diketahui 4% orang yang terkena buta warna adalah laki-laki. Hitunglah probabilitas mendapatkan orang buta warna di antara orang laki-laki!
2. Lemparkan dua koin yang seimbang dan gambarkan kemungkinan hasilnya. Definisikan kejadian-kejadian berikut:
 A = koin pertama muncul ekor
 B = koin kedua muncul muka
 Apakah kedua kejadian A dan B saling *independent*?

E. Referensi

1. Putranto, Leksmono Suryo (2017) Buku Statistika dan Probabilitas. Jakarta: PT. Indeks.
2. Rakhman, A. Probabilitas. <http://repository.ut.ac.id/3980/1/MATA4450-M1.pdf>
3. Safitri Jaya. Modul Statistika Probabilitas.
4. Sudaryono. (2012). Statistika Probabilitas - Teori & Aplikasi. Yogyakarta: Andi.
5. Yakub, S. (2012). Modul Statistika Probabilitas. Medan: Triguna Dharma.



BAB V

POPULASI DAN SAMPEL

Ade Devriany, M.Kes

A. Tujuan Pembelajaran

Mampu memahami perbedaan populasi dan sampel penelitian serta mampu merancang dan mengaplikasikan teknik pengambilan sampel penelitian secara acak (random) dan tidak acak (non-random).

B. MATERI

1. Konsep Dasar Populasi dan Sampel

Pelaksanaan suatu penelitian selalu berhadapan dengan objek yang diteliti atau yang diselidiki. Objek tersebut dapat berupa manusia, hewan, tumbuh-tumbuhan, benda-benda mati lainnya, serta peristiwa dan gejala yang terjadi di dalam masyarakat atau di dalam alam. Pada sebagian besar penelitian, terutama yang menggunakan populasi manusia, tidak semua orang bisa diteliti. Hal ini disebabkan oleh karena populasinya terlalu besar, sehingga tidak mungkin untuk meneliti setiap orang karena keterbatasan waktu, biaya dan kendala-kendala lainnya, atau karena populasinya tidak bisa didefinisikan secara spesifik berdasarkan ruang dan waktu. Dalam keadaan seperti ini hanya sebagian dari populasi, yaitu sampel yang diteliti dan hasilnya diberlakukan secara umum yang mencakup seluruh populasi.

Dengan meneliti sebagian dari populasi, diharapkan bahwa hasil yang diperoleh akan menggambarkan sifat populasi yang bersangkutan. Kesimpulan-kesimpulan penelitian mengenai sampel akan digeneralisasikan terhadap populasi. Generalisasi sampel ke populasi mengandung risiko bahwa akan terdapat kekeliruan atau ketidaktepatan karena sampel tidak akan mencerminkan secara tepat keadaan populasi. Makin tidak sama karakteristik sampel dengan karakteristik populasi, maka makin besar kemungkinan kekeliruan dalam generalisasi. Hal ini menunjukkan pentingnya teknik pengambilan sampel yang tepat untuk memperkecil kekeliruan generalisasi dari sampel ke populasi.

a. Populasi

Populasi atau *universe* adalah jumlah keseluruhan dari unit analisa yang ciri-cirinya akan diduga. Populasi juga diartikan keseluruhan individu yang menjadi acuan hasil-hasil penelitian akan berlaku. Populasi dapat dibagi menjadi dua kelompok yaitu populasi sasaran (target populasi) dan populasi sampling (populasi studi). Populasi sasaran yaitu kumpulan dari satuan atau unit yang ingin dibuat inferensi atau generalisasinya. Populasi sampling adalah kumpulan dari satuan atau unit di mana kita mengambil sampel. Populasi sampling merupakan sebagian dari populasi target. Misalnya akan dilakukan penelitian tentang rata-rata jumlah konsumsi alkohol per minggu di kota A oleh anak remaja usia 15-17 tahun, maka yang menjadi target populasi adalah semua anak remaja yang berusia 15-17 tahun yang ada di kota A, dan populasi sampling adalah sekelompok anak remaja yang dipilih dari sebuah sekolah tertentu yang ada di kota A.

Alasan pertama dilakukannya pengambilan sampel adalah untuk dapat menggambarkan dan menarik kesimpulan tentang populasi. Populasi dapat berupa orang, objek atau karakteristik yang akan diobservasi. Kumpulan dari hasil observasi pada populasi merupakan ringkasan dari gambaran karakteristik populasi yang

disebut “parameter”, karakteristik yang sama dalam sampel disebut “statistik”. Statistik sampel akan menolong untuk menyimpulkan keadaan parameter populasi. Nilai dari parameter populasi bersifat konstan tetapi biasanya tidak diketahui, nilai statistik sampel dapat diketahui dengan mengukur sampel.

Parameter dan statistik adalah besaran yang berupa data ringkasan atau angka ringkasan yang menunjukkan suatu ciri dari populasi dan sampel. Parameter dan statistik merupakan hasil hitungan nilai dari semua unit di dalam populasi dan sampel bersangkutan.

Tabel 1. Lambang Parameter dan Statistik

Besaran	Lambang Parameter (Populasi)	Lambang Statistik (Sampel)
Rata-rata	M	X
Varians	σ^2	S^2
Simpangan Baku	O	S
Jumlah Observasi	N	n
Proporsi	P	p

1) Populasi

Populasi dapat dibagi menjadi dua yaitu populasi *finite* yaitu populasi yang jumlah anggotanya dapat dihitung, sedangkan populasi *infinite* yaitu populasi yang jumlah anggotanya tidak dapat dihitung.

Pengertian dapat dihitung agak abstrak, sebagai pendekatan jika populasi kurang dari 10.000 maka masuk populasi *finite*, sedangkan lebihnya disebut populasi *infinite*. Hasil pengukuran pada sampel akan dijadikan generalisasi untuk menggambarkan populasi, maka hendaknya batasan populasi harus jelas dari sisi dimensi ruang dan waktu

2) Sampel

Sampel adalah bagian populasi yang diambil dengan cara tertentu, dimana pengukuran dilakukan.

3) Unit Sampel

Unit sampel adalah kumpulan individu yang berasal dari populasi yang tidak saling tumpang tindih atau dengan kata lain mempunyai karakteristik atau ciri-ciri tertentu.

a) Elemen Sampel

Elemen sampel adalah individu berasal dari populasi, dimana pengukuran dilakukan kepadanya.

(1) Sampling Frame

Sampling frame adalah daftar populasi yang sangat tergantung dari kelompok atau individu penyusunnya.

(2) Variabel

Variabel adalah ciri-ciri yang melekat pada subyek yang diteliti dan mempunyai variasi dari hasil pengukurannya.

(3) Generalisasi

Generalisasi adalah upaya menarik kesimpulan dari data yang kecil (sampel) untuk menggambarkan keadaan yang ada pada populasi

(4) Random

Random adalah setiap anggota populasi mendapatkan Change (kesempatan) yang sama untuk menjadi anggota sampel.

Guna memberi gambaran yang lebih kongkret dari beberapa pengertian diatas, dapat dibuat ilustrasi sebagai berikut : "suatu penelitian ingin membuat perkiraan status gizi Balita di Puskesmas "X" tahun 2007, ukuran yang dipakai untuk menentukan status gizi adalah berat badan dan umur balita. Ilustrasi diatas dapat digunakan untuk mengidentifikasi pengertian diatas :

a) Populasi

(1) Populasi seluruh Balita yang di Puskesmas "X"

(2) Seluruh Kepala Keluarga yang mempunyai Balita di Puskesmas "X"

Pilihan diatas sangat tergantung pada data awal yang tersedia dilapangan, jika peneliti mendapatkan jumlah seluruh Balita di Puskesmas "X", maka populasinya pada pilihan "a", tetapi jika data dilapangan yang ada adalah daftar KK yang memiliki Balita, maka populasi pilihan "b".

b) Sampel

Sebagian Balita di Puskesmas "X"

c) Unit Sampel

KK yang mempunyai Balita

Pengertian unit sampel adalah sekumpulan individu, sehingga satu KK dengan KK lainnya boleh jadi memiliki Balita tidak sama, ada yang satu ada yang lebih dari satu. Jika KK yang mempunyai Balita lebih dari satu, terpilih jadi anggota sampel maka semua Balita yang ada dalam KK tersebut otomatis menjadi anggota sampel.

d) Elemen Sampel

Individu Balita yang akan dilakukan pengukuran BB dan umurnya.

e) Sampling Frame

Daftar populasi yang diapai ada dua yaitu berdasarkan KK yang punya Balita, berarti daftar populasinya berupa unit sampel, jika nama Balita yang menjadi daftar populasi, maka yang dipakai adalah elemen sampel.

f) Variabel

Ada dua variabel dalam penelitian ini, yaitu :

(1) BB/U yang digunakan ratio keduanya

(2) Status Gizi, hasil kategorisasi atau pengelompokkan ukuran BB/U

b. Sampel

Sampel adalah bagian dari jumlah dan karakteristik yang dimiliki oleh populasi. Bila populasi besar dan peneliti tidak mungkin mempelajari semua yang ada pada populasi, misalnya karena keterbatasan dana, tenaga dan waktu, maka peneliti dapat menggunakan sampel yang diambil dari populasi itu. Apa yang dipelajari dari sampel, kesimpulannya akan dapat diberlakukan untuk populasi. Untuk itu sampel yang diambil dari populasi harus betul-betul representatif (mewakili).

Bila sampel tidak representatif, maka ibarat orang buta disuruh menyimpulkan karakteristik gajah. Satu orang memegang telinga gajah, maka ia menyimpulkan gajah itu seperti kipas. Orang kedua memegang badan gajah, maka ia menyimpulkan gajah itu seperti tembok besar. Satu orang lagi memegang ekornya, maka ia menyimpulkan gajah itu kecil seperti seutas tali. Begitulah kalau sampel yang dipilih tidak representatif, maka ibarat 3 orang buta itu yang membuat kesimpulan salah tentang gajah.

Sampel yang dikehendaki merupakan bagian dari populasi target yang akan diteliti secara langsung yang memenuhi kriteria inklusi dan eksklusi.

1) Kriteria Inklusi

Kriteria inklusi merupakan karakteristik umum subyek penelitian pada populasi target dan sumber. Sering sekali ada kendala dalam memperoleh kriteria inklusi yang sesuai dengan masalah penelitian, biasanya masalah logistik. Dalam hal ini pertimbangan ilmiah sebagian harus dikorbankan karena alasan praktis.

Contoh penelitian : penelitian ingin mengetahui hubungan antara merokok dengan kejadian jantung koroner, maka orang yang boleh dijadikan dalam kelompok kasus pada penelitian ini adalah orang yang tidak menderita penyakit jantung yang lain.

2) Kriteria Eksklusi

Kriteria eksklusi merupakan kriteria dari subjek penelitian yang tidak boleh ada, dan jika subjek mempunyai kriteria eksklusi maka subjek harus dikeluarkan dari penelitian. Hal ini dikarenakan :

- a) Terdapat keadaan yang tidak mungkin dilaksanakan penelitian, misalnya subjek tidak mempunyai tempat tinggal
- b) Terdapat keadaan lain yang mengganggu dalam pengukuran maupun interpretasi

Contoh : penelitian ingin mengetahui hubungan antara merokok dengan kejadian jantung koroner, maka orang yang menderita penyakit jantung lain tidak boleh dijadikan dalam kelompok kasus pada penelitian ini

- c) Adanya hambatan etika
- d) Subjek menolak dijadikan responden

1) Keuntungan Penelitian Dilakukan terhadap Sampel :

- a) Efisien adalah hal biaya, waktu dan tenaga. Dengan meneliti sampel yang jumlah lebih sedikit maka akan lebih murah, lebih cepat mendapatkan data dan tidak memerlukan tenaga yang banyak
- b) Lebih akurat, dalam hal pengukuran jika subjeknya sedikit maka hasil pengukurannya akan lebih akurat dibandingkan dengan mengukur banyak subjek.
- c) Lebih tajam dan mendalam, sebagian penyakit mempunyai manifestasi yang bervariasi, kemudian dengan seleksi sampel akan diperoleh subjek dengan karakteristik tertentu, sehingga akan diperoleh data pada kelompok subjek yang lebih homogen.

c. Teknik Sampling

1) Metode Probability Sampling

a) Pengambilan Sampel Acak Sederhana (*Simple Random Sampling*)

Dalam metode ini pengambilan sampel dilakukan sedemikian rupa sehingga setiap unit dasar penelitian mempunyai kesempatan yang sama untuk diambil sebagai sampel. Terpilihnya setiap unit tersebut kedalam sampel harus benar-benar berdasarkan faktor kebetulan (chance), bebas dari subjektivitas peneliti atau orang lain.

Keuntungan metode ini adalah ketepatan yang tinggi dan setiap unit sampel mempunyai kesempatan yang sama untuk diambil, murah dan mudah dilaksanakan bila populasi sedikit dan homogen, sedangkan kerugiannya adalah bila tidak ada kerangka sampling dan populasi menyebar tidak merata atau populasi yang luas dengan kondisi alam yang tidak menunjang, maka akan sulit dilaksanakan dan bila sifat populasi tidak homogen maka akan terjadi bias.

Contoh :

Misal dari populasi kepala keluarga yang dianggap homogen sebanyak 1000 orang diambil sampel sebanyak 50 orang dengan menggunakan tabel random. Pertama, buat kerangka sampling yaitu daftar nama kepala keluarga yang diberi nomor 0001 hingga 1000. Untuk pemberian nomor, perlu diperhatikan jumlah digit di populasi, karena besar populasi adalah 1000 maka jumlah digit adalah 4. Maka nomor awal dimulai dengan 0001 bukan 1, 01 atau 001. Ini untuk mempertahankan prinsip "equal probability". Selanjutnya peneliti bisa menggunakan tabel bilangan random. Cara menggunakan tabel bilangan random, cari angka di bawah angka seribu dan atau sama dengan 1000 sebanyak 50 angka. Angka yang diambil dua kali yang sama, maka angka yang sama tadi dibuang yang terpilih sebanyak 50 dari tabel random, itulah yang menjadi responden dalam penelitian.

b) Pengambilan Sampel Acak Sistematis (*Systematic Random Sampling*)

Metode pengambilan sampel ini hanya unsur pertama saja dari sampel yang dipilih secara acak, sedangkan unsur-unsur selanjutnya dipilih secara

sistematis menurut pola tertentu. Metode ini dapat dilakukan dengan syarat populasi harus besar, harus tersedia kerangka sampling dan populasi bersifat homogen.

Keuntungan metode ini adalah caranya relatif mudah dan dapat dilakukan oleh petugas lapangan, caranya praktis bila populasi dalam bentuk kartu, variasi akan lebih kecil dibandingkan dengan cara lain dan membutuhkan waktu, biaya yang relatif lebih rendah dibandingkan dengan simple random sampling. Kerugian metode ini adalah setiap unit sampel tidak mempunyai peluang yang sama untuk diambil sebagai sampel dan bila terdapat suatu kecenderungan tertentu maka metode ini menjadi kurang sesuai.

Contoh :

Jika benar populasi adalah $N = 1000$, besar sampel adalah $n = 50$, maka $k = \text{interval} = N/n = 1000/50 = 20$. Untuk pertama kali ambil satu angka secara random antara 1 hingga ke k . Dalam contoh diatas pilih satu angka secara random antara angka 1 hingga angka 20. Misalnya, terpilih angka 4 maka dari daftar kerangka sampling nomor 1 hingga nomor 1000, angka yang terpilih pertama kali adalah angka 4. Angka selanjutnya yang dipilih adalah $4 + 20 = 24$, angka berikutnya adalah $24 + 20 = 54$. Demikian seterusnya hingga sampel sebanyak 50.

c) Pengambilan Sampel Acak Stratifikasi (*Stratified Random Sampling*)

Dalam metode ini populasi yang heterogen dibagi ke dalam lapisan-lapisan (strata) yang seragam. Hal ini dilakukan karena dalam praktek sering dijumpai populasi yang tidak homogen. Makin heterogen suatu populasi, makin besar pula perbedaan sifat antara strata. Syarat penggunaan metode ini adalah harus ada kriteria yang jelas yang akan digunakan sebagai dasar untuk menstratifikasi populasi dalam lapisan-lapisan tersebut. Kemudian harus ada data pendahuluan dari populasi mengenai kriteria yang dipergunakan untuk menstratifikasi dan harus diketahui dengan tepat jumlah satuan-satuan elementer dari tiap stratum dalam populasi tersebut. Keuntungan metode ini adalah semua ciri-ciri populasi yang heterogen dapat terwakili, kemungkinan bagi peneliti untuk dapat meneliti hubungan antara tiap strata dan membandingkannya. Kerugian metode ini adalah harus diketahui terlebih dahulu kondisi populasi yang kadang sering tidak diketahui, agar dapat dilakukan stratifikasi dengan baik dan metode ini sulit untuk membuat kelompok yang homogen.

Contoh :

Suatu populasi Kepala Keluarga dengan $N = 1000$ terbagi menjadi 3 strata, KK Kaya, $N_1 = 100$, KK Cukup Kaya dimana $N_2 = 250$ dan KK Miskin di mana $N_3 = 650$. Dari masing-masing stratum dibuat kerangka sampling. Untuk KK Kaya dari nomor 001 hingga 100, untuk KK Cukup Kaya dari nomor 001 hingga nomor 250, dan KK Miskin dari nomor 001 hingga 650. Misalkan

diambil sampel sebanyak 50 orang kepala keluarga, maka untuk masing-masing stratum diambil sebesar

Untuk KK Kaya :

$$n_1 = N_1/N \times n = 100/1000 \times 50 = 5$$

selanjutnya ambil 5 dari 100 dengan teknik sampling random sederhana dari kerangka sampling yang tersedia.

Untuk KK Cukup Kaya :

$$n_2 = N_2/N \times n = 250/1000 \times 50 = 12,5 = 13$$

selanjutnya diambil 13 dari 250 dengan teknik random sederhana dari kerangka sampling yang tersedia.

Untuk KK Miskin :

$$n_3 = N_3/N \times n = 650/1000 \times 50 = 32,5 = 33$$

selanjutnya diambil 33 dari 650 dengan teknik random sederhana dari kerangka sampling yang tersedia.

d) Pengambilan Sampel Acak Berkelompok (*Cluster Random Sampling*)

Dilapangan sering kita diahapkan dengan kenyataan di mana kerangka sampling digunakan untuk dasar pemilihan sampel tidak tersedia atau tidak lengkap, dan biaya untuk membuat kerangka sampel tersebut terlalu besar. Untuk mengatasi hal tersebut, maka unit-unit analisa dalam populasi digolongkan ke dalam gugus-gugus (*cluster*) dan ini merupakan satuan-satuan di mana sampel akan diambil. Jumlah *cluster* yang diambil sebagai sampel harus secara acak. Kemudian untuk unit penelitian dalam *cluster* tersebut diteliti semua. Pada metode ini unit samplingnya terdiri dari satu elemen populasi, di mana setiap unit sampling adalah gugusan atau group dari elemen populasi.

Keuntungan metode ini adalah tidak diperlukannya daftar kerangka sampling dengan unsur-unsurnya, tetapi kekurangannya sangat sulit untuk menghitung standar kesalahan (*standar error*).

Contoh :

Satuan Blok perumahan merupakan rumpun yang terdiri dari bangunan perumahan yang dialami KK. Umpannya blok perumahan ada 5 blok perumahan. Blok 1 = 50 kk, Blok 2 = 75 kk dan Blok 3 = 65 kk. Untuk mendapatkan sampel yang akan dijadikan objek penelitian maka 3 blok KK ini akan di random secara undian dengan cara membuat gulungan kertas, ada 3 gulungan kertas yang ditandai dengan blok 1, 2, 3. Kemudian gulungan kertas dicabut. Gulungan kertas yang tercabut maka itulah yang akan menjadi sampel penelitian. Misalnya tercabut gulungan kertas yang ditandai Blok 2 dengan KK = 75, maka sampel penelitian adalah blok 2 perumahan dengan jumlah KK 75.

Dapat pula dilakukan dengan cara sejumlah individu yang heterogen dipilih menjadi kelompok (blok) yang anggotanya homogen. Kemudian perlakuan A, B dan C dialokasikan secara random ke individu di dalam blok yang sama.

e) Pengambilan Sampel Acak Bertahap (*Multistage Random Sampling*)

Populasi yang letaknya sangat tersebar geografis, sehingga sangat sulit untuk mendapatkan kerangka sampling dari unit populasi. Untuk mengatasi hal ini mata unit-unit analisa dikelompokkan ke dalam *cluster* yang merupakan satuan di mana sampel akan diambil. Pengambilan sampel dilakukan melalui tahap tertentu. Populasi dalam *cluster* tingkat pertama kemudian *cluster* pertama dibagi menjadi *cluster* tingkat kedua dan gugus tingkat kedua dapat lagi menjadi *cluster* tingkat lebih lanjut.

Masalah pertama dalam pemilihan sampel rumpun dua tahap adalah pemilihan yang sesuai. Ada dua hal yang perlu dipertimbangkan. Pertama, pendekatan geografis dari elemen di dalam suatu rumpun. Kedua, ukuran rumpun yang melegakan bagi pelaksana. Pemilihan rumpun yang sesuai tergantung pada apakah peneliti menginginkan sampel dengan sedikit rumpun namun jumlah elemen di dalam rumpun yang banyak atau banyak rumpun yang diambil dan sesedikit mungkin jumlah elemen di dalam masing-masing rumpun.

f) Metode Non-Probability Sampling

(1) Sampling Kuota (*Quota Sampling*)

Sampling kuota adalah teknik untuk menentukan sampel dari populasi yang mempunyai ciri-ciri tertentu sampai jumlah (kuota) yang diinginkan. Pengambilan sampel secara kuota dilakukan dengan cara menetapkan sejumlah anggota sampel secara *quotum* atau jatah. Teknik sampling ini dilakukan dengan cara : pertama, menetapkan beberapa besar jumlah sampel yang diperlukan atau menetapkan *quotum* (jatah). Kemudian jumlah atau *quotum* itulah yang dijadikan dasar untuk mengambil unit sampel yang diperlukan. Anggota populasi mana pun yang akan diambil tidak menjadi soal, yang penting jumlah *quotum* yang sudah ditetapkan dapat dipenuhi.

Contoh :

Melakukan penelitian tentang pendapat masyarakat terhadap pelayanan masyarakat dalam urusan Ijin Mendirikan Bangunan (IMB). Jumlah sampel yang ditentukan 500 orang. Kalau pengumpulan data belum memenuhi kuota 500 orang tersebut, maka penilaian dipandang belum selsai. Bila pengumpulan data dilakukan secara kelompok yang terdiri atas 5 orang pengumpul data, maka setiap anggota kelompok harus dapat menghubungi 100 orang anggota sampel atau 5 orang tersebut harus dapat mencari data dari 500 anggota sampel.

(2) Sampling Sistematis (*Systematic Sampling*)

Sampling sistematis adalah teknik pengambilan sampel berdasarkan urutan dari anggota populasi yang telah diberi nomor urut.

Contoh :

Anggota populasi yang terdiri dari 100 orang, dari semua anggota itu diberi nomor urut yaitu nomor 1 sampai dengan nomor 100. Pengambilan sampel dapat dilakukan dengan mengambil nomor ganjil saja, genap saja atau kelipatan dan bilangan tertentu. Misalnya kelipatan dari bilangan lima, untuk ini maka yang diambil sebagai sampel adalah nomor 1, 5, 10, 15, 20, 25 dan seterusnya sampai 100.

(3) Sampling Kebetulan (*Accidental Sampling*)

Sampling kebetulan adalah teknik penentuan sampel berdasarkan kebetulan yaitu siapa saja yang secara kebetulan bertemu dengan peneliti dapat digunakan sebagai sampel, bila dipandang orang yang kebetulan ditemui itu cocok sebagai sumber data. Sampling kebetulan berarti sampel yang diambil dari responden atau kasus yang kebetulan ada di suatu tempat atau keadaan tertentu.

Contoh :

Penelitian tentang pemberian ASI oleh ibu-ibu wilayah kerja Puskesmas Pasar Minggu, maka sampel penelitian ini dapat diambil dari ruang KIA tempat pemeriksaan kesehatan ibu dan anak di Puskesmas Pasar Minggu selama periode tertentu misalnya pada bulan Juni 2006 hari senin sampai sabtu atau setiap hari selasa selama bulan Juni 2006. Seberapa banyak pun ibu-ibu yang ditemui pada hari yang ditentukan tersebut menjadi sampel penelitian ini.

(4) Sampling Purposive

Sampling purposive adalah teknik penentuan sampel dengan pertimbangan tertentu yang dibuat oleh peneliti sendiri, berdasarkan ciri atau sifat populasi yang sudah diketahui sebelumnya. Pelaksanaan pengambilan sampel secara *purposive* yaitu peneliti mengidentifikasi semua karakteristik populasi dengan melakukan studi pendahuluan atau dengan mempelajari berbagai hal yang berhubungan dengan populasi. Kemudian peneliti menetapkan berdasarkan petimbangannya sebagian dari anggota populasi menjadi sampel penelitian sehingga teknik pengambilan sampel secara *purposive* ini didasarkan pertimbangan pribadi peneliti sendiri.

Contoh :

Penelitian tentang pola asuh terhadap Balita di Desa "X" kriteria yang dibuat oleh peneliti meliputi : Anak pertama, Mempunyai KMS
Maka Balita yang memenuhi persyaratan tersebut, akan menjadi anggota sampel. Sedikit berbeda dengan kriteria Inklusi, kriterianya

dibuat oleh peneliti dari subyek yang memenuhi kriteria tersebut baru diambil sampel daripadanya sehingga generalisasinya bisa pada semua subyek yang memenuhi kriteria Inklusi.

(5) Sampling Jenuh

Sampling jenuh adalah teknik penentuan sampel bila semua anggota populasi sebagai sampel. Hal ini sering dilakukan bila jumlah populasi relatif kecil kurang dari 30 orang atau penelitian yang ingin membuat generalisasi dengan kesalahan yang sangat kecil. Istilah lain sampel jenuh adalah sensus, dimana semua anggota populasi dijadikan sampel.

(6) Snowball Sampling

Snowball sampling adalah teknik penentuan sampel yang mula-mula jumlahnya kecil, kemudian membesar. Dalam penentuan sampel, pertama dipilih 1 atau 2 orang tetapi karena dengan 2 orang ini belum merasa lengkap terhadap data yang diberikan maka peneliti mencari orang lain yang dipandang lebih tahu dan dapat melengkapi data yang diberikan oleh 2 orang sebelumnya. Pada penelitian kualitatif banyak menggunakan sampel *purposive* dan *snowball*, misalnya akan meneliti siapa provokator kerusuhan atau jaringan teroris maka akan cocok menggunakan *purposive* dan *snowball* sampling.

Jumlah anggota sampel sering dinyatakan dengan ukuran sampel. Jumlah sampel yang diharapkan 100% mewakili populasi adalah sama dengan jumlah anggota populasi itu sendiri. Jadi bila jumlah populasi 1000 dan hasil penelitian itu akan diberlakukan untuk 1000 orang tersebut tanpa ada kesalahan, maka jumlah sampel yang diambil sama dengan jumlah populasi tersebut yaitu 1000 orang. Makin besar jumlah sampel mendekati populasi, maka peluang kesalahan generalisasi semakin kecil dan sebaliknya makin kecil jumlah sampel menjauhi populasi, maka makin besar kesalahan generalisasi (diberlakukan umum).

C. Rangkuman

Populasi juga diartikan keseluruhan individu yang menjadi acuan hasil-hasil penelitian akan berlaku. Populasi dapat dibagi menjadi dua kelompok yaitu populasi sasaran (target populasi) dan populasi sampling (populasi studi). Populasi sasaran yaitu kumpulan dari satuan atau unit yang ingin dibuat inferensi atau generalisasinya. Populasi sampling adalah kumpulan dari satuan atau unit di mana kita mengambil sampel. Populasi sampling merupakan sebagian dari populasi target.

