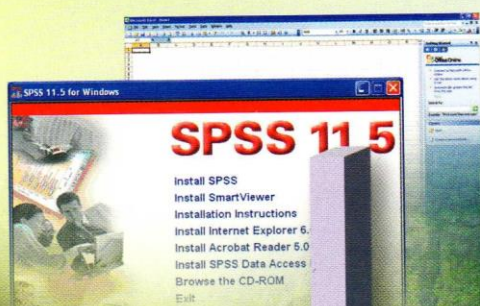


STATISTIKA UNTUK PENELITIAN PENDIDIKAN DAN APLIKASINYA DENGAN SPSS DAN EXCEL



ALI ANWAR

STATISTIKA UNTUK PENELITIAN PENDIDIKAN

DAN APLIKASINYA DENGAN SPSS DAN EXCEL



IAIT PRESS

STATISTIKA UNTUK PENELITIAN PENDIDIKAN

DAN APLIKASINYA DENGAN SPSS DAN EXCEL

**Sanksi Pelanggaran Pasal 44:
Undang-undang No. 7 Tahun 1987
Tentang Hak Cipta**

1. Barang siapa dengan sengaja dan tanpa hak mengumumkan atau memperbanyak suatu ciptaan atau memberi izin untuk itu, dipidana dengan pidana penjara paling lama 7 (tujuh) tahun dan/atau didenda paling banyak Rp.100.000.000,00 (Seratus Juta Rupiah);
2. Barang siapa dengan sengaja menyiarkan, memamerkan, mengedarkan atau menjual kepada umum suatu ciptaan atau barang hasil pelanggaran Hak Cipta sebagaimana dimaksud dalam ayat (1) dipidana dengan pidana penjara paling lama 5 (lima) tahun dan/atau didenda paling banyak Rp.50.000.000,00 (Lima Puluh Juta Rupiah).

ALI ANWAR

**STATISTIKA UNTUK
PENELITIAN PENDIDIKAN
DAN APLIKASINYA DENGAN SPSS DAN EXCEL**



IAIT PRESS

Perpustakaan Nasional: Katalog Dalam Terbitan (KDT)

**STATISTIKA UNTUK PENELITIAN PENDIDIKAN
DAN APLIKASINYA DENGAN SPSS DAN EXCEL
Oleh Ali Anwar**

x + 298 halaman; 15,5 x 23 cm
ISBN: 978-979-18633-2-2

Diterbitkan oleh: IAIT Press
Program Pascasarjana IAIT Kediri
Jl. KH. Wahid Hasyim 62 Kediri 64114
Telp.: 0354 777239, 772879
Fax.: 0354 777239, 772879
E-mail: iait.press@yahoo.co.id

Cetakan pertama: Maret 2009

KATA PENGANTAR

Buku ini merupakan buku statistik terapan yang menyajikan aplikasi dalam menguji instrumen penelitian, menentukan besarnya sampel, mendeskripsikan data dan menganalisisnya serta menguji hipotesis dengan menggunakan SPSS dan Excel. Sebagai buku statistik terapan buku ini tidak menjelaskan bagaimana rumus-rumus statistik itu disusun dan dihasilkan.

Kehadiran buku ini merupakan tanggung jawab akademik penulis karena sejak tahun 2004 penulis mendapat amanat untuk menjadi Ketua Analis Data Kuantitatif STAIN Kediri dan sejak tahun 2005 penulis juga diminta mengampu Mata Kuliah Statistik Pendidikan pada Program Pascasarjana IAIT Kediri. Oleh karena itu, tujuan sederhana dari buku ini adalah untuk memenuhi kebutuhan sivitas akademika pada kedua perguruan tinggi tersebut baik ketika mahasiswa atau dosen mengadakan penelitian dengan pendekatan kuantitatif maupun ketika dosen memberikan bimbingan kepada mahasiswanya.

Buku ini hampir tidak dimungkinkan sampai kepada sidang pembaca jika tidak ada peran dari berbagai pihak. Oleh karena itu penulis mengucapkan terima kasih kepada teman-teman peserta Pelatihan Metodologi Sosial Keagamaan kerjasama Direktorat Pendidikan Tinggi Islam Depag RI dan Sekolah Pascasajana UGM pada tahun 2007 dan 2008 yang memberikan dorongan agar penulis secepatnya merampungkan karya ini. Terima kasih yang sama juga penulis sampaikan kepada mahasiswa Program Pascasarjana IAIT Kediri angkatan 2007 dan 2008 yang senantiasa menagih terselesaikannya buku ini.

Penghargaan yang tinggi juga penulis sampaikan kepada Dr. Kharisudin Aqib, M. Ag. Direktur Program Pascasarjana IAIT Kediri periode 2003-2008 yang telah memberikan kepercayaan penulis untuk mengampu Mata Kuliah Statistik Pendidikan. Penghargaan yang sama juga penulis sampaikan kepada Prof. Dr. Faturahman, Guru Besar Fakultas Psikologi UGM yang telah banyak memberi inspirasi kepada penulis untuk menulis karya ini. Apresiasi serupa juga penulis sampaikan kepada Prof. Fauzan Saleh, Ph. D., Guru Besar Jurusan Ushuluddin STAIN Kediri, Prof. Dr. Irwan Abdullah, Guru Besar Fakultas Ilmu

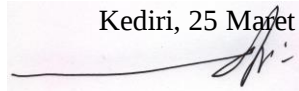
Budaya UGM dan Direktur Sekolah Pascasarjana UGM, dan Prof. Dr. Abuddin Nata, MA. Guru Besar UIN Jakarta yang telah banyak “memprovokasi” penulis untuk senantiasa berkarya. Melalui relasi yang sedemikian bersahabat dan produktivitas akademik yang tinggi, ketiga Guru Besar yang disebutkan terakhir senantiasa dapat dijadikan *modelling* dalam kehidupan ini.

Melalui buku ini perkenankanlah penulis mengucapkan terima kasih kepada Sahabat penulis yang telah berkenan memberikan cobaan dan ujian terhadap perjalanan akademik penulis. Penulis menduga karena pemahaman beliau yang baik tentang al-Qur’an maka beliau berusaha untuk menjadi Khalifatu ‘Ilāh yang hendak mempraktikkan QS al-Baqarah/2: 214, Āli ‘Imrān/3: 142, dan al-‘Ankabūt/29: 2. Hal ini beliau lakukan sejak proses penerimaan mutasi penulis dari IAIN Bandar Lampung ke STAIN Kediri pada tahun 2003, ketika penulis mengajukan kenaikan jabatan fungsional ke Lektor Kepala pada tahun 2006, ketika proses penyelesaian studi penulis pada program doktor, dan terakhir ketika penulis berusaha mencapai karir akademik tertinggi sebagai Guru Besar. Hanya saja, ketiga cobaan pertama ternyata Allah dalam waktu yang tepat memberi jalan keluar, sedangkan yang keempat, yaitu kenaikan jabatan fungsional ke Guru Besar tampaknya untuk sementara kehendak Sahabat penulis tersebut tidak bertentangan dengan kehendak Allah. Tetapi penulis yakin, sebagaimana Allah janjikan dalam Sūrat al-Baqarah ayat 214 di atas bahwa pertolongan Allah amat dekat. Oleh karena itu, semoga penulis dapat berhasil menghadapi cobaan ini dan layak disebut sebagai orang yang beriman dan akhirnya mendapat surga di akhirat kelak.

Tanpa pengertian istri tercinta, Aini Masykurun, dan ketiga anak-anak tersayang, Muhammad Medina Almas Ali, Nayl al-Falāhi, dan Muhammad Fā’iq Ashfa, buku ini hampir tidak mungkin terselesaikan. Oleh karena itu, buku ini kupersembahkan buah segenap keluarga penulis, semoga hal ini menjadikan Allah memudahkan penulis sekeluarga untuk menjadi orang-orang yang bermanfaat. Tidak hanya itu, semoga kehadliran buku merupakan bukti syukur penulis sekeluarga kepada Allah atas segala nikmat yang telah Allah karuniakan.

Dengan senang hati penulis akan menerima saran untuk perbaikan dan kritik jika ada kesalahan. Semoga buku ini bermanfa’at dan kepada semua pihak yang memungkinkan buku ini terselesaikan diberikan pahala oleh Allah yang berlipat ganda dan hidup dan kehidupannya ke depan menjadi lebih barokah dan manfaat, amīn.

Kediri, 25 Maret 2009.



Ali Anwar

E-mail: ali_anwar03@yahoo.co.id

DAFTAR ISI

KATA PENGANTAR	v
DAFTAR ISI	vii
DAFTAR LAMPIRAN	x
 BAB I PENDAHULUAN	 1
A. Latar Belakang	1
B. Manfaat Statistik dalam Penelitian Kuantitatif	2
C. Macam-macam Statistik	2
D. Jenis Data	3
E. Uji Validitas dan Reliabilitas Instrumen	5
F. Pedoman Umum Memilih Teknik Analisis Statistik	21
G. Populasi dan Sampel	23
H. Konsep Pengujian Hipotesis	35
I. Urut-urutan dalam Menginstall Program SPSS Versi 11.5 dan Proses mengaktifkan Menu Data Analysis Microsoft Excel	37
 BAB II STATISTIK DESKRIPTIF	 47
A. Pendahuluan	47
B. Aplikasi untuk Penyajian Data dengan SPSS	50
C. Aplikasi untuk Penyajian Data dengan Microsoft Excel	61
D. Aplikasi untuk Penyajian Pemusatan, Penyebaran, dan Distribusi Data dengan SPSS ...	71
E. Aplikasi untuk Penyajian Pemusatan,	73

	Penyebaran, dan Distribusi Data dengan Microsoft Excel	
BAB III	UJI SATU SAMPUL	85
	A. Pendahuluan	85
	B. Statistik Parametrik: T-test one sample	85
	1. Aplikasi dengan SPSS	86
	2. Aplikasi dengan Microsoft Excel	91
	C. Statistik Non-Parametrik	92
	1. Test Binomial	92
	2. Chi Kuadrat	94
	3. Test Run	98
BAB IV	ANALISIS KORELASI	103
	A. Pendahuluan	103
	B. Arah Korelasi	103
	C. Angka Korelasi	104
	D. Macam-macam Teknik Korelasi	104
	E. Statistik Parametrik	104
	1. Product Moment	104
	2. Korelasi Ganda	115
	3. Korelasi Parsial	120
	F. Statistik Non-parametrik	126
	1. Koefisien Kontingensi	126
	2. Spearman Rank	131
	3. Kendall's tau	135
	G. Koefisien Penentu	139
BAB V	ANALISIS REGRESI	141
	A. Pendahuluan	141
	B. Regresi Linear Sederhana	142
	1. Aplikasi dengan SPSS	142
	2. Aplikasi dengan Microsoft Excel	149
	C. Regresi Ganda Dua Prediktor	152
	1. Aplikasi dengan SPSS	153
	2. Aplikasi dengan Microsoft Excel	161

BAB VI	ANALISIS KOMPARASI	165
	A. Pendahuluan	165
	B. Macam-macam Teknik Analisis Komparasi	165
	C. Komparasi Dua Sampel Berkorelasi	166
	1. Statistik Parametrik: T-test of related	166
	2. Statistik Non Parametrik	172
	a. Mc Nemar	172
	b. Sign Test	176
	c. Wilcoxon Matched Pairs	184
	D. Komparasi Dua Sampel Independen	190
	1. Statistik Parametrik: T-test of Independent ..	190
	2. Statistik Non Parametrik	204
	a. Fisher Exact Probability	204
	b. χ^2 Two Sample	209
	c. Median Test	213
	d. Mann-Whitney U-Test	218
	e. Kolmogorov Semirnov	224
	f. Wald Woldfowitz	228
	E. Komparasi k Sampel Berkorelasi	232
	1. Statistik Parametrik: One-way Anova	239
	2. Statistik Non Parametrik	239
	a. χ^2 k Sample related	239
	b. Cochran's Q	243
	c. Friedman Two-way Anova	247
	F. Komparasi k Sampel Independen	251
	1. Statistik Parametrik: One-way Anova	251
	2. Statistik Non Parametrik	258
	a. χ^2 k Sample Independent.....	258
	b. Median Extention	262
	c. Kruskal-Wallis One-way Anova	265
	Daftar Kepustakaan	271
	Lampiran	273

DAFTAR LAMPIRAN

Tabel I	Luas di bawah Lengkungan Kurva Normal dari 0 s.d. Z	275
Tabel II	Nilai-Nilai dalam Distribusi t	277
Tabel III	Nilai-Nilai r Product Moment	279
Tabel IV	Harga-Harga X dalam Test Binomial (Harga-Harga dalam Tabel Adalah 0,...)	280
Tabel V	Harga Factorial	281
Tabel VI	Nilai-Nilai Chi Kuadrat	282
Tabel VII a	Harga-Harga Kritis r dalam Test Run Satu Sampel, untuk A = 5 %	283
Tabel VII b	Harga-Harga Kritis r dalam Test Run Dua Sampel, Untuk A = 5 %	284
Tabel VIII	Harga-Harga Kritis untuk Test Wilcoxon	285
Tabel IX	Harga-Harga Kritis Man-Whitney U Test	286
Tabel X	Tabel Harga-Harga Kritis dalam Test Kolmogorov-Smirnov	287
Tabel XI	Harga-Harga Z untuk Test Run Wald-Wolfowitz ..	288
Tabel XII	Nilai-Nilai untuk Distribusi F	290
Tabel XIII	Tabel Nilai-Nilai rho	294
Tabel XIV	Tabel Harga-Harga Kritis Z dalam Observasi Distribusi Normal	295
Cara Menghitung Nilai Tabel: T, F, χ^2 , dan Z dengan Aplikasi Microsoft Excel		297

BAB I

PENDAHULUAN

A. Latar Belakang

Selama ini mahasiswa sering merasa bahwa Statistik adalah mata kuliah yang sulit dipahami dan dipraktikkan. Perasaan tersebut tidaklah berlebihan. Hal ini sesuai dengan temuan Singgih Santoso, “Jika orang mendengar kata statistik, maka asosiasi mereka adalah tentang sesuatu yang ruwet, memusingkan, penuh dengan rumus-rumus yang rumit, membosankan, dan sebagainya.”¹ Sebagai akibatnya mahasiswa yang menulis skripsi dengan pendekatan kuantitatif menemukan banyak kesulitan untuk menentukan analisis statistik yang tepat dan mengaplikasikan analisis yang telah dipilih.

Dengan semakin meluasnya ketersediaan komputer pada lembaga pendidikan akhir-akhir ini, metode pengajaran Statistik di beberapa perguruan tinggi telah mengalami perubahan yang dramatis.² Kusnandar selanjutnya mengatakan bahwa ketersediaan perangkat keras dan lunak komputer telah memberikan banyak kemudahan kepada peneliti dalam menganalisis hasil penelitian.³ Bila dahulu ada anggapan bahwa penelitian yang meneliti tiga variabel atau lebih hanya dapat dilakukan oleh peneliti setingkat mahasiswa S2 atau bahkan S3, dikarenakan sulitnya analisis multivariat, sekarang itu hal tersebut dapat dilakukan oleh mahasiswa S1 atau oleh peneliti pemula. Singgih Santoso menganggap bahwa

¹Singgih Santoso, *Statistik Diskriptif: Konsep dan Aplikasi dengan Microsoft Excel dan SPSS*, (Yogyakarta: Penerbit Andi, 2003), hlm. iii.

²Dadan Kusnandar, *Metode Statistik dan Aplikasinya dengan Minitab dan Excel*, (Yogyakarta: Madyan Press, 2004), hlm. 7.

³Kusnandar, *Metode Statistik dan Aplikasinya*, hlm. v.

perkembangan Software Statistik yang pesat membuat penggunaan metode statistik Multivariat yang sangat kompleks menjadi mudah dan praktis.⁴ Kalau terhadap statistik Multivariat saja menjadi mudah dan praktis, apalagi bagi statistik univariat dan bivariat.

Kejadian ini, sejauh pengamatan penulis belum banyak terjadi di berbagai perguruan tinggi di wilayah Kediri dan sekitarnya. Belum dimanfaatkannya teknologi ini untuk analisis statistik boleh jadi disebabkan sivitas akademika belum tahu dan masih meragukan tentang berbagai efektivitas, kemudahan, dan validitas dari teknologi itu atau lembaga tersebut belum memiliki hardware yang dibutuhkan untuk aplikasi tersebut.

Berangkat dari latar belakang masalah di atas, maka penulis melalui buku ini berusaha memberikan kesan mudah terhadap statistik dengan memberikan contoh dan aplikasinya dengan SPSS dan Microsoft Excel.

B. Manfaat Statistik dalam Penelitian Kuantitatif

Setidaknya ada 4 (empat) manfaat Statistik dalam penelitian kuantitatif:

1. Untuk menghitung besarnya anggota sampel yang diambil dari populasi, sehingga sampel dapat lebih representatif dan dapat dipertanggungjawabkan.
2. Untuk menguji validitas dan reliabilitas instrumen.
3. Teknik-teknik untuk mendeskripsikan data sehingga data lebih komunikatif; dan
4. Alat untuk analisis data.⁵

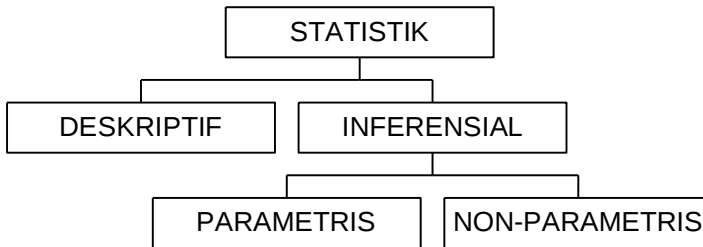
C. Macam-macam Statistik

Statistik dapat dibedakan menjadi dua, deskriptif dan inferensial. Statistik Deskriptif adalah statistik yang digunakan untuk menggambarkan data atau menganalisis suatu hasil penelitian tetapi tidak digunakan untuk generalisasi. Sementara statistik inferensial adalah statistik yang digunakan untuk menganalisis data sampel dan hasilnya akan digeneralisasikan.

⁴Singgih Santoso, *Buku Latihan SPSS Statistik Multivariat*, (Jakarta: Elex Media Komputindo, 2003), hlm. v.

⁵ Sugiyono, *Statistika untuk Penelitian*, (Bandung: CV Alfabeta, 2003), hlm. 12.

Statistik inferensial dapat dibedakan menjadi dua, parametrik dan non-parametrik. Statistik parametrik terutama digunakan untuk menganalisis data interval atau rasio yang diambil dari populasi yang berdistribusi normal. Sedangkan statistik non-parametrik terutama digunakan untuk menganalisis data nominal atau ordinal.⁶ Atau datanya interval atau rasio tetapi tidak berdistribusi normal juga menggunakan statistik non-parametrik. Macam-macam statistik itu dapat digambarkan seperti di bawah ini:



Dikarenakan penentuan teknik analisis ditentukan juga oleh jenis data, maka di bawah ini akan dijelaskan jenis data dan ciri-cirinya.

D. Jenis Data

Data hasil penelitian dapat dikelompokkan menjadi dua: kualitatif dan kuantitatif. Data kualitatif adalah data yang berbentuk kalimat, kata, atau gambar. Sedangkan data kuantitatif adalah data yang berbentuk angka, atau kualitatif yang diangkakan.⁷

Data kuantitatif dapat juga dikelompokkan menjadi dua: diskrit dan kontinum. Data diskrit adalah data yang diperoleh dari menghitung, bukan mengukur. Data ini sering disebut dengan data nominal. Tidak hanya itu, data yang berjenis nominal juga data yang diperoleh dengan cara kotegorisasi atau klasifikasi. Contoh: jenis pekerjaan diklasifikasi sebagai berikut: pegawai negeri diberi tanda 1, pegawai swasta diberi tanda 2, dan wiraswasta diberi tanda 3. Ciri data nominal adalah setara, dalam kasus ini pegawai negeri tidak lebih tinggi dibanding wiraswasta dan data tipe ini tidak dapat dilakukan operasi matematika.

Sementara data kontinum dikelompokkan menjadi tiga, yaitu ordinal, interval, dan rasio. Data ordinal adalah data yang berjenjang

⁶ Sugiyono, *Statistika untuk Penelitian*, hlm 13-14.

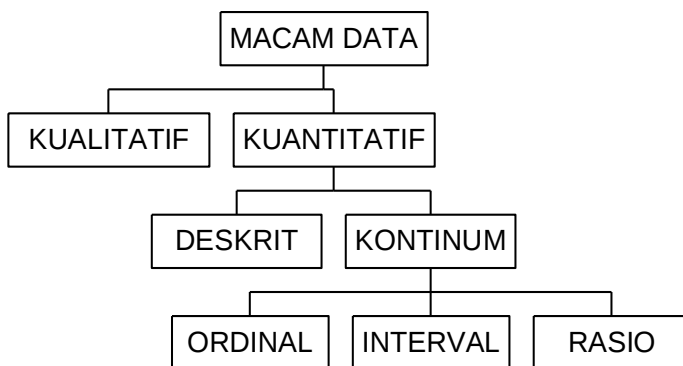
⁷ Sugiyono, *Statistika untuk Penelitian*, hlm 14-15.

dan berbentuk peringkat, atau dengan kata lain data ordinal adalah data yang diperoleh dengan cara kategorisasi atau klasifikasi yang berjenjang dan bertingkat. Contoh: kepuasan mahasiswa diklasifikasi sebagai berikut: sangat puas diberi tanda 5, puas diberi tanda 4, cukup puas diberi tanda 3, tidak puas diberi tanda 2, dan sangat tidak puas diberi tanda 1. Ciri data ordinal adalah posisi data tidak setara (berbentuk peringkat/berjenjang), tidak dapat dilakukan rumus matematika, dan jarak antara dua titik tidak diketahui. Data ordinal ini dapat dibentuk dari data interval atau rasio.

Sedangkan data interval adalah data yang jaraknya sama tetapi tidak mempunyai nol absolut atau dengan kata lain bahwa data interval adalah data yang diperoleh dengan cara pengukuran, di mana jarak antara dua titik diketahui. Contoh: temperatur suatu ruangan: celcius pada 0^0 c sampai 100^0 C. Ciri tipe data interval: tidak ada kategorisasi, bisa dilakukan rumus matematika, dan tidak ada nol absolut.

Sementara data rasio adalah data yang jaraknya sama dan mempunyai nol absolut atau dengan kata lain data rasio adalah data yang diperoleh dengan pengukuran dan ada nol absolut. Contoh: hasil pengukuran panjang atau berat. Ciri tipe data rasio, tidak ada kategorisasi, dapat dilakukan rumus matematika, dan ada nol absolut.⁸

Bermacam data seperti di atas dapat digambarkan seperti di bawah ini:



⁸ Sugiyono, *Statistika untuk Penelitian*, hlm 14-15.

Contoh data dalam skala pengukuran nominal dan ordinal:

No	Skala Penguk.	Data Kualitatif	Kategori
01	Nominal	Suku	1. Sunda 2. Jawa 3. Madura 4. Lainnya
		Kepemilikan motor	1. ya 2. tidak
02	Ordinal	Pendidikan	1. PT 2. SMA 3. SMP 4. SD
		Jabatan Dosen	1. Guru Besar 2. Lektor Kepala 3. Lektor 4. Asisten Ahli
		Nilai Akhir Mata Kuliah	1. A 2. B 3. C 4. D 5. E ⁹

Contoh data dalam skala pengukuran interval dan rasio:

Data Kuantitatif	Skala Pengukuran
Suhu (Celcius atau Fahrenheit)	interval
penanggalan (Masehi atau Hijriah)	interval
tinggi (meter)	rasio
berat (kilogram)	rasio
Umur (tahun atau hari)	rasio ¹⁰

E. Uji Validitas dan Reliabilitas Instrumen

Secara umum, salah satu perbedaan antara skripsi, tesis, dan disertasi dapat dilihat dari aspek metodologi penelitian. Penulis skripsi dituntut untuk menyebutkan adanya upaya untuk memperoleh data penelitian secara akurat dengan menggunakan instrumen pengumpul data yang valid. Penulis tesis harus menyertakan bukti-bukti yang dapat dijadikan pegangan untuk menyatakan bahwa

⁹ Kusnandar, *Metode Statistik dan Aplikasinya*, hlm. 5.

¹⁰ Kusnandar, *Metode Statistik dan Aplikasinya*, hlm. 6.

instrumen pengumpul data yang digunakan cukup valid. Bagi penulis disertasi, bukti-bukti validitas instrumen pengumpul data harus dapat diterima sebagai bukti-bukti yang tepat.¹¹ Dalam skripsi, penyimpangan-penyimpangan yang terjadi dalam pengumpulan data tidak harus dikemukakan, sedangkan dalam tesis terlebih lagi dalam disertasi, penyimpangan itu harus dikemukakan beserta alasan-alasannya, sejauh mana penyimpangan tersebut, dan sejauh mana penyimpangan tersebut dapat ditoleransi.¹²

1. Langkah-Langkah Penyusunan Instrumen

Penyusunan suatu instrumen biasanya dilakukan setelah suatu konsep yang ingin diukur didefinisikan secara jelas. Definisi tersebut sudah harus dapat disampaikan dalam wujud pertanyaan atau pernyataan. Dalam bahasa metodologi definisi yang dimaksudkan sudah harus operasional. Cara untuk mengoperasionalkan definisi, menurut Ancok, paling sedikit ada tiga cara yaitu:

- a. Mencari definisi-definisi tentang konsep yang telah ditulis di dalam literatur oleh para ahli. Kalau sekiranya sudah ada definisi yang cukup operasional untuk ditumpahkan ke dalam suatu instrumen, definisi tersebut dapat langsung digunakan. Tetapi bila definisi yang dikemukakan oleh para ahli belum begitu operasional, maka peneliti harus menjabarkan definisi tersebut seoperasional mungkin. Dengan demikian akan lebih mudahlah kita menyusun pertanyaan.
- b. Kalau sekiranya di dalam literatur tidak diperoleh definisi konsep yang ingin diukur, maka peneliti harus mendefinisikan sendiri dengan menggunakan pemikirannya sendiri. Untuk lebih memantapkan definisi operasional yang dibuat, sebaiknya peneliti mendiskusikan definisi tersebut dengan ahli-ahli lain. Setelah definisi yang dibuat cukup mantap, barulah definisi tersebut diwujudkan dalam bentuk pertanyaan yang akan menjadi komponen instrumen yang akan dibuat. Cara yang kedua ini dapat dilakukan dengan cara kebalikan dengan cara yang di atas. Kita tidak memulai dari menyusun definisi sendiri, tetapi peneliti mengumpulkan lebih dahulu bagaimana pendapat para ahli lewat diskusi.

¹¹Ali Saukah dkk., *Pedoman Penulisan Karya Ilmiah*, (Malang: Universitas Negeri Malang (UM), 2000), hlm. 3.

¹²Ali Saukah dkk., *Pedoman Penulisan Karya Ilmiah*, hlm. 3.

Setelah cukup banyak kesamaan pendapat di kalangan para ahli terhadap definisi suatu konsep, kemudian disusun definisi operasional konsep yang bersangkutan.

- c. Dengan menanyakan langsung kepada para responden. Berdasarkan aspek-aspek yang dijelaskan responden tersebut disusun pertanyaan yang akan menjadi komponen instrumen.¹³

Pendapat Ancok di atas dapat disederhanakan sebagai berikut.

- a. Seluruh variabel, yang dalam istilah Ancok disebut konsep, yang terdapat dalam judul penelitian dijabarkan menjadi sub variabel.
- b. Menentukan indikator dari setiap sub variabel.
- c. Menjabarkan masing-masing indikator menjadi beberapa deskriptor.
- d. Merumuskan setiap deskriptor menjadi butir-butir pertanyaan atau pernyataan.

Agar setiap butir instrumen memenuhi validitas internal dan eksternal, maka penyusunan instrumen tersebut harus mengacu dan berdasarkan pada teori dan fakta empiris. Masing-masing butir instrumen harus diteliti satu persatu untuk melihat apakah pertanyaan atau pernyataan tersebut sudah memenuhi persyaratan dalam penulisannya.

Di dalam menulis pertanyaan atau pernyataan, menurut Ancok, hal-hal berikut ini perlu diperhatikan:

- a. Hindari pertanyaan atau pernyataan yang dapat menimbulkan lebih dari satu pengertian.
- b. Hindari pertanyaan atau pernyataan yang tidak relevan dengan dimensi konsep yang akan diukur.
- c. Hindari pertanyaan atau pernyataan yang diperkirakan orang akan setuju atau tidak setuju.
- d. Gunakan bahasa yang mudah dimengerti, jelas, dan singkat.

¹³Djamaludin Ancok, *Teknik Penyusunan Skala Pengukur*, (Yogyakarta: Pusat Studi Kependudukan dan Kebijakan Universitas Gadjah Mada, 2002), hlm. 8-9.

- e. Setiap pertanyaan atau pernyataan harus berisikan satu hal saja. Hindarkan pertanyaan atau pernyataan yang mengandung dua atau lebih permasalahan.¹⁴

Langkah-langkah tersebut perlu dilakukan agar responden mendapatkan kemudahan untuk mengisi instrumen yang diajukan.

2. Prosedur Uji Validitas dan Reliabilitas Instrumen

Setelah instrumen selesai disusun dan diaudit oleh tenaga ahli, selanjutnya diujicobakan kepada sebagian responden yang menjadi sampel penelitian tersebut yang jumlahnya paling sedikit 30 orang. Jumlah responden yang lebih dari 30 orang biasanya cukup memadai untuk taraf uji-coba. Hal ini disebabkan distribusi skor akan mendekati distribusi normal.

Hasil uji-coba ini kemudian akan digunakan untuk mengetahui sejauh mana instrumen yang telah disusun memiliki validitas dan reliabilitas. Suatu instrumen yang baik harus memiliki validitas dan reliabilitas.

Validitas ialah indeks yang menunjukkan sejauh mana suatu instrumen betul-betul mengukur apa yang perlu diukur. Timbangan hanya valid untuk mengukur berat, tidak valid untuk mengukur panjang. Sebaliknya meteran hanya valid bila digunakan untuk mengukur panjang.

Selama ini uji validitas dilakukan dengan menggunakan korelasi antara skor item dan Skor Total (*Item-Total Correlation*). Korelasi antara skor item dengan skor total haruslah signifikan berdasarkan ukuran statistik tertentu. Bila sekiranya skor semua pertanyaan atau pernyataan yang disusun berdasarkan dimensi konsep berkorelasi dengan skor total, maka dapat dikatakan bahwa instrumen tersebut mempunyai validitas.

Validitas yang seperti ini, menurut Ancok, disebut dengan validitas konstruk (*construct validity*).¹⁵ Bila instrumen telah memiliki validitas konstruk berarti semua item yang ada di dalam instrumen itu mengukur konsep yang ingin diukur.

Bila instrumen yang disusun hanya memiliki jumlah item yang sedikit (kurang dari 30 item), maka harus dilakukan koreksi

¹⁴ Ancok, *Teknik Penyusunan*, hlm. 18.