Sampel adalah bagian dari jumlah dan karakteristik yang dimiliki oleh populasi. Apa yang dipelajari dari sampel, kesimpulannya akan dapat diberlakukan untuk populasi. Untuk itu sampel yang diambil dari populasi harus betul-betul representatif (mewakili). Teknik pengambilan sampel dibagi menjadi. Metode probability sampling yang terdiri atas pengambilan sampel acak sederhana (*Simple Random Sampling*), pengambilan sampel acak sistematis (*Systematic Random Sampling*), pengambilan sampel acak stratifikasi (*Stratified*

Random Sampling), pengambilan sampel acak berkelompok (*Cluster Random Sampling*) dan pengambilan sampel acak bertahap (*Multistage Random Sampling*). Adapun untuk metode non-probability sampling terdiri atas sampling kuota (*Quota Sampling*), sampling sistematis (*Systematic Sampling*), sampling kebetulan (*Accidental Sampling*), *sampling purposive*, *sampling jenuh* dan *snowball sampling*.

D. Tugas

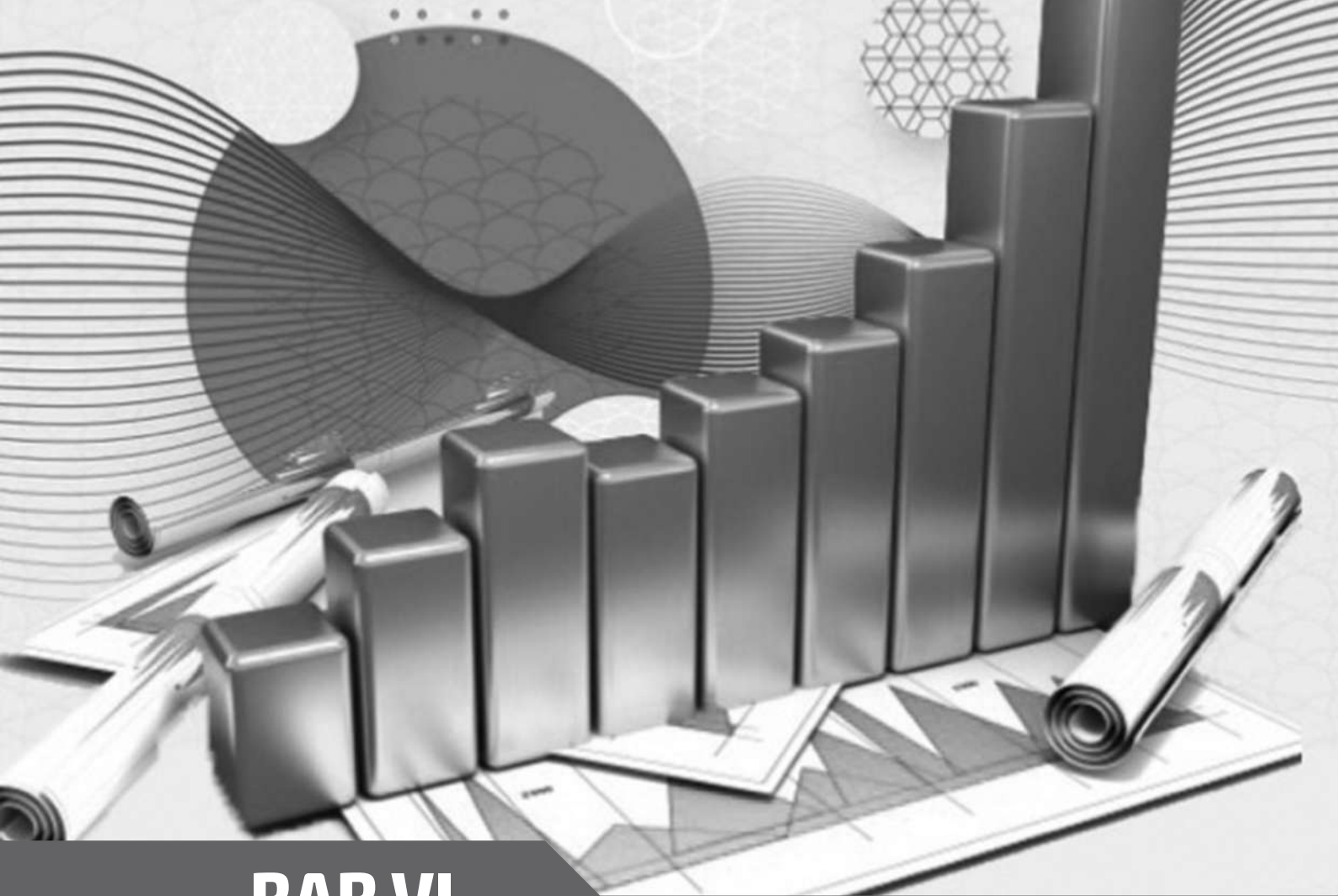
1. Seorang mahasiswa akan menyusun sebuah penelitian sebagai salah satu syarat dalam kelulusannya. Berbagai referensi dibaca dan proses diskusi dilakukan untuk menemukan variabel yang tepat untuk diteliti.
Apakah istilah dari definisi wilayah generalisasi yang terdiri dari objek atau subjek dengan kualitas dan karakteristik tertentu dan ditetapkan peneliti untuk dipelajari kemudian ditarik kesimpulannya?
 - a. Wilayah penelitian
 - b. Populasi
 - c. Sampel
 - d. Tinjauan pustaka
 - e. Definisi operasional
2. Seorang mahasiswa akan menyusun sebuah penelitian sebagai salah satu syarat dalam kelulusannya. Berbagai referensi dibaca dan proses diskusi dilakukan untuk menemukan variabel yang tepat untuk diteliti.
Siapakah yang menjadi sasaran dalam sebuah penelitian ilmiah?
 - a. Wilayah penelitian
 - b. Populasi
 - c. Sampel
 - d. Populasi terbatas
 - e. Sampel terbatas
3. Seorang mahasiswa akan menyusun sebuah penelitian sebagai salah satu syarat dalam kelulusannya. Berbagai referensi dibaca dan proses diskusi dilakukan untuk menemukan variabel yang tepat untuk diteliti.
Dalam suatu penelitian, disarankan untuk menggunakan sampel dibandingkan populasi. Apakah salah satu alasan ilmiah yang paling tepat dari pernyataan tersebut?
 - a. Mudah dilakukan
 - b. Proses penelitian lebih cepat
 - c. Hemat waktu, biaya dan tenaga
 - d. Publikasi dapat dilakukan dengan mudah
 - e. Tidak tersedianya data yang lengkap
4. Seorang mahasiswa akan menyusun sebuah penelitian sebagai salah satu syarat dalam kelulusannya. Berbagai referensi dibaca dan proses diskusi dilakukan untuk menemukan variabel yang tepat untuk diteliti. Diketahui bahwa sebuah penelitian ilmiah akan dilakukan pada lansia di Kota Pangkalpinang Kecamatan X Provinsi Bangka Belitung. Menurut kalian siapakah populasi dalam penelitian tersebut?

- a. Seluruh lansia di Provinsi Bangka Belitung
 - b. Lansia di Kota Pangkalpinang
 - c. Lansia di Kecamatan x Kota Pangkalpinang
 - d. Seluruh masyarakat di Provinsi Bangka Belitung
 - e. Seluruh masyarakat di Kecamatan X
5. Seorang mahasiswa akan menyusun sebuah penelitian sebagai salah satu syarat dalam kelulusannya. Berbagai referensi dibaca dan proses diskusi dilakukan untuk menemukan variabel yang tepat untuk diteliti. Diketahui bahwa sebuah penelitian ilmiah akan dilakukan pada lansia di Kota Pangkalpinang Kecamatan X Provinsi Bangka Belitung. Menurut kalian siapakah sampel dalam penelitian tersebut?
- a. Seluruh lansia di Provinsi Bangka Belitung
 - b. Lansia di Kota Pangkalpinang
 - c. Lansia di Kecamatan X Kota Pangkalpinang
 - d. Seluruh masyarakat di Provinsi Bangka Belitung
 - e. Seluruh masyarakat di Kecamatan X
6. Diketahui bahwa sebuah penelitian dilakukan pada mahasiswa Poltekkes Pangkalpinang untuk mengetahui gambaran status gizinya berdasarkan indikator IMT/U. Apakah salah satu kriteria inklusi dalam penelitian tersebut?
- a. Bertempat tinggal di wilayah Poltekkes Pangkalpinang
 - b. Mahasiswa berumur > 18 tahun
 - c. Mampu baca tulis
 - d. Mahasiswa Jurusan Gizi
 - e. Mampu berdiri tegak
7. Seorang mahasiswa sedang menyusun proposal penelitian tentang sisa makanan pasien di rumah sakit. Proses pengumpulan data akan dilakukan dengan cara menimbang sisa makanan pasien. Apakah yang menjadi kriteria eksklusi dalam kasus tersebut?
- a. dilakukan 3 kali dalam sehari
 - b. pasien kondisi koma
 - c. mau terlibat dalam penelitian
 - d. bisa berkomunikasi baik
 - e. pasien rawat inap
8. Seorang peneliti akan menentukan populasi dan sampel dalam penelitian di sebuah sekolah menengah atas. Sampel yang dipilih adalah remaja putri yang mendapat dan minum tablet tambah darah (TTD). Siapakah populasi dalam penelitian tersebut?
- a. seluruh murid
 - b. seluruh murid dan guru
 - c. semua remaja putri usia > 17 tahun
 - d. semua remaja putri
 - e. semua remaja putri yang mendapat TTD

9. Seorang mahasiswa akan menyusun sebuah penelitian sebagai salah satu syarat dalam kelulusannya. Berbagai referensi dibaca dan proses diskusi dilakukan untuk menemukan variabel yang tepat untuk diteliti. Dalam teknik sampling, diketahui bahwa semua subjek dalam populasi mempunyai kesempatan yang sama untuk dipilih sebagai sampel. Apakah salah satu teknik yang digunakan?
- Cluster sampling
 - Sampling accidental
 - Snowball sampling
 - Purposif sampling
 - Sampling kuota
10. Seorang mahasiswa akan menyusun sebuah penelitian sebagai salah satu syarat dalam kelulusannya. Berbagai referensi dibaca dan proses diskusi dilakukan untuk menemukan variabel yang tepat untuk diteliti. Apakah teknik sampling yang digunakan oleh peneliti yang menggunakan cara membagi jumlah seluruh anggota populasi dengan jumlah sampel yang dibutuhkan sebagai interval sampel?
- Snowball sampling
 - Sistematik random sampling
 - Simpel random sampling
 - Stratified random sampling
 - Sampling kuota

E. Referensi

- Hasmi. (2012). Metode Penelitian Epidemiologi. Jakarta: TIM.
- Kasjono, Heru Subaris ; Yasril. (2009). Teknik Sampling Untuk Penelitian Kesehatan. Yogyakarta: Graha Ilmu.
- Notoatmodjo, Soekidjo. (2012). Metodologi Penelitian Kesehatan. Jakarta: Rineka Cipta.
- Riyanto, Agus. (2011). Aplikasi Metodologi Penelitian Kesehatan. Yogyakarta: Nuha Medika.
- Santjaka, Aris. (2011). Statistik Untuk Penelitian Kesehatan. Yogyakarta: Nuha Medika.
- Sugiyono. (2009). Statistika Untuk Penelitian. Bandung: Cv. Alfabeta.
- Sugiyono. (2019). Metode Penelitian Kuantitatif, Kualitatif dan R&D. Bandung: Cv. Alfabeta.



BAB VI

TEKNIK PEMILIHAN STATISTIK

A. Tujuan Pembelajaran

Setelah mempelajari materi ini mahasiswa diharapkan mampu memahami Teknik p

B. Materi

1. Pengertian Teknik Analisis Data

Penelitian merupakan kegiatan yang terencana untuk mencari jawaban yang obyektif atas permasalahan manusia melalui prosedur ilmiah. Untuk itu di dalam suatu penelitian dibutuhkan suatu proses analisis data yang berguna untuk menganalisis data-data yang telah terkumpul. Data yang terkumpul banyak sekali dan terdiri dari berbagai catatan di lapangan, gambar, foto, dokumen, laporan, biografi, artikel, dan sebagainya. Pekerjaan analisis data dalam hal ini ialah mengatur, mengurutkan, mengelompokkan, memberikan kode, dan mengategorikannya. Pengorganisasian dan pengelolaan data tersebut bertujuan menemukan tema dan hipotesis kerja yang akhirnya diangkat menjadi teori substantif oleh karena itu, analisis data merupakan bagian yang amat penting karena dengan analisislah suatu data dapat diberi arti dan makna yang berguna untuk masalah penelitian. Data yang telah dikumpulkan oleh peneliti tidak akan ada gunanya apabila tidak dianalisis terlebih dahulu.

Data ialah bahan mentah yang perlu di olah sehingga menghasilkan informasi atau keterangan, baik kualitatif maupun kuantitatif yang menunjukkan fakta. Sementara perolehan data seyogyanya relevan, artinya data yang ada hubungannya langsung dengan masalah penelitian. Analisa data berasal dari gabungan dari dua buah kata yaitu "analisis" dan "data".

Analysis is process of resolving data into its constituent component to reveal its characteristic elements and structure (Dey, 1995). Analisis merupakan evaluasi dari sebuah situasi dari sebuah permasalahan yang dibahas, termasuk didalamnya peninjauan dari berbagai aspek dan sudut pandang, sehingga tidak jarang ditemui permasalahan besar dapat dibagi menjadi komponen yang lebih kecil sehingga dapat diteliti dan ditangani lebih mudah, sedangkan data adalah fakta atau bagian dari fakta yang mengandung arti yang dihubungkan dengan kenyataan, simbol-simbol, gambar-gambar, kata-kata, angka-angka atau huruf-huruf yang menunjukkan suatu ide, obyek, kondisi atau situasi dan lain-lain.²

Analisis data adalah proses mencari dan menyusun secara sistematis data yang diperoleh dari hasil wawancara, catatan lapangan, dan dokumentasi, dengan cara mengorganisasikan data ke dalam kategori, menjabarkan ke dalam unit-unit, melakukan sintesa, menyusun ke dalam pola, memilih mana yang penting dan yang akan dipelajari, dan membuat kesimpulan sehingga mudah difahami oleh diri sendiri maupun orang lain.³

Menurut Lexy Y. Moleong (2002) menjelaskan bahwa analisis data adalah proses mengatur urutan data, mengorganisasikannya ke dalam suatu pola, kategori, dan satuan uraian dasar. Analisis data sebagai proses yang merinci usaha secara formal untuk menemukan tema dan merumuskan hipotesis (ide) seperti yang disarankan dan sebagai

usaha untuk memberikan bantuan dan tema pada hipotesis. Jika dikaji, pada dasarnya definisi pertama lebih menitikberatkan pengorganisasian data sedangkan yang ke dua lebih menekankan maksud dan tujuan analisis data. Dengan demikian definisi tersebut dapat disintesiskan bahwa analisis data merupakan proses mengorganisasikan dan mengurutkan data ke dalam pola, kategori dan satuan uraian dasar sehingga dapat ditemukan tema dan dapat dirumuskan hipotesis kerja seperti yang didasarkan oleh data (Ardhana, 2008)

Dalam proses analisis data dimulai dengan menelaah seluruh data yang tersedia dari berbagai sumber, yaitu dari wawancara, pengamatan yang sudah dituliskan dalam catatan lapangan, dokumen pribadi, dokumen resmi, gambar, foto, dan sebagainya. Untuk mengetahui lebih lanjut tentang analisis data, dalam makalah ini akan membahas pengertian analisis data, jenis-jenis analisis data, teknik-teknik analisis data, dan langkah-langkah analisis data.

2. Jenis – Jenis Analisis Data

Analisis dalam penelitian merupakan bagian dalam proses penelitian yang sangat penting, karena dengan analisa inilah data yang ada akan nampak manfaatnya terutama dalam memecahkan masalah penelitian dan mencapai tujuan akhir penelitian (Subagyo, 2011). Data yang belum dianalisis masih merupakan data mentah. Dalam kegiatan penelitian, data mentah akan memberi arti, bila dianalisis dan ditafsirkan.

Dalam rangka analisis dan interpretasi data, perlu dipahami tentang keberadaan data itu sendiri. Secara garis besar, keberadaan data dapat digolongkan ke dalam dua jenis, yaitu :

a) Data Kualitatif

Data kualitatif disebut juga dengan data lunak. Data semacam ini diperoleh melalui penelitian yang menggunakan pendekatan kualitatif, atau penilaian kualitatif. Keberadaan data bermuatan kualitatif adalah catatan lapangan yang berupa catatan atau rekaman kata-kata, kalimat, atau paragraf yang diperoleh dari wawancara menggunakan pertanyaan terbuka, observasi partisipatoris, atau pemaknaan peneliti terhadap dokumen atau peninggalan. Untuk memperoleh arti dari data semacam ini melalui interpretasi data, digunakan teknik analisis data kualitatif.

b) Data Kuantitatif

Data kuantitatif adalah angka-angka (kuantitas), baik diperoleh dari jumlah suatu penggabungan ataupun pengukuran. Data bermuatan kuantitatif yang diperoleh dari jumlah suatu penggabungan selalu menggunakan bilangan cacah. Contoh data seperti ini adalah angka-angka hasil sensus, angka-angka hasil tabulasi terhadap jawaban terhadap angket atau wawancara terstruktur. Adapun data bermuatan kuantitatif hasil pengukuran adalah skor-skor yang diperoleh melalui pengukuran, seperti skor tes prestasi belajar, skor skala motivasi, skor timbangan, dan sebagainya.

3. Teknik Analisis Data

Teknik analisis data merupakan cara menganalisis data penelitian, termasuk alat-alat statistik yang relevan untuk digunakan (Noor, 2012). Dalam hal teknik analisis data, penelitian kualitatif dengan penelitian kuantitatif juga memiliki beberapa perbedaan. Dalam analisis data kuantitatif, teknik analisis datanya sangat bervariasi tergantung kepada tujuan penelitian, hipotesis penelitian, dan jenis data yang diperoleh. Teknik statistik dengan menggunakan peranti lunak komputer sering kali digunakan untuk mempermudah peneliti dalam melakukan analisis data karena bentuk data yang berupa angka, lebih bersifat universal, bebas budaya (culture free), dan lebih objektif serta tidak bermakna ganda (Herdiansyah, 2011). Dalam teknik analisis data menggunakan statistik, terdapat dua macam statistik yang digunakan pada data kuantitatif, yaitu statistik deskriptif dan inferensial.

a) Statistik Deskriptif

Adalah bidang statistik yang berhubungan dengan metode pengelompokan, peringkasan, dan penyajian data dalam cara yang lebih informatif. Pada statistik jenis ini kita melakukan teknik statistik yang berhubungan dengan penyajian data statistik dalam bentuk gambaran angka-angka (Santosa, 2005). Yang termasuk dalam statistik deskriptif antara lain distribusi frekuensi, distribusi persen dan pengukuran tendensi sentral.

- Tabel distribusi frekuensi yaitu menggambarkan pengaturan data secara teratur didalam suatu tabel. Data diatur secara berurutan sesuai besar kecilnya angka atau digolongkan didalam kelas-kelas yang sesuai dengan tingkatan dan jumlah yang sesuai didalam kelas.

Contoh tabel distribusi frekuensi :

Apakah Saudara mendapatkan penyuluhan tentang kesehatan ?

Jawaban	Frekuensi
Pernah	110
Tidak Pernah	90
Jumlah	200

Artinya : ada sebanyak 110 individu yang memilih "pernah" mendapatkan penyuluhan dan 90 yang memilih "tidak pernah" mendapatkan penyuluhan kesehatan.

- Distribusi persen adalah pengaturan data yang dihitung dalam bentuk persen. Cara memperoleh frekuensi relatif ialah :

$$\frac{\text{Frekuensi masing-masing individu} \times 100\%}{\text{jumlah frekuensi}}$$

Umur	Frekuensi	Presentase
<25	121	37%
26 – 30	59	18%
31 – 40	83	25%
>40	66	20%
Jumlah	329	100%

Artinya : ada sebanyak 37% responden berusia <25 tahun, 18% berusia antara 26-30 tahun dan seterusnya.

- Pengukuran tendensi sentral
Cara lain menggambarkan statistik deskriptif ialah dengan menggunakan tendensi sentral. Contoh bilangan tendensi sentral ialah mean (rata-rata), median dan mode. Tendensi sentral berguna untuk menggambarkan bilangan yang dapat mewakili suatu kelompok bilangan tertentu.
- Mean
Dapat dicari dengan menjumlahkan semua nilai kemudian dibagi dengan banyaknya individu. Rumusnya :

Dimana M = mean; X = jumlah data dan N = jumlah individu.

Contoh Ada 5 orang dengan penghasilan sbb:

Individu	Penghasilan Dalam Ribuan (Rp.)
A	100
B	125
C	140
D	150
E	175
$N = 5$	$\sum X = 690$

- Mode
Mode merupakan nilai yang jumlah frekuensinya paling besar. Untuk mencari nilai mode dapat dilihat pada jumlah frekuensi yang paling besar.

Contoh :

Nilai	Frekuensi
60	5
65	6
66	7
70	15
72	2
75	6
80	8
85	10

- Median

Merupakan nilai tengah yang membatasi setengah frekuensi bagian bawah dan setengah frekuensi bagian atas.

Contoh :

Nomor	Nilai
1	60
2	65
3	70
4	75
5	85
6	80
7	81
8	79
9	77

85 adalah *median* yang membagi empat nilai di atasnya dan empat nilai di bawahnya.

b) Statistik Inferensial

Adalah teknik statistik yang berhubungan dengan analisis data untuk penarikan kesimpulan atas data. Teknik statistik inferensial berhubungan dengan pengolahan statistik sehingga dengan menggunakan hasil analisis tersebut kita dapat menarik kesimpulan atas karakteristik populasi (Santosa, 2005).

1) Parametrik

Statistik parametrik adalah cabang ilmu statistik inferensial yang digunakan untuk menganalisis data-data yang memiliki sebaran normal saja. Diartikan pula ilmu statistik yang berhubungan dengan inferensi statistik yang membahas parameter-parameter populasi; jenis data interval atau rasio; distribusi data normal atau mendekati normal (Asep, tt). Statistik parametrik tidak dapat dipergunakan sebagai metode statistik apabila data yang akan dianalisis tidak menyebar secara normal. Dengan kata lain, data yang ingin di analisis harus ditransformasikan terlebih dahulu. Transformasi yang dimaksud adalah data ubah mengikuti sebaran normal. Transformasi dapat dilakukan dengan mengubah data ke dalam bentuk logaritma natural, menggunakan operasi matematik (membagi, menambah, atau mengali dengan bilangan tertentu), dan mengubah skala data dari nominal menjadi interval. Spesifikasi ini disebabkan karena metode statistik parametrik memiliki tingkat akurasi ketepatan yang lebih tinggi dibandingkan statistik non parametrik (akan dijelaskan selanjutnya). Untuk itulah penyajian data dengan sebaran normal harus dilakukan untuk mendapatkan analisis data yang akurat. Contoh statistik parametrik yaitu Normalitas, Homogenitas, Uji T, dan Anova.

2) Non-parametrik

Statistik nonparametrik disebut juga statistik bebas sebaran. Statistik nonparametrik tidak mensyaratkan bentuk sebaran parameter populasi. Statistik nonparametrik dapat digunakan pada data yang memiliki sebaran normal atau tidak. Statistik nonparametrik biasanya digunakan untuk melakukan analisis pada data nominal atau ordinal. Keunggulan dari statistik nonparametrik yaitu, tidak membutuhkan asumsi normalitas; secara umum metode statistik non-parametrik lebih mudah dikerjakan dan lebih mudah dimengerti jika dibandingkan dengan statistik parametrik karena ststistika non-parametrik tidak membutuhkan perhitungan matematik yang rumit seperti halnya statistik parametrik; statistik non-parametrik dapat digantikan data numerik (nominal) dengan jenjang (ordinal); kadang-kadang pada statistik non-parametrik tidak dibutuhkan urutan atau jenjang secara formal karena sering dijumpai hasil pengamatan yang dinyatakan dalam data kualitatif; pengujian hipotesis pada statistik non-parametrik dilakukan secara langsung pada pengamatan yang nyata. Walaupun pada statistik non-parametrik tidak terikat pada distribusi normal populasi, tetapi dapat digunakan pada populasi berdistribusi normal. Contoh statistik nonparametrik yaitu Kolerasi Spearman (*Spearman Rank Order Correlation*) dan Chi Square (Ketut, 2013).

Berbeda halnya dengan analisis data kualitatif. Data yang diperoleh dari berbagai sumber dalam penelitian kualitatif dapat menggunakan teknik pengumpulan data yang bermacam-macam (triangulasi) dan dilakukan secara terus-menerus sampai datanya jenuh (dapat disimpulkan). Pengamatan yang terus-menerus menghasilkan variasi data yang tinggi. Oleh karena itu sering mengalami kesulitan dalam proses menganalisisnya. Analisis data kualitatif adalah bersifat induktif, yaitu suatu analisis berdasarkan data yang diperoleh selanjutnya dikembangkan pola hubungan tertentu atau menjadi hipotesis. Berdasarkan hipotesis yang telah dirumuskan maka selanjutnya mencari data lagi secara terus-menerus agar dapat digeneralisasikan apakah hipotesis diterima atau ditolak berdasarkan data valid yang telah terkumpul. Ketika hipotesis diterima berdasarkan data yang terkumpul maka hipotesis dapat berkembang menjadi teori. Menurut Sugiyono, analisis data dalam penelitian kualitatif dilakukan sejak sebelum memasuki lapangan, selama di lapangan dan setelah selesai di lapangan.

1) Analisis sebelum di lapangan

Penelitian kualitatif telah melakukan analisis data sebelum peneliti memasuki lapangan. Analisis dilakukan terhadap data hasil studi pendahuluan, atau data sekunder, yang akan digunakan untuk menentukan fokus penelitian. Namun demikian fokus penelitian ini masih bersifat sementara, dan akan berkembang setelah peneliti masuk dan selama di lapangan.

2) Analisis selama di lapangan model Miles and Huberman

Analisis data dalam penelitian kualitatif, dilakukan pada saat pengumpulan data berlangsung dan setelah selesai pengumpulan data dalam periode tertentu. Pada saat wawancara, peneliti sudah melakukan analisis terhadap jawaban yang diwawancarai. Bila jawaban yang diwawancarai setelah dianalisis terasa belum memuaskan, maka peneliti akan melanjutkan pertanyaan lagi, sampai tahap tertentu sehingga diperoleh data yang dianggap kredibel. Miles and Huberman, mengemukakan bahwa aktivitas dalam analisis data kualitatif dilakukan secara interaktif dan berlangsung secara terus menerus sampai tuntas, sehingga datanya sudah jenuh. Aktivitas dalam analisis data, yaitu *data reduction*, *data display*, dan *conclusion drawing/verification* (Sugiyono, 2010).

(a) Data *Reduction* (Reduksi Data)

Data yang diperoleh dari lapangan jumlahnya cukup banyak, untuk itu maka perlu dicatat secara teliti dan rinci. Seperti telah dikemukakan, makin lama peneliti ke lapangan, maka jumlah data akan makin banyak, kompleks dan rumit. Untuk itu perlu segera dilakukan analisis data melalui reduksi data. Mereduksi data berarti merangkum, memilih hal-hal yang pokok, memfokuskan pada hal-hal yang penting, dicari tema dan polanya dan membuang yang tidak perlu. Dengan demikian data yang telah direduksi akan memberikan gambaran yang lebih jelas, dan mempermudah peneliti untuk melakukan pengumpulan data selanjutnya, dan mencarinya bila diperlukan.

(b) Data *Display* (Penyajian Data)

Setelah data di reduksi, maka langkah selanjutnya adalah mendisplaykan data. Kalau dalam penelitian kuantitatif penyajian data ini dapat dilakukan dalam bentuk tabel, grafik, pie chart, pictogram dan sejenisnya. Melalui penyajian data tersebut, maka data terorganisasikan, tersusun dalam pola hubungan, sehingga akan semakin mudah difahami.

Dalam penelitian kualitatif, penyajian data bisa dilakukan dalam bentuk uraian singkat, bagan, hubungan antar kategori, *flowchart* dan sejenisnya. Yang paling sering digunakan untuk menyajikan data dalam penelitian kualitatif adalah dengan teks yang bersifat naratif.

(c) Conclusion Drawing/verification

Langkah ke tiga dalam analisis data kualitatif adalah penarikan kesimpulan dan verifikasi. Kesimpulan awal yang dikemukakan masih bersifat sementara, dan akan berubah bila tidak ditemukan bukti-bukti yang kuat yang mendukung pada tahap pengumpulan data berikutnya. Tetapi apabila kesimpulan yang dikemukakan pada tahap awal, didukung oleh bukti-bukti yang valid dan konsisten saat peneliti kembali ke lapangan

mengumpulkan data, maka kesimpulan yang dikemukakan merupakan kesimpulan yang kredibel.

Dengan demikian kesimpulan dalam penelitian kualitatif mungkin dapat menjawab rumusan masalah yang dirumuskan sejak awal, tetapi mungkin juga tidak, karena seperti telah dikemukakan bahwa masalah dan rumusan masalah dalam penelitian kualitatif masih bersifat sementara dan akan berkembang setelah penelitian berada di lapangan (Sugiyono, 2011).

3) Analisis data selama di lapangan model Spradley

Menurut Spradley, proses penelitian kualitatif setelah memasuki lapangan, dimulai dengan menetapkan seseorang informan kunci "*key informant*" yang merupakan informan yang berwibawa dan dipercaya mampu "membukakan pintu" kepada peneliti untuk memasuki obyek penelitian. Setelah itu peneliti melakukan wawancara kepada informan tersebut, dan mencatat hasil wawancara. Selanjutnya, perhatian peneliti pada obyek penelitian dan memulai mengajukan pertanyaan deskriptif, dilanjutkan dengan analisis terhadap hasil wawancara. Berdasarkan hasil dari analisis wawancara selanjutnya peneliti melakukan analisis domain. Pada langkah selanjutnya peneliti sudah menentukan fokus, dan melakukan analisis taksonomi. Berdasarkan hasil analisis taksonomi, selanjutnya peneliti mengajukan pertanyaan kontras, yang dilanjutkan dengan analisis komponensial. Hasil dari analisis komponensial, selanjutnya peneliti menemukan tema-tema budaya. Berdasarkan temuan tersebut, selanjutnya peneliti menuliskan laporan penelitian etnografi.

Jadi proses penelitian berangkat dari yang luas, kemudian memfokus, dan meluas lagi. Terdapat tahapan analisis data yang dilakukan dalam penelitian kualitatif, yaitu analisis domain, taksonomi, dan komponensial, analisis tema kultural (Sugiyono, 2012)

4. Langkah-langkah Analisis Data

Menurut Sukardi (2003), ada beberapa langkah yang perlu dilalui agar proses analisis menjadi lebih terarah, yakni skoring, tabulasi, mendeskripsikan data, dan melakukan uji statistika.

a. Skoring

Skoring adalah pemberian nilai pada setiap jawaban yang dikumpulkan peneliti dari instrumen yang telah disebarkan. Setiap item pertanyaan yang dimunculkan pada instrumen dikuantifikasikan dalam bentuk angka. Misalnya, pada saat angket disebarkan alternatif jawaban yang diberikan masih berupa kualitatif, maka pada tahap ini harus dikuantifikasikan. Pada tahap ini peneliti memberikan nilai atau bobot pada setiap alternatif jawaban.

Contoh alternatif jawaban pada angket.

Selalu : 3

Belum tentu : 2

Tidak : 1

b. *Tabulasi*

Setelah tahap skoring, hasilnya ditransfer dalam bentuk yang lebih ringkas dan mudah dilihat. Mencatat skor secara sistematis akan memudahkan pengamatan data yang diperoleh. Apabila analisis data membandingkan dua kelompok, maka data ditempatkan dalam kolom yang berbeda. Dengan menggunakan prinsip tabulasi ini, seorang peneliti akan dapat menentukan arah selanjutnya teknik analisis apa yang diperlukan, tergantung pada tujuan analisis data yang hendak dicapai.

c. *Mendeskripsikan data*

Mendeskripsikan data adalah menggambarkan data yang ada guna memperoleh bentuk nyata dari responden, sehingga lebih dimengerti oleh peneliti atau seseorang yang tertarik dengan hasil penelitian yang dilakukan. Analisis data yang paling sederhana dan sering digunakan oleh peneliti atau pengembang adalah menganalisis data yang ada dengan menggunakan prinsip-prinsip deskriptif. Dengan menganalisis secara deskriptif dapat mendeskripsikan data secara lebih ringkas, sederhana, dan lebih mudah dimengerti.

d. *Melakukan uji statistika*

Uji statistika atau analisis inferensial merupakan pengolahan data yang diperoleh dengan menggunakan rumus-rumus atau aturan-aturan yang berlaku, sesuai dengan pendekatan penelitian atau desain yang diambil. Penggunaan rumus atau aturan-aturan tersebut hendaknya mampu mengukur dan sesuai dengan tujuan atau hasil penelitian yang ingin peneliti capai (Sugiyono, 2010).

5. ANALISIS

Di dalam sebuah penulisan karya ilmiah yang berdasarkan data-data penelitian, analisis data merupakan suatu hal yang harus dijabarkan oleh penulis. Karena tujuan pokok dari suatu penelitian adalah untuk menjawab pertanyaan-pertanyaan penelitian. Dan untuk mencapai tujuan pokok tersebut, peneliti harus dapat melakukan proses pengolahan data dan menganalisis data tersebut. Analisis data merupakan salah satu langkah terpenting dalam sebuah penelitian, karena merupakan cara berfikir agar memperoleh temuan-temuan yang di hasilkan dari sebuah penelitian.

Dalam melakukan analisis dalam sebuah penelitian, peneliti perlu mengetahui terlebih dahulu mengenai teknik-teknik dan langkah-langkah dalam menganalisis data yang harus di lakukan agar proses analisis lebih terarah. Begitu pula dalam melakukan analisis data, peneliti memerlukan usaha yang sangat perlu untuk di implementasikan

yakni pemikiran para peneliti. Selain melakukan analisis data, peneliti juga perlu menguasai akan kepustakaan yang berguna untuk mengkonfirmasikan adanya teori baru yang barangkali dapat ditemukan. Setelah data dianalisa, hasil-hasilnya harus diinterpretasikan untuk mencari makna dan implikasi yang lebih luas dari hasil-hasil penelitian.

C. TUGAS

1. Jelaskan perbedaan analisis kualitatif dan analisis kuantitatif !
2. Buatlah tabel distribusi frekuensi dan persen dengan menggunakan SPSS berdasarkan, umur, jenis kelamin, pendidikan ?
3. Langkah-langkah apa saja yang harus dilakukan dalam analisis data ?

D. KESIMPULAN

Analisis data merupakan proses mengorganisasikan dan mengurutkan data ke dalam pola, kategori dan satuan uraian dasar sehingga dapat ditemukan tema dan dapat dirumuskan hipotesis kerja seperti yang didasarkan oleh data. Dalam rangka analisis dan interpretasi data, perlu dipahami tentang keberadaan data itu sendiri. Secara garis besar, keberadaan data dapat digolongkan ke dalam dua jenis, yaitu : data bermuatan kualitatif dan data bermuatan kuantitatif.

Teknik analisis data ada dua, yaitu teknik analisis data kuantitatif dan teknik analisis data kualitatif yaitu teknik analisis data kuantitatif dengan menggunakan statistik, meliputi statistik deskriptif dan inferensial. Statistik inferensial meliputi statistik parametris dan non parametris. Teknik analisis data kualitatif dilakukan dari sebelum penelitian, selama penelitian, dan sesudah penelitian yang meliputi analisis sebelum di lapangan, teknik analisis selama di lapangan model Miles dan Huberman dan teknik analisis data menurut Spradley. Dan terdapat pula langkah – langkah yang perlu dilalui agar proses analisis menjadi lebih terarah, yakni skoring, tabulasi, mendeskripsikan data, dan melakukan uji statistika.

E. DAFTAR PUSTAKA

- Dey, Ian. 1995. *Qualitative Data Analysis*, New York: RNY.
- Herdiansyah, Haris. 2011. *Metodologi Penelitian Kualitatif Untuk Ilmu-ilmu Sosial*, Jakarta: Salemba Humanika.
- Noor, Juliansyah. 2012. *Metodologi Penelitian*, Jakarta : Kencana.
- Santosa, Purbayu Budi dan Ashari. 2005. *Analisis Statistik dengan Microsoft Excel dan SPSS*, Yogyakarta : Andi.

- Subagyo, Joko. 2011. *Metode Penelitian Dalam Teori & Praktik*, Jakarta : Rineka Cipta.
- Sugiyono. 2010. *Metode Penelitian Pendidikan Pendekatan Kuantitatif, Kualitatif, dan R&D*, Bandung : Alfabeta
- Sugiyono. 2012. *Memahami Penelitian Kualitatif*, Bandung: AlfaBeta
- Sukmadinata, Nana Syaodih. 2009. *Metodologi Penelitian Pendidikan*. Bandung : Rosda.
- Moleong, Lexy. 2017. *Metodologi Penelitian Kualitatif*, Bandung: Rosda Karya.



BAB VII

UJI STATISTIK PARAMETRIK

A. Tujuan Pembelajaran:

Mampu memahami uji statistik parametrik (uji beda rata-rata yang terdiri satu sampel, dua sampel, lebih dari dua sampel, dan uji regresi linier yang terdiri regresi linier sederhana dan regresi linier berganda) dan mempraktikkan uji statistik dengan software SPSS.

B. Materi

Statistik Parametrik yaitu ilmu statistik yang mempertimbangkan jenis sebaran atau distribusi data, yaitu apakah data menyebar secara normal atau tidak. Dengan kata lain, data yang akan dianalisis menggunakan statistik parametrik harus memenuhi asumsi normalitas. Pada umumnya, jika data tidak menyebar normal, maka data seharusnya dikerjakan dengan metode statistik nonparametrik, atau setidaknya-tidaknya dilakukan transformasi terlebih dahulu agar data mengikuti sebaran normal, sehingga bisa dikerjakan dengan statistik parametrik. Digunakan untuk menguji parameter populasi melalui statistik, atau menguji ukuran populasi melalui data sampel. Statistik parametrik memerlukan terpenuhinya banyak asumsi, antara lain berdistribusi normal, data homogen, harus terpenuhi asumsi linieritas. Statistik parametrik banyak digunakan untuk menganalisis data interval dan rasio, contohnya: uji-z, uji-t, korelasipearson, anova

Keunggulan Parametrik yaitu syarat syarat parameter dari suatu populasi yang menjadi sampel biasanya tidak diuji dan dianggap memenuhi syarat, pengukuran terhadap data dilakukan dengan kuat. Observasi bebas satu sama lain dan ditarik dari populasi yang berdistribusi normal serta memiliki varian yang homogen. Kelemahan Parametrik yaitu populasi harus memiliki varian yang sama. Variabel-variabel yang diteliti harus dapat diukur setidaknya dalam skala interval. Dalam analisis varian ditambahkan persyaratan rata-rata dari populasi harus normal dan bervarian sama, dan harus merupakan kombinasi linear dari efek-efek yang ditimbulkan.

Jenis data yang dianalisis dalam statistik parametrik terutama adalah mempunyai skala interval atau rasio. Dalam analisis menggunakan statistik parametrik, pertimbangan jenis sebaran atau distribusi data yang menyebar secara normal adalah mutlak. Asumsi utama dalam statistik parametrik adalah data yang akan dianalisis harus memenuhi normalitas. Jika distribusi data tidak normal, maka analisis harus dikerjakan dengan teknik statistik non parametrik. Sebagai jalan tengah, dapat dilakukan transformasi lebih dahulu agar data menjadi normal. Asumsi berikutnya yang harus dipenuhi dalam analisis menggunakan statistik parametrik adalah dua data kelompok atau lebih yang akan diuji harus homogen. Sehingga persyaratan homogenitas data-data harus diperhatikan. Asumsi yang harus terpenuhi lainnya adalah data harus linier. Menganalisis hubungan antara dua variabel interal digunakan analisis korelasi product moment dari Pearson. Tetapi dalam teknik non-parametrik analisis korelasi lebih dikenal dengan korelasi tata jenjang (rank order correlation) Spearman. Cara dan langkah-langkah melakukan penghitungan dua teknik korelasi tersebut berbeda. Dalam analisis korelasi tata jenjang, lebih dahulu membuat urutan (ranking) dari data yang akan dikorelasikan, sedangkan dalam product moment dari Pearson tidak perlu dilakukan penyusunan ranking. Dalam teknik analisis statistik parametrik, untuk melakukan uji perbedaan mean antara dua kelompok data dapat digunakan teknik statistik uji t. Teknik statistik uji t dibedakan menjadi uji t untuk sampel yang berkorelasi, dan uji t

untuk sampel yang tidak berkorelasi. Uji beda mean antara lebih dari dua kelompok data dapat digunakan teknik analisis varian (Budiwanto: 2014)

1. Uji Beda Mean Satu Sampel

Pengujian rata-rata satu sampel dimaksudkan untuk menguji nilai tengah atau rata-rata populasi μ sama dengan nilai tertentu μ_o , lawan hipotesis alternatifnya bahwa nilai tengah atau rata-rata populasi μ tidak sama dengan μ_o . Pengujian satu sampel pada prinsipnya ingin menguji apakah suatu nilai tertentu (yang diberikan sebagai pembanding) berbeda secara nyata ataukah tidak dengan rata-rata sebuah sampel. Nilai tertentu di sini pada umumnya adalah sebuah nilai parameter untuk mengukur suatu populasi.

Jadi kita akan menguji:

$H_o : \mu = \mu_o$ lawan $H_1 : \mu \neq \mu_o$

H_o merupakan hipotesa awal sedangkan H_1 merupakan hipotesis alternatif atau hipotesis kerja

- Rumus *one sample t-test* (Uji T):

$$t_{hit} = \frac{\bar{x} - \mu_o}{s/\sqrt{n}}$$

t_{hit} = nilai t hitung

$\bar{x}$ = rata-rata sampel

μ_o = nilai parameter

s = standar deviasi sampel

n = jumlah sampel

Interpretasi :

- Untuk menginterpretasikan t-test terlebih dahulu harus ditentukan:
 - Nilai signifikansi α
 - D_f (*degree of freedom*) = $N-k$, khusus untuk *one sample t-test* $d_f = N - 1$
- Bandingkan nilai t_{hit} dengan t_{tab} , dimana $t_{tab} = t_{\frac{\alpha}{2}; N-1}$
- Apabila:
 - $t_{hit} > t_{tab}$ = ada berbeda secara signifikansi (H_o ditolak)
 - $t_{hit} < t_{tab}$ = tidak ada berbeda secara signifikansi (H_o diterima)

Contoh 1:

Seorang mahasiswa melakukan penelitian mengenai galon susu murni yang rata-rata isinya 10 liter. Telah diambil sampel secara acak dari 10 botol yang telah diukur isinya, dengan hasil sebagai berikut :

Galon ke-	Volume
1	10,2
2	9,7
3	10,1
4	10,3
5	10,1
6	9,8
7	9,9
8	10,4
9	10,3
10	9,8

Dengan taraf signifikasnsi $\alpha = 0,01$. Apakah galon susu murni rata-rata isinya 10 liter?
Penyelesaian :

a. Analisa secara manual:

1. Hipotesis $H_0 : \mu = 10$ lawan $H_1 : \mu \neq 10$
2. Uji statistik t (karena α tidak diketahui atau $n < 30$).
3. $\alpha = 0.01$
4. Wilayah kritik : $t_{hit} < t_{\frac{\alpha}{2}, N-1}$ atau $t_{hit} > t_{\frac{\alpha}{2}, N-1}$ dengan $t_{tab} = 3,259$
5. Perhitungan, dari data: $\bar{x} = 10,06$ dan simpangan baku sampel $s = 0,2459$.

$$t_{hit} = \frac{\bar{x} - \mu_0}{s/\sqrt{n}}$$

$$= \frac{10,06 - 10}{\frac{0,2459}{\sqrt{10}}} = 0,772$$

Karena $t_{hit} = 0,772 < t_{tab} = 3,259$, maka H_0 diterima. Atau untuk menguji Hipotesis nol menggunakan interval Confidence dengan ketentuan apabila terletak diantara -0,1927 dan 0,3127 disimpulkan untuk menerima H_0 artinya pernyataan bahwa rata-rata isi galon susu murni 10 liter dapat diterima.

b. Analisa menggunakan SPSS:

1. Masukkan data diatas pada *Data View*, namun sebelumnya kita harus menentukan nama dan tipe datanya pada *Variable View*.
2. Klik Menu *Analyze* → *Compare Means* → *One Sample T-Test*
3. Masukkan galon susu ke i (X) ke kolom *test variable* dan masukkan nilai rata-rata 10 pada *test value*
4. Klik option dan pada *interval confidence* masukkan 99% (karena $\alpha = 0,01$), kemudian klik continue
5. Kemudian klik OK
6. Sehingga menghasilkan hasil analisa sebagai berikut:

One-Sample Statistics				
	N	Mean	Std. Deviation	Std. Error Mean
galon susu ke-i	10	10.0600	.24585	.07775

One-Sample Test						
	Test Value = 10					
	t	df	Sig. (2-tailed)	Mean Difference	99% Confidence Interval of the Difference	
					Lower	Upper
galon susu ke-i	.772	9	0.460	.06000	-.1927	.3127

Keterangan hasil analisa:

Std error = Standar Error

T = nilai hitung

Df = derajat kebebasan

Sig (2-tailed) = probabilitas ($\alpha/2$)

Mean difference = perbandingan rata-rata

Ho diterima karena sig = 0,46 > 0,01, artinya rata-rata galon susu berisi 10 liter

- Dalam pengujian Uji-Z, data yang diperoleh adalah berdistribusi normal dengan ciri :
 - a. Unimodial, selalu memiliki modus dan hanya satu modus
 - b. Simetrik
 - c. Modus = median = rata-rata
 - d. Asimtotik, kurva distribusi normal tidak akan pernah menyentuh absisnya
 - e. Pengujian Uji-Z dapat dilakukan apabila simpangan baku populasi (σ) diketahui dan n-nya sejumlah lebih dari 30
 - f. Untuk uji perbedaan rata-rata data tunggal dengan Uji-Z, maka diperoleh dari sampel berpopulasi tunggal
 - g. Rumus yang digunakan untuk mengetahui nilai-z adalah :

$$Z = \frac{\bar{x} - \mu}{\sigma / \sqrt{n}}$$

- Dalam penggunaan Uji-Z, derajat kebebasan (df) tidak perlu diperhatikan karena simpangan baku yang diketahui adalah simpangan baku populasi
 - a. Nilai untuk pengujian satu sisi (one tail) pada uji-z dengan $\alpha = 0,01$ maka harga z tabel = 2,33 sedangkan pada $\alpha = 0,05$ harga z tabel = 1,65

- b. Nilai untuk pengujian dua sisi (two tail) pada uji-z dengan $\alpha = 0,01$ maka harga z tabel = 2,58 sedangkan pada $\alpha = 0,05$ harga z tabel = 1,65

Langkah-langkah dalam melaksanakan uji-z :

1. Menyusun formulasi hipotesis nihil dan hipotesis alternatifnya
 - i. Pengujian dua sisi
 $H_0 : \mu = \mu_0$ lawan $H_1 : \mu \neq \mu_0$
 - ii. Pengujian satu sisi kanan
 $H_0 : \mu = \mu_0$ lawan $H_1 : \mu > \mu_0$
 - iii. Pengujian satu sisi kiri
 $H_0 : \mu = \mu_0$ lawan $H_1 : \mu < \mu_0$
2. Menentukan level of significancenya (α)
3. Menentukan peraturan-peraturanpengujiannya/kriterianya/rule of the uji
 - i. Pengujian dua sisi
 Ho diterima apabila : $-z_{\frac{\alpha}{2}} \leq z \leq z_{\frac{\alpha}{2}}$
 Ho ditolak apabila : $z > z_{\frac{\alpha}{2}}$
 - ii. Pengujian satu sisi kanan
 Ho diterima apabila : $z \leq z_{\alpha}$
 Ho ditolak apabila : $z > z_{\alpha}$
 - iii. Pengujian satu sisi kiri
 Ho diterima apabila : $z \geq -z_{\alpha}$
 Ho ditolak apabila : $z < -z_{\alpha}$
4. Dari sampel random yang diambil kemudian dihitung nilai Z, dengan rumus :

$$Z = \frac{\bar{x} - \mu}{\sigma / \sqrt{n}}$$

5. Dengan membandingkan perhitungan pada langkah 4 dengan peraturan pengujian langkah 3 kemudian diambil kesimpulan

Contoh 2 :

Ujilah hipotesis bahwa hasil rata-rata per hari dari suatu pabrik $\mu = 880 \text{ ton}$ dengan alternatif bahwa μ lebih besar atau lebih kecil dari 880 ton per hari. Suatu sampel yang didasarkan pada $n=50$ pengukuran, hasil rata-rata perhari $\bar{x} = 875 \text{ ton}$ dengan simpangan baku $\sigma = 21 \text{ ton}$. Gunakan $\alpha = 5\%$

Penyelesaian :

- a. $H_0 : \mu = 880 \text{ ton}$; $H_1 : \mu \neq 880 \text{ ton}$, digunakan pengujian dua sisi
- b. $\alpha = 5\%$, $Z = 1,96$

c. Kriteria pengujian

Ho diterima apabila : $-1,96 \leq z \leq 1,96$ Ho ditolak apabila : $z > 1,96$ atau $z < -1,96$ d. Perhitungan, dari data: $\bar{x} = 875$ ton dan simpangan baku sampel $\sigma = 21$ ton.

$$Z = \frac{\bar{x} - \mu}{\sigma / \sqrt{n}} = \frac{875 - 880}{\frac{21}{\sqrt{50}}} = -1,6836$$

Karena $t_{hit} = -1,6836 > t_{tab} = -1,96$, maka H_0 diterima. Atau untuk menguji Hipotesis nol menggunakan interval Confidence dengan ketentuan apabila terletak diantara $-1,96$ dan $1,96$ disimpulkan untuk menerima H_0 artinya pernyataan bahwa hasil rata-rata per hari dari suatu pabrik $\mu = 880$ ton dapat diterima.

2. Uji Beda Dua Mean Sampel

Dibidang kesehatan seringkalikita harus menarik kesimpulan apakah parameter dua populasi berbeda atau tidak. Misalnya apakah ada perbedaan tekanan darah penduduk dewasa orang kota dengan orang desa. Atau, apakah ada perbedaan berat badan antara sebelum mengikuti program diet dengan sesudahnya. Uji statistik yang membandingkan mean dua kelompok data ini disebut uji beda dua mean. Sebelum kita melakukan uji statisik dua kelompok data, kita perlu perhatikan apakah dua kelompok data tersebut berasal dari dua kelompok yang independen atau berasal dari dua kelompok yang dependen / pasangan. Dikatakan kedua kelompok data dependen bila data kelompok yang satu tidak tergantung dari data kelompok kedua, misalnya membandingkan mean tekanan darah sistolik orang desa dengan orang kota. \tekanan darah orang kota independen (tidak tergantung) dengan orang desa. Dilain pihak, kedua kelompok data dikatakan dependen f pasangan bila kelompok data yang dibandingkan 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) Berdasarkan karakteristik data tersebut maka uji beda dua mean dibagi dalam dua kelompok, yaitu uji beda dua mean independen dan uji beda mean dependen.

3. Uji Beda Dua Mean Dependen

Uji – t berpasangan (paired t-test) adalah salah satu metode pengujian hipotesis dimana data yang digunakan tidak bebas (dependen). Ciri-ciri yang paling sering ditemui pada kasus yang berpasangan adalah satu individu (objek penelitian) dikenai 2 buah perlakuan yang berbeda. Walaupun menggunakan individu yang sama, peneliti tetap memperoleh 2 macam data sampel, yaitu data dari perlakuan pertama dan data dari perlakuan kedua. Hipotesis dari kasus ini dapat ditulis:

$$H_0 = \mu_1 - \mu_2 = 0 \text{ atau } \mu_1 = \mu_2$$

$$H_a = \mu_1 - \mu_2 \neq 0 \text{ atau } \mu_1 \neq \mu_2$$

H_a berarti bahwa selisih sebenarnya dari kedua rata-rata tidak sama dengan nol.

Rumus *Paired Sample t-test*:

$$t_{hit} = \frac{\bar{D}}{\frac{SD}{\sqrt{n}}}$$

Ingat :

$$Variansi (s^2) = \frac{1}{n-1} \sum_{i=1}^n ((x_j - x_i) - \bar{D})^2$$

t = nilai hitung

$\bar{D}$ = rata-rata selisih pengukuran 1 dan 2

SD = standar deviasi selisih pengukuran 1 dan 2

n = jumlah sampel

Interpretasi :

- Untuk menginterpretasikan t-test terlebih dahulu harus ditentukan:
 - Nilai signifikansi α
 - D_f (*degree of freedom*) = $N-k$, khusus untuk *paired sample t-test* $d_f = N - 1$
- Bandingkan nilai t_{hit} dengan $t_{tab=\alpha;n-1}$
- Apabila:
 - $t_{hit} > t_{tab}$ = berbeda secara signifikansi (H_0 ditolak)
 - $t_{hit} < t_{tab}$ = tidak berbeda secara signifikansi (H_0 diterima)

Contoh :

Peneliti ingin mengetahui penggunaan terapi Antiretroviral terhadap berat badan pasien HIV/AIDS. Rumusan pertanyaannya yaitu "apakah terdapat perbedaan rerata berat badan sebelum dan sesudah diberikan terapi Antiretroviral pada pasien HIV/AIDS.

Data Hasil Penelitian

Reponden	Berat Badan Sebelum Terapi ARV	Berat Badan Sesudah Terapi ARV
1	48	50
2	48	50
3	53	55
4	56	57
5	58	60
6	51	53
7	45	47
8	58	60
9	44	45
10	45	46
11	41	43
12	60	63
13	47	50
14	56	59
15	46	48

Secara manual:

Respons	Sebelum terapi ARV (Xj)	Sesudah Terapi ARV (Xi)	(Xj - Xi)	$\bar{D}$	(Xj - Xi) - $\bar{D}$	$[(Xj - Xi) - \bar{D}]^2$
1	48	50	-2	-2	0	0
2	48	50	-2	-2	0	0
3	53	55	-2	-2	0	0
4	56	57	-1	-2	1	1
5	58	60	-2	-2	0	0
6	51	53	-2	-2	0	0
7	45	47	-2	-2	0	0
8	58	60	-2	-2	0	0
9	44	45	-1	-2	1	1
10	45	46	-1	-2	1	1
11	41	43	-2	-2	0	0
12	60	63	-3	-2	-1	1
13	47	50	-3	-2	-1	1
14	56	59	-3	-2	-1	1
15	46	48	-2	-2	0	0
Total	756	786	-30			6

Dari tabel perhitungan diperoleh:

$$\bar{D} = \frac{\text{jumlah selisih}}{n} = \frac{-30}{15} = -2$$

$$\text{Variansi } (s^2) = \frac{1}{n-1} \sum_{i=1}^n ((x_j - x_i) - \bar{D})^2$$

$$= \frac{6}{14} = 0,4286 \text{ sehingga } SD = \sqrt{0,4286} = 0,6546$$

$$t_{hit} = \frac{\bar{D}}{\frac{SD}{\sqrt{n}}} = \frac{-2}{\frac{0,6546}{\sqrt{15}}} = -11,833$$

Ikuti langkah berikut

- Isikan data tersebut ke dalam SPSS
- Pada Menu SPSS klik *Analyze*
- Klik *Compare Means*
- Klik *Paired Sample t*
- Masukkan data berat badan sebelum dan sesudah ke dalam kotak Paired Variables.
- Klik Continue, diikuti klik OK

Interpretasi Hasil

- a. Kolom Paired Sample Statistics menjelaskan tentang deskripsi masing-masing variable.