¹⁵ Ancok, *Teknik Penyusunan*, hlm. 21.

terhadap angka korelasi yang diperoleh karena angka korelasi tersebut kelebihan bobot. Kelebihan bobot itu terjadi karena skor item yang dikorelasikan dengan skor total ikut sebagai komponen skor total, dan hal ini menyebabkan angka korelasi terjadi lebih besar. Dengan semakin banyak jumlah item, pengaruh skor item terhadap membesarnya angka korelasi dengan skor total akan semakin kecil.

Sedangkan rumus untuk mengoreksi koefisien korelasi skor item dengan skor total adalah berikut ini.

$$r.pq = \frac{(r.tp)(SDy) - (SDx)}{\sqrt{(SDy)^2 + (SDx)^2 - 2(r.tp)(SDx)(SDy)}}$$

Keterangan:

$r.pq$ = angka korelasi setelah dikoreksi

$r.tp$ = angka korelasi sebelum dikoreksi

SDy = standard deviasi skor total

SDx = standar deviasi item

Sebagaimana dijelaskan di atas, bahwa uji validitas konstruk dilakukan dengan menggunakan korelasi antara skor item dan skor total yang dicari dengan rumus Pearson Product Moment.¹⁶

Contoh:

Peneliti sedang mengadakan penelitian tentang Pengaruh Kompetensi Guru terhadap Prestasi belajar Siswa. Peneliti menjabarkan variabel Kompetensi Guru ke dalam 4 (empat) sub variabel, yaitu kompetensi paedagogik, kompetensi profesional, kompetensi sosial, dan kompetensi kepribadian. Instrumen yang digunakan untuk mengumpulkan data tentang variabel kompetensi guru adalah angket sedangkan prestasi belajar siswa dikumpulkan datanya dengan menggunakan dokumentasi. Masing-masing sub variabel dari variabel kompetensi guru dijabarkan dengan panduan teori dan fakta di lapangan. Misalkan, item-item pertanyaan yang dijabarkan dari sub variabel kompetensi paedagogik berjumlah 30. Sedangkan jumlah item untuk sub variabel lainnya menyesuaikan.

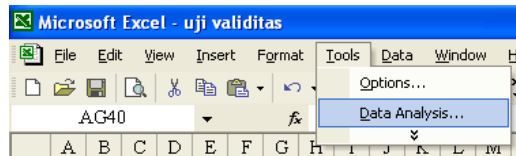
¹⁶ Prosedur analisis Person Product Moment dapat dibaca pada Bab IV, hlm. 104-115.

Setelah instrumen yang berupa angket selesai disusun dan dimintakan persetujuan kepada tenaga ahli, maka instrumen tersebut diujicobakan kepada 30 orang calon responden. Sedangkan data yang didapatkan dari sub variabel kompetensi paedagogik adalah sebagai berikut.

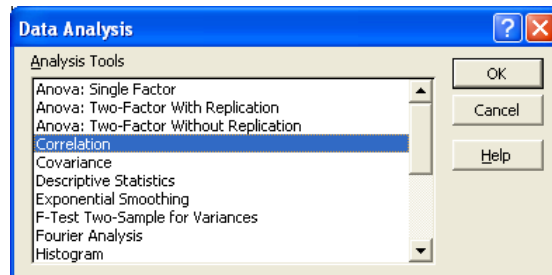
	A	B	C	D	E	F	G	H	I	J	K	L	M	N	O	P	Q	R	S	T	U	V	W	X	Y	Z	AA	AB	AC	AD	AE		
	x ₁	x ₂	x ₃	x ₄	x ₅	x ₆	x ₇	x ₈	x ₉	x ₁₀	x ₁₁	x ₁₂	x ₁₃	x ₁₄	x ₁₅	x ₁₆	x ₁₇	x ₁₈	x ₁₉	x ₂₀	x ₂₁	x ₂₂	x ₂₃	x ₂₄	x ₂₅	x ₂₆	x ₂₇	x ₂₈	x ₂₉	x ₃₀	jml		
1	5	4	5	3	4	3	5	4	4	4	4	4	2	3	2	3	2	3	4	2	3	2	3	4	3	2	3	4	3	2	101		
2	3	4	3	4	3	5	4	3	5	4	5	5	4	5	4	3	5	4	3	4	3	4	3	4	3	5	4	3	3	5	116		
3	4	5	4	4	5	3	5	4	4	3	5	4	4	5	4	5	4	3	5	4	5	4	4	5	4	4	5	4	4	4	127		
4	5	3	5	3	5	4	3	5	3	3	3	4	3	4	3	5	3	5	3	5	3	3	5	3	5	3	4	3	5	3	4	109	
5	4	3	3	4	3	5	4	3	5	4	3	5	4	5	4	5	4	4	3	4	3	4	3	4	3	5	4	3	3	5	113		
6	7	5	4	4	3	4	3	5	4	3	5	4	3	5	3	5	4	3	5	3	5	4	4	5	4	3	5	4	3	5	4	120	
7	3	2	3	4	2	2	4	2	2	4	2	2	4	2	4	2	4	2	4	2	4	2	4	2	3	4	2	3	2	3	2	81	
8	9	5	4	5	4	4	5	4	4	4	4	4	3	2	3	2	3	4	2	3	5	4	5	4	5	4	4	5	4	4	4	118	
9	3	2	2	3	2	4	3	2	2	4	3	2	4	3	4	3	3	4	3	2	3	2	2	3	2	2	4	3	2	2	4	85	
10	11	4	3	3	4	3	4	3	3	5	4	3	5	4	5	4	3	5	4	3	4	3	4	3	4	3	5	4	3	5	5	112	
11	12	5	4	4	5	4	4	5	4	4	5	4	4	5	4	4	4	5	4	5	4	5	4	4	5	4	4	5	4	4	4	130	
12	3	2	2	3	2	4	3	3	2	4	3	2	4	3	4	3	2	4	3	2	3	2	2	3	2	2	4	3	2	2	4	85	
13	4	2	3	5	2	5	3	2	5	3	4	5	3	2	3	2	5	3	2	5	2	3	5	2	5	2	3	2	5	3	3	100	
14	1	2	2	4	3	4	3	2	5	4	4	5	4	4	4	4	2	3	3	3	2	2	2	2	4	2	4	3	2	5	4	95	
15	6	4	3	4	3	4	4	3	4	4	4	4	4	4	4	4	4	3	3	3	3	3	4	3	4	3	4	4	3	4	4	108	
16	17	5	4	4	2	4	5	5	4	5	5	5	5	5	5	5	4	5	5	4	5	4	4	5	4	5	4	5	4	5	5	137	
17	4	3	3	4	3	3	4	3	3	4	3	3	4	3	4	3	4	3	4	3	4	3	3	4	3	4	3	4	3	3	3	100	
18	19	3	2	3	2	3	3	2	4	4	3	4	3	3	4	3	3	4	3	2	3	2	3	2	3	2	4	3	2	4	4	89	
19	20	2	3	3	2	3	5	2	3	5	2	5	5	2	5	2	3	5	2	3	2	3	2	3	3	2	3	5	2	3	5	5	100
20	21	3	2	4	3	2	4	3	4	1	4	1	2	4	3	4	3	2	3	3	2	3	2	4	3	2	4	3	2	1	4	85	
21	22	4	3	3	4	3	2	3	2	2	4	2	2	4	2	4	3	2	4	3	4	3	3	4	3	4	3	2	4	3	2	89	
22	5	4	4	3	4	3	5	4	3	4	5	3	3	4	3	5	2	4	5	4	5	4	4	5	4	5	4	3	5	4	3	117	
23	4	4	5	4	4	5	4	4	4	4	5	4	4	4	4	5	2	4	5	3	5	4	4	4	4	4	5	4	4	4	4	126	
24	5	4	4	5	4	4	5	4	4	4	5	4	4	4	4	4	5	2	4	5	3	5	4	4	5	4	5	4	5	4	5	118	
25	4	3	3	4	3	5	4	3	5	5	4	5	5	4	5	4	3	5	2	3	4	3	3	4	3	4	3	5	4	3	5	5	118
26	3	2	2	3	2	3	3	2	3	3	3	3	3	3	3	3	2	3	2	3	2	2	3	2	3	3	3	2	3	3	3	81	
27	3	2	4	3	2	2	3	2	2	2	3	3	2	3	2	3	2	2	3	2	3	2	4	3	3	3	2	3	2	2	2	76	
28	2	3	2	3	2	3	2	3	3	3	2	3	2	3	2	3	2	3	2	3	2	3	3	2	3	2	3	2	3	2	3	79	
29	3	2	2	3	2	4	3	2	2	4	3	2	4	3	4	3	2	5	3	2	3	2	2	3	2	2	4	3	2	2	4	85	
30	4	3	3	4	3	5	4	3	3	5	4	3	5	2	5	4	3	5	4	3	4	3	3	4	3	4	5	4	3	5	3	5	112
31	3	2	4	3	2	4	3	2	2	4	3	3	4	3	4	3	2	4	3	2	3	2	4	3	2	4	3	2	2	2	4	89	

Data di atas kemudian dicari koefisien korelasinya, yaitu korelasi skor item dengan skor total. Salah satu cara untuk mencari koefisien korelasi antara skor item dengan skor total adalah menggunakan **Data Analysis** yang aplikasinya sebagai berikut:

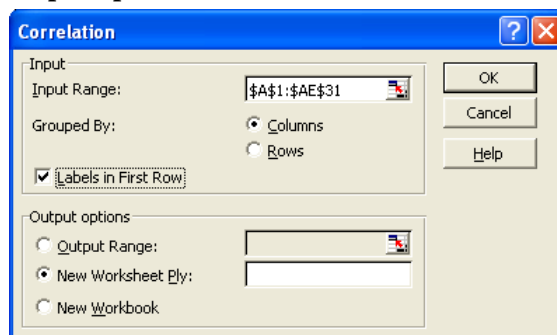
- a. Setelah datanya diinput, dijumlahkan dan diberi label seperti contoh di atas, klik **Tools** pada worksheet Microsoft Excel, lalu seret ke **Data Analysis** sebagai berikut.



- b. Setelah itu pilihlah **Correlation**, dengan cara klik dua kali atau klik **Ok**.



- c. Setelah keluar gambar seperti berikut ini, klik pada kolom yang ada di depan **Input Range**, lalu bloklah seluruh data yang sudah diinput. Aktifkan **Label in First Row**. Karena outputnya juga banyak, maka pilihlah **New Worksheet Ply** untuk **Output options**.



- d. Berikut ini adalah output aplikasi di atas. Karena jumlah item pertanyaan untuk sub variabel kompetensi paedagogik ini

berjumlah 30, maka tidak diperlukan aplikasi rumus untuk mengoreksi koefisien korelasi tersebut.

	A	B	C	D	E	F	G	H	I	J	K	L	M	N	O	P	Q	R	S	T	U	V	W	X	Y	Z	AA	AB	AC	AD	AE	AF
1		x ¹	x ²	x ³	x ⁴	x ⁵	x ⁶	x ⁷	x ⁸	x ⁹	x ¹⁰	x ¹¹	x ¹²	x ¹³	x ¹⁴	x ¹⁵	x ¹⁶	x ¹⁷	x ¹⁸	x ¹⁹	x ²⁰	x ²¹	x ²²	x ²³	x ²⁴	x ²⁵	x ²⁶	x ²⁷	x ²⁸	x ²⁹	x ³⁰	jinl
2	x ¹	1,00																														
3	x ²	0,77	1,00																													
4	x ³	0,42	0,63	1,00																												
5	x ⁴	0,47	0,25	-0,03	1,00																											
6	x ⁵	0,43	0,82	0,75	0,09	1,00																										
7	x ⁶	0,09	0,16	-0,03	0,00	0,12	1,00																									
8	x ⁷	0,83	0,62	0,21	0,61	0,33	0,10	1,00																								
9	x ⁸	0,49	0,76	0,83	0,00	0,88	0,15	0,25	1,00																							
10	x ⁹	0,22	0,55	0,11	0,11	0,50	0,39	0,30	0,27	1,00																						
11	x ¹⁰	0,16	0,17	-0,15	0,03	0,03	0,86	0,18	0,05	0,40	1,00																					
12	x ¹¹	0,69	0,68	0,26	0,52	0,55	-0,03	0,80	0,32	0,46	0,05	1,00																				
13	x ¹²	0,09	0,46	0,27	0,00	0,51	0,41	0,16	0,32	0,86	0,39	0,38	1,00																			
14	x ¹³	0,12	0,17	-0,09	0,09	0,07	0,88	0,17	0,10	0,37	0,90	0,03	0,36	1,00																		
15	x ¹⁴	0,48	0,31	-0,07	0,51	0,08	-0,08	0,77	-0,04	0,23	-0,02	0,64	0,13	0,08	1,00																	
16	x ¹⁵	0,05	0,07	-0,21	0,09	0,00	0,84	0,17	-0,01	0,37	0,90	0,00	0,36	0,96	0,12	1,00																
17	x ¹⁶	0,57	0,40	-0,04	0,58	0,13	0,00	0,87	0,00	0,22	0,09	0,73	0,10	0,15	0,90	0,18	1,00															
18	x ¹⁷	0,19	0,50	0,48	-0,10	0,70	0,10	0,09	0,61	0,35	-0,02	0,30	0,42	0,15	0,12	0,15	0,04	1,00														
19	x ¹⁸	0,22	0,20	-0,14	-0,05	0,04	0,76	0,17	0,04	0,34	0,91	0,08	0,32	0,81	-0,08	0,81	0,04	0,02	1,00													
20	x ¹⁹	0,60	0,42	0,04	0,48	0,15	-0,11	0,83	0,06	0,03	-0,01	0,67	-0,07	0,04	0,80	0,07	0,91	0,07	0,00	1,00												
21	x ²⁰	0,34	0,70	0,65	-0,01	0,87	0,10	0,26	0,77	0,35	0,03	0,46	0,42	0,15	0,16	0,11	0,20	0,82	0,06	0,22	1,00											
22	x ²¹	0,84	0,60	0,22	0,59	0,33	0,01	0,93	0,26	0,18	0,09	0,72	0,05	0,04	0,69	0,08	0,79	0,12	0,12	0,77	0,28	1,00										
23	x ²²	0,74	0,97	0,57	0,28	0,81	0,15	0,67	0,71	0,54	0,17	0,69	0,45	0,16	0,39	0,11	0,48	0,51	0,20	0,49	0,72	0,69	1,00									
24	x ²³	0,40	0,60	0,98	-0,01	0,76	-0,04	0,24	0,81	0,09	-0,16	0,26	0,26	-0,11	-0,01	-0,19	0,01	0,50	-0,16	0,10	0,67	0,29	0,59	1,00								
25	x ²⁴	0,76	0,59	0,14	0,67	0,32	0,04	0,96	0,20	0,33	0,12	0,83	0,17	0,11	0,82	0,11	0,92	0,06	0,08	0,84	0,23	0,90	0,64	0,18	1,00							
26	x ²⁵	0,53	0,82	0,67	0,10	0,85	-0,02	0,42	0,73	0,43	-0,03	0,59	0,50	-0,07	0,13	-0,11	0,19	0,55	0,03	0,23	0,74	0,49	0,85	0,70	0,37	1,00						
27	x ²⁶	0,05	0,07	-0,20	0,13	0,03	0,81	0,20	-0,01	0,35	0,87	0,00	0,35	0,89	0,11	0,96	0,18	0,14	0,78	0,07	0,10	0,18	0,15	-0,14	0,14	-0,02	1,00					
28	x ²⁷	0,86	0,67	0,21	0,61	0,36	0,10	0,96	0,29	0,30	0,18	0,80	0,16	0,17	0,77	0,17	0,87	0,14	0,17	0,83	0,31	0,93	0,72	0,24	0,96	0,46	0,20	1,00				
29	x ²⁸	0,51	0,85	0,79	0,07	0,98	0,11	0,34	0,91	0,43	0,02	0,52	0,45	0,06	0,06	-0,01	0,11	0,73	0,07	0,16	0,89	0,37	0,84	0,79	0,30	0,87	0,02	0,38	1,00			
30	x ²⁹	0,17	0,48	0,05	0,15	0,45	0,33	0,35	0,18	0,97	0,34	0,47	0,82	0,31	0,31	0,34	0,30	0,32	0,28	0,09	0,32	0,25	0,51	0,06	0,37	0,40	0,36	0,31	0,38	1,00		
31	x ³⁰	0,09	0,06	-0,18	0,18	0,04	0,80	0,25	0,00	0,34	0,86	0,06	0,35	0,88	0,16	0,95	0,22	0,14	0,77	0,11	0,10	0,22	0,14	-0,12	0,20	-0,02	0,98	0,25	0,03	0,36	1,00	
32	jinl	0,72	0,84	0,44	0,41	0,71	0,44	0,78	0,59	0,63	0,47	0,75	0,58	0,49	0,53	0,47	0,62	0,50	0,44	0,55	0,65	0,72	0,67	0,46	0,74	0,67	0,49	0,79	0,71	0,61	0,51	1,00

- e. Pengambilan keputusan untuk menentukan item yang valid digunakan r_{hitung} dibandingkan dengan r_{tabel} dengan dk jumlah sampel dikurangi variabel, yang dalam hal ini pasti 2 (dua),

yaitu item dan total. Manakala $r_{hitung} \geq r_{tabel}$ maka item dikatakan valid, akan tetapi kalau $r_{hitung} < r_{tabel}$ maka item tersebut disimpulkan tidak valid. Berdasarkan r tabel untuk dk 28 dan taraf nyata (α): 0,05, didapatkan skornya $r_{tabel: 0,05;28}=0,374$, maka seluruh item adalah valid.

Selanjutnya akan dibicarakan mengenai uji reliabilitas. Reliabilitas adalah indeks yang menunjukkan sejauh mana suatu alat pengukur dapat dipercaya atau dapat diandalkan. Reliabilitas menunjukkan sejauh mana hasil pengukuran tetap konsisten bila dilakukan pengukuran dua kali atau lebih terhadap gejala yang sama, dengan instrumen yang sama.

Seharusnya setiap instrumen memiliki kemampuan untuk menghasilkan pengukuran yang konsisten. Pada instrumen yang dibuat untuk mengukur fenomena fisik seperti berat atau panjang suatu benda, konsistensi hasil pengukuran bukanlah suatu hal yang sulit untuk diperoleh. Tetapi untuk pengukuran gejala non-fisik seperti sikap, opini, dan persepsi tentang perilaku, bukanlah hal yang mudah. Dalam pengukuran gejala sosial boleh dikatakan hampir tidak pernah terjadi hasil pengukuran ulang yang persis sama dengan hasil pengukuran sebelumnya.

Setidaknya ada 4 (empat) teknik untuk menguji reliabilitas instrumen, yaitu test-retest, belah dua, paralel, dan konsistensi internal.

Test-retest dilakukan dengan cara instrumen diujicobakan sebanyak dua kali. Interval waktu antara ujicoba pertama dan kedua sebaiknya sekitar 15 sampai dengan 30 hari. Apabila interval waktu terlalu pendek dikhawatirkan responden masih ingat terhadap jawaban yang pertama, sehingga responden cenderung hanya mengulangi jawaban pertama, tetapi kalau lebih lama dikhawatirkan sudah terjadi perubahan pada diri responden terkait dengan sesuatu yang diteliti.

Sebagai contoh: instrumen yang telah diujicobakan di atas, yang telah dianalisis validitasnya, diujicobakan untuk kedua kalinya. Setelah instrumen diisi responden dan diberi skor, maka data dari ujicoba yang kedua tersebut dapat ditabulasikan. Untuk mengetahui reliabilitas instrumen tersebut, maka dikorelasikan antara skor total ujicoba pertama dan skor total ujicoba kedua.

Sedangkan data hasil ujicoba kedua ditabulasikan berikut ini..

	A	B	C	D	E	F	G	H	I	J	K	L	M	N	O	P	Q	R	S	T	U	V	W	X	Y	Z	AA	AB	AC	AD	AE
1	x ₁	x ₂	x ₃	x ₄	x ₅	x ₆	x ₇	x ₈	x ₉	x ₁₀	x ₁₁	x ₁₂	x ₁₃	x ₁₄	x ₁₅	x ₁₆	x ₁₇	x ₁₈	x ₁₉	x ₂₀	x ₂₁	x ₂₂	x ₂₃	x ₂₄	x ₂₅	x ₂₆	x ₂₇	x ₂₈	x ₂₉	x ₃₀	jml
2	3	4	5	3	3	4	3	3	4	3	3	4	4	3	3	2	3	4	2	3	2	3	4	3	3	2	3	4	3	2	95
3	4	4	3	4	3	5	4	3	3	5	4	5	5	4	4	3	3	4	4	3	4	3	3	4	3	5	4	3	3	5	113
4	5	4	3	5	4	3	5	3	3	5	4	4	5	3	5	3	5	3	5	4	5	4	4	5	4	5	4	5	4	4	124
5	3	3	3	4	5	4	3	5	3	3	3	4	3	4	3	4	3	4	2	3	5	3	2	5	3	4	3	5	3	4	105
6	3	3	3	2	3	5	2	3	3	5	4	3	5	4	5	4	3	4	4	3	4	3	3	4	3	5	4	3	3	5	108
7	5	4	4	3	4	2	4	4	3	2	5	4	3	5	3	5	4	3	5	3	5	4	4	5	4	3	5	4	4	3	116
8	3	2	3	4	2	2	3	2	3	4	2	2	2	2	2	4	2	3	4	2	3	4	2	3	4	2	2	3	2	2	81
9	5	4	5	4	5	4	4	2	4	4	4	4	3	2	2	3	2	3	4	2	3	5	4	5	4	4	4	5	2	3	112
10	3	2	2	3	3	3	3	2	3	3	3	2	4	3	4	3	3	4	3	2	3	2	2	3	2	4	3	2	2	3	83
11	4	3	3	4	3	3	4	3	3	4	4	3	4	3	5	4	3	5	4	3	4	3	3	4	3	5	4	3	3	5	109
12	5	4	4	5	4	4	5	4	4	4	3	4	4	5	4	5	4	4	5	4	5	2	4	5	3	4	5	4	3	4	124
13	3	2	2	3	2	4	3	3	2	4	3	3	4	3	4	3	2	4	3	2	3	2	3	3	2	4	3	2	2	4	87
14	2	3	5	2	5	3	2	5	3	3	4	5	4	2	3	2	5	3	2	5	2	3	3	2	3	2	5	3	2	3	99
15	1	2	2	4	3	4	3	2	5	4	4	5	4	3	4	4	2	3	3	2	3	2	2	4	2	4	3	2	5	4	95
16	4	3	3	4	3	4	4	3	4	3	4	4	4	3	3	4	3	3	3	3	4	3	3	4	3	4	3	4	3	4	105
17	5	4	4	2	4	5	4	5	5	5	3	3	5	5	3	3	4	3	5	4	5	4	4	5	4	5	4	5	5	5	129
18	4	3	3	4	3	3	4	3	3	3	3	3	3	4	3	4	4	3	4	3	4	3	3	4	3	4	3	4	3	3	100
19	3	2	2	3	2	3	3	2	4	4	3	4	3	3	3	3	3	3	3	2	3	2	2	3	2	4	3	2	4	4	87
20	2	3	3	2	3	5	2	3	5	5	2	5	5	2	5	4	3	5	3	4	2	3	3	2	3	2	3	2	3	5	102
21	3	2	4	3	2	4	3	4	2	4	2	2	4	3	4	3	2	3	3	3	3	2	4	3	2	4	3	2	1	4	88
22	4	3	3	4	3	2	3	3	2	2	4	2	2	1	2	4	3	2	4	3	3	2	3	4	3	2	4	3	2	2	84
23	5	4	4	3	4	3	5	4	3	4	5	3	2	3	5	2	4	5	4	5	4	5	3	4	2	4	5	4	3	3	111
24	5	4	4	5	4	4	2	4	4	4	5	2	4	4	4	5	2	4	5	3	5	4	3	5	4	5	4	4	4	4	120
25	4	3	3	4	3	5	4	3	5	2	4	5	2	4	5	4	3	5	2	3	4	3	3	3	3	5	4	3	5	5	111
26	3	2	2	3	2	3	3	2	3	3	3	3	3	3	3	3	3	4	3	2	3	2	2	3	4	3	2	3	3	3	84
27	3	2	4	3	2	2	3	2	2	2	3	3	2	3	2	3	3	4	3	2	3	2	4	3	2	4	3	2	2	2	80
28	2	3	3	2	3	2	3	3	3	2	3	3	2	3	2	3	2	3	2	4	5	3	2	3	2	3	3	3	3	3	85
29	3	2	2	3	2	4	3	2	4	3	2	4	3	4	3	2	5	3	2	5	3	2	4	3	2	4	3	2	4	3	90
30	4	3	3	4	3	5	4	3	3	5	4	3	5	2	5	4	3	3	4	3	3	3	3	4	3	5	4	3	4	5	110
31	3	2	4	3	2	4	3	2	2	4	3	3	4	3	4	3	2	4	3	2	3	2	4	3	2	4	3	2	2	3	88

Skor koefisien korelasi antara skor total ujcoba pertama dan ujcoba kedua seperti disajikan berikut ini.

	A	B	C	D	E	F
1	uji_1	uji_2				
2	101	95				
3	116	113				
4	127	124				
5	100	105				

	uji_1	uji_2
uji_1	1	
uji_2	0,988	1

Berdasarkan hasil analisis tersebut diketahui bahwa r_{hitung} nya sebesar 0,988. Manakala $r_{hitung} \geq r_{tabel}$ maka item dikatakan reliabel, akan tetapi kalau $r_{hitung} < r_{tabel}$ maka item tersebut disimpulkan tidak reliabel. Berdasarkan r tabel untuk dk 28 dan taraf nyata (α): 0,05, didapatkan skornya $r_{tabel: 0,05;28} = 0,374$, maka instrumen tersebut reliabel atau handal.

Sedangkan teknik kedua adalah belah dua. Bila kita ingin menggunakan teknik belah dua sebagai cara untuk menghitung reliabilitas instrumen, instrumen yang disusun, menurut Ancok, harus memiliki cukup banyak item. Jumlah item sekitar 50-60 adalah jumlah yang cukup memadai. Makin besar jumlah item reliabilitas yang diperoleh akan makin bertambah baik.¹⁷

Langkah yang perlu dilakukan adalah sebagai berikut:

- a. Menyajikan instrumen kepada sejumlah responden, kemudian dihitung validitas itemnya. Item-item yang valid dikumpulkan menjadi satu, yang tidak valid dibuang.
- b. Membagi item-item yang valid tersebut menjadi dua belahan. Untuk membelah instrumen menjadi dua dilakukan dengan cara: pertama, membagi item dengan acak (random). Separuh masuk belahan pertama yang separuh lagi masuk belahan kedua; kedua, membagi item berdasarkan nomor genap-ganjil. Item yang bernomor genap dikelompokkan sebagai belahan pertama, sedangkan yang bernomor ganjil dikelompokkan sebagai belahan kedua.
- c. Skor masing-masing item pada tiap belahan dijumlahkan. Langkah ini akan menghasilkan dua skor total untuk masing-masing responden, yakni skor total belahan pertama dan skor total belahan kedua.
- d. Mengkorelasikan skor total belahan pertama dengan skor total belahan kedua dengan menggunakan teknik korelasi Pearson Product Moment.
- e. Selanjutnya skor koefisien korelasi product moment tersebut dimasukkan dalam rumus berikut:

$$r_{tot} = \frac{2r_{tt}}{1 + r_{tt}}$$

r_{tot} = angka reliabilitas keseluruhan item

¹⁷ Ancok, *Teknik Penyusunan*, hlm. 31-32.

r_{tt} = angka korelasi belahan pertama dan belahan kedua

Misalkan didapatkan angka korelasi belahan pertama dan belahan kedua sebesar 0,69. Selanjutnya angka korelasi tersebut dimasukkan ke dalam rumus di atas, hasilnya adalah:

$$0,82 = \frac{2 \times 0,69}{1 + 0,69}$$

Selanjutnya angka korelasi tersebut dibandingkan dengan r_{tabel} untuk derajat kebebasan $\frac{1}{2}$ jumlah item dikurangi 2. Seandainya item yang dibelah semula berjumlah 50, maka derajat kebebasannya ada pada 23. Pada dk 23 dan taraf nyata (α) 5% dari didapati $r_{tabel: 0,03;23}$: 0,413. Karena $r_{hitung} > r_{tabel}$, maka disimpulkan kalau instrumen tersebut reliabel.

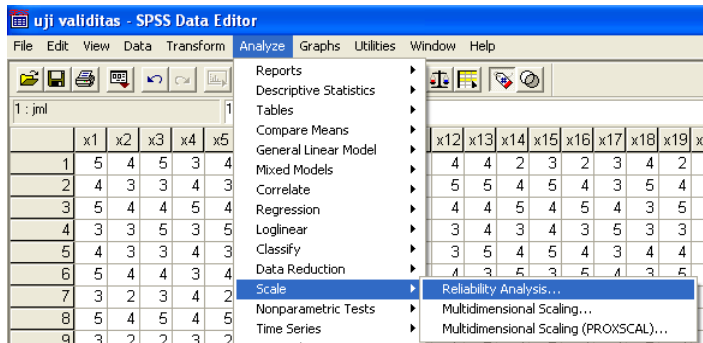
Teknik yang ketiga adalah paralel, yang sering juga disebut teknik ekuivalen. Pada teknik ini, penghitungan reliabilitas dilakukan dengan membuat dua jenis instrumen dengan redaksi yang berbeda tetapi mengukur aspek yang sama. Kedua instrumen tersebut diberikan kepada responden yang sama, kemudian dicari validitasnya untuk masing-masing jenis.

Untuk menghitung reliabilitas perlu mengkorelasikan skor total dari kedua jenis instrumen tersebut. Teknik korelasi yang dipakai ialah teknik korelasi product moment yang rumus serta penghitungannya telah dibahas sebelumnya.

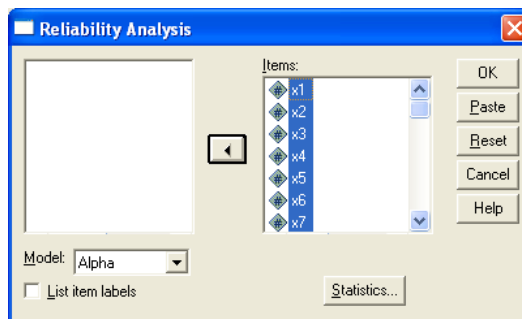
Sedangkan teknik yang keempat adalah teknik uji konsistensi internal. Di antara beberapa uji konsistensi internal, Alpha Cronbach adalah yang paling sering digunakan. Oleh karena itu, pada kesempatan ini akan dicontohkan aplikasi teknik ini baik menggunakan SPSS maupun Microsoft Excel.

Untuk aplikasi SPSS ini sekaligus dapat digunakan untuk menguji validitas data, di samping menguji reliabilitasnya. Sedangkan prosedurnya adalah sebagai berikut.

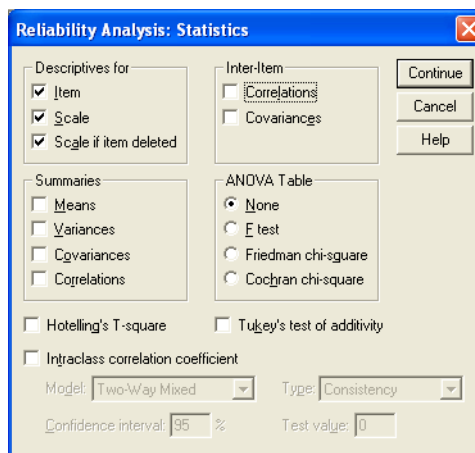
- a. Setelah data diinput, untuk contoh ini digunakan data yang telah diuji validitasnya di atas, lalu klik **Analyze ► Scale ► Reliability Analysis...**, sebagaimana gambar berikut ini.