Paired Samples Statistics				
	Mean	N	Std. Deviation	Std. Error Mean
Pair Sebelum	50.40	15	6.021	1.555
Sesudah	52.40	15	6.266	1.618

- b. Tabel Paired Samples Test menggambarkan hasil uji t berpasangan. Lihat kolom sig. (2 tailed). Diperoleh nilai signifikansi 0,000 ($p < 0,050$). Hal tersebut berarti terdapat perbedaan yang bermakna antara berat badan sebelum dan sesudah menjalani terapi ASV. Nilai IK 95% adalah antara -2,363 sampai dengan -1,637.

Paired Samples Test									
		Paired Differences					t	df	Sig. (2-tailed)
		Mean	Std. Deviation	Std. Error Mean	95% Confidence Interval of the Difference				
					Lower	Upper			
Pair	Sebelum 1 - Sesudah	-2.000	.655	.169	-2.363	-1.637	- 11.832	14	.000

- df = degree of freedom (derajat kebebasan) : Untuk analisis T Paired selalu $N - 1$. Di mana N adalah jumlah sampel.
- T = nilai t hitung: hasil -11,832: Harus dibandingkan dengan t tabel pada df 14. Apabila t hitung > t tabel: signifikan.
- Sig.(2-tailed) = Nilai probabilitas/p value uji T Paired: Hasil = 0,000. Artinya: Tidak ada perbedaan antara sebelum dan sesudah perlakuan. Sebab: Nilai p value > 0,05 (95 % kepercayaan).
- Mean = -2,000. Bernilai Negatif, artinya tidak terjadi kecenderungan kenaikan sesudah terapi ARV. Apabila didapatkan hasil Positif: Artinya terjadi kecenderungan perubahan sesudah perlakuan.

Melaporkan Hasil

Hasil penyajian terbaik dari uji ini adalah dengan mendeskripsikan informasi yang lengkap dengan menampilkan IK serta nilai p. Hal ini dikarenakan informasi nilai IK lebih bermanfaat secara klinis daripada informasi p value. Pada dasarnya pertimbangan tabel manakah yang akan disampaikan dalam laporan penelitian atau jurnal ilmiah bergantung pada pedoman yang dianut oleh institusi atau jurnal tersebut.

Contoh:

Tabel: Hasil uji t berpasangan secara lengkap

	N	Rerata± s.b.	Perbedaan Rerata±s. b.	IK95%	p
BB Sebelum Terapi ARV	15	50,40±6,02	2,00±6,55	1,63±2,36	<0,001
BB Sesudah Terapi ARV	15	52,40±6,26			

Uji t berpasangan

Berdasarkan tabel di atas hasil uji paired t-test didapatkan bahwa rata-rata berat badan pada pasien HIV/AIDS sebelum menjalani terapi ARV sebesar 50,4 kg sedangkan setelah diberikan terapi ARV sebesar 52,4 kg. Hasil uji paired t -test juga didapatkan p value 0,001 (<0,05) dengan demikian dapat disimpulkan bahwa terdapat perbedaan antara berat badan sebelum dan sesudah mendapatkan terapi ARV bagi pasien HIV/AIDS.

4. Uji Beda Dua Mean Independent

Uji ini untuk mengetahui perbedaan rata-rata dua populasi/kelompok data yang independen. Contoh kasus suatu penelitian ingin mengetahui hubungan status merokok ibu hamil dengan berat badan bayi yang dilahirkan. Respondan terbagi dalam dua kelompok, yaitu mereka yang merokok dan yang tidak merokok.

Uji T independen ini memiliki asumsi/syarat yang mesti dipenuhi, yaitu:

- Datanya berdistribusi normal.
- Kedua kelompok data independen (bebas)
- Variabel yang dihubungkan berbentuk numerik dan kategorik (dengan hanya 2 kelompok)

Rumus *Independent Sample t-test*:

$$T_{hit} = \frac{M_1 - M_2}{\sqrt{\frac{SS_1 + SS_2}{n_1 + n_2 - 2} \left(\frac{1}{n_1} + \frac{1}{n_2} \right)}}$$

Keterangan:

M_1 = rata-rata skor kelompok 1

M_2 = rata-rata skor kelompok 2

SS_1 = sum of square kelompok 1

SS_2 = sum of square kelompok 2

n_1 = jumlah subjek/sample kelompok 1

n_2 = jumlah subjek/sample kelompok 2

Dimana :

$$M_1 = \frac{\sum X_1}{n_1} \quad SS_1 = \sum X_1^2 - \frac{(\sum X_1)^2}{n_1}$$

$$M_2 = \frac{\sum X_2}{n_2} \quad SS_2 = \sum X_2^2 - \frac{(\sum X_2)^2}{n_2}$$

Interpretasi :

- Untuk menginterpretasikan t-test terlebih dahulu harus ditentukan:
 - Nilai signifikansi α
 - Interval Confidence = $1 - \alpha$
 - D_f (degree of freedom) = $N - k$, khusus untuk *independent sample t-test* $d_f = N - 2$ atau D_f (degree of freedom) = $(n_1 - n_2) - 2$
- Bandingkan nilai t_{hit} dengan t_{tab}
- Apabila:
 - $t_{hit} > t_{tab}$ = berbeda secara signifikansi (H_0 ditolak)
 - $t_{hit} < t_{tab}$ = tidak berbeda secara signifikansi (H_0 diterima)

Contoh :

Seorang peneliti ingin membandingkan daya analgetik dan antiinflamasi dari indometasin dan paracetamol. Dari hasil percobaan didapatkan hasil berikut

Perlakuan	% Daya Analgetik	% Daya Antiinflamasi
1,00	68,80	30,45
1,00	70,54	34,32
1,00	65,32	32,67
1,00	68,54	35,89
1,00	64,57	33,40
1,00	69,89	34,38
1,00	64,54	32,65
1,00	65,34	35,87
1,00	68,92	34,45
1,00	67,64	33,87
2,00	40,76	78,35
2,00	35,45	76,98
2,00	38,38	79,56
2,00	37,29	78,98
2,00	39,48	75,34
2,00	38,64	80,97
2,00	41,87	79,54
2,00	35,64	76,65
2,00	34,98	77,87
2,00	38,23	75,98

Perlakuan 1 : Indometasin

Perlakuan 2 : Paracetamol

Langkah-langkah :

- Buka lembar kerja baru
- Klik *Variabel View*
- Ketik Perlakuan pada kolom Name dan baris 1
- Ketik *Daya_Analgetik* pada kolom Name dan baris 2
- Ketik *Daya_inflamasi* pada kolom Name dan baris 3
- Klik Data View
- Lalu isi kolom sesuai data
- Klik (...) pada Value baris pertama, akan muncul kotak dialog
- Ketik 1.00 → ketik Indometasin → Klik Add
- Ketik 2.00 → ketik Paracetamol → Klik Add → Ok
- Lalu klik Analyze → Compare Means → Independent Sample T-test
- Masukkan *Daya_analgetik* dan *Daya_antiinflamasi* pada Test Variable List dan masukkan Perlakuan pada Grouping Variable
- Klik Perlakuan → Define Groups
- Klik Use Spesific Values, ketik 1 pada group 1 dan 2 pada group 2. Klik Continue → Ok
- Akan muncul output seperti ini

Group Statistics					
	Perlakuan	N	Mean	Std. Deviation	Std. Error Mean
Daya_An	Indometasin	10	67.4100	2.27168	.71837
	Paracetamol	10	38.0720	2.28611	.72293
Daya_Inf	Indometasin	10	33.7950	1.62169	.51282
	Paracetamol	10	78.0220	1.78699	.56510

Independent Samples Test										
		Levene's Test for Equality of Variances		t-test for Equality of Means						
		F	Sig.	T	df	Sig. (2-tailed)	Mean Difference	Std. Error Difference	95% Confidence Interval of the Difference	
									Lower	Upper
Daya_An	Equal variances assumed	.141	.712	28.787	18	.000	29.33800	1.01916	27.19683	31.47917
	Equal variances not assumed			28.787	17.999	.000	29.33800	1.01916	27.19682	31.47918
Daya_Inf	Equal variances assumed	.353	.560	-57.957	18	.000	-44.22700	.76310	-45.83021	-42.62379
	Equal variances not assumed			-57.957	17.833	.000	-44.22700	.76310	-45.83129	-42.62271

Pada output SPSS (Independent samples test) → check sig (2-tailed) → harganya dibandingkan dengan nilai α (0,05) → kesimpulan

- jika varians sama → pakai nilai sig. pada baris "equal variances assumed"
- jika varians berbeda → "equal variances not assumed".

Sig Daya Analgetik 0,712 > 0,05 → H_0 diterima → identik (memiliki variansi yang sama)
 Sig Daya Analgetik 0,560 > 0,05 → H_0 diterima → identik (memiliki variansi yang sama)
 sehingga hasil uji t menyatakan bahwa ada perbedaan rata-rata antara persentase daya analgetik atau persentase daya antiinflamasi antara paracetamol dan indometasin.

a. Uji Homogenitas Varian

Uji homogenitas dilakukan untuk memberikan keyakinan bahwa sekelompok data yang diteliti dalam proses analisis berasal dari populasi yang tidak jauh berbeda keragamannya. Uji ini dilakukan sebagai prasyarat dalam analisis uji t dan analisis varian sebagai bagian dari statistik parametrik. Asumsi yang mendasari dalam analisis varian adalah bahwa varian dari populasi adalah sama. Pengujian homogenitas adalah pengujian untuk mengetahui sama tidaknya variansi-variansi dua buah distribusi atau lebih. Uji homogenitas yang akan dibahas adalah uji homogenitas menggunakan uji F (Budiwanto: 2017).

Langkah-langkah menghitung uji homogenitas:

- Menghitung varians atau standar deviasi kuadrat variabel X dan Y, dengan menggunakan rumus:

$$S_x^2 = \sqrt{\frac{n \cdot \sum X^2 - (\sum X)^2}{n(n-1)}} \quad S_y^2 = \sqrt{\frac{n \cdot \sum Y^2 - (\sum Y)^2}{n(n-1)}}$$

- Menghitung F hitung dari varians kelompok X dan Y, dengan rumus :

$$F = \frac{S \text{ besar}}{S \text{ kecil}}$$

Pembilang: S besar artinya varians dari kelompok dengan varian terbesar atau lebih banyak. Penyebut: S kecil artinya varian dari kelompok dengan varian terkecil atau lebih sedikit. Jika varians sama pada kedua kelompok, maka bebas menentukan pembilang dan penyebut.

- Membandingkan F hitung dengan F tabel pada tabel distribusi F, dengan memperhatikan beberapa hal sebagai berikut. (1) Untuk varians dari kelompok dengan varians terbesar adalah dk pembilang n-1. (2) Untuk varians dari kelompok dengan varians terkecil adalah dk penyebut n-1. (3) Jika F hitung lebih kecil (<) daripada F tabel, berarti tidak homogeny. Dengan kata lain sekelompok data yang diteliti dalam proses analisis berasal dari populasi yang tidak jauh berbeda keragamannya (4) Jika F hitung lebih besar (>) daripada F tabel, berarti tidak homogen.

Langkah-langkah uji homogenitas menggunakan SPSS :

- Klik Analyze
- Klik Compare Mean
- Klik One Way Anova
- Masukkan data yang akan di uji homogenitas pada kolom dependent list (data variabel 1) dan pada faktor (data variabel 2)
- Klik Options
- Klik Homogeneity of Variance Test
- Klik Continue
- Klik OK

Contoh :

Berikut ini contoh melakukan uji homogenitas dua kelompok eksperimen dan kelompok control, penelitian tentang pengaruh metode latihan keterampilan bulutangkis. Hasil tes akhir keterampilan bulutangkis kelompok eksperimen (x) dan kelompok kontrol (Y) disajikan dalam tabel 8.3 sebagai berikut.

Tabel 8.3. Hasil Tes Keterampilan Bulutangkis Kelompok Eksperimen (X) dan Kelompok Kontrol (Y)

Pasangan	X	Y	X ²	Y ²	XY
1	26	29	676	841	754
2	31	28	961	784	868
3	23	26	529	676	598
4	36	38	1296	1444	1368
5	41	42	1681	1764	1722
6	29	31	841	961	899
7	37	34	1369	1156	1258
8	32	27	1024	729	864
9	26	29	676	841	754
10	38	35	1444	1225	1330
11	42	43	1764	1764	1806
12	35	32	1225	1024	1120
Jumlah	396	394	13486	13209	13341

Selanjutnya dilakukan penghitungan nilai varians setiap kelompok menggunakan rumus sebagai berikut.

$$S_x^2 = \sqrt{\frac{12.13486 - (396)^2}{12(12-1)}} = 6,16$$

$$S_y^2 = \sqrt{\frac{12.13029 - (394)^2}{12(12-1)}} = 29,02$$

Selanjutnya menghitung nilai F menggunakan rumus sebagai berikut

$$F = \frac{2,90}{6,16} = 0,47$$

Hasil penghitungan nilai varian kelompok eksperimen (X) dan kelompok kontrol (Y) diperoleh F hitung 0,47, sedangkan F tabel 5% dengan derajat kebebasan pembilang $12-1=11$ dan derajat kebebasan penyebut $= 12-1 = 11$ diperoleh 2,82. Karena F hitung lebih kecil daripada F tabel, maka disimpulkan bahwa distribusi data kelompok eksperimen (X) dan kelompok kontrol (Y) adalah homogen (Budiwanto: 2017).

5. Uji Beda Lebih dari Dua Mean (Uji Anova/Uji F)

Analisis varians (analysis of variance, ANOVA) adalah suatu metode analisis statistika yang termasuk ke dalam cabang statistika inferensi. Dalam literatur Indonesia metode ini dikenal dengan berbagai nama lain, seperti analisis ragam, sidik ragam, dan analisis variansi. Ia merupakan pengembangan dari masalah Behrens-Fisher, sehingga uji-F juga dipakai dalam pengambilan keputusan. Analisis varians pertama kali diperkenalkan oleh Sir Ronald Fisher, bapak statistika modern. Dalam praktik, analisis varians dapat merupakan uji hipotesis (lebih sering dipakai) maupun pendugaan (estimation, khususnya di bidang genetika terapan).

Secara umum, analisis varians menguji dua varians (atau ragam) berdasarkan hipotesis nol bahwa kedua varians itu sama. Varians pertama adalah varians antarcontoh (among samples) dan varians kedua adalah varians di dalam masing-masing contoh (within samples). Dengan ide semacam ini, analisis varians dengan dua contoh akan memberikan hasil yang sama dengan uji-t untuk dua rerata (mean).

Supaya valid dalam menafsirkan hasilnya, analisis varians menggantungkan diri pada empat asumsi yang harus dipenuhi dalam perancangan percobaan:

- Data berdistribusi normal, karena pengujiannya menggunakan uji FSnedecor
- Varians atau ragamnya homogen, dikenal sebagai homoskedastisitas, karena hanya digunakan satu penduga (estimate) untuk varians dalam contoh
- Masing-masing contoh saling bebas, yang harus dapat diatur dengan perancangan percobaan yang tepat
- Komponen-komponen dalam modelnya bersifat aditif (saling menjumlah).

Analisis varians relatif mudah dimodifikasi dan dapat dikembangkan untuk berbagai bentuk percobaan yang lebih rumit. Selain itu, analisis ini juga masih memiliki keterkaitan dengan analisis regresi. Akibatnya, penggunaannya sangat luas di berbagai bidang, mulai dari eksperimen laboratorium hingga eksperimen psikologi.

6. Uji One Way Anova

Digunakan utk menguji rata-rata/pengaruh perlakuan dari suatu percobaan yang menggunakan 1 faktor, dimana 1 faktor tersebut memiliki 3 atau lebih kelompok. Disebut satu arah karena hanya berkepentingan dengan 1 faktor saja, mengelompokkan data berdasarkan satu kriteria saja.

Asumsi yang harus dipenuhi dalam uji one-way anova adalah:

- a. Data dari populasi (sampel) berjenis interval atau rasio
- b. Sampel yang akan diuji lebih dari 2 kelompok
- c. Populasi yang akan diuji berdistribusi normal
- d. Varian setiap sampel harus sama

Cara menentukan taraf signifikansi:

- $\text{sig} > \alpha \rightarrow \text{Ho diterima}$
- $\text{sig} < \alpha \rightarrow \text{Ho ditolak}$

Cara menentukan kaidah pengujian:

- jika $F_{\text{hit}} \leq F_{\text{tab}} \rightarrow \text{terima Ho}$
- jika $F_{\text{hit}} \geq F_{\text{tab}} \rightarrow \text{tolak Ho}$

Cara melakukan uji One Way Anova dengan SPSS:

- a. Klik Analyze, lalu pilih Compare Means, lalu klik One Way ANOVA.
- b. Lalu pindahkan variabel dependen penelitian ke bagian Dependent list, dan pindahkan variabel yang digunakan sebagai pembanding ke bagian faktor. Klik tanda panah untuk memindahkannya
- c. Setelah itu, klik tombol 'Option', lalu centang pilihan 'descriptive dan homogeneity of variances'. Selain itu abaikan saja, lalu klik continue.
- d. Selanjutnya klik kolom 'Post hoc', centang pilihan 'bonferoni dan tukey'. Lalu klik continue, kemudian yang terakhir klik OK.

7. Uji Two Way Anova

Merupakan pengujian hipotesis komparatif (perbandingan) untuk k sampel (lebih dari 2 sampel) dengan mengukur atau mengelompokkan data berdasarkan dua faktor berpengaruh yang disusun dalam baris dan kolom.

Uji Asumsi yang harus dipenuhi adalah:

- a. Data dari sampel berjenis interval atau rasio
- b. Populasi atau sampel yang akan diuji berdistribusi normal
- c. Varian setiap sampel harus sama
- d. Kelompok data harus memiliki ukuran sampel yang sama

Cara menentukan taraf signifikansi:

- $\text{sig} > \alpha \rightarrow \text{Ho diterima}$
- $\text{sig} < \alpha \rightarrow \text{Ho ditolak}$

Cara menentukan kaidah pengujian:

- jika $F_{\text{hit}} \leq F_{\text{tab}} \rightarrow \text{terima Ho}$
- jika $F_{\text{hit}} \geq F_{\text{tab}} \rightarrow \text{tolak Ho}$

Cara melakukan uji Two Way Anova dengan SPSS:

- a. Klik Analyze - General Linear Model - Univariate.
- b. Masukkan variabel dependen penelitian ke kotak Dependent Variable, masukkan variabel yang digunakan sebagai pembanding (misalnya menggunakan variabel pembanding yaitu gender dan tingkat pendidikan) ke kotak Fixed factor(s).

- c. Klik Plot, maka akan muncul jendela seperti di bawah ini: Masukkan Gender ke kotak Horizontal Axis dan Pendidikan ke kotak Separate Lines. Klik Add
- d. Klik Continue.
- e. Klik Post Hoc, Masukkan variabel pendidikan ke kotak Post Hoc Test, Centang Tukey
- f. Klik Continue
- g. Klik Options, kemudian masukkan Gender, Pendidikan, dan Gender*Pendidikan ke dalam kotak Display Means for. Pada Display centang Descriptive statistics dan Homogeneity test.
- h. Klik Continue

Contoh :

Seorang peneliti ingin mengetahui perbedaan prestasi belajar Statistika antara mahasiswa Akuntansi (X_1), mahasiswa Manajemen (X_2), dengan mahasiswa Matematika (X_3) berdasarkan data seperti di bawah ini.

NO	Akuntansi		Manajemen		Matematika	
	NAMA	PRES	NAMA	PRES	NAMA	PRES
1	Ario	60	Ateew	70	Amro	70
2	Bardhi	75	Bink	50	August	70
3	Baruki	75	Chery	60	Brown	70
4	Cyndi	75	Chian	60	Cherry	90
5	Danur	60	Chow	50	Dino	80
6	Farida	75	Dingo	80	Dorby	65
7	Fasli	65	Feung	50	Forbes	90
8	Haryo	65	Fung	50	Grandy	90
9	Inong	80	Iyon	50	James	90
10	Jarot	60	Jacky	70	Janet	80
11	Kunti	50	Jong	50	John	90
12	Lulu	70	Lee	50	Lennox	40
13	Novia	80	Novie	50	Leroy	80
14	Shinta	50	Simoen	55	Nigel	40
15	Suripto	50	Sontee	60	Roy	70

Penyelesaian :

Mencari nilai statistik dasar

NO	X1	X2	X3	XT	X1 ²	X2 ²	X3 ²	XXT
1	60	70	70	200	3.600	4.900	4.900	13.400
2	75	50	70	195	5.625	2.500	4.900	13.025
3	75	60	70	205	5.625	3.600	4.900	14.125
4	75	60	90	225	5.625	3.600	8.100	17.325
5	60	50	80	190	3.600	2.500	6.400	12.500
6	75	80	65	220	5.625	6.400	4.225	16.250
7	65	50	90	205	4.225	2.500	8.100	14.825
8	65	50	90	205	4.225	2.500	8.100	14.825
9	80	50	90	220	6.400	2.500	8.100	17.000
10	60	70	80	210	3.600	4.900	6.400	14.900
11	50	50	90	190	2.500	2.500	8.100	13.100
12	70	50	40	160	4.900	2.500	1.600	9.000
13	80	50	80	210	6.400	2.500	6.400	15.300
14	50	55	40	145	2.500	3.025	1.600	7.125
15	50	60	70	180	2.500	3.600	4.900	11.000
Σ	990	855	1.115	2.960	66.950	50.025	86.725	203.700

a. Dari tabel di atas diperoleh nilai statistik dasar sbb:

$$\begin{aligned}
 N_1 &= 15; N_2 = 15; N_3 = 15; NT = 45; \\
 \Sigma X_1 &= 990; \Sigma X_2 = 855; \Sigma X_3 = 1.115; \Sigma XT \\
 &= 2.960 \\
 \Sigma X_1^2 &= 66.950; \Sigma X_2^2 = 50.025; \Sigma X_3^2 \\
 &= 86.725 \\
 \Sigma X^2T &= 203.700 \\
 m &= 3
 \end{aligned}$$

b. Mencari Nilai Rata-Rata (M)

$$\begin{aligned}
 M_1 &= \Sigma X_1 : N_1 \\
 M_1 &= 66,000 \\
 \text{analog,} \\
 M_2 &= 57,000 \\
 M_3 &= 74,333 \\
 MT &= 65,778
 \end{aligned}$$

c. Mencari Standar Deviasi Kuadrat (SD^2) $SD_1^2 = (\sum X_1^2 : N_1) - M_1^2$ $SD_1^2 = 107,333$ analog $SD_2^2 = 86,000$ $SD_3^2 = 256,272$ $SDT^2 = 199,922$	d. Mencari Standar Deviasi (SD) $SD_1 = \sqrt{SD_1^2}$ $SD_1 = 10,360$ analog, $SD_2 = 9,274$ $SD_3 = 16,008$ $SDT = 14,139$
e. Mencari Jumlah Kuadrat (JK) $JK_{Tot} = \sum X^2 - \{(\sum XT)^2 : NT\}$ $JK_{Tot} = 8.997,778$ $JK_{Ant} = \{(\sum X_1)^2 : N_1\} + \{(\sum X_2)^2 : N_2\} + \{(\sum X_3)^2 : N_3\} - \{(\sum XT)^2 : NT\}$ $JK_{Ant} = 2.254,445$ $JK_{Dal} = JK_{Tot} - JK_{Ant}$ $JK_{Dal} = 6.743,333$	f. Mencari Derajat Bebas (DB) $DB_{Tot} = NT - 1$ $DB_{Tot} = 44$ $DB_{Ant} = m - 1$ $DB_{Ant} = 2$ $DB_{Dal} = NT - m$ $DB_{Dal} = 42$
g. Mencari Mean Kuadrat (MK) $MK_{Ant} = JK_{Ant} : DB_{Ant}$ $MK_{Ant} = 1.127,222$ $MK_{Dal} = JK_{Dal} : DB_{Dal}$ $MK_{Dal} = 160,555$	h. Mencari Koefisien Varians (F_{Ant}) $F_{Ant} = MK_{Ant} : MK_{Dal}$ $F_{Ant} = 7,021$

i. Menyajikan Hasil Analisis dalam Tabel

Secara ringkas hasil perhitungan statistik tsb dapat disajikan dalam tabel sebagai berikut

Tabel 1:
RINGKASAN HASIL UJI ANALISIS VARIANS (ANAVA)
UNTUK PRESTASI FISIKA MAHASISWA (N=15)

VARIASI	JK	DB	MK	F
Antar	2.254,445	2	1.127,222	7,021
Dalam	6.743,333	42	160,555	--
Total	8.997,780	44	--	--

Jadi, $F_{Ant} = 7,021$

j. Mengkonsultasi F Hitung dengan Nilai F Tabel Di dalam Tabel Nilai F untuk DB = 2/42 ditemukan harga sbb: $F_{(5\%;2;42)} = 3,219$ dan $F_{(1\%;2;42)} = 5,15$. Jadi, FAnt = 7,021 adalah sangat signifikan.	k. Mengkonsultasi F Hitung dengan Nilai F Tabel Di dalam Tabel Nilai F untuk DB = 2/42 ditemukan harga sbb: $F_{(5\%;2;42)} = 3,219$ dan $F_{(1\%;2;42)} = 5,15$. Jadi, FAnt = 7,021 adalah sangat signifikan.
l. Mencari SD Mean Kuadrat (SDm^2) $SDm_1^2 = SD_1^2 : (N_1 - 1)$ $SDm_1^2 = 7,667$ analog, $SDm_2^2 = 6,143$ $SDm_3^2 = 18,305$ $SDmT^2 = 14,280$	m. Mencari SD Beda Mean ($SDbm$) $SDbm_{12} = \sqrt{SDm_1^2 + SDm_2^2}$ $SDbm_{12} = 3,716$ analog, $SDbm_{13} = 5,096$ $SDbm_{23} = 4,944$
n. Mencari Koefisien t $t_{12} = (M_1 - M_2) : SDbm_{12}$ $t_{12} = 2,422$ analog, $t_{13} = -1,635$ $t_{23} = -3,506$	o. Mencari Derajat Bebas (DB) $DB_{12} = (N_1 - 1) + (N_2 - 1)$ $DB_{12} = 28$ analog, $DB_{13} = 28$ $DB_{23} = 28$

Mengkonsultasi t Hitung dengan t Tabel

- Dari Tabel Nilai-Nilai t untuk $DB_{12} = 28$ ditemukan harga sbb:
 $t_{1\%} = 2,763$ dan $t_{5\%} = 2,048$.
Jadi, $t_{12} = 2,422$ adalah signifikan.
- Dari Tabel Nilai-Nilai t untuk $DB_{13} = 28$ ditemukan harga sbb:
 $t_{1\%} = 2,763$ dan $t_{5\%} = 2,048$.
Jadi, $t_{13} = -1,635$ adalah tidak signifikan.
- Dari Tabel Nilai-Nilai t untuk $DB_{23} = 28$ ditemukan harga sbb:
 $t_{1\%} = 2,763$ dan $t_{5\%} = 2,048$.
Jadi, $t_{23} = -3,506$ adalah sangat signifikan.

Menginterpretasi Hasil Analisis

- Dari nilai FAnt = 7,021 berpredikat sangat signifikan dapat diinterpretasi bahwa secara umum terdapat perbedaan prestasi Statistika di antara kelompok mahasiswa Akuntansi, Manajemen dan Matematika.
- Dari nilai $t_{12} = 2,422$ berpredikat signifikan dapat diinterpretasi bahwa secara kasus per kasus prestasi Statistika mahasiswa Akuntansi lebih tinggi daripada mahasiswa Manajemen.
- Dari nilai $t_{tab} = -2,048 < t_{13} = -1,635$ berpredikat tidak signifikan dapat diinterpretasi bahwa secara kasus per kasus tidak terdapat perbedaan prestasi Statistika antara mahasiswa Akuntansi dengan mahasiswa Matematika.

4. Dari nilai $t_{tab} = -2,048 > t_{23} = -3,506$ dengan predikat sangat signifikan dapat diinterpretasi secara kasus per kasus prestasi Statistika mahasiswa Manajemen lebih rendah daripada mahasiswa Matematika.
5. Menyimpulkan Hasil Analisis
Secara umum terdapat perbedaan prestasi Statistika di antara mahasiswa Akuntansi, Manajemen dan Matematika; dan secara kasus per kasus prestasi Statistika mahasiswa Manajemen lebih rendah daripada mahasiswa Akuntansi dan Matematika.

6. Analisis Regresi Sederhana

Istilah regresi pertama kali dalam konsep statistik digunakan oleh Sir Francis Galton pada tahun 1886 di mana yang bersangkutan melakukan kajian yang menunjukkan bahwa tinggi badan anak-anak yang dilahirkan dari para orang tua yang tinggi cenderung bergerak (*regress*) ke arah ketinggian rata-rata populasi secara keseluruhan. Galton memperkenalkan kata regresi (*regression*) sebagai nama proses umum untuk memprediksi satu variabel, yaitu tinggi badan anak dengan menggunakan variabel lain, yaitu tinggi badan orang tua. Dapat disimpulkan bahwa analisis regresi sederhana merupakan instrumen statistik yang digunakan untuk mengerti pengaruh antara satu variabel bebas yang dapat disebut prediktor atau variabel independen (x) terhadap satu variabel terikat yang disebut juga dengan istilah kriterium atau variabel dependen (y).

Penerapan analisis regresi sederhana dapat diketahui secara luas pada berbagai bidang, misalnya pada bidang pertanian, pendidikan, ekonomi, kesehatan dan sebagainya. Contoh penerapan analisis regresi sederhana pada berbagai bidang diantaranya:

- a. Pengaruh berat badan terhadap tinggi badan
- b. Pengaruh penggunaan rokok terhadap paru-paru
- c. Pengaruh iklan terhadap penjualan pada perusahaan
- d. Pengaruh minat belajar siswa terhadap nilai ujian
- e. Pengaruh jumlah pupuk terhadap produksi pertanian

$$\hat{Y} = \hat{\beta}_0 + \hat{\beta}_1 X$$

Dengan:

X = variabel independen

Y = variabel dependen

β_0 = konstanta (*intercept*)

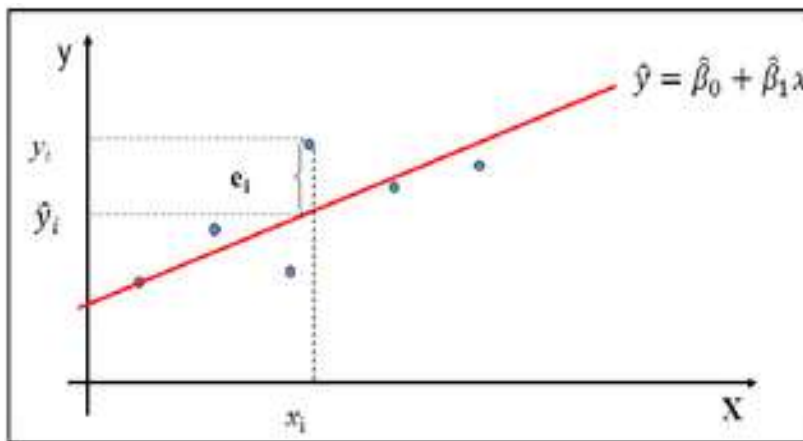
β_1 = koefisien (*slope*)

Untuk β_0 dan β_1 dapat dicari menggunakan persamaan dibawah ini

$$\beta_0 = \frac{(\sum X_i)(\sum X_i^2) - (\sum X_i)(\sum X_i Y_i)}{n \sum X_i^2 - (\sum X_i)^2}$$

$$\beta_1 = \frac{n(\sum X_i Y_i) - (\sum X_i)(\sum Y_i)}{n \sum X_i^2 - (\sum X_i)^2}$$

Persamaan analisis regresi sederhana juga dapat dibuat dalam bentuk grafik dibawah ini:



Dengan:

$\hat{y}$ = persamaan regresi

$\hat{y}_i$ = prediksi y (ke-i = 1, 2, 3, ..., n observasi)

y_1 = variabel dependen (ke-i = 1,2,3...n observasi)

x_1 = variabel independen (ke-i = 1,2,3...n observasi)

β_0 = konstanta (*intercept*)

β_1 = koefisien (*slope*)

e_i = variabel pengganggu atau residual (ke-i = 1, 2, 3, ..., n observasi)

7. Asumsi Untuk Analisis Regresi Sederhana

Sebelum lanjut pada tahap pengerjaan analisis regresi sederhana, ada beberapa asumsi yang harus dipenuhi, diantaranya :

- Model regresi harus linier dalam parameter
- Variabel bebas tidak berkorelasi dengan error
- Nilai error $E\{\varepsilon_i\} = 0$
- Varian untuk masing-masing error term (kesalahan) tetap
- Tidak terjadi auto korelasi
- Model regresi dispesifikasi secara benar

8. Tahap Pengerjaan Analisis Regresi Sederhana

- Menentukan tujuan dari penelitian.
- Mengidentifikasi X (variabel independen) dan y (variabel dependen).
- Menghitung nilai dari X^2 , Y^2 , XY dan total masing-masing sebagai awal untuk menghitung β_0 dan β_1 .
- Menghitung nilai β_0 dan β_1 .
- Membuat model persamaan analisis regresi sederhana
- Melakukan prediksi atau interpretasi dari persamaan analisis regresi sederhana

9. Contoh Kasus Analisis Regresi Sederhana

Misalkan:

Seorang manager rumah sakit ingin mengetahui hubungan antara lamanya tenaga keperawatan pasien UGD dalam satuan tahun (x) dengan banyaknya pasien yang operasi berhasil (y). diperoleh data sebagai berikut.

X	Y
1	2
5	4
4	6
2	4
3	2

Buatlah model regresi hubungan lamanya tenaga keperawatan pasien UGD dengan banyaknya pasien yang operasinya berhasil dan hitung koefisien determinasinya!

Jawab:

Untuk menyelesaikan soal diatas, berikut langkah-langkahnya:

- Menentukan tujuan: Pengaruh lamanya tenaga keperawatan UGD terhadap banyaknya pasien yang operasi berhasil.
- Menentukan variabel X dan Y: variabel X dan Y sudah diketahui dari tabel diatas
- Menghitung nilai dari X^2, Y^2, XY

X	Y	X^2	Y^2	XY
1	2	1	4	2
5	4	25	16	20
4	6	16	36	24
2	4	4	16	8
3	2	9	4	6
$\Sigma X = 15$	$\Sigma Y = 18$	$\Sigma X^2 = 55$	$\Sigma Y^2 = 76$	$\Sigma XY = 60$

- Menghitung nilai β_0 dan β_1

$$\beta_0 = \frac{(\Sigma X_i)(\Sigma X_i^2) - (\Sigma X_i)(\Sigma X_i Y_i)}{n \Sigma X_i^2 - (\Sigma X_i)^2}$$

$$\beta_0 = \frac{(15)(55) - (15)(60)}{(5)(55) - (15)^2}$$

$$\beta_0 = \frac{825 - 900}{275 - 225}$$

$$\beta_0 = \frac{-75}{50}$$

$$\beta_0 = -1,5$$

$$\beta_1 = \frac{n(\sum X_i Y_i) - (\sum X_i)(\sum Y_i)}{n\sum X_i^2 - (\sum X_i)^2}$$

$$\beta_1 = \frac{(5)(60) - (15)(18)}{(5)(55) - (15)^2}$$

$$\beta_1 = \frac{300 - 270}{275 - 225}$$

$$\beta_1 = \frac{30}{50}$$

$$\beta_1 = 0,6$$

2. Membuat model persamaan analisis regresi sederhana

$$\hat{Y} = \hat{\beta}_0 + \hat{\beta}_1 X$$

$$\hat{Y} = -1,5 + 0,6X$$

10. Regresi Linier Berganda

Metode yang juga digunakan adalah regresi berganda. Regresi ini menggunakan beberapa variabel X , misalnya X_1 , X_2 dan seterusnya yang kemudian dianalisis secara bersamaan. Rumus yang digunakan dalam regresi berganda pada prinsipnya sama regresi sederhana, hanya saja regresi berganda ditambahkan variabel lain yang juga disertakan dalam penelitian. Rumus regresi linier berganda adalah sebagai berikut:

$$Y = a + b_1 X_1 + b_2 X_2 + \dots + b_n X_n$$

Keterangan:

Y = Variabel dependen (nilai yang diprediksikan)

X = Variabel independen

a = Konstanta

b = Koefisien regresi

Penelitian yang menggunakan regresi berganda regresi linier berganda yang meliputi:

- a. Analisis hubungan antara populasi tanaman, dosis pemupukan dan hasil menanam.
- b. Analisis hubungan antara waktu tanam dan faktor iklim terhadap dinamika hama menanam.
- c. Menganalisis hubungan antara jumlah ransum dan waktu pemberian peningkatan bobot ternak.
- d. Menganalisis hubungan antara umur, tinggi tanaman, dan hasil.
- e. Analisis hubungan antara gender dan tingkat adopsi teknologi pertanian modern.

CONTOH KASUS:

"Aplikasi Regresi Linier Berganda untuk Mengetahui Pengaruh Umur, Tinggi Tanaman dan Rendemen Terhadap Hasil Jagung."

Sebuah penelitian dilakukan untuk mengkaji hubungan antara ketiga variabel tersebut yakni tinggi tanaman, umur panen dan hasil tanaman jagung dari berbagai varietas. Data yang dikumpulkan sebagai berikut

Nomor	Umur tanaman (hari)	Tinggi tanaman (cm)	Rendemen (%)	Hasil (t/ha)
1	100	203	70	9.5
2	102	206	72	9.8
3	98	200	68	9.1
4	95	198	65	8.6
5	102	204	69	9.7
6	104	210	72	10
7	98	199	69	9
8	92	190	63	8
9	102	204	71	9.7
10	100	202	71	9.6
11	102	205	73	9.8
12	85	190	67	7.8
13	90	193	69	8
14	92	194	64	8.1
15	98	199	69	9
16	102	205	71	9.7

Penyelesaian:

Model yang akan digunakan untuk analisis data adalah regresi linier. Tahapan analisisnya adalah:

1. Buka program Excel Microsoft Office dan lakukan tabulasi seperti berikut simpan dengan nama **regresi berganda.xls**
2. Buka program SPSS pada computer, selanjutnya akan muncul data view pada computer. Impor data dari Excel dengan klik **File > Open > Data**. Selanjutnya pada dialog **File Type** pilih **Excel** dan **File nama** pilih **regresi berganda.xls** dilanjutkan dengan klik **Open**. Klik **Continue** maka data akan ditampilkan di data view spss seperti berikut.
3. Selanjutnya kita akan melakukan analisis regresi, klik **Analyze > Regression > Linear regression**
4. Pilih variabel Hasil dan klik ke **Dependent List**, variabel Hasil akan berpindah ke kanan (lihat gambar 3). Selanjutnya pada **Independent List** pilih variabel **Tinggi**, **umur** dan **rendemen**. Klik tanda panah ke kanan, variabel akan berpindah.
5. Masih pada kotak dialog Linear regression Klik **statistics** dan tandai pada **Estimates**, **Model Fit** dan **Descriptives** dilanjutkan dengan klik **Continue**.
6. Masih pada kotak dialog Linear regression klik **Plots** dan tandai pilihan **Histogram** dan **Normal Probability Plot**. dilanjutkan dengan klik **Continue > OK**.

OUTPUT MODEL

Descriptive Statistics

	Mean	Std. Deviation	N
Hasil	9.0875	.75971	16
Umur	97.62	5.390	16
Tinggi	200.12	5.898	16
Rendemen	68.94	2.932	16

Interpretasi tabel: Tabel ini menjelaskan deskripsi variabel seperti rata-rata (mean), standar deviasi dan jumlah data (N). Nilai rata-rata variabel **Hasil** adalah 9,09 t/ha dengan rata-rata penyimpangan (deviasi mencapai 0,75) dengan jumlah data 16. Demikian pula pada **Umur** dan **Tinggi**, mempunyai nilai rata-rata 97,62 hari dan 200,12 cm dengan penyimpangan 5,39 dan 5,89 dengan jumlah data 16.

Model Summary^b

Model	R	R Square	Adjusted R Square	Std. Error of the Estimate
1	.991 ^a	.982	.977	.11434

a. Predictors: (Constant), Rendemen, Umur, Tinggi

b. Dependent Variabel: Hasil

Interpretasi tabel: Nilai korelasi antara variabel prediktor (umur, tinggi tanaman, rendemen) dengan variabel hasil ($R = 0,991$) sehingga dapat disimpulkan bahwa terdapat hubungan yang sangat erat antara umur dan tinggi tanaman serta rendemen terhadap hasil yang didapatkan.

Nilai R-square atau koefisien determinasi sebesar 0,982. Hal ini menunjukkan bahwa kemampuan variabel umur, tinggi tanaman, rendemen mempengaruhi hasil panen sebesar 98,2% dan masih terdapat $100 - 98,2 = 1,8\%$ variabel lain (selain ketiga variabel tersebut) yang mempengaruhi hasil

Model		Sum of Squares	Df	Mean Square	F	Sig.
1	Regression	8.501	3	2.834	216.741	.000 ^a
	Residual	.157	12	.013		
	Total	8.658	15			

a. Predictors: (Constant), Rendemen, Umur, Tinggi

b. Dependent Variabel: Hasil

Interpretasi tabel: Nilai korelasi antara variabel prediktor (umur, tinggi tanaman, rendemen) dengan variabel hasil (R) = 0,991 sehingga dapat disimpulkan bahwa terdapat hubungan yang sangat erat antara umur dan tinggi tanaman serta rendemen terhadap hasil yang didapatkan.

Nilai R -square atau koefisien determinasi sebesar 0,982. Hal ini menunjukkan bahwa kemampuan variabel umur, tinggi tanaman, rendemen mempengaruhi hasil panen sebesar 98,2% dan masih terdapat $100 - 98,2 = 1,8\%$ variabel lain (selain ketiga variabel tersebut) yang mempengaruhi hasil.

ANOVA^b

Model		Sum of Squares	Df	Mean Square	F	Sig.
1	Regression	8.501	3	2.834	216.741	.000 ^a
	Residual	.157	12	.013		
	Total	8.658	15			

a. Predictors: (Constant), Rendemen, Umur, Tinggi

b. Dependent Variabel: Hasil

Interpretasi: Uji Anova dilakukan untuk mengetahui ada tidaknya pengaruh variabel umur, tinggi dan rendemen terhadap hasil. Apabila nilai Sig atau P-value < 0,05 maka terdapat hubungan yang nyata antara variabel tersebut dengan hasil. Demikian pula apabila Sig > 0,05 maka dapat disimpulkan tidak ada hubungan antara variabel dengan hasil. Seperti terlihat pada tabel Anova, nilai Sig model sebesar 0,000 (<0,05) sehingga dapat disimpulkan terdapat sedikitnya 1 faktor yang berpengaruh secara signifikan dengan hasil jagung.

Coefficients^a

Model		Unstandardized Coefficients		Standardized Coefficients	t	Sig.
		B	Std. Error	Beta		
1	(Constant)	-9.488	1.842		-5.150	.000
	Umur	.079	.019	.558	4.120	.001
	Tinggi	.039	.021	.299	1.861	.087
	Rendemen	.046	.018	.178	2.539	.026

a. Dependent Variabel: Hasil

Interpretasi: Tabel coefficient menampilkan koefisien dari persamaan regresi yang dihasilkan. Berdasarkan tabel di atas, model regresi linier berganda dapat dituliskan sebagai berikut:

$$Y = -9,488 + 0,079 X_1 + 0,039 X_2 + 0,046 X_3$$

Dimana: X_1 = umur tanaman (hari), X_2 = tinggi tanaman (cm), X_3 = Rendemen

C. Rangkuman

1. Uji statistik parametrik adalah teknik pengujian data dalam statistik yang mempertimbangkan jenis sebaran atau distribusi data, yaitu apakah data menyebar secara normal atau tidak
2. Uji beda rata-rata satu sampel (independent) adalah metode pengujian untuk mengetahui perbedaan rata-rata dua populasi/kelompok data yang independen/bebas
3. Uji beda rata-rata dua sampel (dependent) adalah salah satu metode pengujian hipotesis dimana data yang digunakan tidak bebas (dependen)
4. Uji beda rata-rata lebih dari dua sampel (ANOVA) adalah metode untuk menguji dua varians (atau ragam) berdasarkan hipotesis nol bahwa kedua varians itu sama
5. Uji Regresi Linier Sederhana merupakan instrumen statistik yang digunakan untuk mengerti pengaruh antara satu variabel bebas yang dapat disebut prediktor atau variabel independen (x) terhadap satu variabel terikat yang disebut juga dengan istilah kriterium atau variabel dependen (y).
6. Uji Regresi Linier Ganda adalah uji regresi yang menggunakan beberapa variabel X , misalnya X_1 , X_2 dan seterusnya yang kemudian dianalisis secara bersamaan. Rumus yang digunakan dalam regresi berganda pada prinsipnya sama regresi sederhana, hanya saja regresi berganda ditambahkan variabel lain yang juga disertakan dalam penelitian.

D. Tugas

1. Seorang peneliti ingin menerapkan dua metode pengobatan yang berbeda, sebutlah metode A dan metode B. kedua metode pengobatan diterapkan pada sekelompok pasien berjumlah 15 orang. Tentukan metode manakah yang lebih efektif dalam menyembuhkan suatu penyakit. Data lama obat (hari) dalam menyembuhkan suatu penyakit seperti tercantum dibawah ini:

ID	A	B	C	D	E	F	G	H	I	J	K	L	M	N	O
Met.A	6	9	7	6	6	7	5	4	8	7	9	5	4	8	7
Met.B	6	7	8	4	3	9	4	6	7	8	9	4	3		5

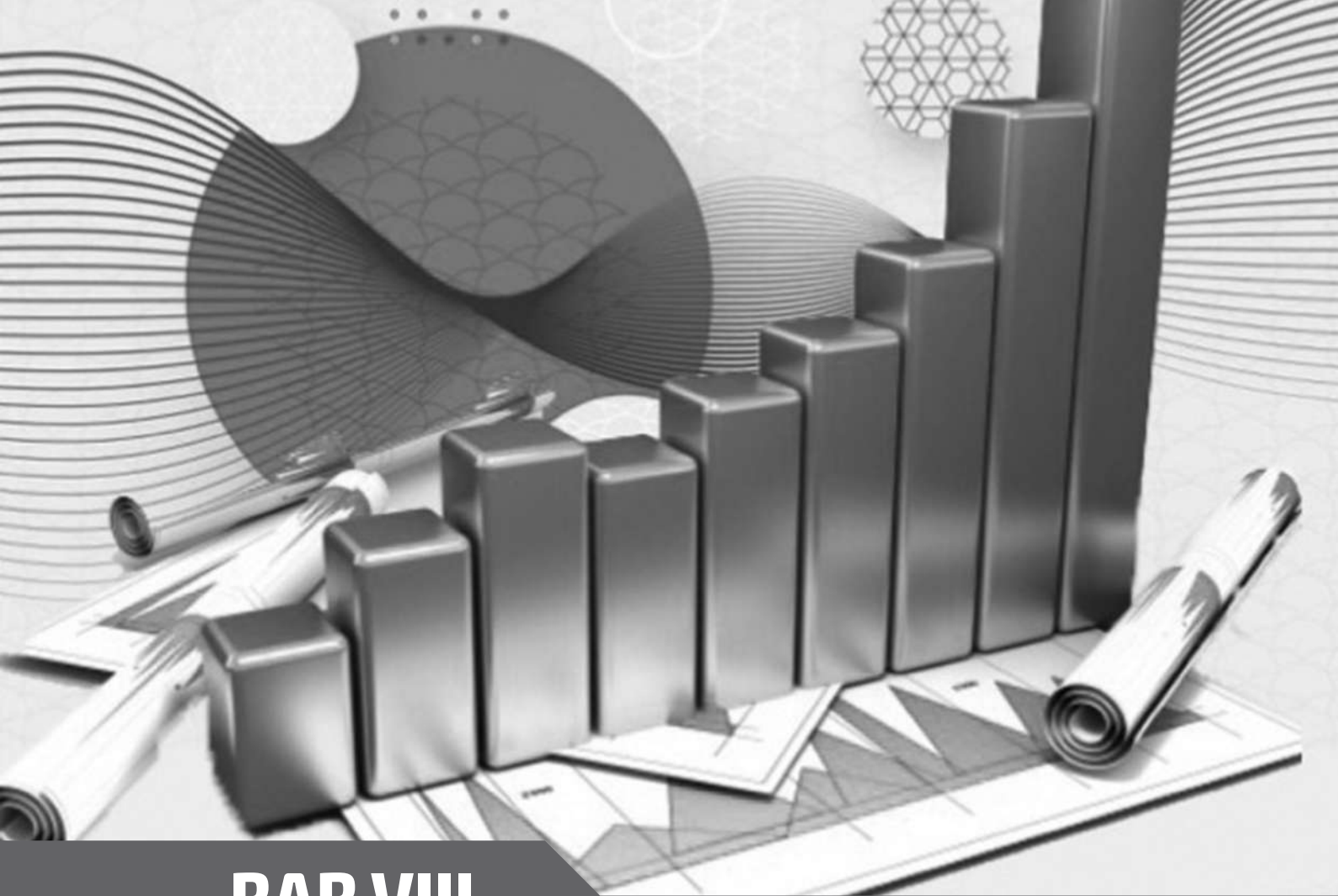
2. Sebuah penelitian dilakukan untuk mengetahui apakah ada pengaruh makanan ikan (tiap hari dalam seminggu) [X_1] dan Panjang ikan (mm) [X_2] terhadap berat ikan (kg) [Y] di Desa Tani Tambak Kalianyar Blitar Sebagai Berikut:

ID	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15
X_1	8	10	7	12	9	10	7	8	11	8	10	8	10	12	15
X_2	125	137	100	122	129	128	98	103	130	95	115	105	120	125	130
Y	37	41	34	39	40	42	38	42	40	36	41	38	40	45	47

- a. Berapa besar persamaan regresi gandanya?
- b. Buktikan apakah terdapat pengaruh yang signifikan makanan ikan dan Panjang ikan terhadap berat ikan di Desa Tani Kalianyar Blitar?

E. Referensi

- Budiwanto, Setyo. 2017. Metode Statistika Untuk Mengolah Data Keolahragaan. Malang: UM Penerbit dan Percetakan.
- Suyanto, Ahmad Ikhlasul Amal, Moh. Arifin Noor, dan Indra Tri Astutik. 2018. ANALISIS DATA PENELITIAN Petunjuk Praktis Bagi Mahasiswa Kesehatan Menggunakan SPSS. Semarang: UNISSULA PRESS.
- Kadek, Luh Pande Ary Susilawati, dkk. 2017. Bahan Ajar PRAKTIKUM STATISTIK. Denpasar.
- Purnomo, Hari, dan Eka Siswanto Syamsul. 2017. STATISTIKA FARMASI (Aplikasi Praktis dengan SPSS). Yogyakarta: Grafika Indah.
- Nuryadi, dkk. 2017. Dasar-Dasar Statistik Penelitian. Yogyakarta: SIBUKU MEDIA.
- Riduwan dan Sunato. 2009. PENGANTAR STATISTKA untuk penelitian Pendidikan, social, ekonomi, komunikasi, dan bisnis. Bandung: Alfabeta



BAB VIII

ANALISIS STATISTIK NON PARAMETRIK

A. Tujuan pembelajaran :

Setelah mempelajari pokok bahasan non parametrik mahasiswa mampu :

1. Memahami pengertian statistika non parametrik
2. Memahami kapan metode statistik non parametrik digunakan
3. Memahami pedoman penggunaan statistik non parametrik
4. Memahami jenis-jenis uji statistik non parametrik pada analisis korelasi (uji rank spearman, uji kendall tau) .
5. Memahami jenis-jenis uji statistik non parametrik pada analisis komparasi (uji chi Square, uji kolmogorov smirnov, uji friedman, uji kruskal-wallis, uji wilcoxon, uji mann whitney)
6. Memahami penggunaan aplikasi SPSS uji statistik non parametrik.
7. Memahami latihan soal dan pembahasan masing-masing jenis uji statistik non parametrik.

B. Materi

1. Statistik Non-parametrik

Statistik Non-parametrik didasarkan pada model yang tidak berdasarkan pada bentuk khusus dari distribusi data. Asumsi yang berhubungan dengan uji statistik Non-parametrik meliputi observasi harus independen, pengukuran variabel dengan skala ordinal dan skala nominal (kategorikal), data tidak berdistribusi normal dan jumlah sampel kecil (<20). Berikut beberapa keunggulan uji statistik Non-parametrik:

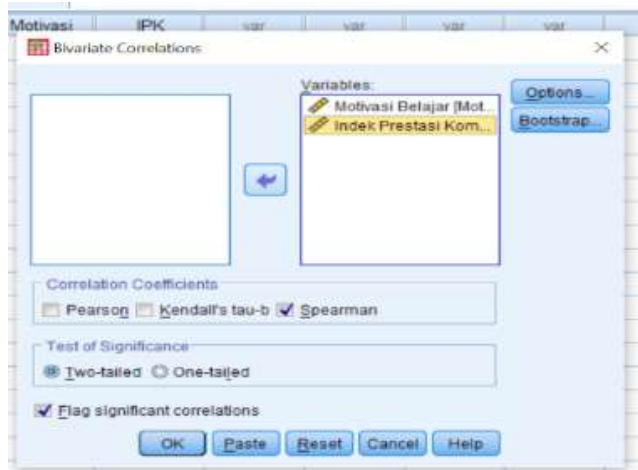
- a. Jika jumlah sampel terlalu kecil, maka tidak ada alternatif lain menggunakan uji statistik Non-parametrik, kecuali distribusi populasi diketahui dengan pasti.
- b. Uji Non-parametrik memiliki asumsi yang lebih sedikit berkaitan dengan data dan mungkin lebih relevan pada situasi tertentu. Hipotesis yang diuji dengan Non-parametrik mungkin lebih sesuai dengan tujuan penelitian.
- c. Uji Non-parametrik digunakan untuk menganalisis data yang secara inheren adalah data dalam bentuk ranking. Jadi si peneliti hanya dapat mengatakan terhadap subyek penelitian bahwa yang satu memiliki lebih atau kurang karakteristik dibandingkan dengan lainnya, tanpa dapat mengatakan seberapa besar atau kecil itu.
- d. Uji Non-parametrik cocok untuk menguji data yang bersifat klasifikasi atau kategorikal (skala nominal). Tidak ada uji parametrik yang cocok untuk menguji data seperti ini.
- e. Ada uji statistik Non-parametrik yang cocok untuk menguji sampai yang berasal dari observasi yang diambil dari populasi yang berbeda. Uji parametrik sering kesulitan menguji data seperti ini.