- b. Setelah keluar gambar seperti berikut ini, destinasikan seluruh skor ke **Item**, dan **Model** masih dibiarkan **Alpha**, lalu klik **Statistics**.



- c. Setelah keluar gambar seperti berikut ini, aktifkan **Item**, **Scale**, dan **Scale if item deleted**, lalu klik **Continue**, lalu klik **Ok**.



d. Berikut ini adalah output analisis **Alpha Cronbach**.

RELIABILITY ANALYSIS - SCALE (ALPHA)				
		Mean	Std Dev	Cases
1.	X1	3,6333	1,0981	30,0
2.	X2	2,9333	,7849	30,0
3.	X3	3,3667	,9643	30,0
4.	X4	3,4333	,8584	30,0
5.	X5	3,1333	,9732	30,0
6.	X6	3,7000	,9154	30,0
7.	X7	3,6000	,9322	30,0
8.	X8	3,2333	1,0063	30,0
9.	X9	3,2333	1,0726	30,0
10.	X10	3,8000	,9248	30,0
11.	X11	3,7000	,9879	30,0
12.	X12	3,4667	1,0417	30,0
13.	X13	3,7667	,9353	30,0
14.	X14	3,4333	,9714	30,0
15.	X15	3,7667	,9353	30,0
16.	X16	3,5667	1,0063	30,0
17.	X17	2,9333	,8683	30,0
18.	X18	3,7000	,9523	30,0
19.	X19	3,4333	1,0400	30,0
20.	X20	2,9333	,8683	30,0
21.	X21	3,6000	1,0372	30,0
22.	X22	2,9000	,7589	30,0
23.	X23	3,3333	,9223	30,0
24.	X24	3,6667	,9223	30,0
25.	X25	3,0667	,8683	30,0
26.	X26	3,7667	,9714	30,0
27.	X27	3,6000	,9322	30,0
28.	X28	3,1000	,9948	30,0
29.	X29	3,2333	1,0400	30,0
30.	X30	3,7333	1,0148	30,0
Statistics for	Mean	Variance	Std Dev	N of
SCALE	102,7667	298,8747	17,2880	Variables 30

Output pertama ini menunjukkan tentang rata-rata skor per-item, dan standard deviasinya.

RELIABILITY ANALYSIS - SCALE (ALPHA)

Item-total Statistics

	Scale Mean if Item Deleted	Scale Variance if Item Deleted	Corrected Item- Total Correlation	Alpha if Item Deleted
X1	99,1333	272,6023	,6913	,9360
X2	99,8333	276,6954	,8258	,9353
X3	99,4000	285,1448	,3930	,9394

X4	99,3333	287,4713	,3665	,9395
X5	99,6333	276,0333	,6770	,9363
X6	99,0667	285,6506	,4003	,9392
X7	99,1667	274,6954	,7544	,9355
X8	99,5333	279,3609	,5500	,9377
X9	99,5333	276,6023	,5920	,9373
X10	98,9667	284,7920	,4238	,9390
X11	99,0667	274,3402	,7199	,9358
X12	99,3000	279,1138	,5366	,9379
X13	99,0000	283,7931	,4509	,9387
X14	99,3333	282,0920	,4854	,9384
X15	99,0000	284,6207	,4240	,9390
X16	99,2000	278,3034	,5825	,9373
X17	99,8333	284,6264	,4606	,9386
X18	99,0667	285,3747	,3914	,9394
X19	99,3333	280,2989	,5024	,9383
X20	99,8333	280,0747	,6209	,9370
X21	99,1667	274,1437	,6887	,9361
X22	99,8667	276,5333	,8624	,9351
X23	99,4333	284,9437	,4201	,9390
X24	99,1000	276,1621	,7132	,9360
X25	99,7000	279,3897	,6453	,9368
X26	99,0000	283,5172	,4406	,9389
X27	99,1667	274,1437	,7730	,9353
X28	99,6667	275,5402	,6766	,9363
X29	99,5333	277,9126	,5733	,9375
X30	99,0333	281,9644	,4660	,9387

Reliability Coefficients

N of Cases = 30,0

N of Items = 30

Alpha = ,9395

Output kedua ini menunjukkan tentang rata-rata skor jumlah seluruh item apabila jumlah item yang ada pada baris tersebut dihapus. Sebagai contoh, bahwa skor 99,1333 yang ada pada baris X1 adalah hasil dari jumlah seluruh skor (3083) dikurangi dengan jumlah skor dari X1 (109) dibagi 30. Pada kolom ketiga menampilkan varians dari skor seluruh item apabila jumlah salah satu item dihilangkan. Sedangkan kolom keempat adalah koefisien korelasi antara item dan skor total yang telah diberikan koreksi dengan rumus.

$$r.pq = \frac{(r.tp)(SDy) - (SDx)}{\sqrt{(SDy)^2 + (SDx)^2 - 2(r.tp)(SDx)(SDy)}}$$

Keterangan:

$r.pq$ = angka korelasi setelah dikoreksi

$r.tp$ = angka korelasi sebelum dikoreksi

SDy = standard deviasi skor total

SDx = standar deviasi item

Skor pada kolom keempat (Corrected item-total corralation) digunakan untuk menguji validitas instrumen. Sebagaimana sudah dijelaskan di atas bahwa pengambilan keputusan untuk menentukan item yang valid digunakan r_{hitung} dibandingkan dengan r_{tabel} dengan dk jumlah sampel dikurangi variabel, yang dalam hal ini pasti 2 (dua), yaitu item dan total. Manakala $r_{hitung} \geq r_{tabel}$ maka item dikatakan valid, akan tetapi kalau $r_{hitung} < r_{tabel}$ maka item tersebut disimpulkan tidak valid. Berdasarkan r_{tabel} untuk dk 28 dan taraf nyata (α): 0,05, didapatkan skornya $r_{tabel: 0,05;28}=0,374$, maka seluruh item adalah valid kecuali item nomor 4. Hanya saja karena jumlah item instrumen berjumlah 30, maka semestinya tidak diperlukan adanya aplikasi koreksi dengan rumus di atas. Padahal skor koefisien korelasi antara item 4 dengan skor jumlah sebelum koreksi adalah 0,4091 yang berarti juga lebih tinggi dibanding 0,374.

Sedangkan kolom kelima adalah skor koefisien alpha apabila skor itemnya dihapus. Apabila masing-masing skor koefisien alpha dibandingkan dengan $r_{tabel: 0,05;28}=0,374$, maka seluruh item adalah reliabel. Kesimpulan dikuatkan dengan koefisien alpha secara keseluruhan sebesar 0,9394 yang lebih besar dibanding $r_{tabel: 0,05;28}=0,374$, maka seluruh item adalah reliabel.

Sedangkan prosedur menguji Alpha Cronbach dengan Microsoft Excel adalah sebagai berikut:

- Setelah data diinput, lalu dihitung jumlah skor item yang diperoleh dari masing-masing responden.
- Menghitung kuadrat jumlah skor item yang diperoleh dari masing-masing responden.
- Menghitung varians masing-masing item.
- Mengaplikasikan rumus Alpha Cronbach, seperti di bawah ini.

$$r_{11} = \left[\frac{k}{k-1} \right] \left[1 - \frac{\sum \sigma_i^2}{\sigma_t^2} \right]$$

di mana

$$\sigma^2 = \frac{\sum x^2 - \frac{(\sum x)^2}{N}}{N}$$

Keterangan

r_{11} = Reliabilitas instrumen

k = Banyaknya item instrumen

$\sum \sigma_i^2$ = Jumlah varians item

σ_t^2 = Varians total

N = Jumlah responden.

Sedangkan aplikasi penjelasan di atas dapat dilihat pada aplikasi Microsoft Excel berikut ini.

	A	AE	AF	AG	AH	AI
1	x_1	jml				
2	5	101	10201			
33	1,17	26,53				
34						
35	25		288,9122		Formula A33	= (A65 - ((A32^2)/30))/30
36	16		0,9395		Formula AE33	= SUM(A33:AD33)
37	25				Formula AF35	= (AF32 - ((AE32^2)/30))/30
38	9				Formula AF36	= (30/29) * (1 - (AE33/AF35))
65	431	11431				
66						

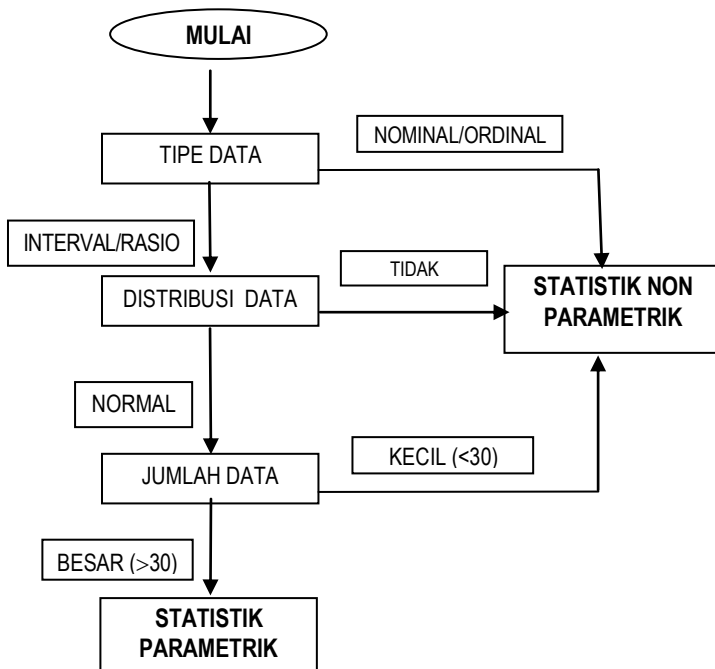
Hasil analisis dengan Microsoft excel juga menghasilkan skor koefisien alpha yang sama dengan hasil penghitungan SPSS yaitu 0,9395.

F. Pedoman Umum Memilih Teknik Statistik

Terdapat bermacam-macam teknik statistik yang dapat digunakan menguji hipotesis. Teknik statistik mana yang akan digunakan tergantung pada interaksi tiga hal: macam data, terpenuhi tidaknya asumsi, dan bentuk hipotesisnya. Sedangkan prosedur penentuan teknik analisis tersebut dapat dilihat pada bagan berikut ini.

BAGAN

PROSEDUR PENENTUAN TEKNIK ANALISIS



Berdasarkan bagan di atas dapat dipahami bahwa data penelitian yang akan dianalisis dengan menggunakan analisis parametrik disyaratkan bertipe interval atau rasio, distribusi datanya normal, dan jumlah sampelnya 30 atau lebih. Data penelitian yang bertipe nominal atau ordinal, teknik analisis yang dapat digunakan adalah analisis non-parametrik. Data interval atau rasio yang tidak berdistribusi normal juga menggunakan analisis non-parametrik. Demikian juga data interval atau rasio yang berdistribusi normal tetapi jumlah sampelnya kurang dari 30 juga menggunakan analisis non-parametrik.

Secara singkat dapat dilihat pada tabel di bawah ini:

MACAM DATA	BENTUK HIPOTESIS					
	Deskriptif (Satu Variabel)	Komparatif (Dua Sampel)		Komparatif (Lebih dari dua sampel)		Asosiatif (hubungan)
		Related	Independen	Related	Independen	
NOMINAL	Binomial	Mc Nemar	Fisher Exact Probability	χ^2 for k Sample	χ^2 for k	Contingency Coefficient

	χ^2 One Sample		χ^2 Two Sample	Cochran Q	Sample	
ORDINAL	Run Test	Sign Test Wilcoxon Matched Pairs	Median Test Mann-Whitney U Test Kolmogorov-Smirnov Wald-Wolfowitz	Friedman Two-way Anova	Median Extention Kruskal-Wallis One Way Anova	Spearman Rank Correlation Kendall Tau
INTERVAL RASIO	t-test	T-test of Related	T-test of independent	One-way Anova Two-way Anova	One-way Anova Two-way Anova	Pearson Product Moment Partial Correlation Multiple Correlation

Buku ini akan menjelaskan secara keseluruhan teknik statistik di atas dengan menggunakan aplikasi SPSS maupun Microsoft Excel.

G. Populasi dan Sampel

Istilah populasi dan sampel tepat digunakan jika penelitian yang dilakukan mengambil sampel sebagai subjek penelitian. Akan tetapi jika sasaran penelitiannya adalah seluruh anggota populasi, akan lebih cocok digunakan istilah subyek penelitian, terutama dalam penelitian eksperimental. Dalam survai, sumber data lazim disebut responden¹⁸ dan dalam penelitian kualitatif disebut informan¹⁹ atau subyek tergantung pada cara pengambilan datanya.

1. Pengertian Populasi dan Sampel

Populasi, menurut Nazir, adalah ”kumpulan dari individu dengan kualitas dan ciri-ciri yang ditetapkan.”²⁰ Sedangkan Surakhmad kelihatannya mendefinisikan populasi sebagai sekelompok subyek, baik manusia, gejala, nilai test, benda-

¹⁸Responden adalah seseorang yang merespons pertanyaan atau pernyataan dari peneliti terkait dengan keadaan diri responden sendiri.

¹⁹Informan adalah seseorang yang dapat memberikan informasi tentang keadaan yang terkait dengan dirinya sendiri dan juga keadaan orang atau hal lain.

²⁰Moh. Nazir, *Metode Penelitian*, (Bogor: Ghalia Indonesia, 2005), hlm. 271.

benda, atau peristiwa yang diberlakukan generalisasi dari sebuah penelitian.²¹ Senada dengan definisi Surakhmad, Fraenkel dan Wallen, sebagaimana dikutip Riyanto, juga mendefinisikan populasi sebagai "kelompok yang menarik peneliti, di mana kelompok tersebut oleh peneliti dijadikan sebagai obyek untuk menggeneralisasikan hasil penelitian."²² Sedangkan Sugiyono mendefinisikan populasi dengan "wilayah generalisasi yang terdiri atas obyek/subyek yang mempunyai kualitas dan karakteristik tertentu yang ditetapkan oleh peneliti untuk dipelajari dan kemudian ditarik kesimpulannya."²³

Berangkat dari empat definisi di atas, setidaknya ada tiga hal yang perlu diperhatikan terkait populasi. Pertama, populasi menjadi wilayah generalisasi dari kesimpulan analisis data sampel; kedua, jumlah populasi, dan ketiga, karakteristik dari populasi.

Sedangkan sampel, menurut Nazir, adalah "bagian dari populasi."²⁴ Surakhmad memberi batasan sampel dengan "bagian dari populasi yang dipandang representatif terhadap populasi."²⁵ Sedangkan Riyanto menganggap bahwa sampel adalah "sembarang himpunan yang merupakan bagian dari suatu populasi."²⁶ Sedangkan Sugiyono membatasi sampel dengan "bagian dari jumlah dan karakteristik yang dimiliki oleh populasi tersebut."²⁷

Penulis berpendapat bahwa sampel tidak hanya bagian dari populasi, tetapi yang terpenting sampel harus mencerminkan karakteristik dari populasi. Oleh karena itu, penjelasan yang akurat tentang karakteristik populasi penelitian perlu diberikan agar besarnya sampel dan cara pengambilannya dapat ditentukan secara tepat. Tujuannya adalah agar sampel yang dipilih benar-

²¹Winarno Surakhmad, *Pengantar Penelitian Ilmiah: dasar, Metode, dan Teknik*, (Bandung: Tarsito, 1990), hlm. 93.

²²Yatim Riyanto, *Metodologi Penelitian Pendidikan*, (Surabaya: SIC, 2001), hlm. 63.

²³Sugiyono, *Statistika untuk Penelitian*, hlm. 72.

²⁴ Nazir, *Metode Penelitian*, hlm. 271.

²⁵ Surakhmad, *Pengantar Penelitian Ilmiah*, hlm. 93.

²⁶ Riyanto, *Metodologi Penelitian Pendidikan*, hlm. 63.

²⁷ Sugiyono, *Statistika untuk Penelitian*, hlm. 73.

benar representatif, dalam arti dapat mencerminkan keadaan populasinya secara cermat. Kerepresentatifan sampel merupakan kriteria terpenting dalam pemilihan sampel dalam kaitannya dengan maksud menggeneralisasikan hasil-hasil penelitian sampel terhadap populasinya. Jika keadaan sampel semakin berbeda dengan karakteristik populasinya, maka semakin besar kemungkinan kekeliruan dalam generalisasinya.

Jadi, hal-hal yang dibahas dalam bagian Populasi dan Sampel adalah (a) identifikasi dan batasan-batasan tentang populasi atau subyek penelitian, (b) besarnya sampel; serta (c) prosedur dan teknik pengambilan sampel.

2. Teknik Penentuan Besarnya Sampel

Selama ini banyak mahasiswa S1 ketika memaparkan tentang sub bab populasi dan sampel hanya menjelaskan tentang jumlah populasi tanpa menjelaskan karakteristiknya. Sedangkan ketika menentukan besarnya sampel mahasiswa sering mengutip pendapat Suharsimi Arikunto berikut ini.

Untuk sekedar ancer-ancer maka apabila subjeknya kurang dari 100, lebih baik diambil semua sehingga penelitiannya merupakan penelitian populasi. Selanjutnya, jika jumlah subjeknya besar dapat diambil antara 10-15% atau 20-25% atau lebih, tergantung setidaknya setidaknya-tidaknya:

- a. Kemampuan peneliti dilihat dari waktu, tenaga, dan dana.
- b. Sempit luasnya wilayah pengamatan dari setiap subjek, karena hal ini menyangkut banyak sedikitnya data.
- c. Besar kecilnya resiko yang ditanggung peneliti. Untuk penelitian yang risikonya besar, tentu saja jika sampel besar, hasilnya akan lebih baik. (Sic!)²⁸

Ketentuan di atas semestinya sudah direvisi oleh Suharsimi Arikunto sendiri pada buku tersebut pada halaman berikutnya,

Penentuan besarnya sampel dengan persentase seperti yang dahulu banyak digunakan tampaknya kini sudah harus ditinggalkan. Agar diperoleh hasil penelitian yang lebih baik, diperlukan sampel yang baik pula, yakni betul-betul mencerminkan populasi. Supaya perolehan sampel

²⁸Suharsimi Arikunto, *Prosedur Penelitian: Suatu Pendekatan Praktek*, (Jakarta: Rineka Cipta, 2002), hlm. 112.

lebih akurat, diperlukan rumus-rumus penentuan besarnya.²⁹

Oleh karena itu, yang perlu diperhatikan ketika menentukan besarnya sampel adalah jumlah populasi, karakteristik populasi, dan tingkat kesalahan yang ditoleransi.

Berikut ini akan diberikan contoh aplikasi dua rumus untuk menentukan besarnya sampel apabila populasinya heterogen. Rumus pertama digunakan apabila jumlah populasi diketahui sedangkan rumus kedua digunakan apabila jumlah populasi tidak diketahui.

Rumus pertama adalah rumus Issac and Michael sebagai berikut:

$$s = \frac{\chi^2 \cdot N \cdot p \cdot q}{d^2 \cdot (N-1) + \chi^2 \cdot p \cdot q}$$

Keterangan:

s : Jumlah sampel

χ^2 : Diambilkan dari χ^2_{tabel} untuk tingkat kesalahan (α) 1%: 6,634891; untuk 5%: 3,841455, dan untuk 10%: 2,705541.

N : Jumlah populasi

p : Jumlah proporsi populasi; misalkan dari 1000 kali pelemparan koin yang jatuh burung sebanyak 597, maka $p = 597/1000$. Akan tetapi kalau proporsi tidak diketahui, maka digunakan angka 0,5.

q : 1 dikurangi nilai proporsi. Seandainya nilai proporsi $597/1000$, maka nilai q adalah $403/1000$.

d : Kesalahan yang ditoleransi.

Apabila rumus di atas diaplikasikan untuk jumlah populasi = 1000, $p = 0,5$, $q = 0,5$ dan kesalahan yang ditoleransi = 0,05, maka caranya sebagai berikut:

$$258 = \frac{3,481 \times 1100 \times 0,5 \times 0,5}{0,05^2 \times (1000 - 1) + 3,481 \times 0,5 \times 0,5}$$

²⁹Suharsimi Arikunto, *Prosedur Penelitian*, hlm. 113.

Untuk mempermudah pembaca ketika menentukan besarnya sampel, berikut ini disajikan tabel yang menyajikan jumlah populasi, jumlah sampel sebagai aplikasi rumus Issac and Michael di atas yang diperbandingkan dengan jumlah sampel menurut Krejcie.

TABEL
JUMLAH SAMPEL

N	S ₁	S ₂
10	10	10
15	14	14
20	19	19
25	23	24
30	28	28
35	32	32
40	36	36
45	40	40
50	44	44
55	48	48
60	51	52
65	55	56
70	58	59
75	62	63
80	65	66
85	68	70
90	72	73
95	75	76
100	78	80
110	84	86
120	89	92
130	95	97
140	100	103
150	105	108
160	110	113
170	114	118
180	119	123
190	123	127
200	127	132
210	131	136

N	S ₁	S ₂
220	135	140
230	139	144
240	142	148
250	146	152
260	149	155
270	152	159
280	155	162
290	158	165
300	161	169
320	167	175
340	172	181
360	177	186
380	182	191
400	186	196
420	191	201
440	195	205
460	198	210
480	202	214
500	205	217
550	213	226
600	221	234
650	227	242
700	233	248
750	238	254
800	243	260
850	247	265
900	251	269
950	255	274
1000	258	278
1100	265	285

N	S ₁	S ₂
1200	270	291
1300	275	297
1400	279	302
1500	283	306
1600	286	310
1700	289	313
1800	292	317
1900	294	320
2000	297	322
2200	301	327
2400	304	331
2600	307	335
2800	310	338
3000	312	341
3500	317	346
4000	320	351
4500	323	354
5000	326	357
6000	329	361
7000	332	364
8000	334	367
9000	335	368
10000	336	370
15000	340	375
20000	342	377
30000	344	379
40000	345	380
50000	346	381
75000	346	382
100000	347	384

Keterangan:

- N** : Jumlah populasi
- S₁** : Jumlah sampel, aplikasi rumus Issac and Michael, untuk tingkat kesalahan (α): 0,05, dan proporsi: 0,5.
- S₂** : Jumlah sampel menurut Krejcie untuk tingkat kesalahan (α) 0,05.

Apabila jumlah populasi tidak diketahui/tidak terbatas, maka dapat digunakan rumus berikut ini:

$$n = \left(\frac{Z}{e}\right)^2 \cdot p \cdot q$$

Keterangan:

- n** : Jumlah sampel
- z** : Diambilkan dari z_{tabel} untuk tingkat kesalahan (α) 1%: 2,58; untuk 5%: 1,96, dan untuk 10%: 1,645.
- e** : Kesalahan yang ditoleransi.
- p** : Jumlah proporsi populasi.
- q** : 1 dikurangi nilai proporsi.

Apabila rumus di atas diaplikasikan untuk tingkat kepercayaan = 95%, $p = 0,5$, $q = 0,5$ dan kesalahan yang ditoleransi = 0,05, maka caranya sebagai berikut:

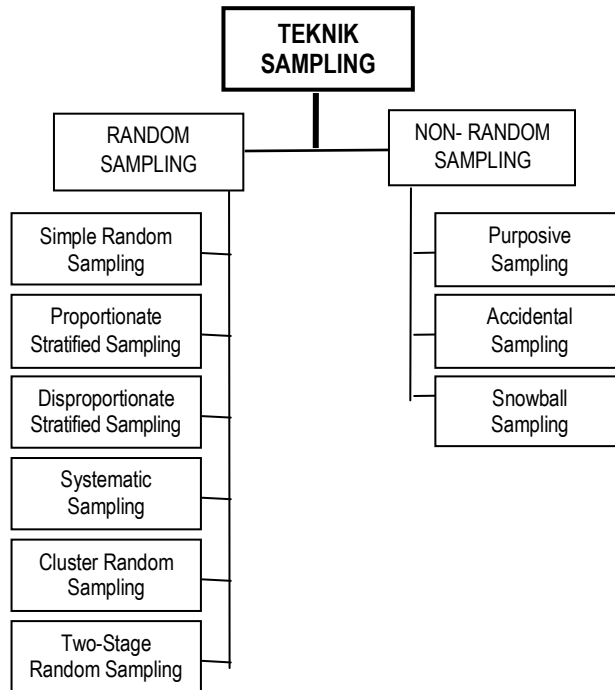
$$384 = \left(\frac{1,96}{0,05}\right)^2 \times 0,5 \times 0,5$$

3. Prosedur dan Teknik Pengambilan Sampel

Di samping itu, ketika jumlah sampel sudah dapat ditentukan, maka yang juga perlu ditentukan adalah teknik sampling dengan mendasarkan pada katakteristik dari populasi tersebut.

Pada dasarnya teknik sampling dapat dikelompokkan menjadi dua, yaitu random sampling dan non-random sampling. Teknik sampling yang dikategorikan sebagai random sampling adalah simple random sampling, proportionate stratified random sampling, disproportionate stratified random sampling, systematic sampling, cluster random sampling, dan two-stage random sampling. Sedangkan yang termasuk non-random

sampling adalah purposive sampling, accidental sampling, dan snowball sampling. Untuk lebih mudahnya dapat dilihat pada bagan berikut ini.



4. Prosedur dan Teknik Pengambilan Sampel

a) Random Sampling

Random sampling, yang sering juga disebut probability sampling, adalah teknik pengambilan sampel yang memberikan peluang yang sama bagi setiap unsur dalam populasi untuk menjadi sampel. Yang termasuk teknik random sampling adalah beberapa teknik sampling berikut ini.

1) Simple Random Sampling

Teknik simple random sampling adalah teknik pengambilan sampel yang memberikan peluang yang sama bagi setiap unsur dalam populasi untuk menjadi sampel. Teknik ini dilakukan bila anggota populasi dianggap homogen.

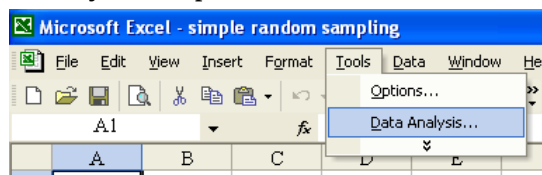
Aplikasi simple random sampling ini dapat dilakukan dengan berbagai cara, misalnya seluruh anggota populasi dicatat dan diambil satu persatu seperti dalam arisan. Selain itu, dapat menggunakan nomor random yang biasanya disertakan dalam buku-buku statistik atau metodologi penelitian. Dengan perkembangan software komputer, aplikasi simple random sampling ini menjadi lebih cepat.

Contoh:

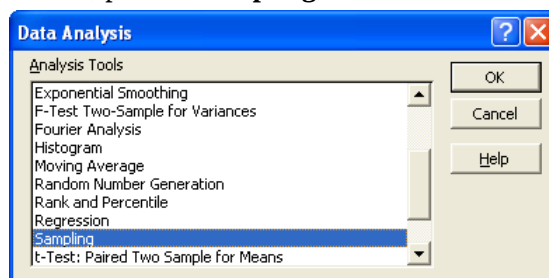
Seorang peneliti berkeinginan mengetahui motivasi belajar siswa kelas I MTs Madrasah Hidayatul Muftadi'in Lirboyo sebanyak 729. Jumlah yang akan diambil untuk dijadikan sampel berdasarkan rumusnya Issac and Michael adalah 236. Dikarenakan populasinya dianggap homogen, maka digunakan simple random sampling.

Caranya:

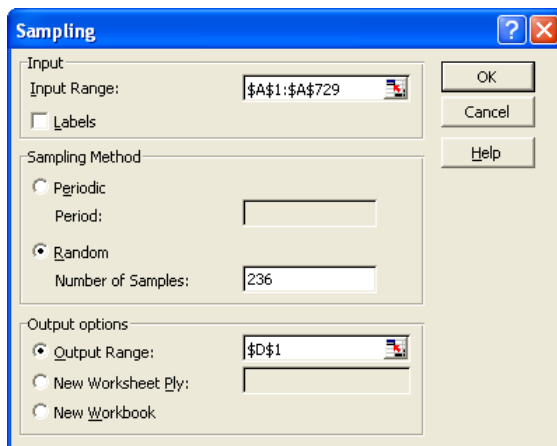
- a. Setelah nomor dan nama siswa diinput, klik **Tools** ➤ **Data Analysis...**, seperti berikut ini.



- b. Setelah itu, pilihlah **sampling**, lalu klik **Ok**.



- c. Setelah itu, klik kotak **Input Range** lalu bloklah seluruh nomor populasi yang berjumlah 729 yang sudah diinput. Lalu ketiklah angka **236 (jumlah sampel)**. Pilihlah tempat untuk output nomor siswa yang akan dijadikan sampel dengan cara memilih salah satu pilihan di **Output options**.



- d. Setelah keluar hasilnya, maka nomor sampel tersebut dapat diurutkan, dengan submenu yang ada di Microsoft Excel, yaitu melalui **Data ► Sort**. Apabila ada nomor yang sama, maka salah satu nomor dapat dihapus dan nomor tersebut dalam daftar populasi juga dihapus lalu diulangi lagi untuk mengambil sampel seperti dijelaskan di atas sampai jumlah sampel yang ditentukan terpenuhi.

Microsoft Excel - simple random sampling

	A	B	C	D	E
1	1	Ahmad		305	
2	2	Sodiq		673	
3	3	Ali		562	
4	4	Faiq		86	
5	5	Muhammad		408	
6	6	Almas		678	
7	7	Anam		258	
		...			

2) Proportionate Stratified Random Sampling

Teknik ini digunakan bila populasi mempunyai anggota/unsur yang berstrata secara proporsional.

Contoh:

Seorang peneliti akan meneliti tentang motivasi berpretasi siswa MTs MHM Lirboyo Kediri. Jumlah populasi adalah 2.055 siswa. Jumlah siswa kelas I

sebanyak 729 (35,47%), kelas II sebanyak 681 (33,14), dan kelas III sebanyak 645 (31,39%).

Berdasarkan aplikasi rumus penentuan besarnya sampel dari Issac and Michael diketahui bahwa manakala jumlah populasinya 2.055, maka didapatkan jumlah sampelnya 298. Dari 298 sampel tersebut, 35,47% diambilkan dari kelas I, yaitu sebanyak 106; 33,14% diambilkan dari kelas II, yaitu sebanyak 99 siswa, dan 31,39% dari kelas III, yaitu sebanyak 94 siswa. Selanjutnya, untuk memilih 106 siswa dari 729 siswa kelas I, 99 siswa dari 681 siswa kelas II, dan 94 siswa dari 645 siswa kelas III dilakukan secara random yang penjelasannya telah diberikan di atas.

3) Disproportionate Stratified Random Sampling

Teknik ini digunakan bila populasi mempunyai anggota/unsur yang berstrata secara tidak proporsional.