2. Analisis Korelasi dalam Uji Statistika Non Parametrik

a. Analisis Korelasi Rank Spearman

Uji Korelasi Spearman digunakan untuk mencari hubungan atau untuk menguji signifikansi hipotesis asosiatif antara variabel bebas dan variabel terikat tidak berdistribusi normal atau non-parametrik, dengan menggunakan skala ukur paling tidak skala ordinal

- 1) Langkah Analisis
 - a) Buka file **Spearman.xls**
 - b) Klik **variable View** dan **data View**.
 - c) Pada menu utama SPSS, lalu pilih **Analyze**, kemudia pilih **Correlate**, lalu pilih **Bivariate**
 - d) sikan variabel yang akan dianalisis pada kotak **variabel** dalam hal ini adalah variabel Motivasi dan IPK
 - e) Pada **Corerelation Corfficient** pilih **Sperman**
 - f) Tampak dilayar tampilan *Windows Bivariate Correlation*



- g) Abaikan lainnya lalu tekan **OK**
- 2) Tampilan output SPSS

Correlations				
			Motivasi Belajar	Indek Prestasi Kumulatif
Spearman's rho	Motivasi Belajar	Correlation Coefficient	1.000	.928**
		Sig. (2-tailed)	.	.000
		N	15	15
	Indek Prestasi Kumulatif	Correlation Coefficient	.928**	1.000
		Sig. (2-tailed)	.000	.
		N	15	15

** . Correlation is significant at the 0.01 level (2-tailed).

- 3) Latihan Soal

Gunakan langkah seperti yang tertera diatas,

Misalkan: Sekolah Insani Husada mendapatkan gambaran motivasi siswa dan hasil belajar. Kemudian para siswa memberikan penilaian terhadap motivasi belajar yang dimiliki dengan ketentuan, 1. Kurang, 2. Sedang, 3. Baik. Tujuannya untuk melihat hubungan signifikan antara motivasi belajar dan hasil belajar.

Hasil penilaian masing-masing bimbingan sebagai berikut:

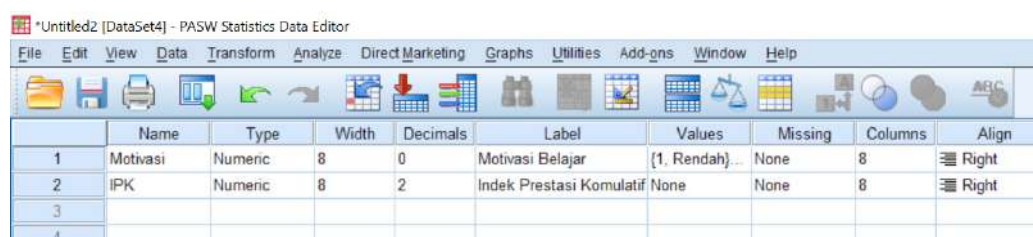
Tabel 8.1. Uji Spearman

Sampel	Motivasi Belajar	Hasil Belajar (IPK)
1.	1	2.40
2.	2	3.00
3.	3	3.55
4.	2	2.90
5.	2	3.05
6.	1	2.42
7	3	3.55
8	2	3.10
9	3	3.45
10	1	2.52
11	1	2.30
11	2	3.00
12	2	3.30
13	2	3.05
15	1	2.32

Langkah Uji

- Buka file **Spearman.xls**
- Klik **variable View** dan **data View**.

Pada variable View pada baris pertama nama , ketikkan "Motivasi" , pada baris kedua ketikkan "IPK"



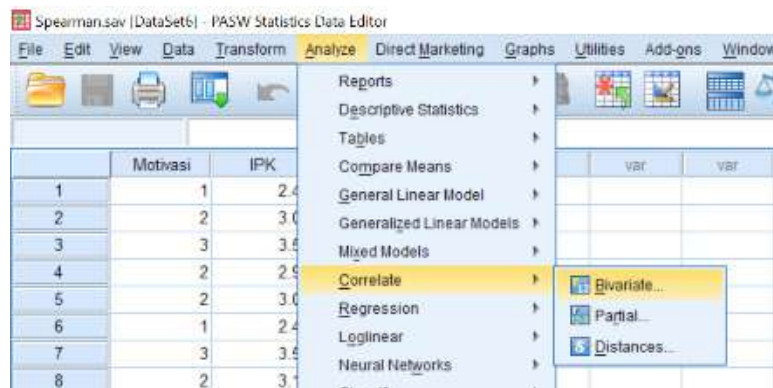
Pada data View ketik data berbentuk interval seperti pada table 8.1 jika telah selesai akan terlihat seperti ini

SPSS Statistics Data Editor - *Untitled2 [DataSet4]

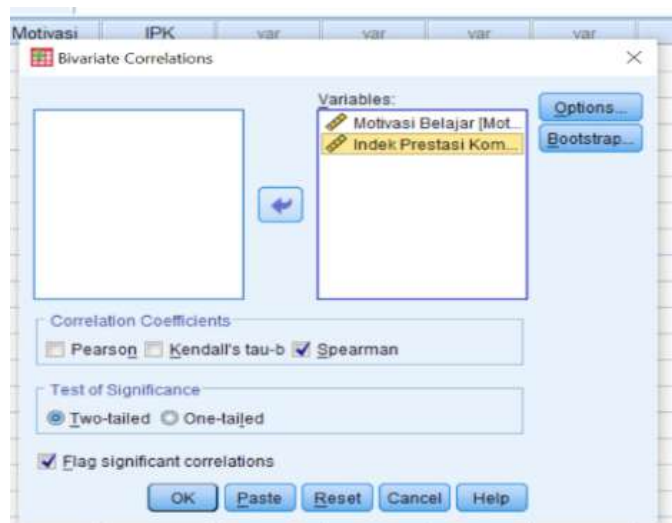
	Motivasi	IPK	var
1	1	2.40	
2	2	3.00	
3	3	3.55	
4	2	2.90	
5	2	3.05	
6	1	2.42	
7	3	3.55	
8	2	3.10	
9	3	3.45	
10	1	2.52	
11	1	2.30	

Dsb...

- Pada menu utama SPSS, lalu pilih **Analyze**, kemudia pilih **Correlate**, lalu pilih **Bivariate**



- Isikan variabel yang akan dianalisis pada kotak **variabel** dalam hal ini adalah variabel Motivasi dan IPK
- Pada **Correlation Coefficient** pilih **Spearman** Tampak dilayar tampilan *Windows Bivariate Correlation*



- Abaikan lainnya lalu tekan **OK**

Tampilan output SPSS

Correlations				
			Motivasi Belajar	Indek Prestasi Kumulatif
Spearman's rho	Motivasi Belajar	Correlation Coefficient	1.000	.928**
		Sig. (2-tailed)	.	.000
		N	15	15
	Indek Prestasi Kumulatif	Correlation Coefficient	.928**	1.000
		Sig. (2-tailed)	.000	.
		N	15	15

** . Correlation is significant at the 0.01 level (2-tailed).

Kesimpulannya :

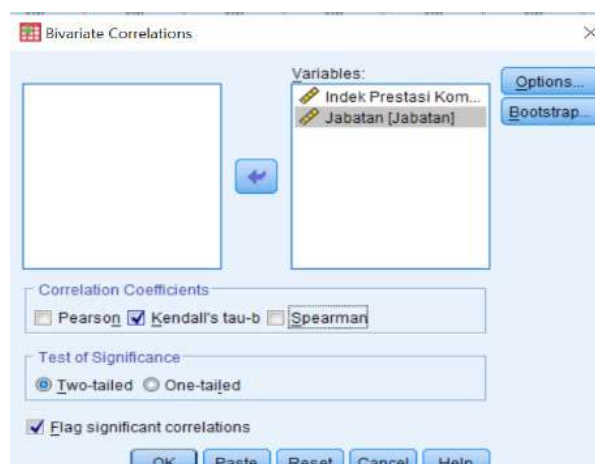
Berdasarkan hasil olah data terlihat Asymp Sig sebesar 0,000, karena nilai Asymp Sig lebih kecil dari α 5 %, maka tolak H_0 , terima H_1 , dapat disimpulkan terdapat hubungan positif atau signifikan antara motivasi belajar dengan hasil belajar (IPK).

b. Analisis Koefisien Korelasi Kendall Tau

Uji Korelasi Kendall Tau digunakan untuk mengukur seberapa kuat (derajat kedekatan) suatu hubungan (asosiasi) yang terjadi antara variabel bebas dan variabel terikat tidak berdistribusi normal atau non-parametrik, dengan menggunakan skala ukur paling tidak skala ordinal

1) Langkah Analisis

- Buka file **Kendal Tau.xls**
- Klik **variable View** dan **data View**.
 - Pada menu utama SPSS, lalu pilih **Analyze**, kemudia pilih **Correlate**, lalu pilih **Bivariate**
- Isikan variabel yang akan dianalisis pada kotak **variabel**
- Pada **Correlation Coefficient** pilih **Kendal Tau**
- Tampak dilayar tampilan *Windows Bivariate Correlation*



- f) Abaikan lainnya lalu tekan **OK**
- g) Tampilan output SPSS

Correlations			Indek Prestasi Kumulatif	Jabatan
Kendall's tau_b	Indek Prestasi Kumulatif	Correlation Coefficient	1.000	.830**
		Sig. (2-tailed)	.	.000
		N	30	30
	Jabatan	Correlation Coefficient	.830**	1.000
		Sig. (2-tailed)	.000	.
		N	30	30

** . Correlation is significant at the 0.01 level (2-tailed).

2) Latihan Soal

Gunakan langkah seperti yang tertera diatas,

Peneliti ingin mengetahui apakah ada hubungan kaitannya antara IPK, jabatan dan gaji 30 sampel karyawan Rumah Sakit Panti Husada dengan masa kerja sekian tahun, dimana sampel diambil berupa 1. Staff 2. Superfisor, 3. Manager. Data sampel disajikan sebagai berikut :

Tabel 8.2. Uji Kendal Tau

Sampel	IPK	Jabatan	Sampel	IPK	Jabatan
1.	2.00	1	16	3.00	2
2.	3.00	2	17	3.65	3
3.	3.55	3	18	2.90	2
4.	2.90	2	19	3.05	2
5.	2.55	2	20	2.12	1
6.	2.12	1	21	2.00	1
7	3.55	3	22	2.90	2
8	2.90	2	23	3.05	2
9	3.05	2	24	2.12	1
10	2.12	1	25	2.00	1
11	2.00	1	26	3.00	2
11	3.00	2	27	3.75	3
12	2.90	2	28	2.00	1
13	3.05	2	29	2.90	2
15	2.12	1	30	3.05	2

Langkah Uji

- Buka file **Kendal Tau.xls**
- Klik **variable View** dan **data View**.

Pada variable View pada baris pertama nema , ketikan "IPK", pada baris kedua ketikkan " Jabatan"

*Kendal Tau.sav [DataSet4] - PASW Statistics Data Editor

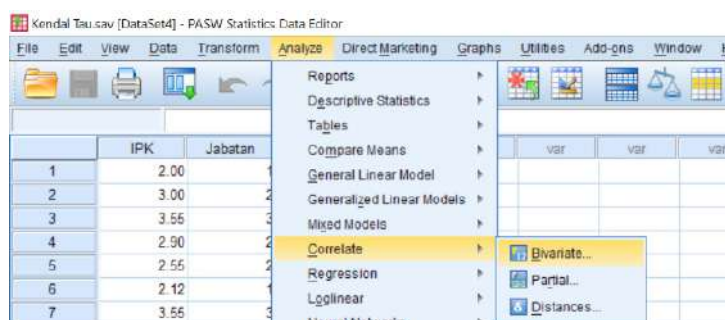
	Name	Type	Width	Decimals	Label	Values	Missing	Columns	
1	IPK	Numeric	8	2	Indek Prestasi Kumulatif	None	None	8	Rig
2	Jabatan	Numeric	8	0	Jabatan	{1, Staff}...	None	8	Rig

Pada data View ketik data berbentuk interval seperti pada table 8.2. jika telah selesai akan terlihat seperti ini

	IPK	Jabatan	var	var
1	2.00	1		
2	3.00	2		
3	3.55	3		
4	2.90	2		
5	2.55	2		
6	2.12	1		
7	3.55	3		
8	2.90	2		
9	3.05	2		
10	2.12	1		
11	2.00	1		

Dst.....

- Pada menu utama SPSS, lalu pilih **Analyze**, kemudia pilih **Correlate**, lalu pilih **Bivariate**



- Isikan variabel yang akan dianalisis pada kotak **variabel** pada **Corerelation Corfficient** pilih **Kendal Tau**.
- Tampak dilayar tampilan *Windows Bivariate Correlation*



- Abaikan lainnya lalu tekan **OK**

- Tampilan output SPSS

Correlations			Indek Prestasi Kumulatif	Jabatan
Kendall's tau_b	Indek Prestasi Kumulatif	Correlation Coefficient	1.000	.830**
		Sig. (2-tailed)	.	.000
		N	30	30
	Jabatan	Correlation Coefficient	.830**	1.000
		Sig. (2-tailed)	.000	.
		N	30	30

** . Correlation is significant at the 0.01 level (2-tailed).

Kesimpulannya :

Berdasarkan hasil olah data terlihat Asymp Sig sebesar 0,000, karena nilai Asymp Sig lebih kecil dari α 5 %, maka tolak H_0 , terima H_1 , dapat disimpulkan terdapat hubungan positif atau signifikan antara Indek Prestasi Kumulatif dengan Jabatan

3. Analisis Komparasi dalam Uji Statiska Non Parametrik

a. Analisis Komparasi Chi Square

Uji Chi Square merupakan salah satu metode statistik non parametrik, antara lain digunakan untuk uji kesesuaian, uji independensi, dan uji homogenitas. Uji Chi Square cocok untuk menganalisis data dengan jumlah kategori dua atau lebih. Teknik yang digunakan adalah *goodness of-fit* dan dapat digunakan untuk menguji apakah terdapat perbedaan signifikansi antara jumlah obyek atau response yang diobservasi yang jatuh pada setiap kategori dan jumlah obyek yang diharapkan (*expected*) berdasarkan pada hipotesis nol. Teknik yang dapat digunakan untuk mengetahui kuat lemah hubungan, demikian juga dapat digunakan untuk menentukan homogenitas . Data yang digunakan dengan data berdistribusi tidak normal dan dengan skala ukuran nominal dan ordinal.

b. Metode

Untuk membandingkan antara frekuensi hasil observasi dengan frekuensi grup yang diharapkan (*expected*), kita harus mampu menyatakan frekuensi yang diharapkan. Hipotesis nol H_0 menyatakan bahwa proposi obyek masuk dalam setiap kategori pada populasi yang diasumsikan. Dari hipotesis nol, kita dapat menarik kesimpulan frekuensi apa yang diharapkan. Teknik Chi-Square memberikan probabilitas bahwa frekuensi yang diobservasi telah dipilih dari populasi dengan nilai *expected* tertentu.

Aturan yang berlaku pada Chi Square :

- Bila pada table 2x2 dijumpai nilai Expected kurang dari 5, maka digunakan adalah **Fisher's exact test**
- Bila pada table 2x2 dan tidak ada nilai E, maka digunakan adalah Continuity correction

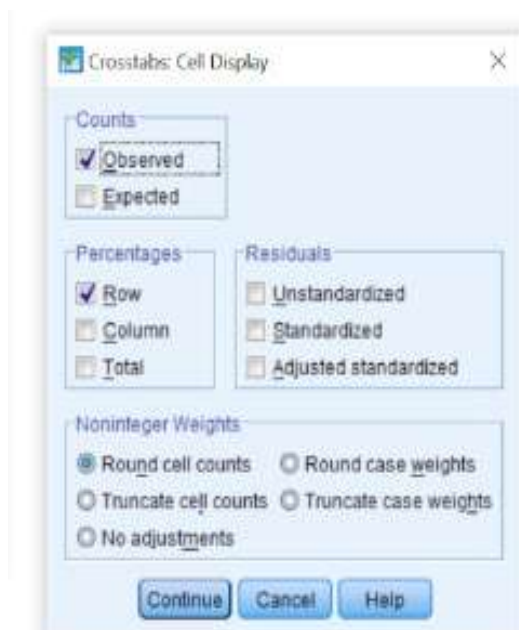
- Bila table lebih dari 2x2, misalnya 3x2, 3x3 maka digunakan adalah **Pearson chi square**

c. Langkah Analisis

- Buka **file KS2.xls**
- Klik **variable View** dan **data View**.
- Dari menu utama SPSS, pilih menu **analyze**, kemudian pilih **Discriptif statistik**, lalu pilih **crosstabs**.
- Dari menu crosstab ada 2 dua kotak yang harus diisi, pada kotak Row, diisi variable independent, variable pekerjaan dimasukan, sedang pada colum diisi variable dependen, variable menyusui dimasukan
- Klik **option statistic** , klik pilihan **Chisure** dan klik pilihan Risk
- Klik Continue



- Klik option **Cells** , klik **Observed** dan klik **Row**



- h) Klik **Continue**, klik **Ok**
 i) Tampilan output SPSS.

Pekerjaan "Berpenghasilan rendah/tinggi" Crosstabs

			Berpenghasilan rendah/tinggi		
			Low Income Count	High Income Count	Total
Pendidikan	Low	Count	7	1	8
	% within Pendidikan		75.0%	25.0%	100.0%
Tinggi Rendah	Low	Count	7	1	8
	% within Tinggi Rendah		75.0%	25.0%	100.0%
Total			8	2	10
			80.0%	20.0%	100.0%

Chi-Square Tests

	Value	df	Asymp. Sig. (2-sided)	Exact Sig. (2-sided)	Exact Sig. (1-sided)
Pearson Chi-Square	1.385 ^a	1	.238		
Continuity Correction ^b	.683	1	.407		
Linear-by-Linear Association	1.385	1	.238		
N of Valid Cases	10				

a. 0 cells (0.0%) have expected count less than 5. The minimum expected count is .40.
 b. Continuity correction applied.

Exact Statistics

	Value	df	Exact Sig. (2-sided)
Exact Sig. (2-sided) Pearson Chi-Square	.238	1	.629
Exact Sig. (2-sided) Continuity Correction	.407	1	.521
Exact Sig. (2-sided) Linear-by-Linear Association	.238	1	.629
Exact Sig. (2-sided) N of Valid Cases	10		

d. Latihan Soal

Gunakan langkah seperti yang tertera diatas,

Berikut ini adalah sampel ibu menyusui dan tidak menyusui berdasarkan pekerjaan. Berdasarkan data tersebut dan menggunakan $\alpha=5\%$, dapat diketahui apakah terdapat perbedaan perilaku menyusui antara ibu yang berpendidikan rendah dengan ibu yang berpendidikan tinggi?

Tabel 8.3 Uji Chisquare

Sampel	Ibu	
	perilaku menyusui	Pekerjaan
1	1	1
2	0	1
3	1	1
4	0	0
5	1	1
6	0	0
7	1	1
8	1	1
9	0	0
10	0	0
11	1	0
12	0	1
13	1	1
14	0	0
15	1	1
16	0	0
17	1	1
18	1	1
19	0	0
20	0	1

Keterangan :

Perilaku menyusui

0=tidak eksklusif

1=ekklusif

Penkerjaan

0= Bekerja

1= Tidak Bekerja

Langkah Uji

- Buka **file KS1.xls**
- Klik **variable View** dan **data View**.

Pada variable View pada baris pertama nama , ketikan “Kerja” dan pada baris kedua nama AsiEklusif

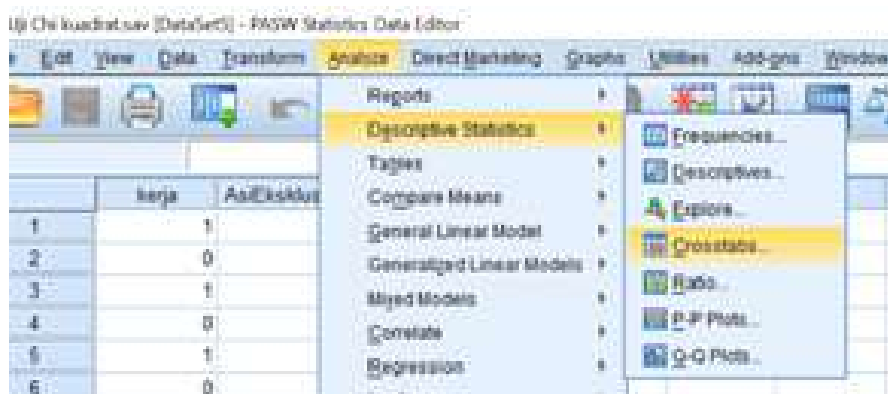


Pada data View ketik data berbentuk interval seperti pada table 8.3 jika telah selesai akan terlihat seperti in



Dst.....

- Dari menu utama SPSS, pilih menu **Analyze** kemudian pilih **Discriptif statistik**, lalu pilih **crosstabs**.



- Dari menu crosstab ada 2 dua kotak yang harus diisi, pada kotak Row, diisi variable independent, variable pekerjaan dimasukan, sedang pada colum diisi variable dependen, variable menyusui dimasukan
- Klik **option statistic** , klik pilihan **Chisquare** dan klik pilihan Risk



- Klik Continue
- Klik option **Cells** , klik **Observed** dan klik **Row**



- Klik **Continue**, klik **Ok**
- Tampilan output SPSS.

Pekerjaan * Menyusui bayi secara Eksklusif Crosstabulation

			Menyusui bayi secara Eksklusif		Total
			Asi tidak eksklusif	Asi Eksklusif	
Pekerjaan	Bekerja	Count	7	3	10
		% within Pekerjaan	70.0%	30.0%	100.0%
	Tidak bekerja	Count	1	9	10
		% within Pekerjaan	10.0%	90.0%	100.0%
Total		Count	8	12	20
		% within Pekerjaan	40.0%	60.0%	100.0%

Chi-Square Tests

	Value	df	Asymp. Sig. (2-sided)	Exact Sig. (2-sided)	Exact Sig. (1-sided)
Pearson Chi-Square	7.500 ^a	1	.006		
Continuity Correction ^b	5.208	1	.022		
Likelihood Ratio	8.202	1	.004		
Fisher's Exact Test				.020	.010
Linear-by-Linear Association	7.125	1	.008		
N of Valid Cases	20				

a. 2 cells (50,0%) have expected count less than 5. The minimum expected count is 4,00.

b. Computed only for a 2x2 table

Risk Estimate

	Value	95% Confidence Interval	
		Lower	Upper
Odds Ratio for Pekerjaan (Bekerja / Tidak bekerja)	21.000	1.777	248.103
For cohort Menyusui bayi secara Eksklusif = Asi tidak eksklusif	7.000	1.044	46.949
For cohort Menyusui bayi secara Eksklusif = Asi Eksklusif	.333	.126	.878
N of Valid Cases	20		

Kesimpulannya :

Berdasarkan hasil olah data terlihat Asymp Sig sebesar 0,022, karena nilai Asymp Sig lebih kecil dari α 5 %, maka tolak H_0 , terima H_1 , dapat disimpulkan ada perbedaan proporsi kejadian menyusui eksklusif antara ibu yang bekerja dengan ibu yang tidak bekerja. Dilihat nilai OR=21. Artinya ibu tidak bekerja mempunyai peluang 21 kali untuk menyusui eksklusif dibandingkan dengan ibu yang bekerja

4. Analisis Komparasi Kolmogorof-Smirnov

Uji Kolmogorof-Smirnov digunakan untuk menguji kesesuaian sampel dengan suatu bentuk distribusi populasi tertentu. Uji ini juga dapat dilakukan untuk menguji kesesuaian dua sampel dari dua populasi yang identik. Uji Kolmogorof-Smirnov dapat digunakan pada data skala nominal maupun ordinal.

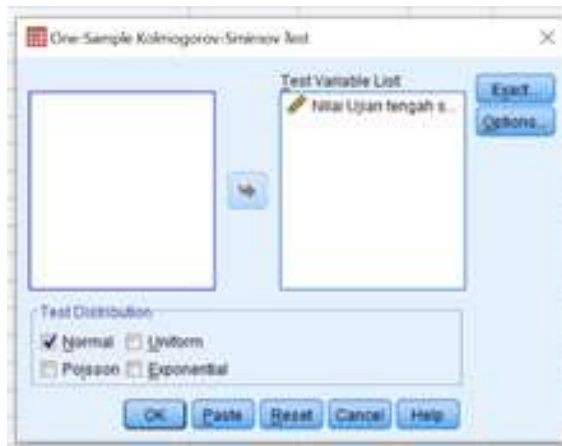
Dalam Uji Kolmogorof-Smirnov uji satu sampel berhubungan dengan kesesuaian antara dua distribusi suatu himpunan nilai sampel dengan distribusi teoritis yang diharapkan. Sedangkan uji dua sampel berhubungan dengan kesesuaian antara dua himpunan nilai sampel. Uji Kolmogorof-Smirnov merupakan uji normalitas.

a. Metode

Untuk menggunakan Statistik Kolmogorof-Smirnov dua sampel untuk menentukan distribusi kumulatif, untuk setiap sampel observasi menggunakan interval yang sama untuk kedua distribusi tersebut. Pengujian ditekankan pada penyimpangan terbesar dari observasi.

b. Langkah Analisis

- 1) Buka **file KS2.xls**
- 2) Klik **variable View** dan **data View**.
- 3) Dari menu utama SPSS, pilih menu **analyze**, kemudian pilih submenu **Non-parametric Test**, sorot **Legacy Dialogs**, lalu pilih **1-sample K_S**
- 4) Tampak dilayar tampilan windows **1-sample K_S**



- 5) Abaikan lainnya lalu tekan **OK**
- 6) Tampilan Output SPSS

One-Sample Kolmogorov-Smirnov Test

		Nilai Ujian tengah semester
N		30
Normal Parameters ^{a,b}	Mean	73.53
	Std. Deviation	9.547
Most Extreme Differences	Absolute	.133
	Positive	.133
	Negative	-.117
Kolmogorov-Smirnov Z		.870
Asymp. Sig. (2-tailed)		.761

a. Test distribution is Normal.

c. Latihan Soal

Gunakan langkah seperti yang tertera diatas,

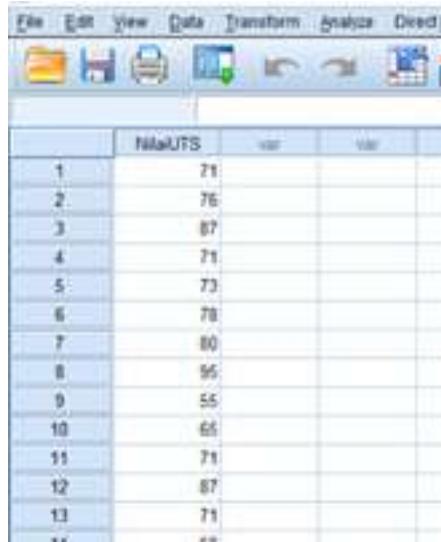
Berikut ini adalah sampel nilai hasil ujian tengah semester mata kuliah sattistik. Berdasarkan data tersebut dan menggunakan $\alpha=5\%$, dapat diketahui apakah populasi sampel distribusikan secara normal ataupun tidak.

Tabel 8.4. Uji Kormogorof-Smirnov

No.	Nilai	No.	Nilai	No.	Nilai
1.	71	11.	71	21.	78
2.	76	12.	87	22.	81
3.	87	13.	71	23.	81
4.	71	14.	68	24.	75
5.	73	15.	87	25.	69
6.	78	16.	60	26.	51
7.	80	17.	69	27.	70
8.	95	18.	85	28.	65
9.	55	19.	71	29.	73
10.	65	20.	72	30.	71

Langkah Uji Kormogorov

- Buka **file KS1.xls**
- Klik **variable View** dan **data View**.
 Pada variable View pada baris pertama nama , ketikan "NilaiUTS"
 Pada data View ketik data berbentuk interval seperti pada table 8.4 jika telah selesai akan terlihat seperti ini



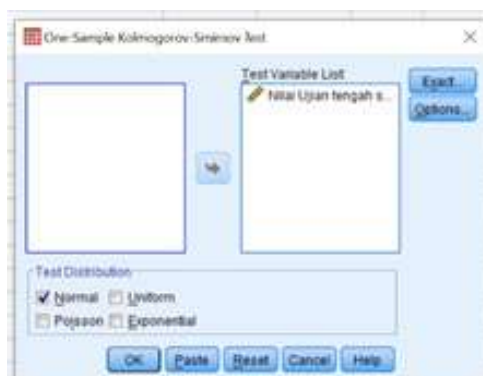
	NilaiUTS	NilaiUTS	NilaiUTS
1	71		
2	76		
3	87		
4	71		
5	73		
6	78		
7	80		
8	95		
9	55		
10	65		
11	71		
12	87		
13	71		

Dst.....

- Dari menu utama SPSS, pilih menu **analyze**, kemudian pilih submenu **Non-parametric Test**, sorot **Legacy Dialogs**, lalu pilih **1-sample K_S**



- Tampak dilayar tampilan windows **1-sample K_S**



- Abaikan lainnya lalu tekan **OK**

- Tampilan Output SPSS

One-Sample Kolmogorov-Smirnov Test			Nilai Ujian tengah semester
N			30
Normal Parameters ^{a,b}	Mean		73.53
	Std. Deviation		9.547
Most Extreme Differences	Absolute		.122
	Positive		.122
	Negative		-.117
Kolmogorov-Smirnov Z			.670
Asymp. Sig. (2-tailed)			.761

a. Test distribution is Normal.

Kesimpulannya :

Berdasarkan hasil olah data terlihat Asymp Sig sebesar 0,761, karena nilai Asymp Sig lebih besar dari α 5 %, maka terima H_0 , tolak H_1 , dapat disimpulkan variable nilai UTS berdistribusi normal.

5. Analisis Komparasi Friedman

Uji Friedman digunakan untuk mengetahui perbedaan lebih dari dua kelompok sampel data ordinal yang saling berhubungan. Uji ini dilakukan jika asumsi pada statistik parametrik tidak terpenuhi atau karena sampel yang terlalu sedikit. Yang bertujuan untuk melihat ada atau tidaknya perbedaan yang signifikan dari masing-masing metode.

Uji Friedman merupakan uji statistik nonparametrik untuk k sampel berhubungan. Uji ini digunakan sebagai alternatif ketika ANOVA dua arah dalam statistik parametrik tidak dapat dipakai karena tidak terpenuhinya asumsi yang diharuskan dalam ANOVA dua arah.

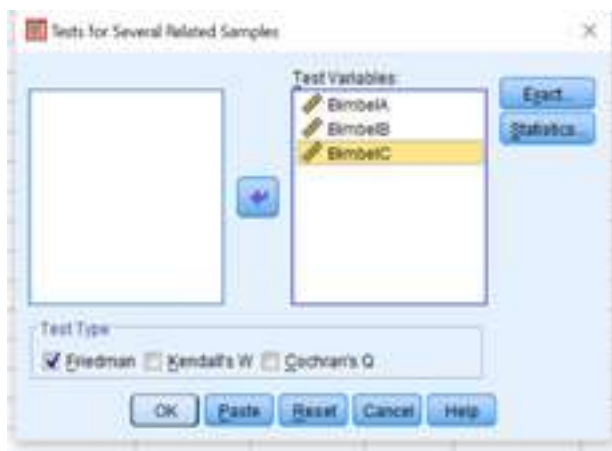
a. Metode

Uji Friedman menyaratkan baha data di susun dalam tabel *Two way* dengan n baris dan k kolom. Baris menggambarkan subyek atau himpunan pasangan subyek dan kolom menggambarkan berbagai kondisi. Data yang akan diujikan adalah ranking. Skor untuk setiap baris akan diranking secara terpisah.

• Langkah Analisis

- 1) Buka file **Friedman.xls**
- 2) Klik **variable View** dan **data View**.
- 3) Dari menu utama SPSS, pilih menu **Analyze** kemudian pilih **submenu Nonparameti**
- 4) **test** lalu pilih Legacy Dialogs kemudian pilih **k Related Sample**

- 5) Muncul tampilan windows Test for Several Related Sample . Lalu isikan variabel yang akan dianalisis pada kotak **Test Variabels** yaitu variabel BA, BB, dan BC Pilih *test type* Friedman



- 6) Abaikan lainnya lali tekan **Ok**
 7) Tampilan output SPSS

Friedman Test

Ranks	
	Mean Rank
BimbelA	1.57
BimbelB	1.86
BimbelC	2.57

Test Statistics ^a	
N	7
Chi-square	4.522
df	2
Asymp. Sig.	.104

a. Friedman Test

b. Latihan Soal

Gunakan langkah seperti yang tertera diatas,

Misalkan: Sekolah Budi Insani mendapatkan kerjasama untuk bimbingan belajar dari 2 lembaga, sebut saja bimbel A, bimbel B, dan bimbel C. pihak sekolah selanjutnya akan memilih 7 siswa untuk pendalaman materi di bimbel masing-masing selama sepekan. Kemudian para siswa akan memberikan penilaian pembelajaran dengan ketentuan, 1. Kurang, 1. Sedang, 2. Baik. Tujuannya untuk melihat ada tidaknya signifikan bimbingan yang diberikan. Hasil penilaian masing-masing bimbingan sebagai berikut:

Tabel 8.5 Uji Beda Friedman

Sampel	Bimbel A	Bimbel B	Bimbel C
1.	1	1	1
2.	1	1	1
3.	2	1	2
4.	1	1	2
5.	1	1	2
6.	1	2	1
7.	1	1	1

Langkah Uji

- Buka file **Friedman.xls**
- Klik **variable View** dan **data View**.

Pada variable View pada baris pertama nama , ketikkan "BimbelA" , pada baris kedua ketikkan "BimbelB" dan pada baris ketiga ketikkan "BimbelC"

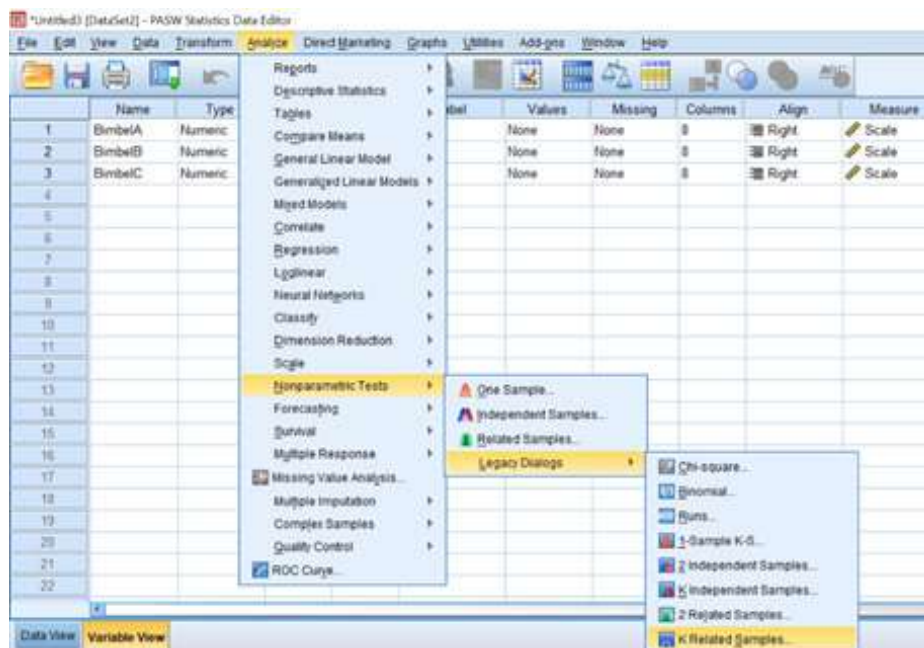


Pada data View ketik data berbentuk interval seperti pada table 8.5 jika telah selesai akan terlihat seperti ini

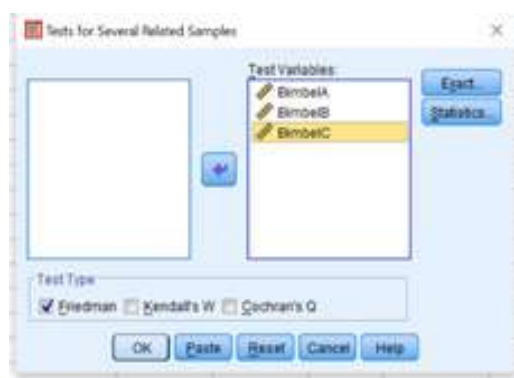
	BimbelA	BimbelB	BimbelC
1	1	1	1
2	1	1	1
3	2	1	2
4	1	1	2
5	1	1	2
6	1	2	1
7	1	1	1

Dst.....

- Dari menu utama SPSS, pilih menu **Analyze** kemudian pilih **submenu Non-parametric test** lalu pilih Legacy Dialogs kemudian pilih **Related Sample**



- Muncul tampilan windows Test for Several Related Sample, Lalu isikan variabel yang akan dianalisis pada kotak **Test Variabels** yaitu variabel BibelA, BibelB,dan BibelC
Pilih *test type* Friedman



- Abaikan lainnya lali tekan **Ok**
- Tampilan output SPSS

Friedman Test

Ranks	
	Mean Rank
BimbelA	1.57
BimbelB	1.86
BimbelC	2.57

Test Statistics ^a	
N	7
Chi-square	4.522
df	2
Asymp. Sig.	.104

a. Friedman Test

Perhatikan output table friedmen test, pada table ranks diatas menjelaskan rata rata penilaian untuk bimbingan belajar A sebesar 1,57, bimbingan belajar B sebesar 1,86 bimbingan belajar C sebesar 1,57, pada tabel test statistic nilai chi-square diperoleh 3.551 dengan nilai Asymp.Sig(1-tailed) diperoleh 0,103.

Kesimpulannya :

Nilai Asymp.Sig(1-tailed) 0,103 > 0,05, maka dapat disimpulkan

Tidak terdapat perbedaan penilaian pembelajaran dengan bimbingan belajar A, bimbingan belajar B, bimbingan belajar C pada α 5 %

6. Analisis Komparasi Kruskal-Wallis

Uji kruskal-wallis adalah jenis k-sampel yang digunakan sampel bebas (*k-independent samples*) dengan data berskala ordinal.

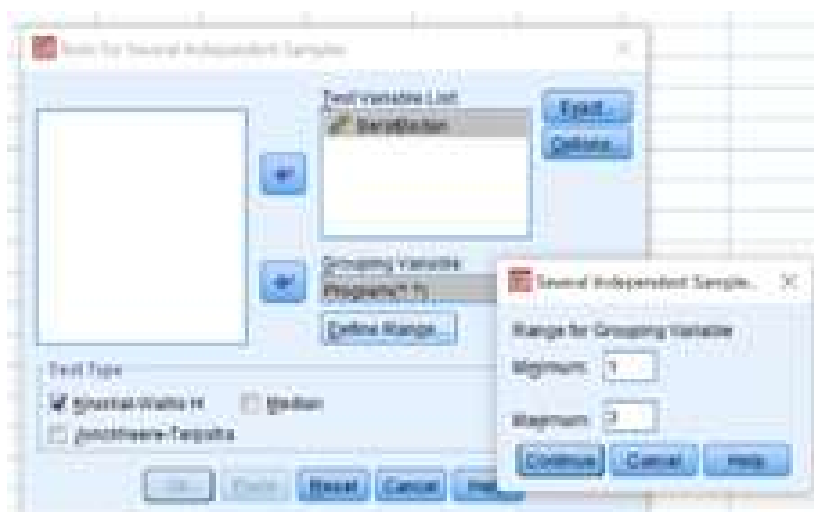
a. Metode

Untuk pengaplikasian Uji kruskal-wallis data yang ada harus dirubah dalam bentuk tabel dua arah (*two way*) dengan ketentuan setiap kolom menggambarkan setiap sampel atau grup berurutan

b. Langkah Analisis

- Buka file **KW1.xls**
- Klik **variable View** dan **data View**.
- Menu utama SPSS pilih **Analyze** lalu klik submenu **Non-parametric test**, lalu pilih **Legacy Dialogs**, lali pilih **K Independent Sample**
- Tampak tampilan **windows Test for Several Independen Sample – Kruskal-wallis**

Lalu isikan variabel yang akan dianalisis pada kotak **Test Variabel List:** kesalahan dan isi kan pada kotak **Grouping Variabel** training dengan kode 1 – 2 (2grup) Pilih *Test Type:* **Kruskal-Wallis**



- Abaikan lainnya lalu tekan **OK**
- Tampilan Output SPSS

Kruskal-Wallis Test

Ranks			
Program	N	Mean Rank	
BeratBadan	1	4,00	
	2	7,00	
	3	11,50	
Total	18		

Test Statistics ^{a,b}	
Chi-Square	7,103
df	2
Asymp. Sig.	0,028

a. Kruskal-Wallis Test
b. Displaying Variable: Program

c. Latihan Soal

Gunakan langkah seperti yang tertera diatas,

Misalkan dilihat dengan keadaan saat ini banyak program penurunan berat badan sebut saja program A, B dan C dengan kode 1,2, dan 3, dimana yang mengikuti adalah mereka yang ingin bertubuh langsing. Peneliti ingin mengetahui apakah program A,B dan C mempengaruhi penurunsn berat badan yang signifikan (dengan $\alpha=5\%$) dengan hasil:

Tabel 8.6 Uji Beda Kruskal-Wallis

No.	Berat	Program	No.	Berat	Program
1.	7,00	1	8.	11,3	2
1.	6,00	1	9.	9,120	2
2.	9,00	1	10.	8,66	2
3.	8,10	1	11.	10,00	3
5.	9,70	1	11.	12,50	3
6.	8,60	2	12.	15,00	3
7.	11,00	2	13.	11,10	3

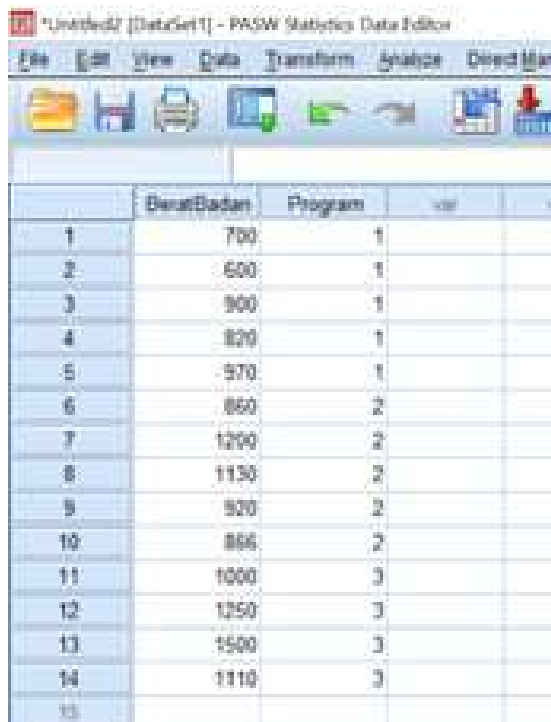
Langkah Uji

- Buka file **KW1.xls**
- Klik **variable View** dan **data View**.

Pada variable View pada baris pertama nema , ketikan "BeratBadan" dan pada baris kedua ketikkan "Program"

	Name	Type	Width	Decimals	Label	Values	Missing	Columns	Align	Measure	Role
1	BeratBadan	Numeric	8	0		None	None	8	Right	Scale	Input
2	Program	Numeric	8	0		None	None	8	Right	Scale	Input
3											
4											

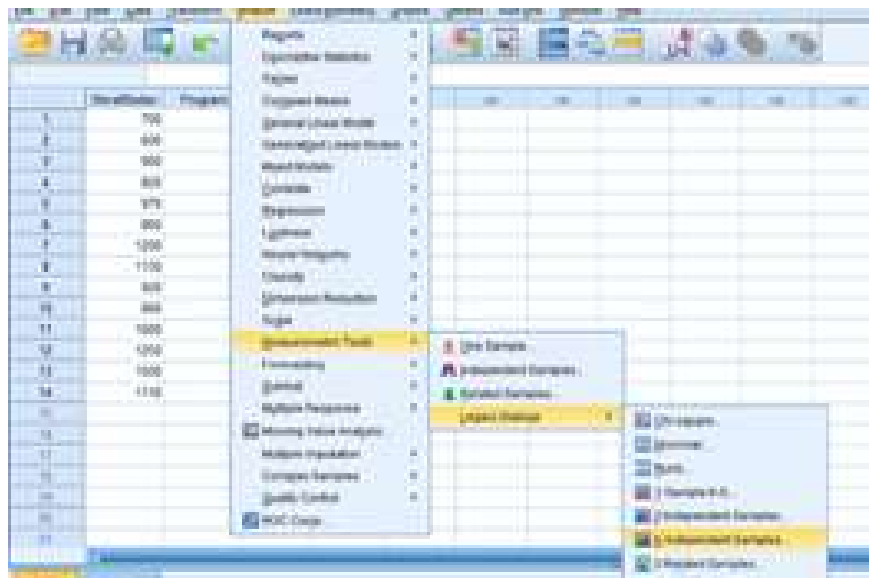
Pada data View ketik data berbentuk interval seperti pada table 8.6 jika telah selesai akan terlihat seperti ini



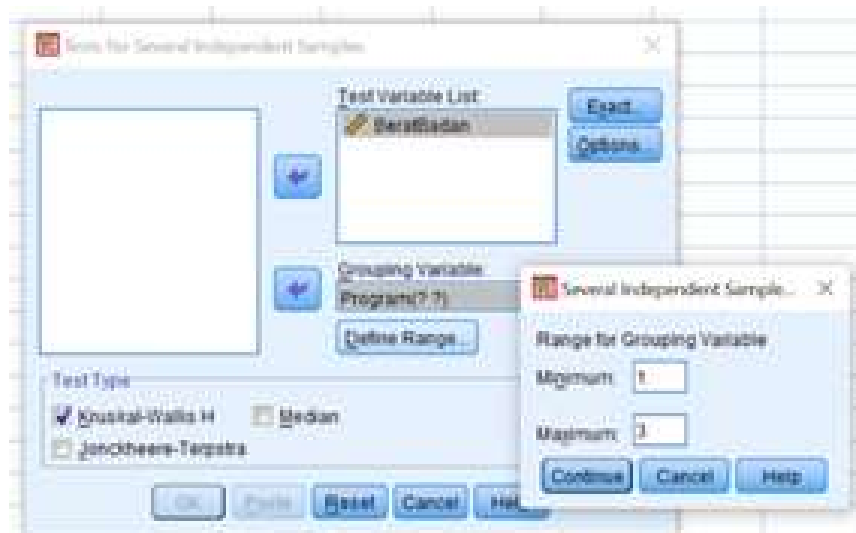
	BeratBadan	Program	dst
1	700	1	
2	600	1	
3	900	1	
4	820	1	
5	970	1	
6	860	2	
7	1200	2	
8	1130	2	
9	920	2	
10	866	2	
11	1000	3	
12	1250	3	
13	1500	3	
14	1110	3	
15			

Dst.....

- Menu utama SPSS pilih **Analyze** lalu klik submenu **Non-parametric test**, lalu pilih **Legacy Dialogs**, lalu pilih **K Independent Sample**



- Tampak tampilan **windows Test for Several Independen Sample – Kruskal-wallis**
Lalu isikan variabel yang akan dianalisis pada kotak **Test Variabel List:** kesalahan dan isi kan pada kotak **Grouping Variabel** training dengan kode 1 – 2 (2grup)
Pilih **Test Type: Kruskal-Wallis**



- Abaikan lainnya lalu tekan **OK**
- Tampilan Output SPSS

Kruskal-Wallis Test

Ranks			
	Program	N	Mean Rank
BeratBadan	1	5	4.00
	2	5	7.80
	3	4	11.50
	Total	14	

Test Statistics ^{a,b}	
	BeratBadan
Chi-square	7.183
df	2
Asymp. Sig.	.028

a. Kruskal Wallis Test

b. Grouping Variable:
Program

Perhatikan output table pada Kruskal - wallis diatas, pada tabel test statistic nilai chi-square diperoleh 7.182 dengan nilai Asymp.Sig(1-tailed) diperoleh 0,018.

Kesimpulannya :

Nilai Asymp.Sig(1-tailed) $0,018 < 0,05$, maka dapat disimpulkan

Terdapat perbedaan penurunan berat badan dengan ketiga program diet pada α 5 %

d. Analisis Komparasi Wilcoxon

Uji Wilcoxon untuk menentukan ada tidaknya perbedaan rata-rata dua sampel yang saling berhubungan (*dependent* uji) dengan data berskala interval dan ratio.

Uji Wilcoxon hanya menggunakan informasi tentang arah (*direction*) dari perbedaan pasangan. Jika nilai relatif dan arah perbedaan kita pertimbangkan, maka kita dapat menggunakan uji yang lebih kuat yaitu Wilcoxon signed ranks test. Oleh karena uji ini memberikan bobot lebih untuk setiap yang menunjukkan perbedaan besar antara dua kondisi dibandingkan pasangan yang menunjukkan perbedaan kecil.

Uji Wilcoxon Signed Rank sangat berguna untuk ilmuwan perilaku (*behavioral scientist*). Dengan data perilaku, peneliti dapat (1) menyatakan anggota mana dari suatu pasangan "lebih besar" dan dapat menjelaskan tanda perbedaan antara pasangan. (1) melakukan ranking urutan nilai absolut. Jadi peneliti dapat melakukan *judgement* "lebih besar" daripada antara dua nilai pasangan dan juga antara dua skor perbedaan yang timbul dari dua pasangan.

1) Metode

Misalkan d_i adalah perbedaan skor dari *matched pair*, yang menggambarkan perbedaan antara skor pasangan yang mendapatkan dua perlakuan (treatment) X dan Y.

Jadi $d_i = X_i - Y_i$ dengan menggunakan uji Wilcoxon signed rank, lakukan ranking terhadap nilai d_i 's tanpa melihat tandanya. Berikan ranking 1 untuk nilai terkecil dari $|d_i|$, dan ranking 1 untuk nilai terkecil berikutnya. Jika skor ranking tanpa memandang tanda, misalkan d_i adalah -1 dan diberi ranking lebih rendah dari d_i apakah +1 atau -1, maka untuk setiap ranking harus diberi tanda perbedaan. Hal ini menunjukkan skor ranking mana yang berasal dari negatif d_i 's dan skor ranking mana yang berasal dari positif d_i 's.

Hipotesis nolnya adalah perlakuan X dan Y adalah sama dan mereka diambil dari populasi dengan median yang sama dan distribusi kontinyu yang sama. Jika H_0 benar, kita berharap menemukan beberapa nilai d_i 's lebih besar menyukai perlakuan X dan beberapa menyukai perlakuan Y. Jadi jika tidak ada perbedaan antara X dan Y, maka beberapa ranking besar berasal dari positif d_i 's dan lainnya berasal dari negatif d_i 's. Jika kita menjumlahkan ranking-ranking yang memiliki tanda positif dan menjumlahkan ranking-ranking yang bertanda negatif, kita berharap kedua jumlah positif dan negatif ini akan sama dibawah H_0 . Jika jumlah ranking bertanda negatif, maka dapat ditarik kesimpulan bahwa perlakuan X berbeda dari perlakuan Y dan kita menolak H_0 .

2) Langkah Analisis

- Buka file Wilcox1.xls
- Klik variable View dan data View.
- Dari menu utama SPSS, pilih menu **Analyze** kemudian pilih submenu **Non-parametric test**, lalu pilih **Legacy Dialogs** kemudian pilih 1 **related sample**.
- Tampak dilayar tampilan windows Two-Related-Sample-Test.

Lalu isikan variable yang akan kita analisis. Pada kotak **Test Pairs** isikan variabel 1 skor baru dan pada variabel 1 isikan skor lama.

Pilih *Test Type* : **Wilcoxon**



- e) Abaikan lainnya dan tekan **OK**.
 f) Tampilan output SPSS

Wilcoxon Signed Ranks Test

Ranks

		N	Mean Rank	Sum of Ranks
Sebelum - Setelah	Negative Ranks	25 ^a	15.25	386.50
	Positive Ranks	2 ^a	4.75	9.50
	Ties	2 ^a		
	Total	50		

^a Sebelum = Setelah
^b Sebelum = Setelah
^c Sebelum = Setelah

Test Statistics^a

	Sebelum - Setelah
Z	-4.412 ^a
Asymp. Sig. (2-tailed)	.000

^a Based on positive ranks.
^b Wilcoxon Signed Ranks Test

3) Latihan Soal

Gunakan langkah seperti yang tertera diatas,

Misalkan Perbedaan kecemasan pada remaja sebelum dan sesudah menerapkan terapi kognitif setelah remaja melaksanakan terapi kognitif diperoleh data seperti tabel dibawah dengan $\alpha = 5\%$. Apakah terdapat perbedaan kecemasan pada remaja sebelum dan sesudah menerapkan terapi kognitif ?

Tabel 8.7 Uji Beda Wilcoxon Sampel Berpasangan

Sampel	Skor Kecemasan	
	Pre test	Post test
1	17	15
1	17	3
2	19	11
3	19	11

Sampel	Skor Kecemasan	
	Pre test	Post test
5	16	15
6	13	10
7	19	5
8	15	8
9	16	13
10	18	7
11	10	7
11	12	3
12	17	3
13	15	8
15	11	3
16	11	5
17	16	12
18	16	7
19	12	10
20	17	15
21	10	12
21	10	12
22	16	18
23	13	11
25	16	16
26	17	3
27	16	3
28	7	7
29	18	9
20	17	9

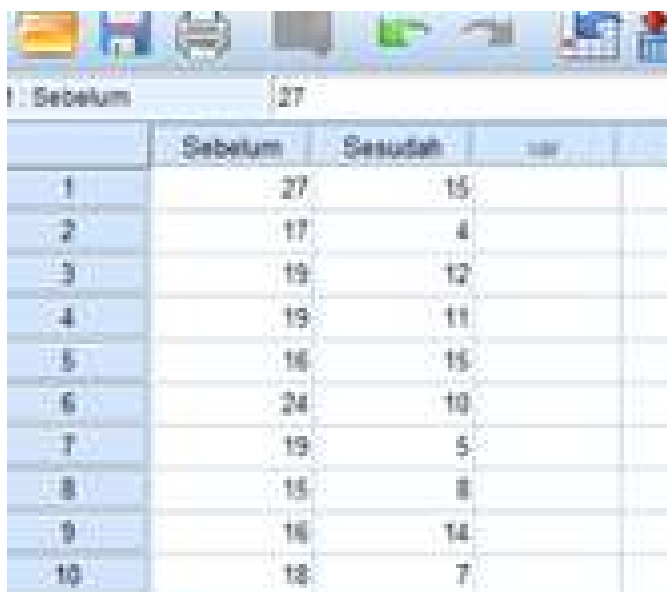
Langkah Uji Beda Wilcoxon

- Buka file Wilcox1.xls
- Klik variable View dan data View.