Contoh:

Diadakan penelitian tentang frekwensi penulisan karya ilmiah. Di lembaga, tempat penelitian dilakukan, memiliki dosen yang berpendidikan S3 sebanyak 7 orang, S2 sebanyak 87 orang, dan S1 sebanyak 5 orang.

Berdasarkan aplikasi rumus penentuan besarnya sampel dari Issac and Michael diketahui bahwa manakala jumlah populasinya 99, maka didapatkan jumlah sampelnya 77. Dikarenakan jumlah dosen setiap strata tidak proporsional, maka sampel untuk dosen yang berpendidikan strata S3 dan S1 diambil semua, sedangkan sisanya, yaitu sebanyak 65 dosen dipilih secara random dari 87 dosen yang berpendidikan S2.

4) Systematic Sampling

Teknik systematic sampling merupakan pengambilan sampel berdasarkan nomor urut ke k dalam populasi. Pengambilan sampel secara acak hanya dilakukan pada awalnya saja, sedangkan pengambilan sampel kedua dan seterusnya digunakan interval tertentu sebesar k.

Contoh:

Seorang peneliti berkeinginan mengetahui motivasi belajar siswa kelas I MTs Madrasah Hidayatul Mubtadi'in Lirboyo sebanyak 729.

Berdasarkan aplikasi rumus penentuan besarnya sampel dari Issac and Michael diketahui bahwa manakala jumlah populasinya 729, maka didapatkan jumlah sampelnya 236. Untuk mendapatkan sampel sebanyak 236 dari 729 dapat digunakan kelipatan 3. Hanya saja nomor pertama yang harus digunakan apakah 1, 2, atau 3 ditentukan secara acak.

Teknik systematic sampling ini dianggap lebih mudah dibandingkan dengan simple random sekalipun. Oleh karena itu, teknik ini sering dijadikan alternatif dari teknik simple random sampling.

5) Cluster Random Sampling

Teknik ini digunakan apabila kerangka sampling yang memuat elemen populasi tidak tersedia. Cluster Random Sampling yaitu teknik sampling yang menggunakan kumpulan atau kelompok (cluster) elemen populasi sebagai dasar penarikan sampel.

Contoh:

Diadakan penelitian tentang pengaruh sertifikasi terhadap produktifitas kerja guru.

Karena peneliti tidak memiliki data tentang jumlah guru yang sudah mendapat sertifikat professional pada setiap sekolah, maka peneliti menentukan sampel berdasarkan sekolahnya.

6) Two-Stage Random Sampling

Teknik ini merupakan penggabungan dari dua teknik sampling.

Contoh:

Peneliti menggabungkan teknik Cluster Random Sampling dan Simple Random Sampling.

Cara yang ditempuh adalah menentukan cluster yang dijadikan sampel secara random, setelah itu elemen dari

cluster yang akan dijadikan sampel juga ditentukan secara random.

Untuk mendapatkan sampel yang representatif, peneliti diperkenankan menggabungkan berbagai teknik sampling sepanjang teknik tersebut memungkinkan didapatkannya sampel yang sesuai dengan karakteristik populasi.

b) Non-Random Sampling

Non-random sampling, yang sering juga disebut non-probability sampling, adalah teknik pengambilan sampel yang tidak memberikan peluang yang sama bagi setiap unsur dalam populasi untuk menjadi sampel. Yang termasuk teknik non-random sampling adalah beberapa teknik sampling berikut ini.

1) Purposive Sampling

Purposive sampling, dikenal juga dengan judgement sampling, adalah teknik penarikan sampel yang didasarkan pada tujuan penelitian. Berdasarkan tujuan penelitian tersebut, peneliti menentukan kriteria sampel yang akan diambilnya.

Contoh:

Dilakukan penelitian tentang kualitas karya ilmiah guru yang akan mengajukan kenaikan pangkat ke IV/b, maka sampel sumber datanya adalah dosen dan peneliti yang produktif menulis karya ilmiah. Teknik sampling ini lebih cocok digunakan dalam penelitian kualitatif.

2) Accidental Sampling

Teknik ini sering juga disebut convenience sampling. Peneliti yang menggunakan teknik ini akan menjadikan sampel siapa saja yang ketemu dengannya apabila orang tersebut dianggap cocok sebagai sumber data. Teknik ini digunakan apabila peneliti tidak memungkinkan untuk mengidentifikasi karakteristik elemen populasi secara jelas. Data yang didapatkan dengan teknik ini mempunyai kelemahan untuk digunakan generalisasi terhadap populasi.

3) Snowball Sampling

Snowball sampling adalah teknik penentuan sampel secara berantai. Mula-mula peneliti memilih satu dua

orang, dan setelah mendapatkan data dari orang tersebut, peneliti meminta informasi kepada siapa lagi harus menggali data untuk mendapatkan data yang lebih spesifik dan detail. Teknik ini banyak digunakan dalam penelitian kualitatif.

Berdasarkan penjelasan di atas, ketika akan memilih sampel dari populasi, peneliti harus memperhatikan bahwa teknik sampling yang dipilih memungkinkan dapat menghasilkan data yang dapat digunakan membuat generalisasi, dapat menghasilkan presisi yang tinggi, sederhana, mudah dilaksanakan, dan memberikan sevalid mungkin data tentang populasi dengan biaya, tenaga, dan waktu yang minimal.

H. Konsep Pengujian Hipotesis

Pada dasarnya menguji hipotesis itu adalah menaksir parameter populasi berdasarkan data sampel. Terdapat dua cara menaksir, *a point estimate* dan *interval estimate*. Yang disebutkan pertama adalah suatu taksiran parameter berdasarkan suatu nilai data sampel sedangkan yang kedua adalah suatu taksiran parameter populasi berdasarkan nilai interval data sampel.³⁰

Menaksir parameter populasi menggunakan nilai tunggal mempunyai resiko kesalahan yang lebih tinggi dibandingkan dengan menggunakan *interval estimate*. Makin besar taksiran maka akan semakin kecil kesalahannya tetapi tingkat ketelitian taksiran semakin rendah. Untuk selanjutnya kesalahan taksiran ini dinyatakan dalam peluang yang berbentuk prosentase. Biasanya dalam penelitian kesalahan taksiran ditetapkan lebih dahulu, yang digunakan adalah 5% dan 1%. Tingkat kesalahan ini selanjutnya dinamakan tingkat signifikansi.³¹

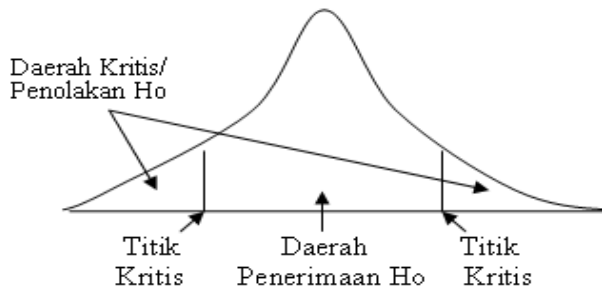
Dalam setiap pengujian hipotesis, kita harus selalu memutuskan apakah menerima ataupun menolak H_0 dan selalu ada kemungkinan bahwa kita membuat kesalahan dalam pengambilan keputusan tersebut. Kesalahan tersebut terjadi ketika kita menolak suatu hipotesis yang benar atau menerima hipotesis yang salah. Kedua jenis kesalahan ini diberi nama secara khusus dalam pengujian hipotesis:

³⁰Kusnandar, *Metode Statistik dan Aplikasinya*, hlm. 86-88.

³¹Kusnandar, *Metode Statistik dan Aplikasinya*, hlm. 86-88.

1. Salah jenis I. Kesalahan ini terjadi ketika kita menolak H_0 padahal H_0 benar. Peluang terjadinya kesalahan ini dinyatakan dengan α dan disebut taraf signifikansi.
2. Salah tipe II. Kesalahan ini terjadi ketika kita menerima H_0 padahal H_0 salah. Peluang terjadinya kesalahan ini dinyatakan dengan β , yang disebut dengan *power of statistical test*.³²

Dalam pengujian hipotesis kebanyakan digunakan kesalahan tipe I yaitu berapa persen kesalahan untuk menolak hipotesis nol (H_0) yang benar. Sebagai ilustrasi dapat dilihat pada gambar di bawah ini:



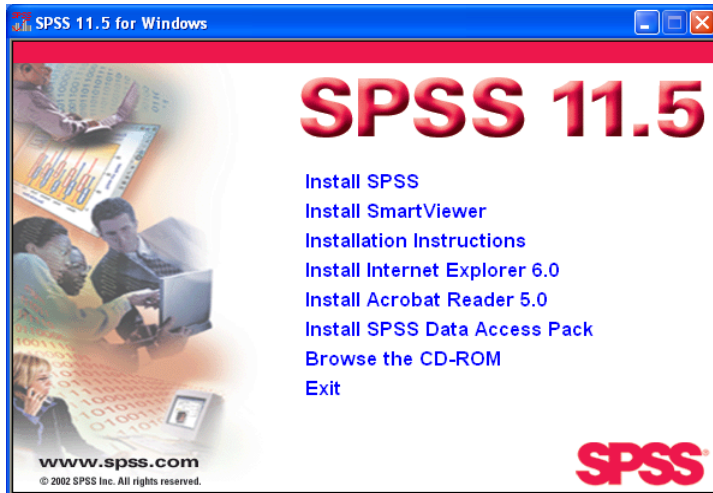
Sebagaimana dijelaskan bahwa software yang akan digunakan untuk menganalisis data penelitian dalam buku ini adalah SPSS dan Microsoft Excel. Oleh karena itu, dalam sub bab berikut ini akan dijelaskan urutan-urutan dalam menginstall SPSS dengan menggunakan versi 11.5 dan bagaimana cara menginput data. Setelah itu akan dijelaskan juga bagaimana mengaktifkan sub menu Data Analysis pada Program Microsoft Excel.

³² Dadan Kusnandar, *Metode Statistik dan Aplikasinya*, hlm. 150.

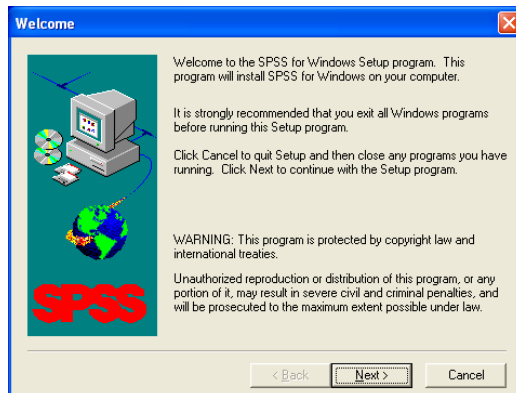
I. Urut-Urutan dalam Menginstall Program SPSS Versi 11.5 dan Proses Mengaktifkan Menu Data Analysis Microsoft Excel

Proses menginstall SPSS versi 11.5 adalah sebagai berikut:

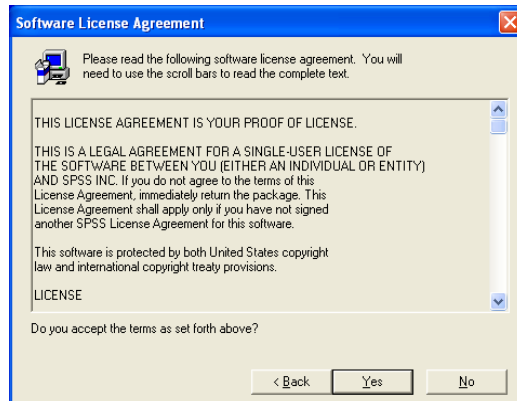
1. Masukkan CD SPSS pada CD ROM, maka secara otomatis akan keluar sebagaimana gambar di bawah ini. **Klik Install SPSS** untuk menginstallnya.



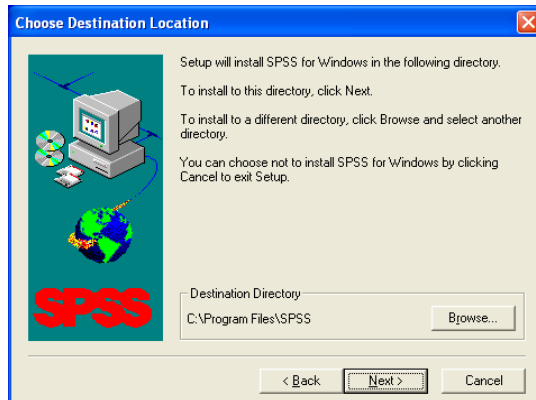
2. Setelah keluar gambar seperti di bawah ini **klik Next**.



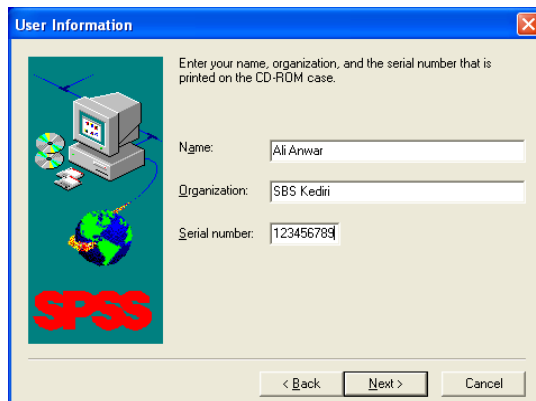
3. Setelah keluar gambar seperti di bawah ini **klik Yes**.



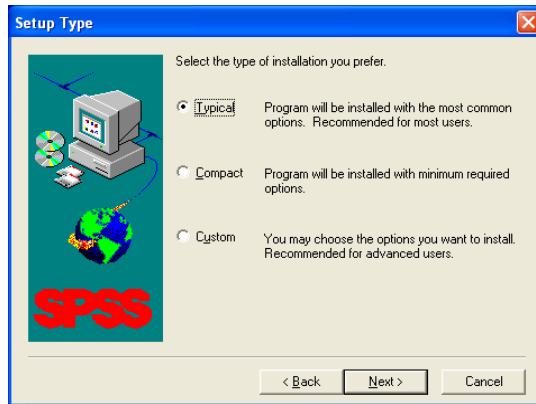
4. Setelah keluar gambar seperti di bawah ini **klik Next**.



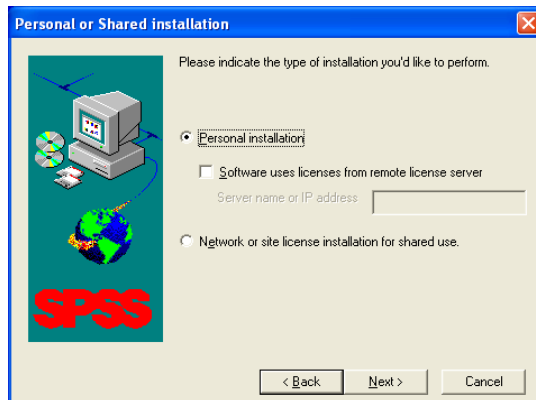
5. Setelah keluar gambar seperti di bawah ini ketikkan **Serial Number** dari program tersebut lalu klik **Next**.



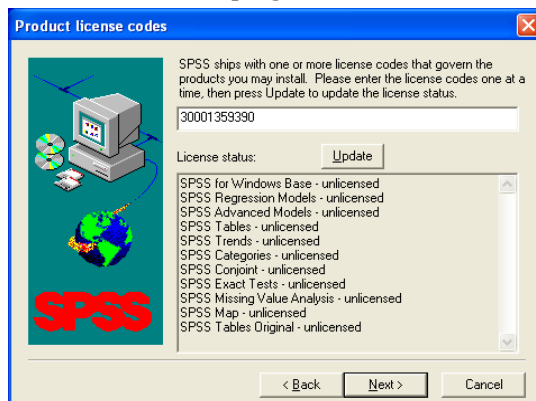
6. Setelah keluar gambar seperti di bawah ini **klik Next**.



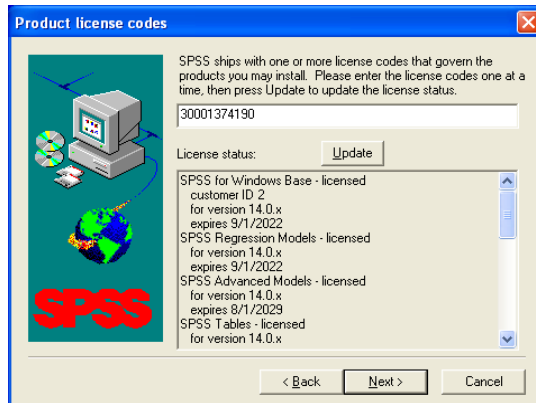
7. Setelah keluar gambar seperti di bawah ini **klik Next**.



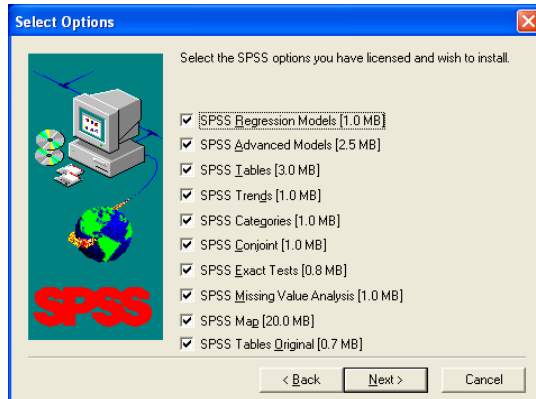
8. Setelah keluar gambar seperti di bawah ini ketik **Lisensi I**: yang biasanya disertakan dalam CD program tersebut, lalu **Klik Update**.



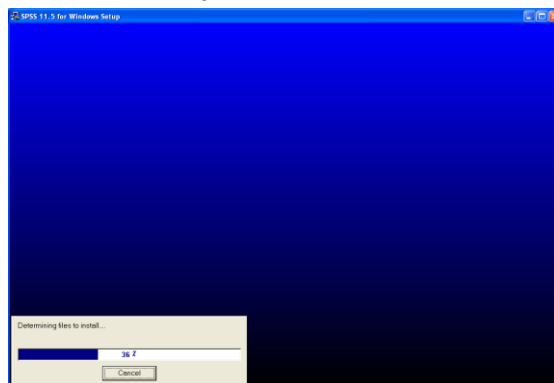
9. Setelah keluar gambar seperti di bawah ini ketik **Lisensi II**, lalu Klik **Update** selanjutnya klik **Next**.



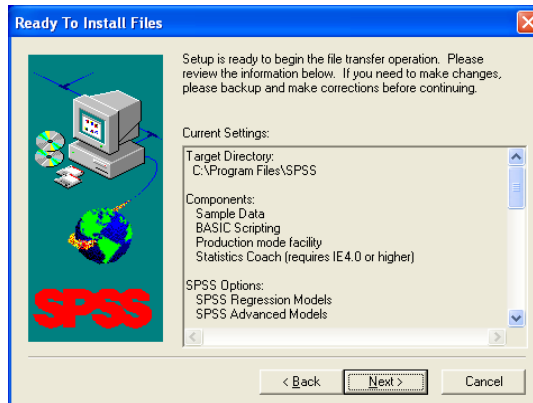
10. Setelah keluar gambar seperti di bawah ini **klik Next**.



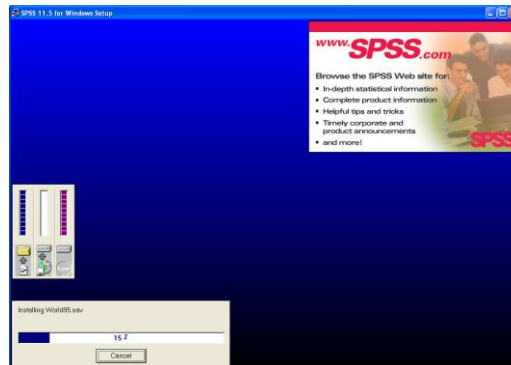
11. Setelah keluar gambar seperti di bawah ini **tunggu sampai keluar gambar berikutnya**.



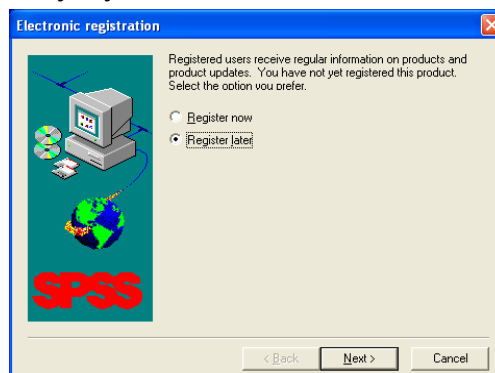
12. Setelah keluar gambar seperti di bawah ini **klik Next**.



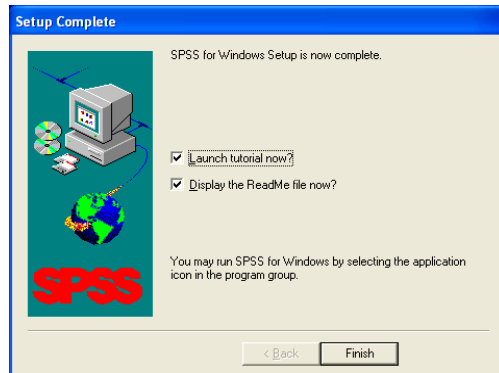
13. Setelah keluar gambar seperti di bawah ini berarti proses instalasi sedang berlangsung. Tunggu sampai keluar gambar berikutnya.



14. Setelah keluar gambar seperti di bawah ini klik **Next** kalau ingin mendaftarkan program tersebut kepada **SPSS Programmer**, akan tetapi kalau tidak diperlukan pendaftaran tersebut klik **Register Later** dan selanjutnya klik Next.

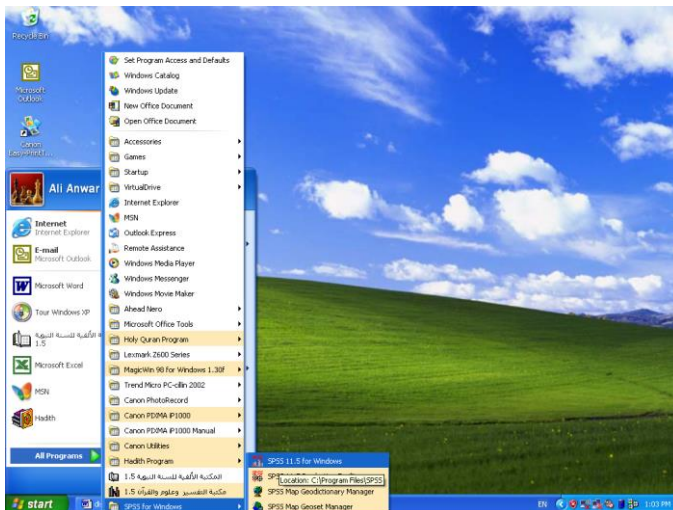


15. Setelah keluar gambar seperti di bawah ini klik **Launch ...** dan klik juga **Display...** untuk menghilangkan tanda contengan dan selanjutnya Finish berarti instalasi selesai.

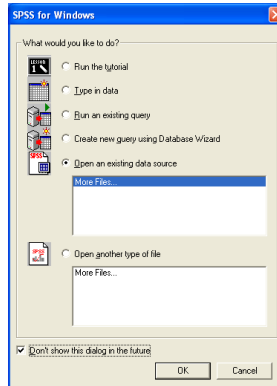


Sedangkan untuk aplikasi program SPSS versi 11.5 dan cara menginput datanya adalah sebagai berikut:

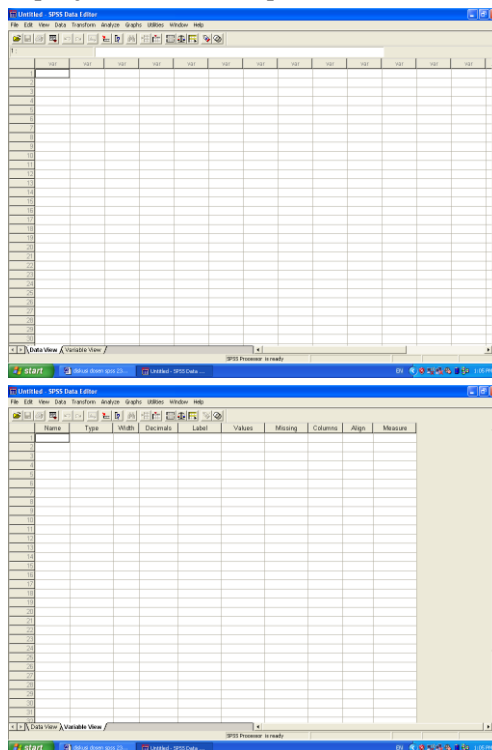
1. Klik Start, lalu
2. seret ke SPSS for windows, lalu
3. seret ke SPSS 11.5 for windows, lalu klik; sebagaimana gambar di bawah ini.



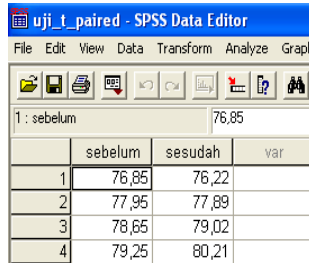
4. Setelah keluar gambar seperti berikut ikon SPSS yang ada di depan **Open...**, kalau anda akan membuka data SPSS yang telah dibuat, tetapi kalau tidak klik **Cancel**.



5. Setelah keluar gambar seperti berikut berarti lembar SPSS siap untuk digunakan mengetik data penelitian. Dalam lembar ini ada dua sheet: yaitu Data View untuk pengetikan Data, dan Variable View untuk untuk memberi penjelasan variabel penelitian.

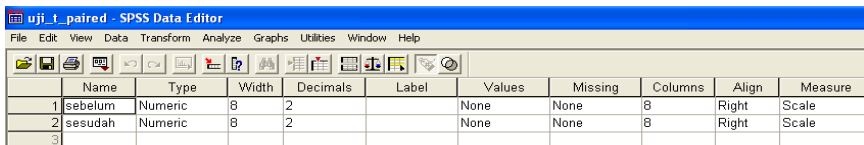


6. Misalkan kita akan memasukkan data penelitian tentang produktifitas kerja pegawai sebelum dan setelah mendapatkan kendaraan dinas seperti di bawah ini.



	sebelum	sesudah	var
1	76,85	76,22	
2	77,95	77,89	
3	78,65	79,02	
4	79,25	80,21	

- Setelah data dimasukkan pada data view, maka klik **variable view** untuk memberikan penjelasan tentang berbagai variabel itu.

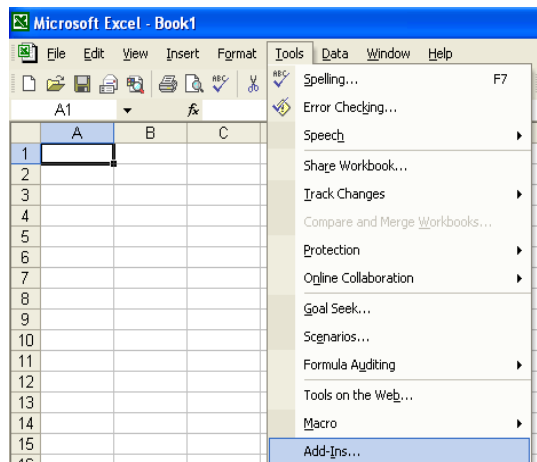


	Name	Type	Width	Decimals	Label	Values	Missing	Columns	Align	Measure
1	sebelum	Numeric	8	2		None	None	8	Right	Scale
2	sesudah	Numeric	8	2		None	None	8	Right	Scale

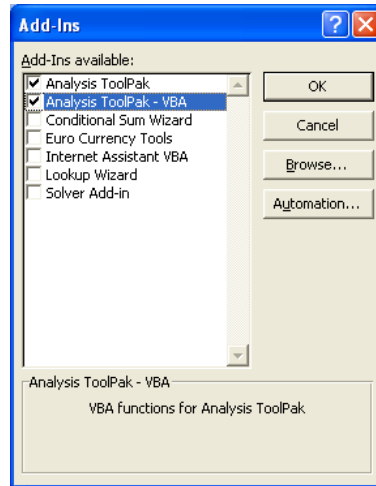
- Setelah selesai, lalu simpan data itu dan siap untuk dianalisis.

Di samping menggunakan SPSS, buku ini juga berusaha menggunakan Microsoft Excel untuk memberikan contoh aplikasi dari analisis statistik yang dipaparkan. Aplikasi Microsoft Excel sedapat mungkin diaplikasikan dengan tiga cara yaitu melalui Data Analysis, aplikasi function, maupun aplikasi rumus-rumus statistik dengan berbagai prosedurnya. Dari tiga cara tersebut, sub menu Data Analysis belum terinstall ketika program Microsoft Excel itu diinstall. Cara mengaktifkan program tersebut adalah sebagai berikut:

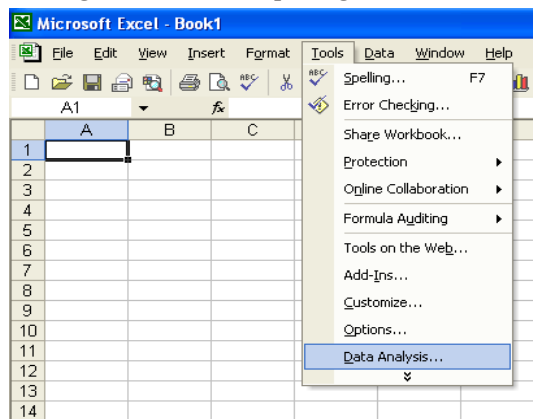
- Klik **Tools** ► **Add-Ins...**, sebagaimana gambar berikut ini.



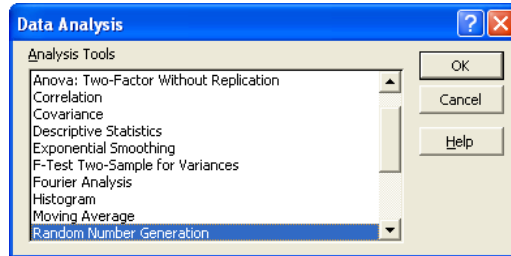
- Setelah keluar gambar berikut ini klik pada kotak **Analysis ToolPak** dan **Analysis ToolPak – VBA** lalu klik **Ok**, sebagaimana gambar berikut ini.



- Setelah selesai, maka akan keluar sub menu **Data Analysis** pada menu **Tools**, sebagaimana terlihat pada gambar di bawah ini.



- Ketika sub menu **Data Analysis** itu diklik, maka akan keluar gambar yang berisi berbagai teknik analisis, sebagaimana terlihat pada gambar berikut ini.



Aplikasi dari berbagai teknik analisis yang menggunakan cara di atas, maupun menggunakan function, dan breakdown dari rumus-rumus statistik akan diberikan ketika menjelaskan berbagai teknik analisis pada bab-bab mendatang.

BAB II

STATISTIK DESKRIPTIF

A. Pendahuluan

Statistik deskriptif, menurut Kusnandar, adalah cabang ilmu statistik yang berkaitan dengan prosedur-prosedur yang digunakan untuk menjelaskan karakteristik data secara umum.¹ Sedangkan Sugiyono mendefinisikannya sebagai statistik yang berfungsi untuk mendeskripsikan atau memberi gambaran terhadap obyek yang diteliti melalui data sampel atau populasi sebagaimana adanya, tanpa melakukan analisis dan membuat kesimpulan yang berlaku untuk umum.² Dari dua definisi di atas dapat dipahami bahwa tujuan utama dari statistik deskriptif adalah menggambarkan data, baik dengan tabel, grafik, maupun ringkasan data. Berangkat dari kesimpulan tersebut, maka berlaku prinsip dasar dalam penyajian data yaitu komunikatif dan lengkap, dalam arti data yang disajikan dapat menarik perhatian pembaca dan mudah dipahami isinya.

Sebagaimana dijelaskan, bahwa ada tiga cara untuk mendeskripsikan data, yaitu:

1. Tabel, yaitu penyajian data yang bertujuan untuk mengelompokkan data ke dalam beberapa kelompok yang masing-masing mempunyai karakteristik yang sama. Bentuk tabel yang sering digunakan adalah tabel distribusi frekuensi dan tabel distribusi frekuensi relatif.

¹Dadan Kusnandar, *Metode Statistik dan Aplikasinya dengan Minitab dan Excel*, (Yogyakarta: Madyan Press, 2004), hlm. 10.

²Untuk elaborasi lebih jauh baca Sugiyono, *Statistika untuk Penelitian*, (Bandung: CV Alfabeta, 2003), hlm. 21.

2. Grafik atau diagram, yaitu penyajian data yang bertujuan memvisualisasikan data secara keseluruhan dengan menonjolkan karakteristik tertentu dari data tersebut.
3. Ringkasan Data Statistik, yaitu penyajian data yang digunakan untuk menjelaskan pemusatan dan penyebaran data.

Sebelum *direlease* berbagai software komputer untuk aplikasi statistik, mendeskripsikan data adalah pekerjaan yang melelahkan dan menjemukan. Saat ini, berbagai pekerjaan itu dapat dikerjakan dengan lebih mudah dan menghasilkan output yang lebih baik.

Contoh:

Peneliti yang sedang mengadakan penelitian di salah satu sebuah Perguruan Tinggi Agama Islam dan mendapatkan data penelitian tentang tenaga edukatif sebagai berikut.

DATA TENAGA EDUKATIF SEBUAH PERGURUAN TINGGI AGAMA ISLAM

NO	KEPANGKATAN	JABATAN FUNGSIONAN	MASA KERJA	IJAZAH	JENIS KELAMIN	UMUR
1	IV/c	Lektor	35	S2	Laki-laki	62
2	IV/c	Lektor Kepala	30	S2	Laki-laki	61
3	IV/b	Lektor	31	S2	Laki-laki	56
4	IV/b	Lektor Kepala	26	S2	Laki-laki	55
5	IV/b	Lektor	30	S2	Laki-laki	56
6	IV/b	Lektor Kepala	22	S2	Laki-laki	58
7	IV/b	Lektor Kepala	27	S1	Laki-laki	60
8	IV/b	Guru Besar	20	S3	Laki-laki	54
9	IV/b	Lektor Kepala	14	S2	Laki-laki	45
10	IV/a	Lektor Kepala	20	S2	Perempuan	47
11	IV/a	Lektor Kepala	18	S2	Laki-laki	46
12	IV/a	Lektor	28	S2	Laki-laki	55
13	IV/a	Lektor Kepala	16	S2	Laki-laki	44
14	IV/a	Lektor Kepala	15	S3	Laki-laki	41
15	IV/a	Lektor Kepala	33	S2	Laki-laki	60
16	IV/a	Lektor Kepala	14	S2	Laki-laki	44
17	IV/a	Lektor Kepala	13	S2	Perempuan	37
18	IV/a	Lektor Kepala	15	S2	Laki-laki	48
19	IV/a	Lektor Kepala	13	S2	Laki-laki	48
20	III/d	lektor	16	S1	Laki-laki	41
21	III/d	lektor	10	S2	Laki-laki	39
22	III/d	lektor	9	S2	Perempuan	45
23	III/d	lektor	9	S2	Perempuan	38
24	III/d	lektor	9	S2	Laki-laki	43
25	III/d	lektor	9	S2	Perempuan	39

26	III/d	lektor	9	S2	Laki-laki	32
27	III/d	lektor	8	S2	Laki-laki	34
28	III/d	lektor	9	S2	Laki-laki	36
29	III/d	lektor	8	S2	Perempuan	35
30	III/d	lektor	7	S2	Laki-laki	38
31	III/d	Lektor Kepala	11	S2	Laki-laki	43
32	III/c	lektor	8	S2	Laki-laki	38
33	III/c	lektor	9	S2	Laki-laki	38
34	III/c	lektor	8	S2	Laki-laki	35
35	III/c	lektor	8	S2	Laki-laki	34
36	III/c	lektor	8	S2	Laki-laki	33
37	III/c	lektor	7	S2	Laki-laki	43
38	III/c	lektor	8	S2	Perempuan	35
39	III/c	lektor	7	S2	Laki-laki	34
40	III/c	lektor	7	S2	Laki-laki	33
41	III/c	lektor	6	S2	Laki-laki	40
42	III/c	lektor	12	S2	Laki-laki	43
43	III/c	lektor	7	S2	Laki-laki	35
44	III/c	lektor	7	S1	Laki-laki	35
45	III/b	Asisten Ahli	10	S2	Laki-laki	34
46	III/b	Asisten Ahli	4	S2	Laki-laki	33
47	III/b	Asisten Ahli	4	S2	Laki-laki	33
48	III/b	Asisten Ahli	4	S2	Laki-laki	32
49	III/b	Asisten Ahli	4	S2	Laki-laki	31
50	III/b	Asisten Ahli	4	S2	Laki-laki	38
51	III/b	Asisten Ahli	4	S2	Laki-laki	37
52	III/b	Asisten Ahli	4	S2	Perempuan	34
53	III/b	Asisten Ahli	4	S2	Perempuan	33
54	III/b	Asisten Ahli	3	S2	Laki-laki	34
55	III/b	Asisten Ahli	7	S2	Laki-laki	42
56	III/b	Asisten Ahli	3	S2	Laki-laki	36
57	III/b	Asisten Ahli	3	S2	Laki-laki	35
58	III/b	Asisten Ahli	3	S2	Perempuan	34
59	III/b	Asisten Ahli	3	S2	Perempuan	34
60	III/b	Asisten Ahli	3	S2	Perempuan	30
61	III/b	Asisten Ahli	3	S2	Perempuan	29
62	III/b	Asisten Ahli	3	S2	Perempuan	30
63	III/b	Asisten Ahli	3	S2	Laki-laki	29
64	III/b	Asisten Ahli	7	S2	Laki-laki	42
65	III/b	Asisten Ahli	8	S2	Perempuan	41
66	III/b	Asisten Ahli	6	S2	Laki-laki	38
67	III/b	calon dosen	6	S2	Laki-laki	41
68	III/b	calon dosen	1	S2	Perempuan	35
69	III/b	calon dosen	1	S2	Laki-laki	31
70	III/b	calon dosen	6	S2	Laki-laki	39
71	III/b	Asisten Ahli	4	S2	Laki-laki	30

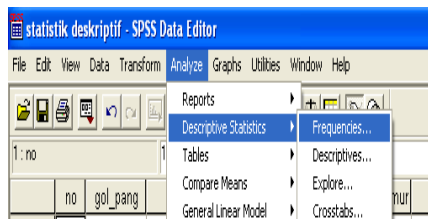
72	III/b	Asisten Ahli	7	S2	Laki-laki	36
73	III/a	Asisten Ahli	7	S1	Perempuan	33
74	III/a	Asisten Ahli	7	S2	Perempuan	32
75	III/a	calon dosen	4	S1	Perempuan	30
76	III/a	calon dosen	4	S2	Perempuan	27
77	III/a	calon dosen	3	S1	Perempuan	27
78	III/a	calon dosen	1	S2	Laki-laki	36
79	III/a	calon dosen	1	S2	Perempuan	35
80	III/a	calon dosen	1	S2	Perempuan	34
81	III/a	calon dosen	1	S1	Laki-laki	35

Agar data di atas mudah untuk dipahami, maka peneliti dapat menyajikannya dalam bentuk tabel maupun grafik, dan untuk data kuantitatif dapat dideskripsikan dengan ringkasan data statistik. Berikut ini akan disajikan aplikasinya dengan SPSS dan Microsoft Excel.

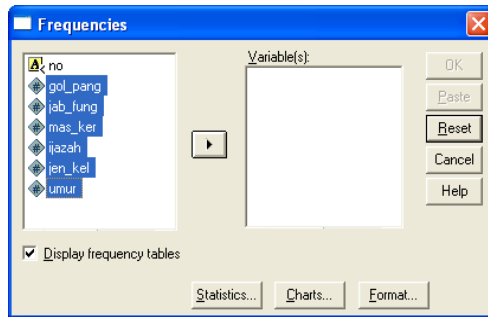
B. Aplikasi untuk Penyajian Data dengan SPSS

Untuk menampilkan data di atas dalam bentuk tabel distribusi frekuensi, distribusi frekuensi kumulatif, dan grafik dengan SPSS dapat dilakukan dengan cara sebagai berikut:

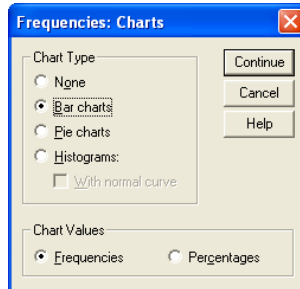
1. Setelah data diinput pada lembar SPSS Data Editor, maka klik **Analyze ► Descriptive Statistics ► Frequencies**, sebagaimana gambar di bawah ini.



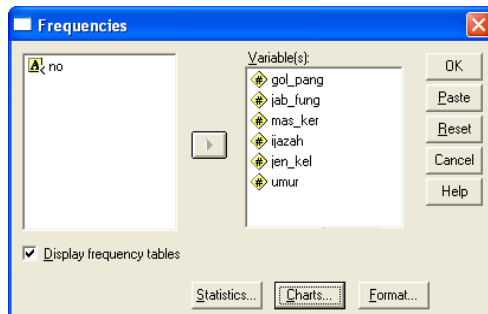
2. Ketika keluar gambar berikut, bloklah variabel-variabel yang akan disajikan lalu masukkan pada kolom **Variable(s)** dengan mengeklik tanda gambar ►. Apabila data tersebut juga ingin ditampilkan dalam bentuk grafik, maka klik **Charts**.



3. Ketika keluar gambar berikut, Pilihlah salah satu dari tiga pilihan grafik, yaitu grafik batang (Bar Charts), grafik serabi (Pie charts), atau histogram. Setelah dipilih, klik Continue.



4. Ketika keluar gambar berikut, klik OK.



5. Berikut ini adalah output dari aplikasi di atas.

Frequencies

Statistics

		GOL PANG	JAB FUNG	MAS KER	IJAZAH	JEN KEL	UMUR
N	Valid	81	81	81	81	81	81
	Missing	0	0	0	0	0	0

Frequency Table

GOL_PANG

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	III/a	9	11,1	11,1	11,1
	III/b	28	34,6	34,6	45,7
	III/c	13	16,0	16,0	61,7
	III/d	12	14,8	14,8	76,5
	IV/a	10	12,3	12,3	88,9
	IV/b	7	8,6	8,6	97,5
	IV/c	2	2,5	2,5	100,0
	Total	81	100,0	100,0	

JAB_FUNG

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	calon dosen	11	13,6	13,6	13,6
	Asisten Ahli	26	32,1	32,1	45,7
	lektor	28	34,6	34,6	80,2
	Lektor Kepala	15	18,5	18,5	98,8
	Guru Besar	1	1,2	1,2	100,0
	Total	81	100,0	100,0	

MAS_KER

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	1,00	6	7,4	7,4	7,4
	3,00	10	12,3	12,3	19,8
	4,00	11	13,6	13,6	33,3
	6,00	4	4,9	4,9	38,3
	7,00	11	13,6	13,6	51,9
	8,00	8	9,9	9,9	61,7
	9,00	7	8,6	8,6	70,4
	10,00	2	2,5	2,5	72,8
	11,00	1	1,2	1,2	74,1
	12,00	1	1,2	1,2	75,3
	13,00	2	2,5	2,5	77,8
	14,00	2	2,5	2,5	80,2
	15,00	2	2,5	2,5	82,7
	16,00	2	2,5	2,5	85,2
	18,00	1	1,2	1,2	86,4
	20,00	2	2,5	2,5	88,9
	22,00	1	1,2	1,2	90,1
	26,00	1	1,2	1,2	91,4
	27,00	1	1,2	1,2	92,6
	28,00	1	1,2	1,2	93,8
	30,00	2	2,5	2,5	96,3
	31,00	1	1,2	1,2	97,5
	33,00	1	1,2	1,2	98,8
	35,00	1	1,2	1,2	100,0
	Total	81	100,0	100,0	

IJAZAH

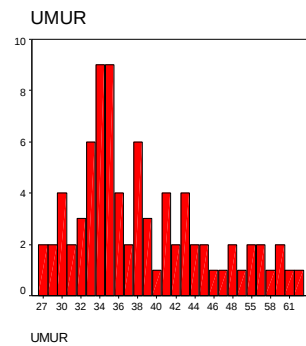
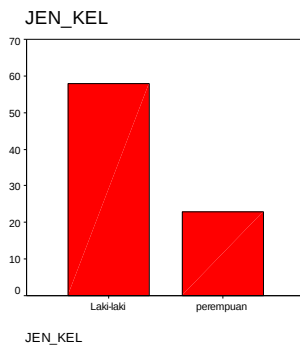
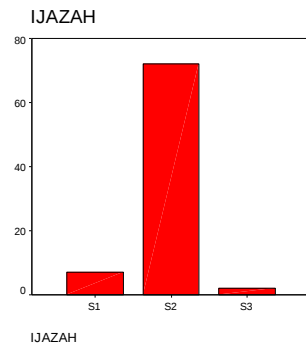
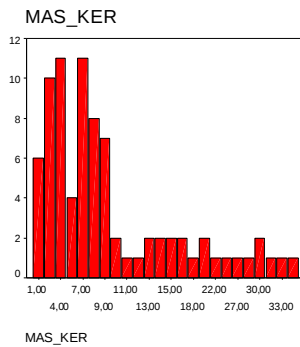
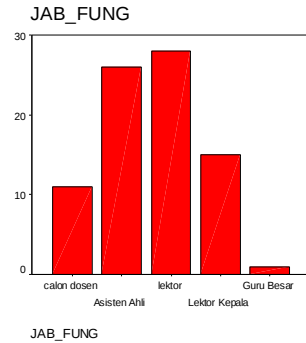
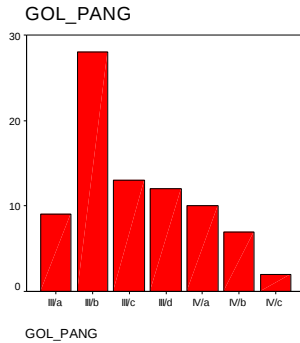
		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	S1	7	8,6	8,6	8,6
	S2	72	88,9	88,9	97,5
	S3	2	2,5	2,5	100,0
	Total	81	100,0	100,0	

JEN_KEL

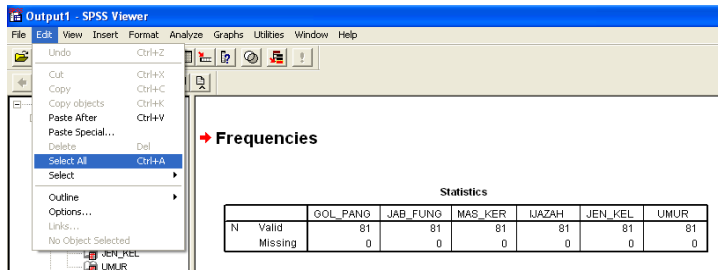
		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	Laki-laki	58	71,6	71,6	71,6
	perempuan	23	28,4	28,4	100,0
	Total	81	100,0	100,0	

UMUR

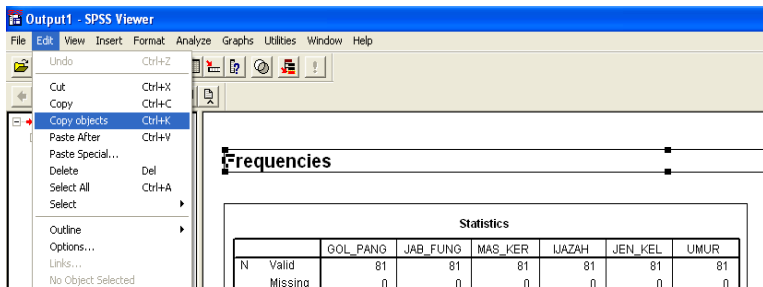
		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	27	2	2,5	2,5	2,5
	29	2	2,5	2,5	4,9
	30	4	4,9	4,9	9,9
	31	2	2,5	2,5	12,3
	32	3	3,7	3,7	16,0
	33	6	7,4	7,4	23,5
	34	9	11,1	11,1	34,6
	35	9	11,1	11,1	45,7
	36	4	4,9	4,9	50,6
	37	2	2,5	2,5	53,1
	38	6	7,4	7,4	60,5
	39	3	3,7	3,7	64,2
	40	1	1,2	1,2	65,4
	41	4	4,9	4,9	70,4
	42	2	2,5	2,5	72,8
	43	4	4,9	4,9	77,8
	44	2	2,5	2,5	80,2
	45	2	2,5	2,5	82,7
	46	1	1,2	1,2	84,0
	47	1	1,2	1,2	85,2
	48	2	2,5	2,5	87,7
	54	1	1,2	1,2	88,9
	55	2	2,5	2,5	91,4
	56	2	2,5	2,5	93,8
	58	1	1,2	1,2	95,1
	60	2	2,5	2,5	97,5
	61	1	1,2	1,2	98,8
	62	1	1,2	1,2	100,0
	Total	81	100,0	100,0	



6. Dikarenakan peneliti ketika menulis laporan penelitiannya biasanya menggunakan program Microsoft Word, maka output SPSS harus dicopy ke Microsoft Word, dengan cara sebagai berikut: Klik **Edit** ➤ **Select All**, sebagaimana gambar berikut ini.



7. Setelah keluar berikut Klik **Copy Object**. Setelah itu buka Microsoft Word, tempatkan kursor pada tempat yang diinginkan dan Klik **Edit ► paste**.



8. Dikarenakan SPSS itu menggunakan *basic language* Bahasa Inggris, maka outputnya juga menggunakan Bahasa Inggris. Untuk mengeditnya dilakukan dengan cara sebagaimana tata cara mengedit gambar. Sebagai contoh kita akan mengedit tabel GOL_PANG, yaitu tabel tentang golongan kepangkatan dosen maka klik gambar tersebut, lalu klik tombol kanan mouse, lalu klik sub menu **Edit Picture...**, maka akan keluar gambar berikut ini:

	Frequency	Percent	Valid Percent	Cumulative Percent
1102	28	34,6	34,6	34,6
1106	28	34,6	34,6	69,2
1109	12	14,8	14,8	84,0
1125	12	14,8	14,8	98,8
1126	3	3,7	3,7	100,0
1172	0	0,0	0,0	
TOTAL	81	100,0	100,0	

Setelah keluar gambar seperti di atas, klik pada tulisan yang akan diedit dan gantilah sesuai dengan kebutuhan.

Apabila dicermati output SPSS di atas, ada penyajian data tentang Masa Kerja dan Umur yang terlalu panjang dan hal itu mengakibatkan pemahaman terhadap data tersebut agak sulit. Untuk mengatasi hal tersebut, maka tabel distribusi frekuensi tersebut perlu dibuat berdasarkan kelas data. Salah satu cara untuk menentukan jumlah kelas adalah dengan menggunakan rumus Struges:

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

K = Jumlah klas interval

N = Jumlah data observasi

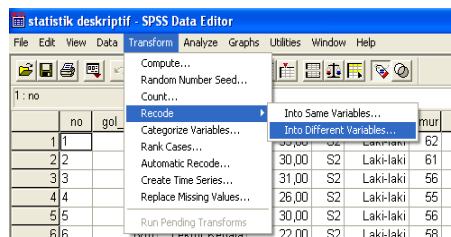
log = Logaritma

Dikarenakan jumlah data di atas berjumlah 81, maka jumlah kelasnya adalah: $1 + 3,3 \log .81 = 1 + (3,3 \cdot 1,91) = 7,30$ dan dapat dibulatkan menjadi 7 kelas.

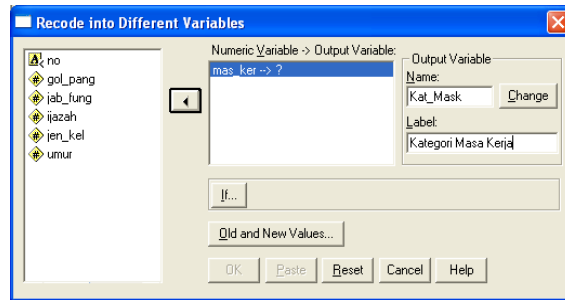
Selanjutnya rentang datanya dihitung dengan cara mengurangi data terbesar dengan data terkecil. Untuk kasus data masa kerja paling sedikit 1 tahun dan paling lama 35 tahun, sedangkan data umur paling muda 27 tahun dan paling tua 62 tahun. Dari hasil pengurangan data terbesar dengan data terkecil dan selanjutnya hasil pengurangan tersebut dibagi dengan jumlah kelas, maka panjang kelas sekitar 5.

Dalam rangka memproses aplikasi selanjutnya dapat digunakan sub menu **recode** dalam SPSS sebagai berikut:

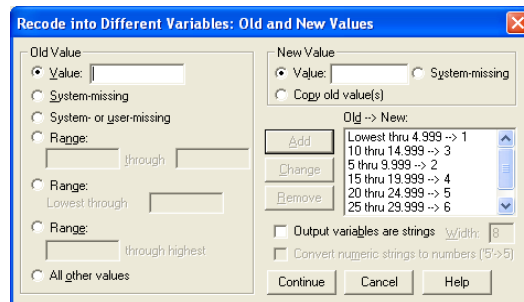
1. Klik **Transform ► Recode ► Into Different Variables**. Pilihan **Different Variables** dimaksudkan akan ada variabel baru sebagai hasil recode tersebut, tetapi kalau yang diinginkan hasil recode itu menggantikan variabel yang akan dibuat kategorinya, maka pilihlah **Into Same Variables**.



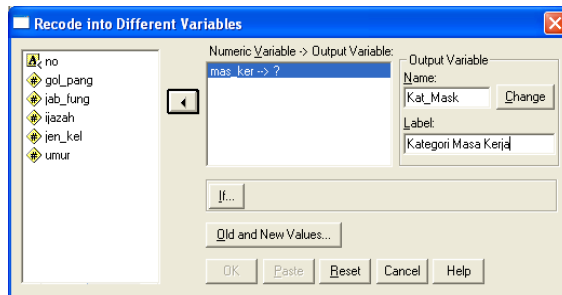
2. Setelah keluar gambar berikut, bloklah variabel yang akan direcode, lalu masukkan dalam **Numeric Variable ► Output Variable** dengan cara mengeklik tanda gambar ►. Setelah itu berilah Nama variabel baru itu pada kolom di bawah **Name** dan labelnya dari variabel tersebut pada kolom di bawah **Label**. Setelah itu klik **Old and New Values**.



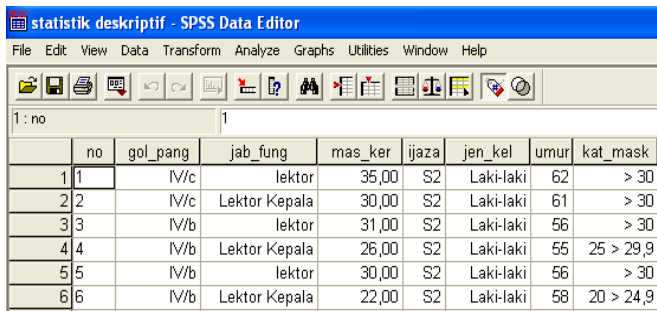
3. Setelah keluar gambar berikut, maka kita dapat membagi kelas data menjadi 7 dan panjang kelasnya adalah 5, sebagaimana dijelaskan di atas. Oleh karena itu kita dapat menggunakan batas bawah dan batas atas terbuka. Dalam hal ini kita dapat mengisi dengan cara mengeklik **Range** yang ada **lowest through** dan pada kotak tersebut diisi dengan 4,999. Setelah itu, Isikan angka **1** pada kolom **Value** yang terdapat dalam **New Value**. Selanjutnya klik **Range** yang isilah kotak yang ada di bawahnya dengan angka 5 dan pada kotak disebelah kanan **through** diisi dengan 9,999. Setelah itu, Isikan angka **2** pada kolom **Value** yang terdapat dalam **New Value**. Dan seterusnya sampai selesai. Sebagaimana tertera dalam gambar berikut. Selanjutnya klik **Continue**.



4. Setelah keluar gambar berikut ini, klik **Change** lalu klik **Ok**.

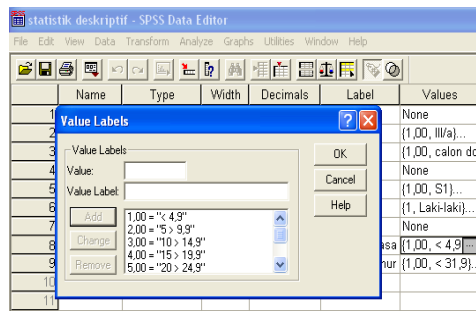


5. Dikarenakan dalam recode tersebut, dipilih **Into Different Variables**, maka akan ada tambahan satu variabel dalam lembaran SPSS Data Editor sebagaimana terlihat pada gambar berikut.



	no	gol_pang	jab_fung	mas_ker	ijaza	jen_kel	umur	kat_mask
1	1	IV/c	lektor	35,00	S2	Laki-laki	62	> 30
2	2	IV/c	Lektor Kepala	30,00	S2	Laki-laki	61	> 30
3	3	IV/b	lektor	31,00	S2	Laki-laki	56	> 30
4	4	IV/b	Lektor Kepala	26,00	S2	Laki-laki	55	25 > 29,9
5	5	IV/b	lektor	30,00	S2	Laki-laki	56	> 30
6	6	IV/b	Lektor Kepala	22,00	S2	Laki-laki	58	20 > 24,9

6. Dalam recode sudah dibagi menjadi 7 kelas dengan diberi atribut 1 sampai dengan 7. Agar program SPSS mengenali yang dimaksud dengan 1 sampai dengan 7, maka dibuat **values**, dengan cara buka **Variabel View**, lalu klik pada **value** di baris variabel yang diberi skor kategori/valuenya, dan buatlah value tersebut dengan cara ketik 1 pada kolom **value** dan pada kolom **Value Label** ketik < 4,9, lalu klik **add**, begitu seterusnya.



Name	Type	Width	Decimals	Label	Values
1					None
2					(1,00, III/a)...
3					(1,00, calon do)
4					None
5					(1,00, S1)...
6					(1, Laki-laki)...
7					None
8					sa (1,00, < 4,9)...
9					hur (1,00, < 31,9)...

7. Ketika variabel baru tersebut akan dibuat tabel distribusi frekuensi caranya adalah klik **Analyze** ► **Descriptive Statistics** ► **Frequencies**. Setelah itu bloklah variabel yang telah dikategorisasi lalu masukkan pada kolom **Variable(s)** dengan mengeklik tanda gambar ►. Setelah itu klik **Ok**.
8. Berikut ini adalah Output dari tabel distribusi frekwensi tersebut.

Frequencies

Statistics

		Kategori Masa Kerja	Kategori Umur
N	Valid	81	81
	Missing	0	0

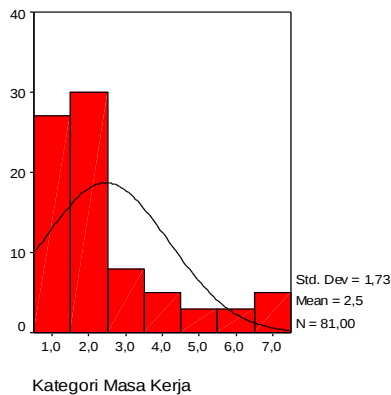
Frequency Table

Kategori Masa Kerja

	Frequency	Percent	Valid Percent	Cumulative Percent
Valid < 4,9	27	33,3	33,3	33,3
5 > 9,9	30	37,0	37,0	70,4
10 > 14,9	8	9,9	9,9	80,2
15 > 19,9	5	6,2	6,2	86,4
20 > 24,9	3	3,7	3,7	90,1
25 > 29,9	3	3,7	3,7	93,8
> 30	5	6,2	6,2	100,0
Total	81	100,0	100,0	

Histogram

Kategori Masa Kerja



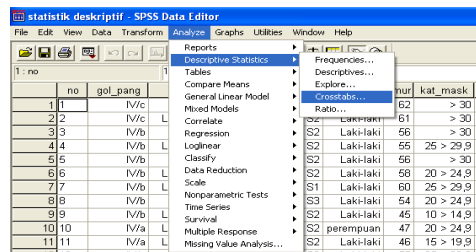
Walaupun deskripsi data di atas sudah memberi pemahaman yang mudah terhadap data, tetapi penyajian data seperti itu masih sangat umum. Memang, selama ini kebanyakan peneliti pemula tidak mengolah data mentah yang didapatkan dari lokasi penelitian untuk disesuaikan dengan topik penelitian. Seharusnya data tersebut disajikan dengan memperhatikan topik penelitian.

Contoh:

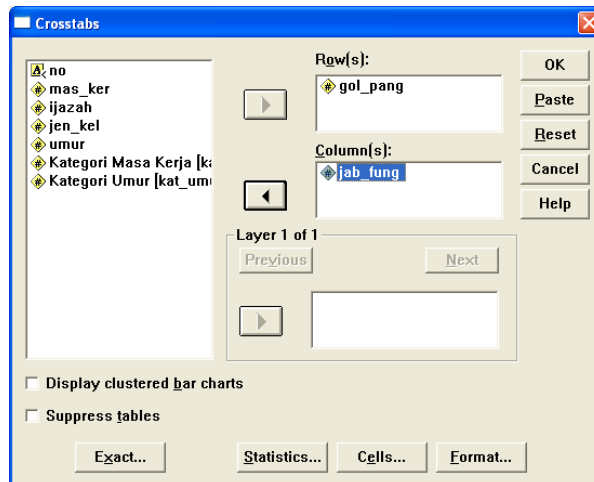
Seorang peneliti yang mempunyai topik penelitian tentang pengaruh gaya kepemimpinan terhadap produktifitas kerja dosen, maka peneliti perlu mentabulasikan data dosen yang terkait dengan pangkat dan jabatan fungsionalnya. Hal ini diperlukan mengingat salah satu cara untuk mengukur produktifitas kerja dosen adalah dengan melihat apakah jabatan fungsional dosen itu lebih tinggi dibandingkan dengan posisi pangkat yang melingkupinya.

Untuk memperjelas hal ini dapat ditabulasikan data dosen di atas. Sedangkan cara sebagai berikut:

1. Setelah data diinput pada lembar SPSS Data Editor, maka klik **Analyze ► Descriptive Statistics ► Crosstabs**, sebagaimana gambar berikut.



2. Masukkan variabel **gol_pang** pada kolom Row(s), dan variabel **jab_fung** pada kolom Column(s) dengan mengeblok variabel tadi dan mengklik tanda gambar ►, setelah itu klik OK.