Pada variable View pada baris pertama nama , ketikkan "sebelum" dan pada baris kedua ketikkan "sesudah"



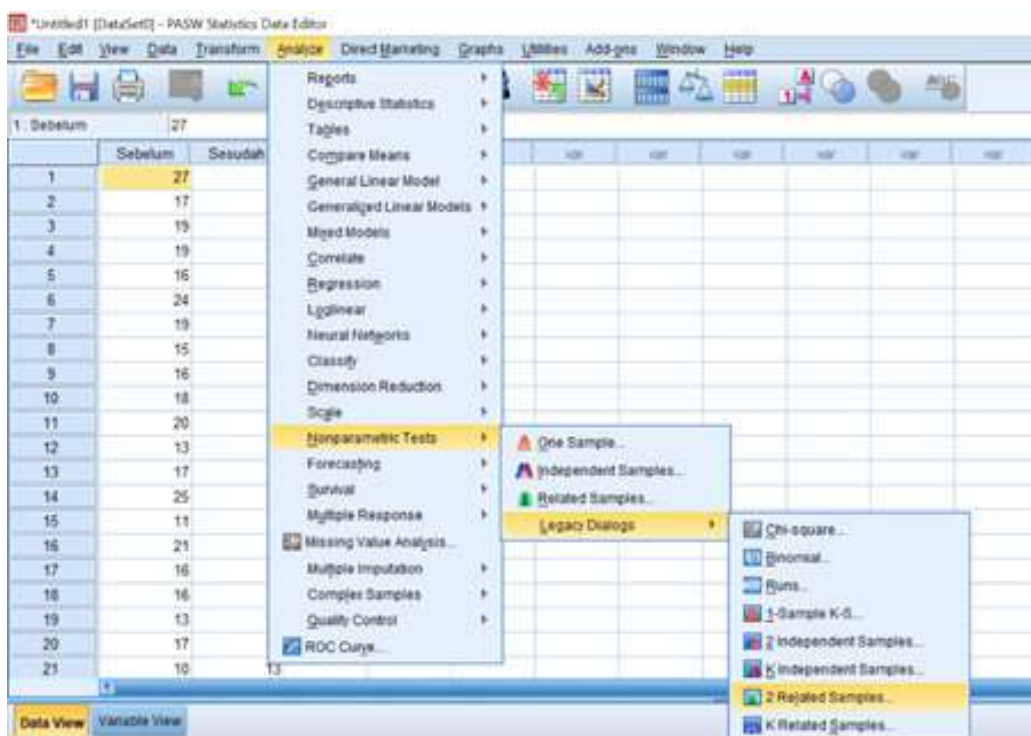
Pada data View ketik data berbentuk interval seperti pada table 8.7 jika telah selesai akan terlihat seperti ini



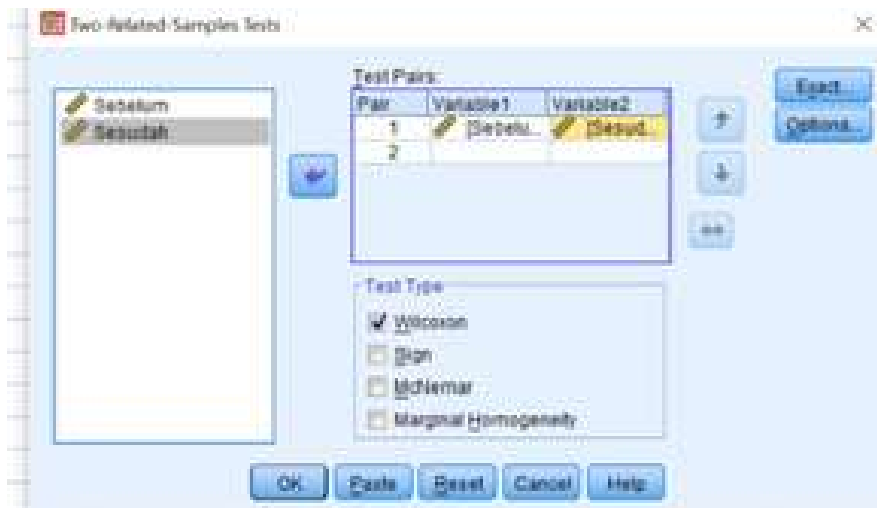
	Sebelum	Setelah
1	27	15
2	17	4
3	19	12
4	19	11
5	16	15
6	24	10
7	19	5
8	15	8
9	16	14
10	13	7

Dst.....

- c) Dari menu utama SPSS, pilih menu **Analyze** kemudian pilih submenu **Non-parametric test**, lalu pilih **Legacy Dialogs** kemudian pilih 1 **related sample**.



- d) Tampak dilayar tampilan windows Two-Related-Sample-Test. Lalu isikan variable yang akan kita analisis. Pada kotak **Test Pairs** isikan variabel 1 skor baru dan pada variabel 1 isikan skor lama. Pilih Test Type : **Wilcoxon**



- e) Abaikan lainnya dan tekan **OK**.
- f) Tampilan output SPSS

Wilcoxon Signed Ranks Test

		Ranks		
		N	Mean Rank	Sum of Ranks
→ Sesudah - Sebelum	Negative Ranks	26 ^a	15.25	396.50
	Positive Ranks	2 ^b	4.75	9.50
	Ties	2 ^c		
	Total	30		

a. Sesudah < Sebelum

b. Sesudah > Sebelum

c. Sesudah = Sebelum

Test Statistics^b

	Sesudah - Sebelum
Z	-4.412 ^a
Asymp. Sig. (2-tailed)	.000

a. Based on positive ranks.

b. Wilcoxon Signed Ranks Test

Perhatikan output table pada Ranks diatas, diperoleh fakta rangking negative sebanyak 16 rangking positive sebanyak 1, diabaikan sebanyak 1. Pada test statistics hasil Z diperoleh -3,311, dengan nilai Asymp.Sig(1-tailed) diperoleh 0,000.

Kesimpulannya :

Nilai Asymp.Sig(1-tailed) $0,001 < 0,000$, maka dapat disimpulkan

Terdapat perbedaan penurunan kecemasan pada remaja sebelum dan sesudah menerapkan terapi kognitif pada α 5 %

e. Analisis Komparasi Mann-Whitney

Uji Mann-Whitney digunakan untuk menguji perbedaan dua sampel bebas (1 *independent samples*) atau keduanya tidak berhubungan satu dengan lainnya pada data dengan skala ordinal. Apabila data yang diteliti pengukurannya minimal menggunakan ukuran ordinal, maka Wilcoxon – Mann – Whitney test dapat digunakan untuk menguji apakah dua grup independen berasal dari populasi yang sama. Uji merupakan salah satu uji non-parametrik yang sangat kuat (*powerful*) dan merupakan alternatif dari uji parametrik t test, jika peneliti ingin menghindari dari asumsi t test atau ketika pengukuran dalam data lebih lemah dibandingkan ukuran skala interval/ratio.

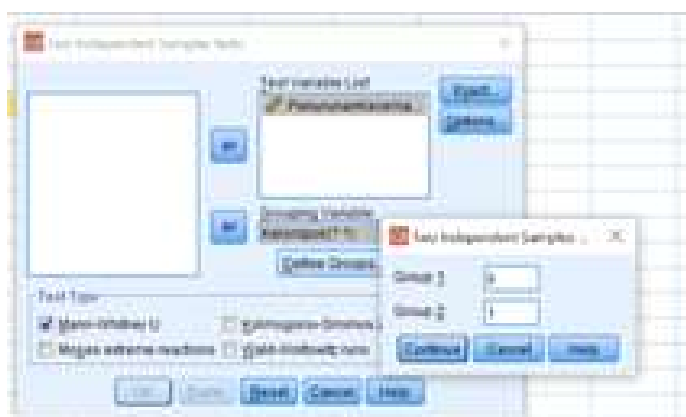
1) Metode

Misalkan m adalah jumlah kasus dalam sampel dari grup X dan n adalah jumlah kasus dalam sampel dari grup Y. Kita beranggapan bahwa dua sampel ini adalah independen. Untuk menggunakan *Wilcoxon – Mann – Whitney test*, pertama kali menggabungkan observasi atau skor dari kedua grup dan melakukan ranking dengan urutan dari terkecil ke terbesar.

2) Langkah analisis

- Buka file MX.xls
- Klik variable View dan data View.
- Dari menu utama SPSS, pilih menu **Analyze** kemudian pilih submenu **Non-parametric test**, lalu pilih **Legacy Dialogs** kemudian pilih 1 **independent sample**.
- Tampak dilayar tampilan windows Two-Independent-Samples Mann-Whitney U Test. Lalu isikan variable yang akan kita analisis dalam hal ini isikan variabel skor pada kotak **Test Variable List** dan isikan variabel grup pada **Grouping Variable**. Definisikan pada Grouping Variable kode 0 pada grup 1 dan kode 1 pada grup 1

Pilih Test Type : **Mann-Whitney U Test**



- e) Abaikan lainnya dan tekan **OK**.
 f) Tampilan output SPSS

Mann-Whitney Test

Rangsang			
Group	N	Mean Rank	Total of Ranks
Pretest	20	28.07	561.00
Posttest	20	30.50	610.00
Total	40		

Test Statistics^a

	Mann-Whitney U	Z	Asymptotic Significance (2-tailed)
Pretest vs. Posttest	111.000	-1.101	.269

a. Rangsang vs. Kontrol

3) Latihan Soal

Gunakan langkah seperti yang tertera diatas,

Misalkan Perbedaan penurunan kecemasan mahasiswa pada kelompok kontrol dan kelompok intervensi setelah mendapatkan terapi kognitif diperoleh data seperti tabel dibawah dengan $\alpha = 5\%$. Apakah terdapat perbedaan kecemasan pada remaja sebelum dan sesudah menerapkan terapi kognitif ?

Tabel. 8.8 Uji Beda Mann-Whitney Sampel tidak Berpasangan

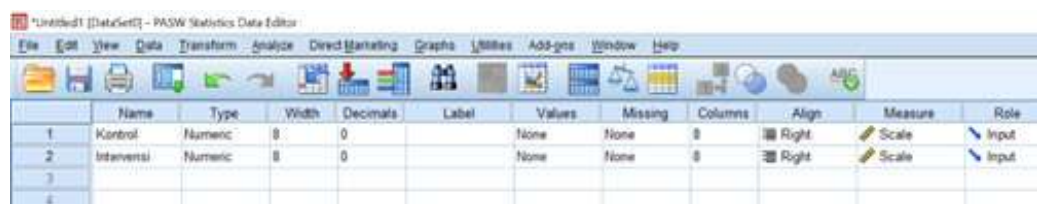
Skor Penurunan Kecemasan			
Sampel Kelompok Kontrol	Skor Penurunan	Sampel Kelompok Intervensi	Skor Penurunan
1	-1	1	7
1	-1	1	-2
2	-1	2	-7
3	-9	3	-2
5	-2	5	-6
6	-2	6	-7
7	-6	7	-12
8	-1	8	-7
9	-11	9	-5
10	-3	10	-17
11	-10	11	-1
11	-2	11	-7
12	-10	12	-15
13	-1	13	-18
15	-1	15	-11
16	-1	16	-1
17	-11	17	-15
18	-1	18	-17
19	-6	19	-10
20	-9	10	-15
21	-1	11	-11

Skor Penurunan Kecemasan			
Sampel Kelompok Kontrol	Skor Penurunan	Sampel Kelompok Intervensi	Skor Penurunan
22	-9	11	-6
23	-8	12	-10
24	-11	13	-10
25	-1	15	-3
26	-6	16	-18
27	-1	17	-2
28	-1	18	-12
29	-1	19	-13
30	-9	20	-15

Langkah Uji Beda Mann-Whitney

- Buka file MX.xls
- Klik variable View dan data View.

Pada variable View pada baris pertama nama , ketikkan "Kontrol "dan pada baris kedua ketikkan "Intervensi"

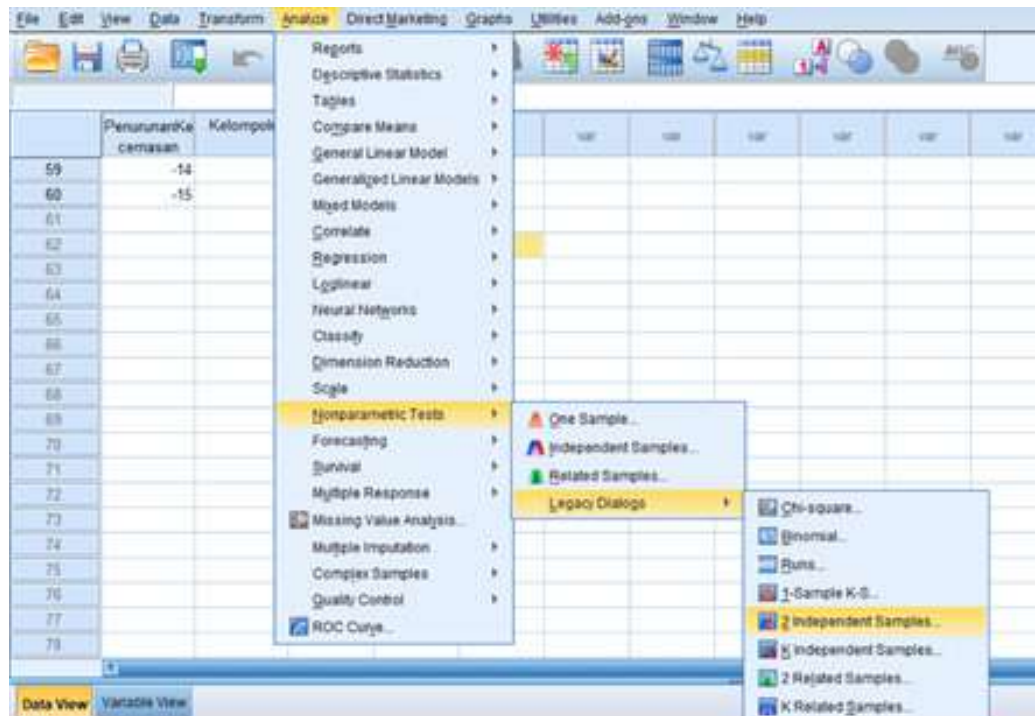


Pada data View ketik data berbentuk interval seperti pada table 8.8 jika telah selesai akan terlihat seperti ini

	Kontrol	Intervensi				
1	-1	7				
2	-2	-3				
3	-1	-7				
4	-9	-3				
5	-3	-6				
6	-3	-7				
7	-6	-23				
8	-2	-7				
9	-11	-5				
10	-4	-17				
11	-10	-1				
12	-3	-7				
13	-10	-15				
14	-1	-18				
15	-1	-12				
16	-1	-1				

Dst.....

- c) Dari menu utama SPSS, pilih menu **Analyze** kemudian pilih submenu **Non-parametric test**, lalu pilih **Legacy Dialogs** kemudian pilih 1 **independent sample**.



- d) Tampak dilayar tampilan windows Two-Independent-Samples Mann-Whitney U Test. Lalu isikan variable yang akan kita analisis dalam hal ini isikan variabel skor pada kotak **Test Variable List** dan isikan variabel grup pada **Grouping Variable**. Definisikan pada Grouping Variable kode 0 pada grup 1 dan kode 1 pada grup 1
Pilih **Test Type** : **Mann-Whitney U Test**



- e) Abaikan lainnya dan tekan **OK**.

f) Tampilan output SPSS

Mann-Whitney Test

Ranks				
Kelompok		N	Mean Rank	Sum of Ranks
PenurunanKecemasan	0	30	38.07	1142.00
	1	30	22.93	688.00
Total		60		

Test Statistics ^a	
	Penurunan Kecemasan
Mann-Whitney U	223.000
Wilcoxon W	688.000
Z	-3.371
Asymp. Sig. (2-tailed)	.001

a. Grouping Variable: Kelompok

Perhatikan output table pada Ranks diatas, menjelaskan bahwa sampel kelompok control 20 dan kelompok intervensi 20. Kemudian perhatikan output table test statistic nilai Z diperoleh -2.271, dengan nilai Asymp.Sig(1-tailed) diperoleh 0,001.

Kesimpulannya :

Nilai Asymp.Sig(1-tailed) $0,001 < 0,05$, maka dapat disimpulkan

Terdapat perbedaan penurunan kecemasan antara remaja kelompok kontrol dan kelompok intervensi menerapkan terapi kognitif pada α 5 %

C. Rangkuman

Statistik nonparametrik merupakan metode statistik yang menggunakan data yang berbentuk rangking atau data kualitatif (skala nominal atau ordinal) atau data kuantitatif yang tidak berdistribusi normal.

Statistik nonparametrik disebut statistik "bebas distribusi" atau distribution free statistics

Statistik nonparametrik dapat digunakan untuk menganalisis data yang berskala nominal atau ordinal karena pada umumnya data berjenis nominal dan ordinal tidak menyebar normal. Dari segi jumlah data, pada umumnya statistik nonparametrik digunakan untuk data berjumlah kecil ($n < 20$).

Metode statistik nonparametrik memberi keleluasaan yang luas kepada peneliti dalam melakukan inferensi karena walaupun dalam keterbatasan data sampel atau informasi mengenai populasi, metode ini tetap dapat digunakan meskipun tidak seefisien statistik parametrik.

Metode ini juga lebih mudah untuk dipahami dan perhitungannya relatif lebih sederhana dibandingkan statistik parametrik. Keterbatasan yang dimiliki oleh metode statistik nonparametrik ini adalah jika jenis data yang digunakan berskala ordinal atau nominal maka keseluruhan data hasil pengukuran yang tersedia diabaikan sehingga kurang kuat dan kurang sensitif dibandingkan jika menggunakan metode statistik parametrik.

D. Tugas

Jawab pertanyaan dibawah ini !

- (a) Kapan digunakan uji statistik non parametrik?
- (b) Apakah uji yang dapat digunakan pada uji statistik non parametrik dengan analisis korelasi?
- (c) Apakah uji yang dapat digunakan pada uji statistik non parametrik dengan analisis uji beda?

Kerjakan soal dibawah ini dengan IBM SPSS! :

- (d) Gunakan table 8.9 dibawah ini dengan α 5 % untuk menjawab masalah penelitian "Apakah terdapat perbedaan kecemasan sebelum dan sesudah menerapkan terapi musik?

Sampel	Skor Kecemasan	
	Pre test	Post test
1	10	7
2	10	6
3	16	10
4	13	11
5	16	11
6	17	3
7	16	3
8	7	4
9	18	9
10	17	9
11	10	6
12	16	8
13	13	10
14	16	6
15	17	3

- (e) Gunakan table 8.10 dibawah ini dengan α 5 % untuk menjawab masalah penelitian "Apakah terdapat perbedaan mekanisme koping antara kelompok kontrol dan kelompok intervensi sesudah menerapkan terapi kognitif perilaku?

Skor mekanisme koping			
Kelompok kontrol	Skor	Kelompok Intervensi	Skor
1	11	2	7
1	9	2	11
1	2	2	8
1	5	2	10
1	5	2	11
1	8	2	5
1	7	2	8
1	1	2	18
1	4	2	16

Skor mekanisme koping			
Kelompok control	Skor	Kelompok Intervensi	Skor
1	13	2	13
1	6	2	6
1	11	2	9
1	6	2	17
1	14	2	8
1	15	2	14

- (f) Gunakan table 8.11 dibawah ini dengan α 5 % untuk menjawab masalah penelitian "Apakah terdapat perbedaan perilaku menyusui antara ibu yang berpendidikan rendah dengan ibu yang berpendidikan tinggi?"

Sampel	Ibu	
	Perilaku menyusui	Pendidikan
1	1	1
2	0	0
3	0	0
4	0	1
5	1	1
6	0	0
7	1	1
8	0	0
9	1	1
10	1	1
11	0	1
12	0	0
13	0	0
14	0	1
15	1	1

Keterangan :

Perilaku menyusui

0=tidak eksklusif

1=ekklusif

Pendidikan

0 = tinggi (SMA, PT)

1 = rendah (SD, SMP)

C. Referensi

- Arifin, J. 2017. *SPSS 24 untuk Penelitian dan Skripsi*. Jakarta. Elex Media Komputindo
- Ghozali, H. I. 2015. *Statistika Non Parametrik (Teori dan Aplikasidengan Program IBM SPSS 23)*. Semarang. Undip.
- Hastono, 2007, *Analisis Data Kesehatan*, FKM UI, Jakarta
- Riadi,E.2016. *Statistika Penelitian (Analisis Manual dan IBM SPSS)* . Yogyakarta. Andi

RIWAYAT HIDUP PENULIS

Ig. Dodiet Aditya Setyawan, SKM.,MPH., lahir di Sragen, 12 Januari 1974. Penulis bertempat tinggal di Jalan Sukowati No. 164, Sragen Kulon, Sragen, Jawa Tengah. Mendapatkan gelar *Master of Public Health* (M.P.H) pada Program Studi S2 Ilmu Kesehatan Masyarakat di Fakultas Kedokteran Universitas Gadjah Mada Yogyakarta, pada Tahun 2014.

Berkarir sebagai Dosen di Politeknik Kesehatan Kementerian Kesehatan Surakarta (Polkesta) dengan Jabatan Lektor sampai dengan saat ini. Selain sebagai Dosen, penulis juga merangkap jabatan sebagai Sekretaris Jurusan Terapi Wicara sejak Tahun 2014 sampai sekarang.

Mata kuliah yang diampu oleh penulis pada saat ini diantaranya adalah Metodologi Penelitian, Statistika, Biostatistika, Ilmu Kesehatan Masyarakat dan Sistem Informasi Kesehatan (SIK). Selain mengampu mata kuliah tersebut, penulis juga sangat tertarik dengan pemanfaatan aplikasi Sistem Informasi Geografis (SIG) pada Kesehatan Masyarakat. Hal ini ditunjukkan dengan dihasilkannya berbagai artikel ilmiah hasil penelitian terkait dengan SIG yang dimuat pada Jurnal Internasional Bereputasi maupun Jurnal-Jurnal Nasional Terakreditasi.

Karya-karya ilmiah dari penulis juga sudah mendapatkan HKI baik dalam bentuk Artikel Ilmiah, Laporan Kegiatan Pengabdian Masyarakat, Poster Ringkasan Hasil Penelitian, Buku Petunjuk Praktikum dan Modul.



Ade Devriany, SKM., M.Kes. Lulus S1 di Jurusan Epidemiologi Fakultas Kesehatan Masyarakat Universitas Hasanuddin tahun 2010, lulus S2 di Program Magister Kesehatan di Program Pasca Sarjana Universitas Hasanuddin. Saat ini adalah dosen tetap Program Studi D-III Gizi di Poltekkes Kemenkes Pangkalpinang. Mengampu mata kuliah Ilmu Kesehatan Masyarakat, Epidemiologi Gizi dan Metodologi Penelitian. Aktif menulis artikel di berbagai jurnal ilmiah baik jurnal nasional maupun internasional dan menjadi narasumber dalam beberapa seminar dan workshop terkait Breastfeeding.



Nuril Huda lahir pada 07 Juli 1987 di Sragen, Jawa Tengah. Berasal dari keluarga petani, ia berprinsip kerja keras dan ulet. Selepas meraih sarjana Pendidikan Matematika dari UIN Sunan Kalijaga lulus 2010. Aktif mengajar di sekolah swasta dan lembaga bimbingan belajar di Yogyakarta, serta menjadi tim penyusun soal matematika jenjang SMA di bimbingan belajar tersebut.

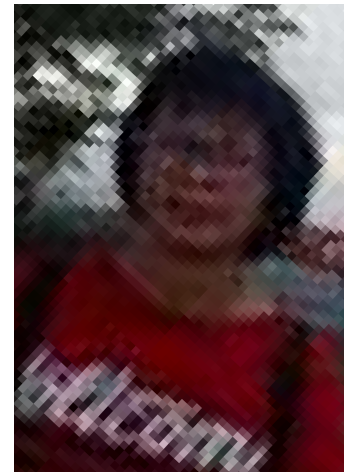


Tahun 2015 lulus dari Pascasarjana UNY dari Program Studi Penelitian dan Evaluasi Pendidikan (konsentrasi pada Pengukuran dan Pengujian bidang Matematika). Tahun 2016 menjadi dosen luar biasa (DLB) di FITK IAIN Tulungagung. Tahun 2019 diangkat menjadi Pengawai Negeri Sipil di UIN Maulana Malik Ibrahim Malang sebagai dosen Tadris Matematika.

Nina Rahmadiliyani, S.Kep., MPH. Penulis adalah dosen di STIKes Husada Borneo sejak tahun 2009. Penulis juga Lulus S1 Keperawatan tahun 2006 di Universitas Muhamadiyah Surakarta dan S2 Kesehatan Masyarakat tahun 2009 di Universitas Gadjah Mada Yogyakarta. Aktif menulis buku ajar, jurnal nasional terakreditasi. Penulis sering mengisi berbagai seminar dan workshop tingkat nasional serta oral presentasi tingkat seminar internasional



Ros Endah Happy Patriyani, S.Kp., Ns., M.Kep, lahir di Surakarta, 18 Mei 1973. Memulai pendidikan di Akademi Keperawatan Panti Kosala Surakarta (1991-1994), melanjutkan pendidikan Sarjana Keperawatan dan Profesi Ners di PSIK FK UNDIP Semarang (1998-2002). Pada tahun 2009 lulus dari Program Studi Magister Keperawatan FIK. Bekerja sebagai perawat pelaksana di Rumah Sakit Telogorejo Semarang selama 3 tahun. Mulai tahun 1998 sampai sekarang sebagai dosen Jurusan Keperawatan Poltekkes Kemenkes Surakarta dan aktif mengajar mata kuliah Keperawatan Komunitas, Keperawatan Keluarga, Keperawatan Gerontik, dan Promosi Kesehatan. Saat ini sedang menempuh Program Doktorat Di Universitas Sebelas Maret Surakarta.



Endang Caturini Sulistyowati, lahir pada 20 April 1970 di Surakarta, adalah dosen Keperawatan di Politeknik Kesehatan Surakarta. Ia menyelesaikan pendidikan Sarjana Keperawatan di Universitas Gadjah Mada tahun 2002. Pendidikan Magister Keperawatan di Universitas Indonesia tahun 2009. Bekerja sejak tahun 1998 sampai sekarang. Hasil penelitiannya dipublikasikan di jurnal ilmiah nasional terakreditasi.



[illegible]

BUKU AJAR STATISTIKA



Puji syukur kami panjatkan kehadiran Allah SWT atas karunia dan hidayah-Nya, kami dapat menyusun Buku Ajar Statistika untuk Mahasiswa, yakni mata kuliah Statistika. Buku Ajar ini disusun berdasarkan RPS Statistika. Buku Statistika terdiri dari beberapa penulis/dosen Perguruan tinggi ternama. Isi Buku membahas mengenai Distribusi Frekuensi, Ukuran Pemusatan, Dispersi, Probabilitas, Populasi dan Sampel, Teknik Pemilihan Analisis Statistik, Analisis Statistik Parametrik dan Analisis Statistik Non Parametrik.

Dengan dibuatnya Buku Ajar ini penulis berharap agar dapat bermanfaat dan membantu dalam memahami materi Statistika Selanjutnya, rasa terima kasih yang penulis ucapkan kepada semua pihak yang membantu dalam penyelesaian Buku Ajar ini.

Penulis menyadari bahwa Buku Ajar ini masih jauh dari kesempurnaan, maka dari itu penulis mengharapkan kritik dan saran pembaca demi kesempurnaan Buku Ajar ini kedepannya. Akhir kata penulis ucapkan terima kasih, mudah-mudahan bermanfaat bagi para pembaca.



 Penerbit Adab
 @penerbitadab
 www.PenerbitAdab.id

Penerbit Adab - Indramayu - Jawa Barat
Telp. 081221151025 | penerbitadab@gmail.com