3. Berikut ini adalah output aplikasi di atas.

Case Processing Summary

	Cases					
	Valid		Missing		Total	
	N	Percent	N	Percent	N	Percent
GOL_PANG * JAB_FUNG	81	100,0%	0	,0%	81	100,0%

GOL_PANG * JAB_FUNG Crosstabulation

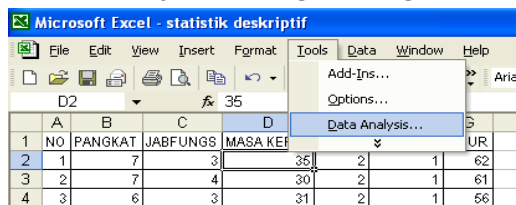
Count		JAB_FUNG					Total
		calon dosen	Asisten Ahli	lektor	Lektor Kepala	Guru Besar	
GOL_PANG	III/a	7	2	0	0	0	9
	III/b	4	24	0	0	0	28
	III/c	0	0	13	0	0	13
	III/d	0	0	11	1	0	12
	IV/a	0	0	1	9	0	10
	IV/b	0	0	2	4	1	7
	IV/c	0	0	1	1	0	2
Total		11	26	28	15	1	81

Tabel di atas memperlihatkan ada 4 orang dosen yang produktifitas akademiknya rendah mengingat mereka sudah berpangkat IV/a, IV/b, dan IV/c tetapi masih mempunyai jabatan fungsional lektor. Di samping itu ada 2 orang yang mempunyai produktifitas akademik tinggi karena salah satunya mempunyai jabatan fungsional lektor kepala padahal pangkatnya baru III/d, sedangkan satu lainnya sudah mempunyai jabatan fungsional Guru Besar padahal ia baru mempunyai pangkat IV/b.³

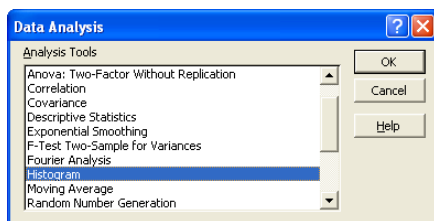
C. Aplikasi untuk Penyajian Data dengan Microsoft Excel

Untuk menyajikan data dengan program Microsoft Excel, baik dalam bentuk tabel maupun grafik tidak selengkap SPSS di atas. Dalam aplikasi ini akan dicontohkan penyajian data tentang masa kerja dosen di sebuah Perguruan Tinggi Agama Islam di atas. Salah satu langkah yang dapat ditempuh adalah sebagai berikut:

1. Input data tentang masa kerja dosen tersebut.
2. Setelah itu, buka sub menu **Data Analysis...**, dengan cara klik **Tools ► Data Analysis...**, sebagaimana gambar berikut ini.



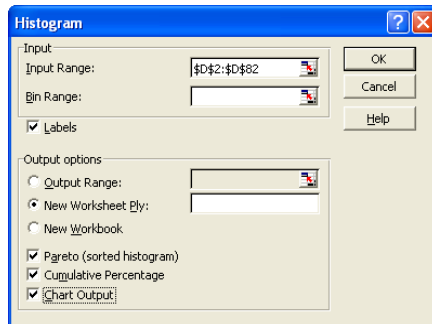
3. Setelah keluar gambar berikut ini pilihlah **Histogram** lalu klik **Ok**.



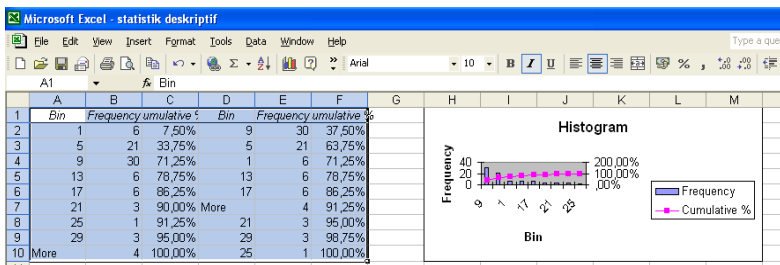
4. Setelah keluar gambar berikut ini bloklah data yang telah diinput tentang masa kerja dosen. Setelah itu aktifkan/pilihlah **Label**,

³Untuk mendapatkan penjelasan tentang kepengkatan dosen dan jabatan fungsionalnya baca Keputusan Mendiknas nomor: 36/D/O/2001 tertanggal 4 Mei 2001 tentang Juklak. Baca Suwito dkk., *Buku Pedoman Tenaga Akademik Perguruan Tinggi Agama Islam dan PAI pada PTU*, (Jakarta: Dirjen Binbaga Islam Depag RI, 2003), hlm. 70-93.

Pareto (sorted histogram), Cumulative Percentage, dan Chart Output. Setelah itu klik **Ok**.



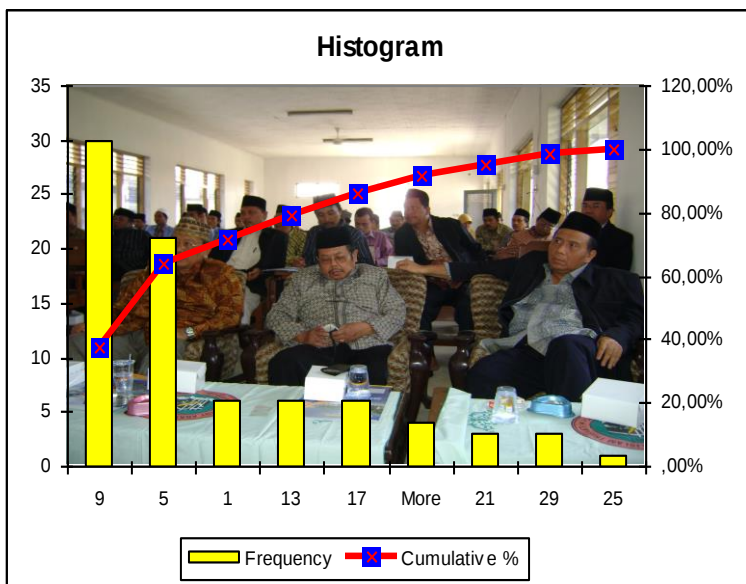
5. Gambar berikut ini adalah hasil dari aplikasi tersebut.



6. Hasil sebagaimana tertera di atas dapat diedit sesuai dengan kebutuhan. Sebagai contoh tabel di atas dapat diedit menjadi seperti berikut ini.

Interval	Frekwensi	Kumulatif %	Interval	Frekwensi	Kumulatif %
1	6	7,50%	9	30	37,50%
5	21	33,75%	5	21	63,75%
9	30	71,25%	1	6	71,25%
13	6	78,75%	13	6	78,75%
17	6	86,25%	17	6	86,25%
21	3	90,00%	More	4	91,25%
25	1	91,25%	21	3	95,00%
29	3	95,00%	29	3	98,75%
More	4	100,00%	25	1	100,00%

7. Sedang cara untuk mengedit gambar adalah dengan mengeklik pada gambar tersebut, lalu klik kanan pada mouse. Selanjutnya, editlah sesuai dengan keinginan. Sebagai contoh grafik di atas dapat diedit menjadi seperti berikut ini.



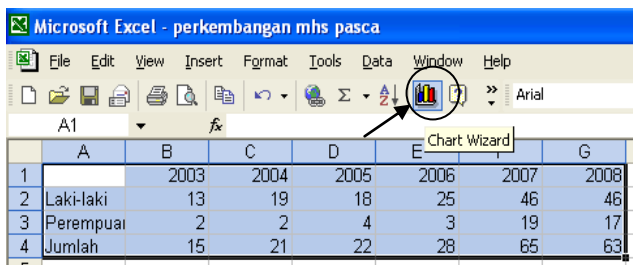
Selain itu, untuk membuat grafik, Microsoft Excel juga menyediakan ikon grafik.

Contoh 1:

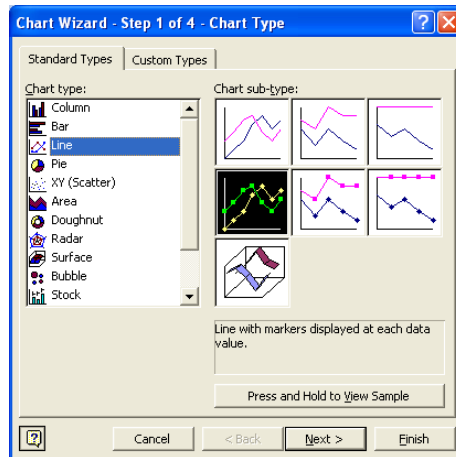
Peneliti akan menampilkan perkembangan mahasiswa baru Program Pascasarjana IAIT Kediri dengan tabel sebagai berikut:

	2003	2004	2005	2006	2007	2008
Laki-laki	13	19	18	25	46	46
Perempuan	2	2	4	3	19	17
Jumlah	15	21	22	28	65	63

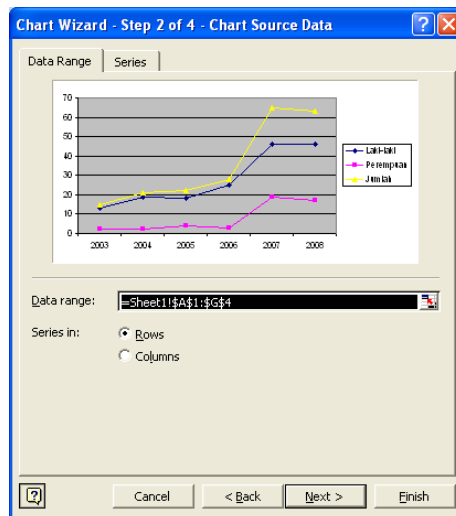
1. Setelah data di atas diinput pada lembar kerja Microsoft Excel, lalu diblok. Setelah itu, klik ikon grafik seperti gambar di bawah ini.



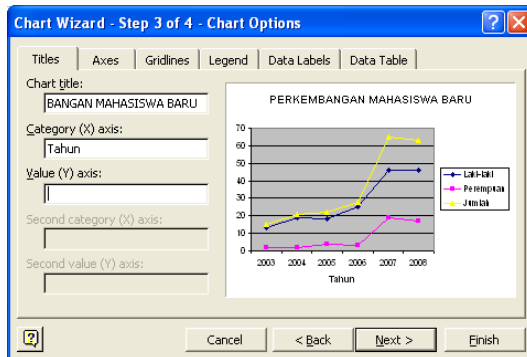
- Setelah keluar seperti berikut ini, pilihlah grafik sesuai dengan kebutuhan. Setelah itu klik **next**.



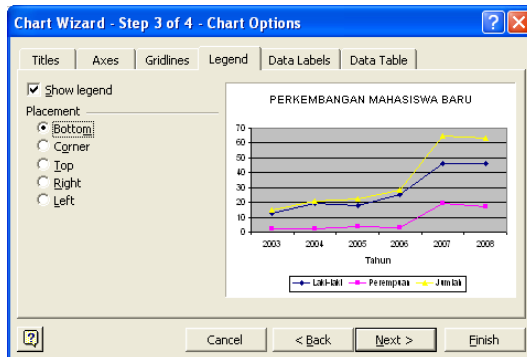
- Setelah keluar seperti berikut ini, klik **next**.



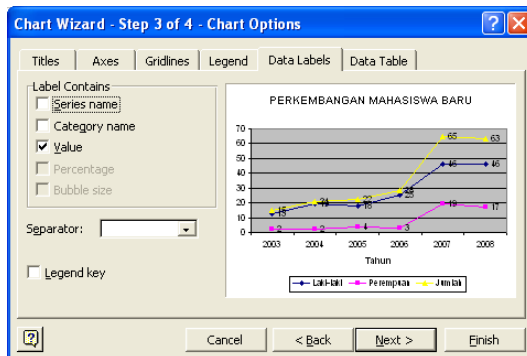
- Setelah keluar seperti berikut ini, klik **Titles**, lalu pada kolom **Chart Title** ketiklah judul grafik yang diinginkan, lalu pada **Category (X) axis (horizontal)** ketiklah sesuai kebutuhan, dan pada **Value (Y) axis (vertikal)** ketiklah sesuai dengan kebutuhan.



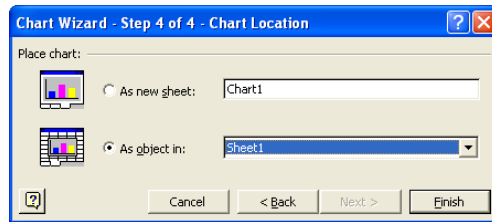
5. Apabila keterangan tentang laki-laki dan seterusnya ingin diletakkan di bawah grafik, maka klik **Legend**, lalu pilihlah **Bottom**.



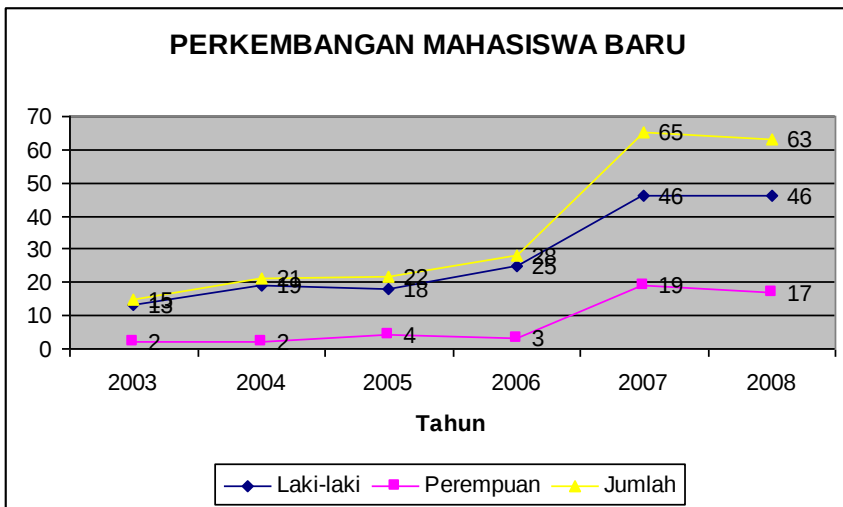
6. Apabila skor/angka jumlah mahasiswa ingin disajikan pada grafik, maka klik **Data Labels** lalu klik **Value**. Setelah itu klik **next**.



7. Setelah keluar gambar berikut ini klik **Finish**.



8. Berikut ini adalah hasil pembuatan grafik tersebut. Grafik tersebut agar lebih indah dapat diedit sesuai dengan kebutuhan, baik terkait dengan background maupun warnanya.

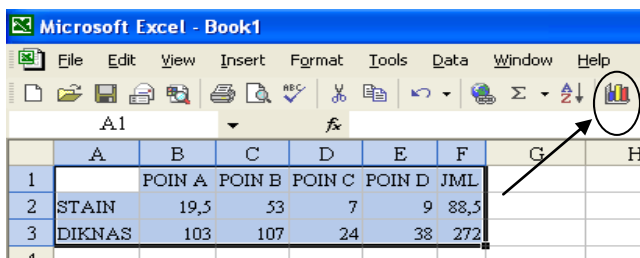


Contoh 2:

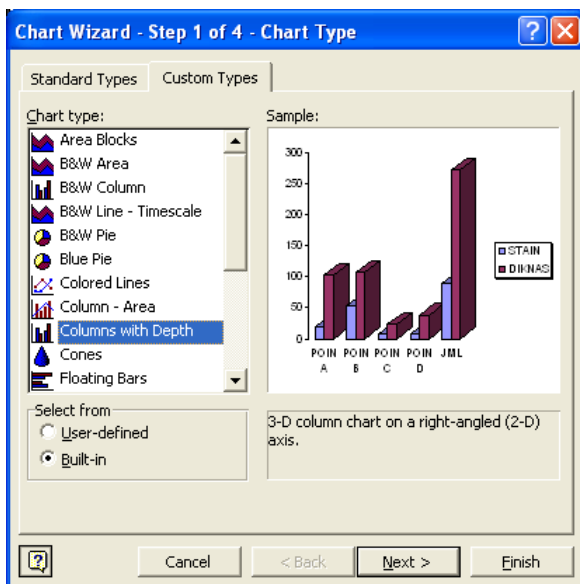
Seorang dosen yang merasa mendapatkan perlakuan tidak adil dari atasannya ketika mengajukan kenaikan jabatan fungsionalnya ke Lektor Kepala karena permohonannya tidak mendapatkan tanggapan semestinya bahkan angka kreditnya mendapatkan penilaian yang sangat rendah dan sebagian di antaranya melanggar aturan, maka dosen tersebut mengajukan kepada instansi yang ada di atasnya untuk diberikan penilaian ulang. Akhirnya berkas yang sudah dinilai oleh Sekolah Tinggi tersebut dinilai ulang dan mendapatkan Penetapan Angka Kredit dari Diknas.

Cara untuk menampilkan data di atas dalam bentuk grafik adalah sebagai berikut:

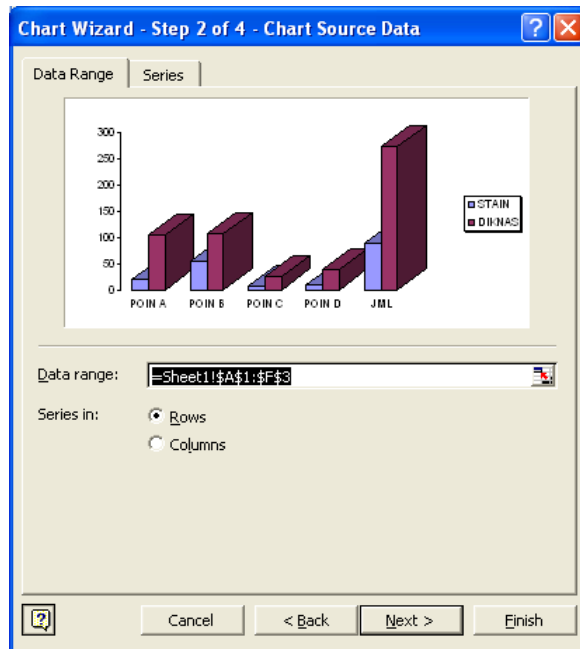
1. Setelah data di atas diinput pada lembar kerja Microsoft Excel, lalu diblok. Setelah itu, klik ikon grafik seperti gambar di bawah ini.



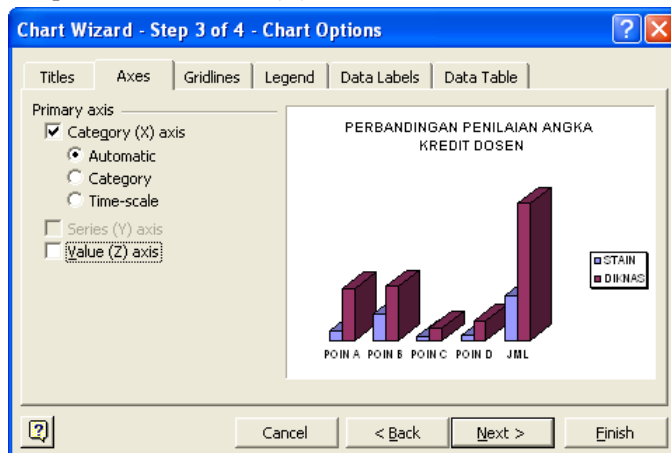
2. Setelah keluar seperti berikut ini, model grafik yang diinginkan, sebagai contoh dipilih **Custom Types** lalu diklik **Columns with Depth**. Setelah itu klik **next**.



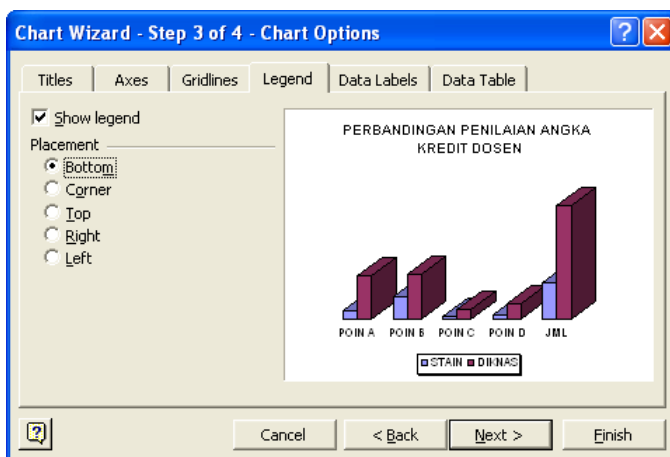
3. Setelah keluar seperti berikut ini, klik **next**.



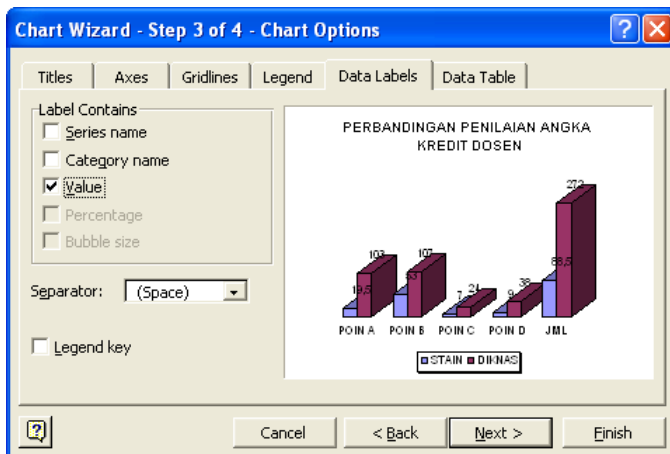
4. Setelah keluar seperti berikut ini, klik **Titles**, lalu ketik **PERBANDINGAN PENILAIAN ANGKA KREDIT DOSEN**, kakalu direncanakan angka kredit langsung tertera di atas grafik, maka pada **Axes, Value (Z) axis** dinonaktifkan.



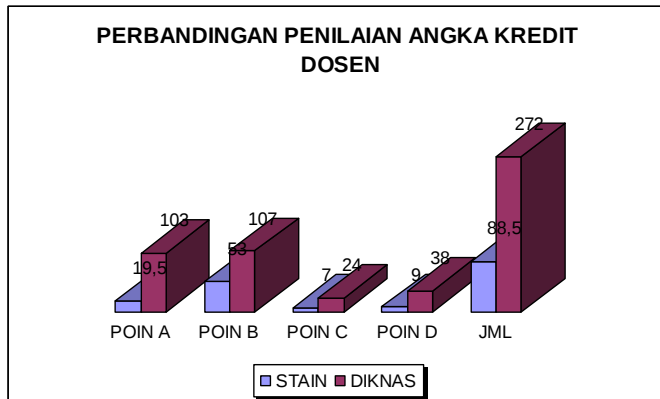
5. Apabila keterangan tentang STAIN dan DIKNAS ingin diletakkan di bawah grafik, maka klik **Legend**, lalu pilihlah **Bottom**.



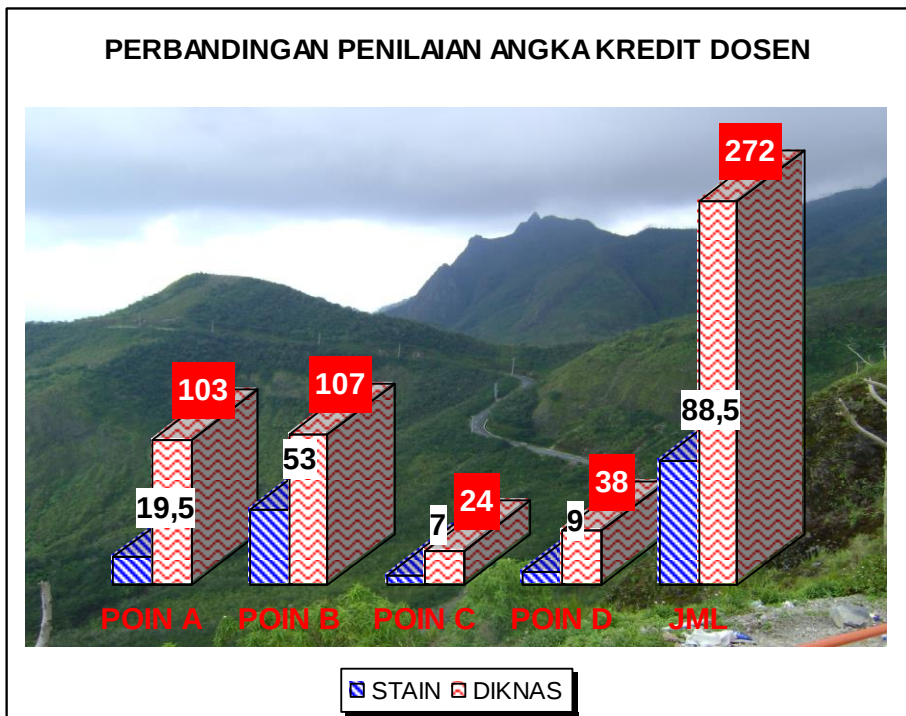
6. Apabila skor/angka jumlah mahasiswa ingin disajikan pada grafik, maka klik **Data Labels** lalu klik **Value**. Setelah itu klik **next**.



7. Setelah keluar gambar berikut ini klik **Finish**. Berikut ini adalah hasil pembuatan grafik tersebut. Grafik tersebut agar lebih indah dapat diedit sesuai dengan kebutuhan, baik terkait dengan background maupun warnanya.



8. Berikut ini adalah hasil suntingan output di atas.

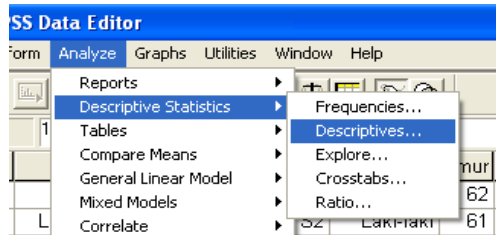


Demikianlah langkah-langkah dalam pembuatan tabel maupun grafik dengan menggunakan SPSS maupun Microsoft Excel. Selanjutnya akan dipaparkan cara untuk menyajikan ringkasan data Statistik, yaitu penyajian data yang digunakan untuk menjelaskan pemusatan dan penyebaran data.

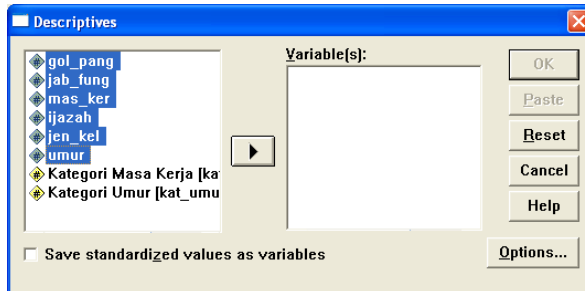
D. Aplikasi untuk Penyajian Pemusatan, Penyebaran, dan Distribusi Data dengan SPSS

Untuk contoh aplikasi ini digunakan data tenaga edukatif di sebuah Perguruan Tinggi Agama Islam yang sudah ditampilkan tabel dan grafiknya di atas.

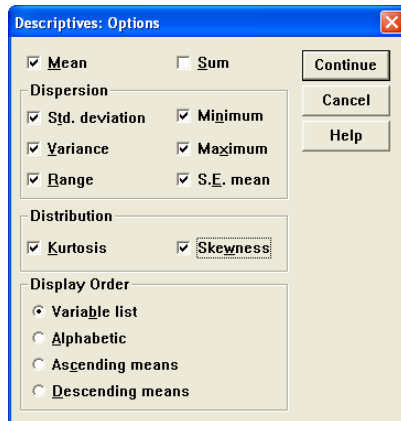
1. Setelah file data tersebut di buka, maka klik **Analyze** ➤ **Descriptive Statistics** ➤ **Descriptives**, sebagaimana gambar berikut ini.



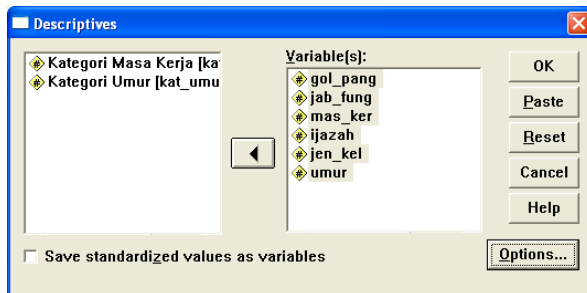
2. Setelah keluar gambar sebagaimana berikut ini, bloklah variabel-variabel yang akan dicari pemusatan, penyebaran, dan distribusi datanya; lalu klik anak panah sehingga variabel-variabel tersebut berpindah ke kolom **Variable(s)**. Setelah itu klik **Options...**



3. Setelah keluar gambar sebagaimana berikut ini, Klik pada kotak yang ada di depan pilihan-pilihan yang akan ditampilkan dalam ringkasan data. Misalnya dalam latihan ini dipilih yang akan ditampilkan dalam ringkasan data seperti di bawah ini. Setelah itu klik **Continue**.



4. Setelah keluar gambar sebagaimana berikut ini, Klik **Ok**.



5. Berikut ini adalah **output** dari aplikasi di atas.

Descriptive Statistics								
		GOL PANG	JAB FUNG	MAS KER	IJAZAH	JEN KEL	UMUR	Valid N (listwise)
N	Statistic	81	81	81	81	81	81	81
Range	Statistic	6,00	4,00	34,00	2,00	1	35	
Minimum	Statistic	1,00	1,00	1,00	1,00	1	27	
Maximum	Statistic	7,00	5,00	35,00	3,00	2	62	
Mean	Statistic	3,1852	2,6173	9,7160	1,9383	1,28	39,21	
	Std. Error	,1793	,1091	,9143	,0366	,05	,94	
Std. Deviation	Statistic	1,61331	,98194	8,22836	,32961	,454	8,504	
Variance	Statistic	2,603	,964	67,706	,109	,206	72,318	
Skewness	Statistic	,588	,028	1,537	-1,207	,976	1,161	
	Std. Error	,267	,267	,267	,267	,267	,267	
Kurtosis	Statistic	-,655	-,700	1,708	5,980	-1,074	,691	
	Std. Error	,529	,529	,529	,529	,529	,529	

Perlu diketahui bahwa tipe data *gol_pang*, *jab_fung*, *ijazah*, dan *jen_kel* adalah data kategori, sedangkan *mas_ker* dan *umur* adalah data rasio.

Kategorisasi yang dimaksud adalah sebagai berikut:

GOL_PANG	JAB_FUNG	IJAZAH	JEN_KEL
1 = III/a	1 = Calon Dosen	1 = S1	1 = Laki-laki
2 = III/b	2 = Asisten Ahli	2 = S2	2 = Perempuan

3	= III/c	3	= Lektor	3	= S3		
4	= III/d	4	= Lektor Kepala				
5	= IV/a	5	= Guru Besar				
6	= IV/b						
7	= IV/c						

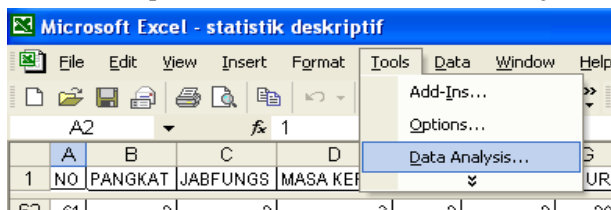
Output di atas dapat menampilkan tentang pemusatan data (mean), penyebaran data (range, minimum, maximum, standard error of mean, standard deviasi, dan variance), dan distribusi data (skewness, kurtosis, dan standard errornya).

Untuk memperjelas prosedur statistik yang menghasilkan output di atas akan ditampilkan dengan aplikasi Microsoft excel dengan cara menjabarkan rumus-rumus yang terkait. Sedangkan aplikasi **Data Analysis** dalam Microsoft Excel juga akan langsung menampilkan output tanpa diketahui prosedurnya.

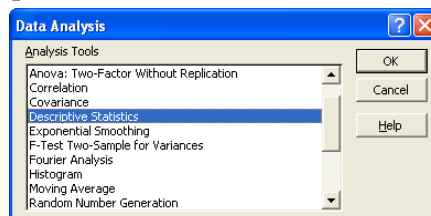
E. Aplikasi untuk Penyajian Pemusatan, Penyebaran, dan Distribusi Data dengan Microsoft Excel

Microsoft Excel juga mempunyai menu yang dapat menampilkan statistik deskriptif dan statistik inferensial secara cepat tanpa melalui aplikasi rumus-rumus statistik yaitu dengan menggunakan menu **Data Analysis**. Sebagai contoh untuk aplikasi ini adalah penyajian statistik deskriptif tentang masa kerja tenaga edukatif yang juga ditampilkan dengan SPSS di atas. Sedangkan langkah-langkahnya adalah sebagai berikut:

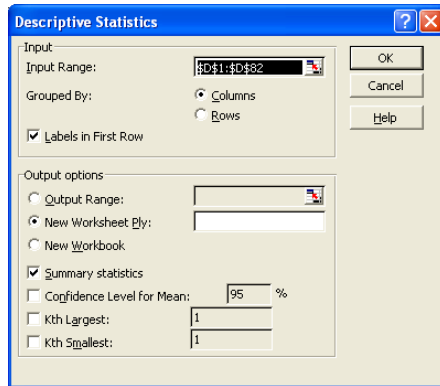
1. Setelah data diinput, lalu klik **Tools ► Data Analysis...**



2. Setelah keluar gambar seperti di bawah ini, klik dua kali secara cepat **Descriptive Statistics**.



3. Setelah keluar gambar seperti berikut ini, ketika kursor berada pada kolom Input Range, maka bloklah seluruh data yang berada pada kolom Masa Kerja. Jikalau nama variabel (Masa Kerja) akan ditampilkan data output, maka nama variabel itu harus juga diblok, dan pada kotak **Label in First Row** harus diklik. Selanjutnya klik kotak **Summary Statistics** lalu klik **Ok**.



4. Berikut ini adalah output dari aplikasi di atas. Apabila dibandingkan dengan aplikasi SPSS di atas ternyata hasilnya sama. Bedanya kalau dalam aplikasi SPSS standard error dari skewness dan kurtosis ditampilkan sedangkan dalam Microsoft Excel kedua hal itu tidak ditampilkan.

MASA KERJA	
Mean	9,716049383
Standard Error	0,914261925
Median	7
Mode	7
Standard Deviation	8,228357321
Sample Variance	67,7058642
Kurtosis	1,708296047
Skewness	1,537190627
Range	34
Minimum	1
Maximum	35
Sum	787
Count	81

Dalam rangka menjelaskan prosedur yang harus dilalui sampai keluarnya hasil berbagai ringkasan data di atas, maka akan diaplikasikan berbagai rumus statistik yang terkait.

Rumus mean (rata-rata):

$$\text{Me} = \frac{\sum X_i}{n}$$

Rumus varians untuk populasi:

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

Rumus standard deviasi untuk populasi:

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

Rumus varians untuk sampel:

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

Rumus standard deviasi untuk sampel:

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

Rumus Standard error of mean:

$$E = \frac{s}{\sqrt{n}}$$

Rumus Skewness:

$$\alpha_3 = \left(\frac{n}{[n-1][n-2]} \right) \left(\frac{\sum_{i=1}^n [X_i - \bar{X}]^3}{s^3} \right)$$

Rumus Kurtosis

$$\alpha_4 = \left[\left(\frac{n[n+1]}{[n-1][n-2][n-3]} \right) \left(\frac{\sum_{i=1}^n [X_i - \bar{X}]^4}{s^4} \right) \right] - \left[\frac{3[n-1]^2}{[n-2][n-3]} \right]$$

Keterangan:

Me= Mean (rata-rata)

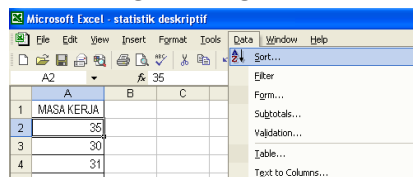
Σ = Epsilon (baca jumlah)

- X_i = Nilai X ke i sampai ke n
 n = Jumlah individu
 σ^2 = Varians populasi
 σ = Simpangan baku populasi
 s^2 = Varians sampel
 s = Simpangan baku sampel

Cara untuk mencari mode adalah menghitung nilai yang sering muncul dalam kelompok data tersebut. Sebagaimana digambarkan pada halaman 24 diketahui bahwa nilai yang paling sering muncul dari data masa kerja adalah 4 dan 7 yang keduanya muncul sebanyak 11 kali. Artinya, dari 81 tenaga edukatif ada 11 orang yang mempunyai masa kerja 4 tahun dan ada 11 orang juga yang mempunyai masa kerja 7 tahun.

Sedangkan cara untuk mencari median adalah dengan mengurutkan nilai dari terkecil kepada nilai terbesar atau sebaliknya, selanjutnya ditentukan nilai tengahnya. Apabila jumlah datanya genap, maka dua data yang ada ditengah dijumlahkan lalu dibagi dua. Karena data masa kerja di atas karena belum urut, maka perlu diurutkan dengan cara:

1. Klik **Data** ➤ **Sort**, sebagaimana gambar berikut ini:



2. Apabila keluar gambar sebagaimana berikut ini dan data yang akan diurutkan sudah tertera **Sort by**, maka selanjutnya pilihlah **ascending** apabila akan diurutkan dari data terkecil ke terbesar, dan pilihlah **descending** apabila akan diurutkan dari data terbesar ke data terkecil. Setelah itu klik **Ok**.



Selanjutnya data masa kerja tersebut akan dihitung rata-rata, standard error of mean, varians, standard deviasi, skewness, dan kurtosisnya. Aplikasi dari penghitungan tersebut adalah dengan cara menjabarkan rumus-rumus yang sudah dipaparkan di atas. Adapun langkah-langkahnya adalah sebagai berikut:

1. Menjumlah seluruh nilai dari data masa kerja.
2. Membagi hasil penjumlahan tersebut dengan jumlah data (81). Hasil pembagian ini disebut rata-rata (mean)
3. Mengurangi masing-masing nilai dengan rata-rata.
4. Mengkwadratkan hasil pengurangan yang ada di nomor 3.
5. Menjumlah hasil pengkwadratan yang ada di nomor 4.
6. Hasil penjumlahan yang ada di nomor 5 dibagi jumlah data dikurangi 1 ($81-1$). Hasil pembagian ini disebut varians sampel.
7. untuk mengetahui standard deviasi sampel, carilah akar dari varians sampel di atas.
8. Untuk mencari skewness dimulai dengan cara masing-masing skor hasil pengurangan yang ada di nomor 3 dipangkatkan tiga. Selanjutnya dijumlahkan hasil pengkuadratan tersebut. Berdasarkan hasil penjumlahan tersebut dan skor standard deviasi dapat diaplikasikan rumus skewness.
9. Untuk mencari kurtosis dimulai dengan masing-masing skor hasil pengurangan yang ada di nomor 3 dipangkatkan empat. Selanjutnya dijumlahkan hasil pengkuadratan tersebut. Berdasarkan hasil penjumlahan tersebut dan skor standard deviasi yang sudah diketahui dapat diaplikasikan rumus kurtosis.

Untuk lebih jelasnya aplikasi tersebut dapat disimak pada contoh aplikasi pada halaman berikut.

Dari seluruh hasil deskripsi SPSS, hanya standard error of skewness dan standard error of kurtosis yang tidak diketahui rumusnya. Dua standar kesalahan yang disebutkan terakhir memang tidak ditampilkan ketika menggunakan aplikasi **Data Analysis**. Dari seluruh buku statistik yang penulis miliki tidak ada satupun yang menjelaskan hal tersebut. Sedangkan cara untuk mengetahui range adalah dilakukan dengan cara mengurangi nilai tertinggi dengan nilai terendah. Untuk contoh data di atas adalah $35 - 1 = 34$.

	A	B	C	D	E	F	G
1	MASA KERJA	xi - mean	Bi^A 2	Bi^A 3	Bi^A 4		
2	35	25,28395062	639,2781588	16163,4774	408676,5843		Formula sel A83
3	30	20,28395062	411,4386526	8345,601312	163281,7649		=SUM(A2:A82)
4	31	21,28395062	453,0065539	9641,769122	205214,9379		
5	26	16,28395062	285,1670477	4317,96711	70313,56319		Formula sel A84 (rumus rata-rata)
6	30	20,28395062	411,4386526	8345,601312	163281,7649		=A83*81
7	22	12,28395062	150,8964428	1853,592167	22769,43465		
8	27	17,28395062	298,7349489	5163,320105	89242,56972		Formula sel B2
9	20	10,28395062	105,7596403	1087,626918	11185,10152		=A2-\$A\$84
10	14	4,283950617	18,35223289	78,62005942	336,8044521		
11	20	10,28395062	105,7596403	1087,626918	11185,10152		Formula sel C2
12	18	8,283950617	68,62383783	568,4764837	4709,231118		=B2^2
13	28	18,28395062	334,3028502	6112,376804	111758,3956		
14	16	6,283950617	39,48903536	248,1408642	1559,304937		Formula sel D2
15	15	5,283950617	27,92013413	147,5286099	779,5338896		=B2^3
16	33	23,28395062	542,1423563	12623,21585	293918,3345		
17	14	4,283950617	18,35223289	78,62005942	336,8044521		Formula sel E2
18	13	3,283950617	10,78433166	35,4152126	116,3018093		=B2^4
19	15	5,283950617	27,92013413	147,5286099	779,5338896		
20	13	3,283950617	10,78433166	35,4152126	116,3018093		Formula sel C84 (rumus varian)
21	16	6,283950617	39,48903536	248,1408642	1559,304937		=C83/(81-1)
22	10	0,283950617	0,090627953	0,022894357	0,006500867		
23	9	-0,716049383	0,512726718	-0,36713765	0,262888688		Formula sel C85 (rumus standar deviasi)
24	9	-0,716049383	0,512726718	-0,36713765	0,262888688		=SQRT(C84)
25	9	-0,716049383	0,512726718	-0,36713765	0,262888688		
26	9	-0,716049383	0,512726718	-0,36713765	0,262888688		Formula sel B88 (rumus standar error of mean)
27	9	-0,716049383	0,512726718	-0,36713765	0,262888688		=C85/(SQRT(81))
28	8	-1,716049383	2,944825484	-5,063465954	8,671997131		
29	9	-0,716049383	0,512726718	-0,36713765	0,262888688		rumus skewness
30	8	-1,716049383	2,944825484	-5,063465954	8,671997131		Formula sel D86
31	7	-2,716049383	7,376924249	-20,03609055	54,41901138		=81/(80^79)

A1		M WINDA KURNIA									
A		B	C	D	E	F	G				
32	11	1,283950617	1,648529188	2,116630068	2,717648482						
33	8	-1,716049383	2,944825484	-5,063465954	8,671997131						Formula sel D87 =D83/(C85*3)
34	9	-0,716049383	0,512726718	-0,36713765	0,262888888						Formula sel D88 =D86*D87
35	8	-1,716049383	2,944825484	-5,063465954	8,671997131						
36	8	-1,716049383	2,944825484	-5,063465954	8,671997131						
37	8	-1,716049383	2,944825484	-5,063465954	8,671997131						rumus kurtosis
38	7	-2,716049383	7,376924249	-20,03609055	54,41901138						Formula sel E86 =(81*82)/(80*79*78)
39	8	-1,716049383	2,944825484	-5,063465954	8,671997131						Formula sel E87 =E83/(C85*4)
40	7	-2,716049383	7,376924249	-20,03609055	54,41901138						Formula sel E88 =E86*E87
41	7	-2,716049383	7,376924249	-20,03609055	54,41901138						Formula sel E89 =(3*(80*2))/(79*78)
42	6	-3,716049383	13,90902301	-51,31501145	190,6891166						Formula sel E90 =E88-E89
43	12	2,283950617	5,216430422	11,91406948	27,21114635						
44	7	-2,716049383	7,376924249	-20,03609055	54,41901138						
45	7	-2,716049383	7,376924249	-20,03609055	54,41901138						
46	10	0,283950617	0,080627953	0,022894357	0,006500867						
47	4	-5,716049383	32,67322055	-186,7617421	1067,539341						
48	4	-5,716049383	32,67322055	-186,7617421	1067,539341						
49	4	-5,716049383	32,67322055	-186,7617421	1067,539341						
50	4	-5,716049383	32,67322055	-186,7617421	1067,539341						
51	4	-5,716049383	32,67322055	-186,7617421	1067,539341						
52	4	-5,716049383	32,67322055	-186,7617421	1067,539341						
53	4	-5,716049383	32,67322055	-186,7617421	1067,539341						
54	4	-5,716049383	32,67322055	-186,7617421	1067,539341						
55	3	-6,716049383	45,10631931	-302,9295519	2034,48983						
56	7	-2,716049383	7,376924249	-20,03609055	54,41901138						
57	3	-6,716049383	45,10631931	-302,9295519	2034,48983						
58	3	-6,716049383	45,10631931	-302,9295519	2034,48983						
59	3	-6,716049383	45,10631931	-302,9295519	2034,48983						
60	3	-6,716049383	45,10631931	-302,9295519	2034,48983						
61	3	-6,716049383	45,10631931	-302,9295519	2034,48983						
62	3	-6,716049383	45,10631931	-302,9295519	2034,48983						
mhs baru		histogram		ringkasan data		aplikasi rumus		data / Tabulasi silang			

	A	B	C	D	E	F	G
63	3	-6,716049383	45,10631931	-302,9295519	2034,48983		
64	3	-6,716049383	45,10631931	-302,9295519	2034,48983		
65	7	-2,716049383	7,376924249	-20,03609055	54,41901138		
66	8	-1,716049383	2,944825484	-5,063465954	8,671997131		
67	6	-3,716049383	13,80902301	-51,31501145	190,6891166		
68	6	-3,716049383	13,80902301	-51,31501145	190,6891166		
69	1	-8,716049383	75,96951684	-662,1540604	5771,367489		
70	1	-8,716049383	75,96951684	-662,1540604	5771,367489		
71	6	-3,716049383	13,80902301	-51,31501145	190,6891166		
72	4	-5,716049383	32,67322055	-186,7617421	1067,539341		
73	7	-2,716049383	7,376924249	-20,03609055	54,41901138		
74	7	-2,716049383	7,376924249	-20,03609055	54,41901138		
75	7	-2,716049383	7,376924249	-20,03609055	54,41901138		
76	4	-5,716049383	32,67322055	-186,7617421	1067,539341		
77	4	-5,716049383	32,67322055	-186,7617421	1067,539341		
78	3	-6,716049383	45,10631931	-302,9295519	2034,48983		
79	1	-8,716049383	75,96951684	-662,1540604	5771,367489		
80	1	-8,716049383	75,96951684	-662,1540604	5771,367489		
81	1	-8,716049383	75,96951684	-662,1540604	5771,367489		
82	1	-8,716049383	75,96951684	-662,1540604	5771,367489		
83	787		5416,469136	66818,88371	1641299,214		
84	9,716049383		6,71058642				
85			8,228357321				
86				0,012816456	0,01347371		
87				119,9388243	358,043002		
88		0,914261925		1,537190627	4,824167518		
89					3,11587147		
90					1,708296047		

Dari ketiga cara di atas (SPSS, Data Analysis dari Microsoft Excel, maupun menjabarkan rumus-rumus dengan Microsoft Excel) ternyata mempunyai hasil yang sama.

Data masa kerja tenaga edukatif di atas memperlihatkan bahwa rata-rata masa kerjanya adalah 9,7 tahun. Sedangkan rentang rata-rata

tersebut dapat diketahui dengan output dari standard error of mean yang bernilai 0,9 tahun. Hal ini berarti rentang rata-rata masa kerjanya adalah antara 8,8 tahun sampai dengan 10,6 tahun. Cara untuk mengetahui rentang rata-rata tersebut adalah rata-rata dikurangi dan ditambahkan dengan standar kesalahan rata-ratanya. Sedangkan median dari data tersebut adalah 7 tahun. Median ini memperlihatkan bahwa 50% tenaga edukatif di sekolah tinggi tersebut mempunyai masa kerja kurang dari 7 tahun dan 50% lainnya mempunyai masa kerja lebih dari 7 tahun. Sedangkan data masa kerja tenaga edukatif tersebut mempunyai bimode, karena nilai data yang paling sering muncul adalah 4 dan 7. Karena nilai 4 dalam urutan data muncul lebih dahulu, maka aplikasi SPSS menyajikan 4 sebagai mode. Dan karena nilai 7 muncul mendekati median, maka Microsoft Excel menyajikan 7 sebagai mode.

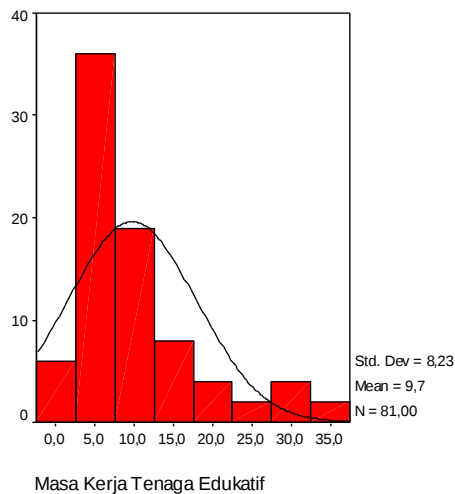
Seluruh penyajian di atas sering dikategorikan dengan penyajian pemusatan data. Sedangkan penyajian tentang penyebaran, keragaman, dan bentuk distribusi data akan dijelaskan kemudian.

Keragaman data tersebut dapat diketahui dari skor varians, yaitu 67,7, sedangkan simpangan baku dapat diketahui dari skor standard deviasi yang merupakan akar dari skor varians, yang dalam data ini adalah 8,2. Simpangan baku ini dapat digunakan untuk mengetahui posisi nilai individu bila dibandingkan dengan nilai rata-ratanya. Bila dikaitkan dengan prosentase luas kurva normal dalam rangka mengetahui normalitas distribusi data, maka jumlah individu yang berada antara mean ditambah dan dikurangi 1 standard deviasi adalah 68%, dan 95% berada antara mean ditambah dan dikurangi 2 standard deviasi, dan 99% berada antara mean ditambah dan dikurangi 3 standard deviasi.

Di samping diketahui dengan cara di atas, normalitas data juga dapat diketahui dengan rumus skewness (kemencengan) dan kurtosis (keruncingan) kurva. Apabila nilai mean lebih besar dengan mode, berarti bentuk kurva tersebut menceng ke kanan, apabila mean, median, dan mode mempunyai skor yang sama, maka bentuk kurva tersebut simetris, selanjutnya apabila nilai mean lebih kecil dari modus, maka kurva menceng ke kiri. Apabila digunakan aplikasi SPSS atau Microsoft Excel dapat dijelaskan manakala skor skewness positif (tidak ada tanda negatif) maka kurva menceng ke kanan, tetapi manakala skornya bertanda negatif, berarti kurva tersebut menceng ke kiri. Sedangkan untuk skor kurtosis, manakala skornya positif (tidak ada tanda negatif), maka kurvanya berbentuk runcing (Leptokurtik).

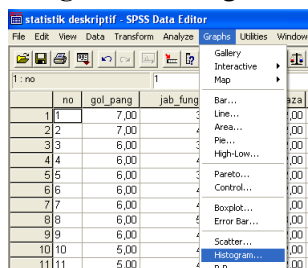
Manakala skor kurtosis bertanda negatif berarti kurvanya berbentuk flat (platikurtik).

Contoh untuk data masa kerja tenaga edukatif di atas, karena meannya (9,7) lebih besar dibanding dengan modenya (7), dan skor skewness adalah 1,5 (yang tidak ada tanda negatif) menunjukkan bahwa kurva menceng ke kanan. Bila digunakan skor kurtosis (1,7) yang juga tidak ada tanda negatif menunjukkan bahwa bentuk kurva tersebut adalah runcing (leptokurtik). Untuk membuktikan data tersebut dapat ditampilkan dengan grafik histogram dengan disertai garis kurva normalnya, sebagai berikut:

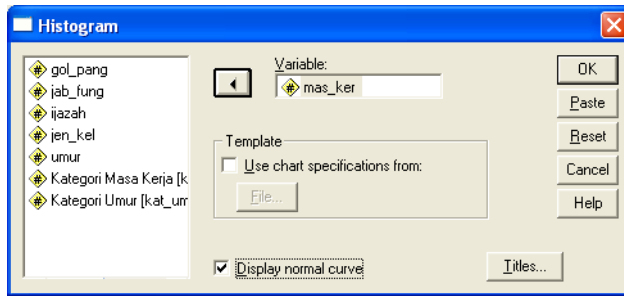


Sedangkan langkah-langkah untuk menampilkan histogram dengan kurva normal dengan SPSS adalah sebagai berikut:

1. Klik **Graphs ► Histogram...**, sebagaimana gambar berikut ini:



2. Setelah keluar gambar seperti berikut ini, masukkan data masa kerja tenaga edukatif pada kolom **variable**. Lalu klik kotak yang ada di depan **Display normal curve**, lalu klik **Ok**. Maka akan keluar sebagaimana tampilan di atas.



Berbagai penjelasan di atas adalah termasuk statistik Deskriptif. Sebagaimana dijelaskan pada Bab I bahwa statistik deskriptif adalah statistik yang digunakan untuk menggambarkan data atau menganalisis suatu hasil penelitian tetapi tidak digunakan untuk generalisasi. Pada Bab III sampai dengan Bab VI akan dipaparkan tentang statistik inferensial. Statistik inferensial adalah statistik yang digunakan untuk menganalisis data sampel dan hasilnya akan digeneralisasikan.

Dalam setiap Bab kecuali Bab yang menjelaskan regresi akan dipaparkan kedua macam statistik inferensial, yaitu parametris dan non-parametris. Statistik parametris terutama digunakan untuk menganalisis data interval atau rasio yang diambil dari populasi yang berdistribusi normal. Sementara non-parametris terutama digunakan untuk menganalisis data nominal atau ordinal. Atau datanya interval atau rasio tetapi tidak berdistribusi normal juga menggunakan statistik non-parametris.

BAB III

UJI SATU SAMPEL

A. Pendahuluan

Pada dasarnya, pengujian hipotesis satu sampel adalah menguji apakah hipotesis yang diajukan dapat digeneralisasikan untuk populasi atau tidak. Atau apakah dapat dilakukan generalisasi untuk populasi berdasarkan hasil kesimpulan dari sampel atau tidak.

Setidaknya ada 4 teknik analisis yang dapat digunakan untuk menguji hipotesis satu sampel:

Jenis Data	Deskriptif (Satu Variabel)
Nominal	Binomial χ^2 One Sample
Ordinal	Run Test
Interval/Rasio	T-Test

B. Statistik Parametris: T-test one Sampel

Statistik parametris yang digunakan untuk menguji hipotesis yang datanya satu sampel adalah t-test one sample.

Contoh:

Seorang peneliti merasa tertarik untuk meneliti mengapa Rasulullah diutus ke muka bumi ini untuk menyempurnakan akhlak, tidak mencerdaskan intelektual. Hal ini ternyata terbukti dari suatu penelitian psikologis di Amerika Serikat bahwa IQ ternyata rata-rata hanya berpengaruh 6% terhadap kesuksesan seseorang. Berangkat dari hasil penelitian di Amerika Serikat tersebut, peneliti akan membuktikan apakah pengaruh IQ terhadap kesuksesan seseorang di Indonesia juga rata-rata = 6%.

Proses Pengambilan Keputusan.

Hipotesis:

Ho Pengaruh IQ terhadap kesuksesan seseorang di Indonesia = 6%

Ha Pengaruh IQ terhadap kesuksesan seseorang di Indonesia $\neq$ 6%

Dasar Pengambilan Keputusan:

Dengan membandingkan t_{hitung} dengan t_{tabel} dengan ketentuan:

Ho diterima $t_{hitung} < t_{tabel}$

Ho ditolak $t_{hitung} \geq t_{tabel}$

Dengan menggunakan angka probabilitas, dengan ketentuan:

Ho diterima Probabilitas > taraf nyata (α)

Ho ditolak Probabilitas $\leq$ taraf nyata (α)

1. Aplikasi dengan SPSS

Untuk membuktikan apakah hasil penelitian psikologis di Amerika Serikat itu juga berlaku di Indonesia atau tidak, maka diadakan penelitian serupa di beberapa kota dan kabupaten. Sedangkan hasil penelitian tersebut adalah sebagai berikut:

NO	Koefisien Korelasi	Koefisien determinasi
1	0,25	6,25
2	0,26	6,76
3	0,29	8,41
4	0,25	6,25
5	0,25	6,25
6	0,23	5,29
7	0,24	5,76
8	0,27	7,29
9	0,25	6,25
10	0,27	7,29
11	0,22	4,84
12	0,21	4,41
13	0,22	4,84
14	0,28	7,84
15	0,27	7,29
16	0,25	6,25
17	0,25	6,25
18	0,26	6,76
19	0,26	6,76
20	0,22	4,84
21	0,25	6,25

22	0,26	6,76
23	0,28	7,84
24	0,27	7,29
25	0,28	7,84
26	0,21	4,41
27	0,23	5,29
28	0,25	6,25
29	0,25	6,25
30	0,24	5,76
31	0,23	5,29
32	0,21	4,41
33	0,23	5,29
34	0,23	5,29
35	0,22	4,84

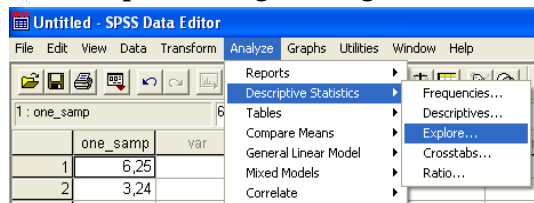
Keterangan:

- Yang dimaksud dengan koefisien korelasi adalah hasil hitung tentang korelasi skor IQ dan skor kesuksesan.
- Sedangkan koefisien determinasi adalah pengkuadratan skor koefisien korelasi dikalikan 100.

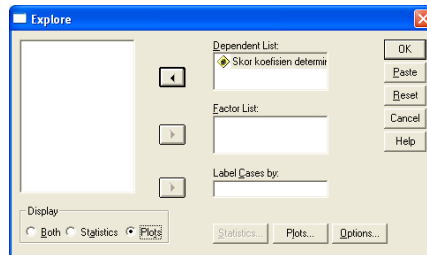
Karena t-test one sample termasuk parametrik, maka skor yang akan dianalisis harus diuji normalitasnya. Cara termudah untuk menguji ini adalah dengan program SPSS.

Prosedur untuk menguji distribusi data adalah sebagai berikut:

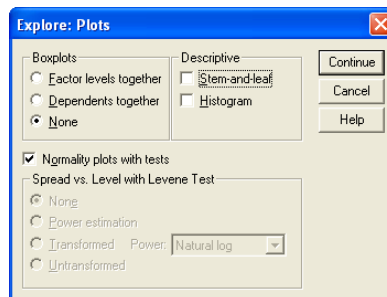
- Setelah data diinput, lalu klik **Analyze ► Descriptive Statistics ► Explore**, sebagaimana gambar berikut ini:



- Setelah itu akan keluar gambar berikut ini, destinasikan variabel itu ke dalam kotak **Dependent List**, lalu pada bagian **Display**, klik kotak **Plots**, Kemudian klik kotak **Plots**.



- Setelah itu akan keluar gambar berikut ini, aktifkan kotak **Normality Plots with tests**, lalu nonaktifkan pilihan **stem and leaf**, kemudian pilih **None** pada bagian **Boxplot**, Setelah itu klik **Ok**.



- Berikut ini adalah hasil analisis normalitas data tersebut.

Case Processing Summary

	Cases					
	Valid		Missing		Total	
	N	Percent	N	Percent	N	Percent
Skor koefisien determinasi	35	50,0%	35	50,0%	70	100,0%

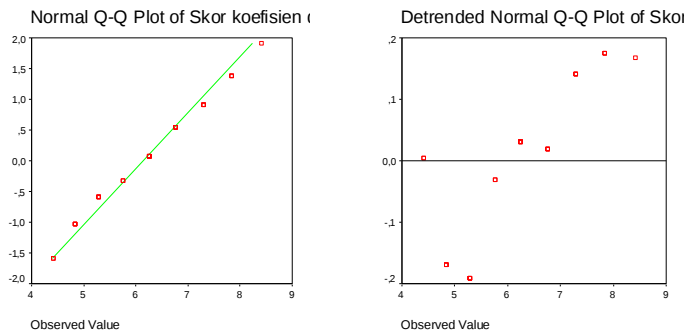
Tests of Normality

	Kolmogorov-Smirnov ^a			Shapiro-Wilk		
	Statistic	df	Sig.	Statistic	df	Sig.
Skor koefisien determinasi	,140	35	,081	,954	35	,151

a. Lilliefors Significance Correction

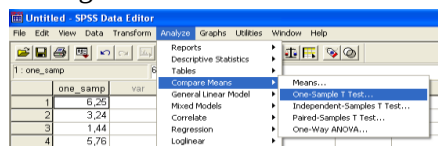
Untuk menguji normalitas dapat digunakan skor **Sig.** Yang ada pada hasil penghitungan **Kolmogorov-Smirnov**. Bila angka **Sig.** Lebih besar atau sama dengan 0,05, maka data tersebut berdistribusi normal, tetapi apabila kurang, maka data itu tidak berdistribusi normal. Karena **Sig.** Untuk variabel tersebut (0,810), itu lebih besar dari 0,05, maka data variabel itu berdistribusi normal.

Hal ini juga dapat dilihat pada grafik Normal Q-Q Plot maupun Detrended Normal Q-Q Plot. Untuk Normal Q-Q Plot itu apabila sebaran data dari variabel itu bergerombol di sekitar garis uji yang mengarah ke kanan atas dan tidak ada yang terletak jauh dari sebaran data, maka data tersebut berdistribusi normal. Sementara untuk Detrended Normal Q-Q Plot, apabila datanya tidak membentuk suatu pola tertentu atau menyebar secara acak, maka data itu berarti berdistribusi normal. Sesuai dengan nilai **Sig.** di atas, maka hasil uji dengan dua model grafik di bawah ini juga menunjukkan data dari variabel itu berdistribusi normal.

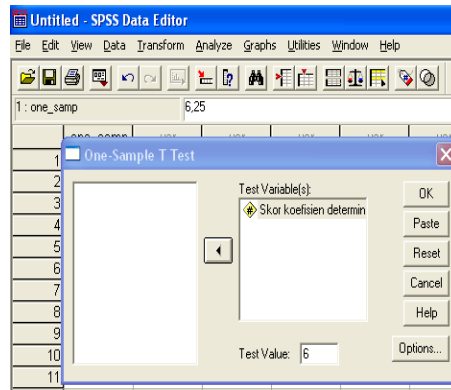


Dikarenakan data variabel itu distribusi datanya normal, maka dapat dilakukan analisis dengan t-test one sampel.

1. Untuk menganalisis menggunakan t-test one sample klik **Analyze ► Compare Means ► One-sample T Test**, sebagaimana gambar berikut ini:



2. Setelah keluar seperti berikut ini destinasikan variabelnya ke kotak **Test Variable(s)**, dan pada **Test Value** isilah angka **6** sebagaimana skor data dari hasil penelitian psikologis di Amerika bahwa rata-rata hanya berpengaruh 6% dan maksimal 20% terhadap kesuksesan seseorang, sebagaimana gambar berikut ini:



3. Berikut ini adalah hasil analisis di atas:

One-Sample Statistics

	N	Mean	Std. Deviation	Std. Error Mean
Skor koefisien determinasi	35	6,1411	1,09005	,18425

Output di atas memperlihatkan bahwa jumlah data yang dianalisis sebanyak 35. Rata-rata pengaruh IQ terhadap kesuksesan seseorang sebanyak 6,1411% dengan standard deviasi sebanyak 1,09005, dan standard error of mean sebanyak 0,18425. Dengan kata lain interval rata-ratanya antara 5,95985% sampai dengan 6,32835.

One-Sample Test

	Test Value = 6					
	t	df	Sig. (2-tailed)	Mean Difference	95% Confidence Interval of the Difference	
					Lower	Upper
Skor koefisien determinasi	,766	34	,449	,1411	-,2333	,5156

Output kedua ini memperlihatkan bahwa skor t-test_{hitung}nya sebesar 0,766. Bila dibandingkan dengan $t_{\text{tabel}; 0,05; 34}$ sebesar 2,032243, maka H_0 diterima dan H_a ditolak. Demikian juga kalau digunakan skor Signifikansinya sebesar 0,449 yang lebih tinggi dari tingkat kesalahan yang ditoleransi yaitu 0,05, maka H_0 diterima dan H_a ditolak. Oleh karena itu Hipotesis nihil yang berbunyi pengaruh IQ terhadap kesuksesan seseorang rata-rata hanya 6% berlaku untuk populasi. Atau dengan kata lain bahwa kecenderungan di Amerika bahwa pengaruh IQ sebesar 6% terhadap kesuksesan seseorang itu juga berlaku di Indonesia.

2. Aplikasi dengan Microsoft Excel

Sedangkan aplikasi Microsoft Excel dapat dilakukan dengan menggunakan rumus berikut ini.

$$t = \frac{\bar{X} - \mu_o}{\frac{s}{\sqrt{n}}}$$

t = Nilai t yang dihitung, selanjutnya disebut t hitung

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

μ_o = Nilai yang dihipotesiskan

s = Simpangan baku

n = Jumlah anggota sampel

Untuk dapat mengaplikasikan rumus tersebut, maka perlu diketahui skor standard deviasi. Langkah-langkah untuk mencari standar deviasi dapat dibaca pada Bab II halaman 73-80.

	A	B	C	D	E	F	G	H
1	NO	Koefisien Korelasi	Koefisien determinasi	$x_i - \text{mean}$	$(x_i - \text{mean})^2$			
2	1	0,25	6,25	0,10886	0,01185		aplikasi rumus	
3	2	0,26	6,76	0,61886	0,38298			
4	3	0,29	8,41	2,26886	5,14771		0,141143	
5	4	0,25	6,25	0,10886	0,01185		1,09005	
6	5	0,25	6,25	0,10886	0,01185		5,91608	
7	6	0,23	5,29	-0,85114	0,72444			
8	7	0,24	5,76	-0,38114	0,14527		0,766029	
9	8	0,27	7,29	1,14886	1,31987			
10	9	0,25	6,25	0,10886	0,01185			
11	10	0,27	7,29	1,14886	1,31987		aplikasi rumus	
12	11	0,22	4,84	-1,30114	1,69297		formula di sel G4	
13	12	0,21	4,41	-1,73114	2,99686		=C37-6	
14	13	0,22	4,84	-1,30114	1,69297		formula di sel G5	
15	14	0,28	7,84	1,69886	2,88612		=+E39	
16	15	0,27	7,29	1,14886	1,31987		formula di sel G6	
17	16	0,25	6,25	0,10886	0,01185		=SQRT(35)	
18	17	0,25	6,25	0,10886	0,01185		formula di sel G8	
19	18	0,26	6,76	0,61886	0,38298		=G4/(G5/G6)	
20	19	0,26	6,76	0,61886	0,38298			
21	20	0,22	4,84	-1,30114	1,69297			
22	21	0,25	6,25	0,10886	0,01185			
23	22	0,26	6,76	0,61886	0,38298			
24	23	0,28	7,84	1,69886	2,88612			
25	24	0,27	7,29	1,14886	1,31987			
26	25	0,28	7,84	1,69886	2,88612			
27	26	0,21	4,41	-1,73114	2,99686			
28	27	0,23	5,29	-0,85114	0,72444			
29	28	0,25	6,25	0,10886	0,01185			
30	29	0,25	6,25	0,10886	0,01185			
31	30	0,24	5,76	-0,38114	0,14527			
32	31	0,23	5,29	-0,85114	0,72444			
33	32	0,21	4,41	-1,73114	2,99686			
34	33	0,23	5,29	-0,85114	0,72444			
35	34	0,23	5,29	-0,85114	0,72444			
36	35	0,22	4,84	-1,30114	1,69297			
37			6,141142857		40,39935			
38					1,18822			
39					1,09005			

Seluruh hasil penghitungan dengan Microsoft Excel ternyata sama dengan hasil penghitungan SPSS. Bedanya SPSS mencantumkan skor signifikansi, sementara Microsoft Excel tidak. Oleh karena itu, untuk memutuskan apakah menerima H_0 ataukah H_a , peneliti yang menganalisis dengan menggunakan Microsoft Excel harus membandingkan t_{hitung} dengan t_{tabel} . Untuk mendapatkan t_{tabel} dapat dilakukan melalui menu function sebagai berikut:

=TINV(0,05,34)
2,032243

Keterangan:

- =TINV : formula untuk pencarian tabel t
- 0,05 : taraf nyata (α)
- 34 : derajat kebebasan (jumlah sampel (35) dikurangi dengan variabel (1))

Dikarenakan t_{hitung} (0,766029) lebih kecil dibanding t_{tabel} (2,032243), maka H_0 diterima dan H_a ditolak. Kesimpulan terakhir juga sama bila digunakan skor signifikansi yang ada pada output SPSS.

C. Statistik Non-parametris

Setidaknya ada 3 (tiga) teknik analisis yang dapat digunakan untuk menguji satu sampel, yaitu Test Binomial, Chi Kuadrat, dan Run Test.

1. Test Binomial

Test ini dikatakan sebagai test binomial karena distribusi data dalam populasi itu berbentuk binomial. Distribusi binomial adalah suatu distribusi yang terdiri dari dua kelas. Analisis binomial ini digunakan manakala datanya maksimal 25.

Contoh:

Seorang peneliti ingin mengetahui apakah anak kyai yang tidak memiliki pesantren ketika studi lebih memilih sekolah ataukah madrasah. Berdasarkan 25 sampel yang dipilih secara random ternyata 16 memilih lembaga pendidikan yang menggunakan nama sekolah sementara sisanya, yaitu 9 memilih madrasah.

Proses Pengambilan Keputusan.

Hipotesis:

- Ho Pilihan anak kyai yang tidak memiliki pesantren ketika studi memilih sekolah atau madrasah secara seimbang/sama.
- Ha Pilihan anak kyai yang tidak memiliki pesantren ketika studi memilih sekolah atau madrasah secara tidak seimbang/tidak sama.

Dasar Pengambilan Keputusan:

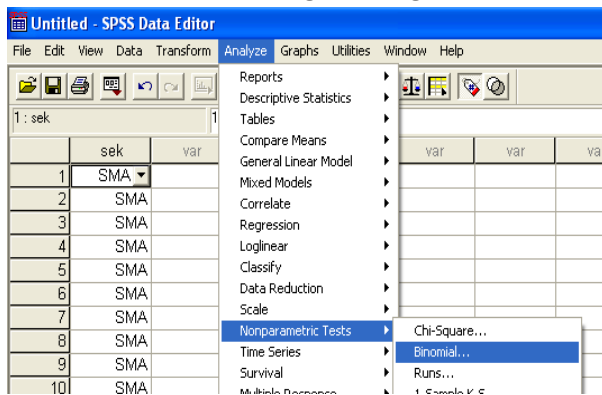
Dengan menggunakan angka probabilitas, dengan ketentuan:

- Ho diterima Probabilitas $>$ taraf nyata (α)
- Ho ditolak Probabilitas $\leq$ taraf nyata (α)

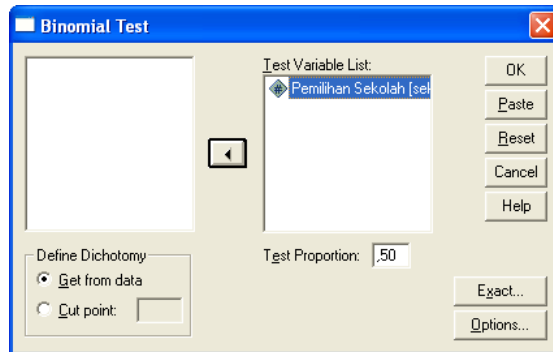
a. Aplikasi dengan SPSS

Untuk analisis binomial dapat dilakukan dengan langkah-langkah sebagai berikut:

1. Setelah data diinput, maka klik **Analyxe ► Nonparametric Test ► Binomial** sebagaimana gambar di bawah ini.



2. Setelah keluar seperti gambar berikut ini, destinasikan variabel tersebut menuju **Test Variable List**, lalu klik **Ok**.



3. Berikut ini adalah hasil analisis binomial.

Binomial Test						
		Category	N	Observed Prop.	Test Prop.	Exact Sig. (2-tailed)
Pemilihan Sekolah	Group 1	SMA	16	,64	,50	,230
	Group 2	MA	9	,36		
	Total		25	1,00		

Output di atas memperlihatkan bahwa anak kyai yang tidak memiliki pesantren mempunyai kecenderungan yang sama antara studi pada sekolah dan madrasah. Walaupun data dari sampel itu berbeda, yaitu 16 memilih SMA dan 9 memilih Madrasah, karena exact sig mempunyai skor 0,230 yang lebih tinggi dari 0,05, maka yang diterima berarti H_0 . Ini artinya kesamaan pilihan itu berlaku untuk populasi.

b. Aplikasi dengan Microsoft Excel

Analisis binomial ini tidak perlu menggunakan aplikasi Microsoft Excel. Untuk membuktikan H_0 dilakukan dengan cara membandingkan Harga x dalam test binomial dengan taraf kesalahan yang ditetapkan 1% (0,01) atau 5% (0,05). Dalam kasus ini dari 25 sampel, 16 memilih SMA dan 9 memilih MA. Frekwensi terkecilnya adalah (x) adalah 9. Berdasarkan tabel yang menjelaskan Harga-harga x dalam Test Binomial didapatkan skornya 0,115. Skor ini ternyata lebih tinggi dibandingkan 0,05 maupun 0,01. Oleh karena itu, maka H_0 diterima dan H_a ditolak. Kesimpulannya adalah Anak kyai yang tidak memiliki pesantren untuk studi memilih belajar di SMA dan MA adalah sama yaitu 50%.

2. Chi Kuadrat (χ^2)

Chi Kuadrat (χ^2) satu sampel adalah teknik analisis yang digunakan untuk menguji hipotesis bila dalam populasi terdiri

dari dua kelas atau lebih, data berbentuk nominal dan sampelnya besar.

Contoh:

Seorang peneliti ingin mengetahui apakah anak kyai yang memiliki pesantren ketika studi akan lebih memilih madrasah diniyah, madrasah sebagai sekolah umum yang berciri khas Islam, atau memilih sekolah. Berdasarkan 36 sampel yang dipilih secara random ternyata 15 memilih madrasah diniyah, 7 memilih madrasah sebagai sekolah umum yang berciri khas Islam, dan 14 memilih sekolah.

Proses Pengambilan Keputusan.

Hipotesis:

- | | |
|----|---|
| Ho | Anak kyai yang memiliki pesantren ketika studi akan lebih memilih madrasah diniyah, madrasah sebagai sekolah umum yang berciri khas Islam, atau sekolah secara seimbang/sama. |
| Ha | Anak kyai yang memiliki pesantren ketika studi akan lebih memilih madrasah diniyah, madrasah sebagai sekolah umum yang berciri khas Islam, atau sekolah secara tidak seimbang/tidak sama. |

Dasar Pengambilan Keputusan:

Dengan membandingkan χ^2_{hitung} dengan χ^2_{tabel} dengan ketentuan:

Ho diterima $\chi^2_{hitung} < \chi^2_{tabel}$

Ho ditolak $\chi^2_{hitung} \geq \chi^2_{tabel}$

Dengan menggunakan angka probabilitas, dengan ketentuan:

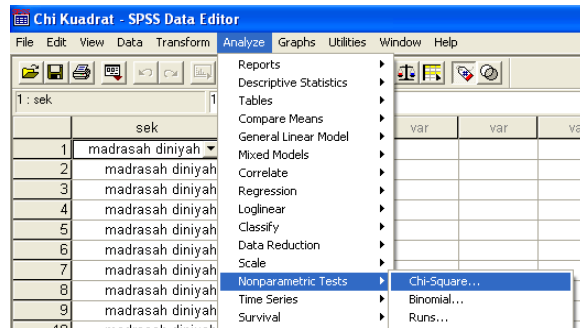
Ho diterima Probabilitas > taraf nyata (α)

Ho ditolak Probabilitas $\leq$ taraf nyata (α)

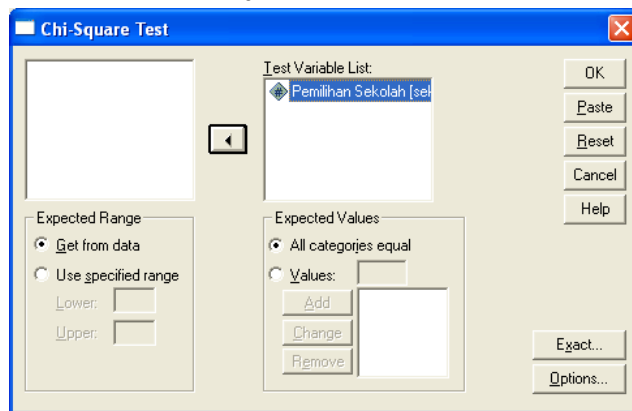
a. Aplikasi dengan SPSS

Untuk analisis Chi Kuadrat (χ^2) satu sampel dapat dilakukan dengan langkah-langkah sebagai berikut:

1. Setelah data diinput, maka klik **Analyze ► Nonparametric Test ► Chi-Square...** sebagaimana gambar di bawah ini.



2. Setelah keluar seperti gambar berikut ini, destinasikan variabel tersebut menuju **Test Variable List**, lalu klik **Ok**.



3. Berikut ini adalah hasil analisis Chi Kuadrat (χ^2) satu sampel.

Pemilihan Sekolah			
	Observed N	Expected N	Residual
madrasah diniyah	15	12,0	3,0
Madrasah (SUCI)	7	12,0	-5,0
sekolah	14	12,0	2,0
Total	36		

Test Statistics	
	Pemilihan Sekolah
Chi-Square ^a	3,167
df	2
Asymp. Sig.	,205

a. 0 cells (.0%) have expected frequencies less than 5. The minimum expected cell frequency is 12,0.

Output pertama memperlihatkan bahwa dari 36 anak kyai yang memiliki pesantren terdapat 15 anak yang memilih

madrasah diniyah sebagai tempat studi, 7 anak memilih madrasah sebagai sekolah umum yang berciri khas Islam, dan 14 di antaranya memilih sekolah.

Output kedua memperlihatkan skor chi kuadrat sebesar 3,167. Apabila skor ini dibandingkan dengan chi kuadrat tabel dengan dk 2 dan tingkat signifikansi 5% ternyata skornya 5,591. Dikarenakan chi kuadrat hitung lebih kecil dibanding chi kuadrat tabel, maka H_0 diterima dan H_a ditolak. Ini sesuai dengan Asymp. Sig sebesar 0,205 yang berarti lebih besar bila dibandingkan dengan peluang kesalahan yang 0,05. Berdasarkan hasil analisis tersebut maka dapat disimpulkan bahwa anak kyai pesantren memilih ketiga tipe pendidikan, madrasah diniyah, madrasah sebagai sekolah umum yang berciri khas Islam, dan sekolah secara berimbang atau sama.

b. Aplikasi dengan Microsoft Excel

Rumus Chi Kuadrat (χ^2) satu sampel adalah sebagai berikut:

$$\chi^2 = \sum_{i=1}^k \frac{(f_o - f_h)^2}{f_n}$$

χ^2 = Chi kuadrat

f_o = Frekuensi yang diobservasi

f_n = Frekuensi yang diharapkan

Untuk mengaplikasikan rumus di atas perlu dibuatkan tabel penolong sebagai berikut:

	A	B	C	D	E	F
1	Alternatif	f_o	f_h	$f_o - f_h$	$(f_o - f_h)^2$	$(f_o - f_h)^2 / f_h$
2	Madrasah Diniyah	15	12	3	9	0,750
3	Madrasah SUKI	7	12	-5	25	2,083
4	Sekolah	14	12	2	4	0,333
5	JUMLAH	36	36	0	38	3,167
6						
7	formula di C2	=36/3				
8	formula di D2	=B2-C2				
9	formula di E2	=D2^2				
10	formula di F2	=E2/C2				
11	formula di F5	=SUM(F2:F4)				
12						

Hasil hitung Chi Kuadrat (χ^2) satu sampel dengan Microsoft excel ternyata sama dengan hasil analisis yang menggunakan SPSS, yaitu 3,167.

3. Run Test

Run Test digunakan untuk menguji hipotesis deskriptif (satu sampel) bila datanya ordinal. Tujuan analisis ini adalah menguji kerandoman populasi yang didasarkan atas data sampel.

Contoh:

Seorang peneliti ingin mengetahui apakah kunci jawaban test yang menggunakan pilihan jawaban Benar (B) dan Salah (S) telah disusun secara random atau tidak. Peneliti mendapatkan data terhadap kunci jawaban sebagai berikut:

1	2	3	4	5	6	7	8	9	10
B	S	B	B	S	S	S	B	B	B

11	12	13	14	15	16	17	18	19	20
B	B	S	S	S	B	S	B	B	B

21	22	23	24	25	26	27	28	29	30
S	S	S	B	B	B	S	S	S	S

31	32	33	34	35	36	37	38	39	40
B	B	S	S	B	B	B	S	S	S

41	42	43	44	45	46	47	48	49	50
S	S	B	B	B	S	S	B	B	S

Proses Pengambilan Keputusan.

Hipotesis:

Ho Kunci jawaban test yang menggunakan pilihan jawaban Benar (B) dan Salah (S) telah disusun secara random.

Ha kunci jawaban test yang menggunakan pilihan jawaban Benar (B) dan Salah (S) tidak disusun secara random.

Dasar Pengambilan Keputusan:

Dengan membandingkan z_{hitung} dengan z_{tabel} dengan ketentuan:

Ho diterima $z_{hitung} < z_{tabel}$

Ho ditolak $z_{hitung} \geq z_{tabel}$

Dengan menggunakan angka probabilitas, dengan ketentuan:

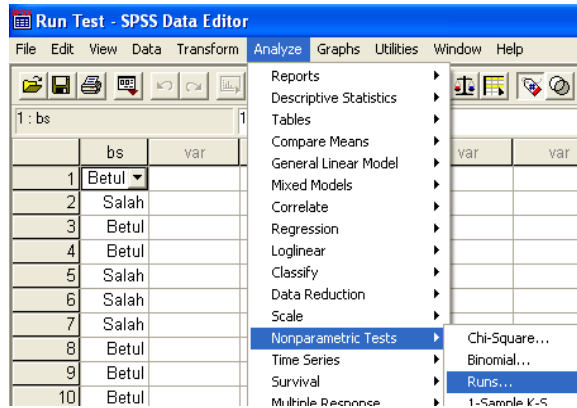
Ho diterima Probabilitas > taraf nyata (α)

Ho ditolak Probabilitas $\leq$ taraf nyata (α)

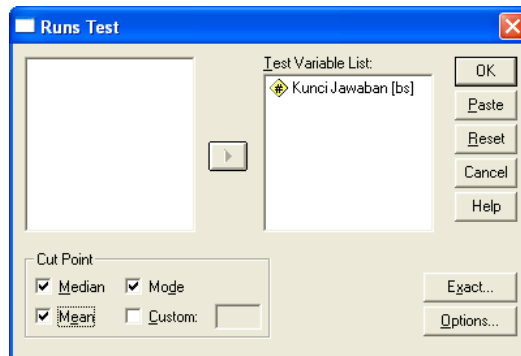
a. Aplikasi dengan SPSS

Untuk analisis Run Test dapat dilakukan dengan langkah-langkah sebagai berikut:

1. Setelah data diinput, maka klik Analyze ► Nonparametric Test ► Runs... sebagaimana gambar di bawah ini.



2. Setelah keluar seperti gambar berikut ini, destinasikan variabel tersebut menuju Test Variable List, lalu pada **Cut Point** dapat dipilih **Median, Mode, dan Mean**. Lalu klik **Ok**.



3. Berikut ini adalah hasil analisis Run Test satu sampel.

Runs Test

	Kunci Jawaban
Test Value ^a	2
Cases < Test Value	25
Cases >= Test Value	25
Total Cases	50
Number of Runs	20
Z	-1,715
Asymp. Sig. (2-tailed)	,086

a. Median

Runs Test 2

	Kunci Jawaban
Test Value ^a	1,50
Cases < Test Value	25
Cases >= Test Value	25
Total Cases	50
Number of Runs	20
Z	-1,715
Asymp. Sig. (2-tailed)	,086

a. Mean

Runs Test 3

	Kunci Jawaban
Test Value ^a	2 ^b
Cases < Test Value	25
Cases >= Test Value	25
Total Cases	50
Number of Runs	20
Z	-1,715
Asymp. Sig. (2-tailed)	,086

a. Mode

b. There are multiple modes. The mode with the largest data value is used.

Output di atas memperlihatkan bahwa test value (diambilkan dari nilai pemusatan data) apabila dihitung dengan median maka nilainya 2, kalau menggunakan mean maka nilainya 1,5, sedangkan kalau dihitung dengan mode maka nilainya 2. Skor yang lebih kecil dari test value 25 sedangkan yang sama dengan atau lebih besar dari test value juga 25. Sedangkan jumlah data ada 50 yang berganti antara kelompok satu dengan lainnya sebanyak 20 kali. Sedangkan cara menghitung jumlah run adalah sebagai berikut:

1	2	3	4	5	6	7	8	9	10
B	S	B	B	S	S	S	B	B	B
1	2	3		4			5		

11	12	13	14	15	16	17	18	19	20
B	B	S	S	S	B	S	B	B	B
		6			7	8	9		

21	22	23	24	25	26	27	28	29	30
S	S	S	B	B	B	S	S	S	S
10			11			12			

31	32	33	34	35	36	37	38	39	40
B	B	S	S	B	B	B	S	S	S
13		14		15			16		

41	42	43	44	45	46	47	48	49	50
S	S	B	B	B	S	S	B	B	S
		17			18		19		20

Sedangkan z hitung adalah -1,715. Apabila z hitung dibandingkan dengan z tabel pada Tabel I untuk uji dua pihak untuk tingkat kesalahan 5%, maka didapatkan z tabel 1,96. Cara untuk mengetahui z tabel adalah tingkat kepercayaan yang digunakan, misalkan 95% (0,95) dibagi dua menjadi 0,475. Skor tersebut ternyata terdapat pada 1,9 dan 0,06. Jadi kalau digabungkan menjadi 1,96. Karena z hitung lebih kecil dari z tabel, maka H_0 diterima dan H_a ditolak. Kesimpulan ini juga sama dengan skor Asymp. Sig. sebesar 0,086 yang lebih tinggi dari kesalahan yang ditoleransi yaitu paling tinggi 0,05. Dari analisis ini dapat disimpulkan bahwa pilihan jawaban **Benar (B)** dan **Salah (S)** telah disusun secara random.

b. Aplikasi dengan Microsoft Excel

Rumus Run Test satu sampel adalah sebagai berikut:

$$z = \frac{r - \mu_r}{\sigma_r} = \frac{r - \left(\frac{2n_1n_2}{n_1 + n_2} + 1 \right) - 0,5}{\sqrt{\frac{2n_1n_2(2n_1n_2 - n_1 - n_2)}{(n_1 + n_2)^2(n_1 + n_2 - 1)}}}$$

Harga (mean) μ_r dan simpangan baku σ_r dapat dihitung dengan rumus berikut :

$$\mu_r = \left(\frac{2n_1n_2}{n_1 + n_2} + 1 \right) - 0,5$$

$$\sigma_r = \sqrt{\frac{2n_1n_2(2n_1n_2 - n_1 - n_2)}{(n_1 + n_2)^2(n_1 + n_2 - 1)}}$$

Sedangkan aplikasi rumus di atas dengan Microsoft Excel sebagai berikut ini:

	A	B	C	D	E
1	run	20			
2	n1	25			
3	n2	25			
4					
5		26	formula di B5	=((2*B2*B3)/(B2+B3))+1	
6		-6,5	formula di B6	=B1-B5-0,5	
7					
8		1500000	formula di B8	=(2*B2*B3)*(2*B2*B3-B2-B3)	
9		2500	formula di B9	=(B2+B3)^2	
10		49	formula di B10	=B2+B3-1	
11		122500	formula di B11	=B9*B10	
12		12,245	formula di B12	=B8/B11	
13		3,499	formula di B13	=SQRT(B12)	
14					
15		-1,858	formula di B15	=B6/B13	

Hasil penghitungan dengan menggunakan Microsoft Excel ternyata tidak sama persis.¹ Walaupun tidak sama, tetapi selisihnya sangat sedikit sehingga tidak merubah kesimpulan akhirnya yaitu menerima H_0 dan menolak H_a . Jadi urutan jawaban antara yang Betul (B) dan Salah (S) telah tersusun secara random.

¹ Ketidaksamaan hasil hitung juga terjadi dalam bukunya Sugiyono, *Statistika untuk Penelitian*, (Bandung: CV Alfabeta, 2003), hlm. 112; bandingkan dengan bukunya Sugiyono dan Eri Wibowo, *Statistika Penelitian dan Aplikasinya dengan SPSS 10.00 for Windows*, (Bandung: CV Alfabeta, 2002), hlm. 94.

BAB IV

ANALISIS KORELASI

A. Pendahuluan

Dalam ilmu statistik istilah korelasi berarti hubungan antardua variabel atau lebih. Hubungan antardua variabel disebut *bivariate correlation*, sementara hubungan antarlebih dua variabel disebut *multivariate correlation*.

Hubungan antardua variabel misalnya korelasi antara intensitas mengikuti diskusi dosen (variabel x) dengan produktifitas kerja (variabel y). Hubungan antarlebih dari dua variabel misalnya korelasi antara kualitas layanan kepada peserta didik (variabel x_1), kompetensi tutor (variabel x_2), dengan banyaknya jumlah peserta didik (variabel y).

B. Arah Korelasi

Hubungan antardua variabel atau lebih itu bila dilihat dari arahnya dapat dibagi menjadi dua, yaitu hubungan yang sifatnya searah dan berlawanan arah. Hubungan searah disebut korelasi positif, sementara yang berlawanan arah disebut korelasi negatif.

Jadi, jika variabel x mengalami kenaikan, maka akan diikuti kenaikan variabel y . Itulah korelasi positif. Contohnya bila layanan terhadap mahasiswa naik (variabel x) maka naik pula jumlah mahasiswa perguruan tinggi itu (variabel y). Sementara korelasi negatif adalah apabila variabel x mengalami peningkatan mengakibatkan variabel y mengalami penurunan dan sebaliknya. Contohnya bila dominasi pimpinan kepada dosen meningkat (variabel x) maka produktivitas kerja dosen akan mengalami penurunan (variabel y).

C. Angka Korelasi

Besar angka korelasi itu berkisar antara 0 sampai 1, baik positif maupun negatif. Bila dalam penghitungan diperoleh angka korelasi lebih dari 1 berarti telah terjadi kesalahan penghitungan. Bila angka korelasi itu bertanda negatif menunjukkan korelasi antarvariabel itu negatif. Interpretasi kasar terhadap angka korelasi sebagai berikut:

Interval Koefisien	Tingkat Hubungan
0,00 – 0,199	Sangat rendah
0,20 – 0,399	Rendah
0,40 – 0,599	Sedang
0,60 – 0,799	Kuat
0,80 – 1,000	Sangat kuat

D. Macam-macam Teknik Korelasi

Setidaknya ada 8 teknik analisis korelasi sebagai berikut.

TIPE DATA	TEKNIK ANALISIS
NOMINAL	Contingency Coefficient
ORDINAL	Spearman Rank Correlation Kendall Tau
INTERVAL RASIO	Pearson Product Moment Partial Correlation Multiple Correlation

E. Statistik Parametrik

1. Product moment

Teknik korelasi ini digunakan untuk mencari hubungan dan membuktikan hipotesis hubungan dua variabel bila data kedua variabel berbentuk interval atau rasio. Karena product moment termasuk parametrik, maka harus memenuhi uji asumsi yaitu kedua variabel itu berdistribusi normal.

Contoh:

Seorang pemilik Kursus yang memiliki cabang di seluruh Indonesia berkeinginan untuk mengetahui apakah

Kompetensi Tutor berkorelasi dengan jumlah peserta didik. Oleh karena itu diambil sampel sebanyak 35 cabang kursus untuk ditelitinya.

Proses Pengambilan Keputusan.

Hipotesis:

Ho Kompetensi Tutor tidak berkorelasi dengan jumlah peserta didik

Ha Kompetensi Tutor berkorelasi dengan jumlah peserta didik

Dasar Pengambilan Keputusan:

Dengan membandingkan r_{hitung} dengan r_{tabel} dengan ketentuan:

Ho diterima $r_{hitung} < r_{tabel}$

Ho ditolak $r_{hitung} \geq r_{tabel}$

Dengan menggunakan angka probabilitas, dengan ketentuan:

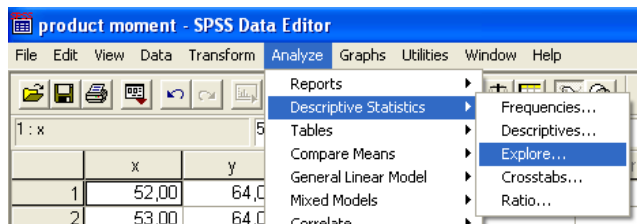
Ho diterima Probabilitas $>$ taraf nyata (α)

Ho ditolak Probabilitas $\leq$ taraf nyata (α)

Karena product moment termasuk parametrik, maka analisis ini dapat digunakan apabila tipe datanya interval atau rasio, distribusi datanya normal, dan jumlah sampelnya lebih dari 30. Dalam contoh ini, tipe datanya interval dan berjumlah 35, maka asumsi yang harus diuji adalah distribusi datanya. Cara termudah untuk menguji ini adalah dengan program SPSS.

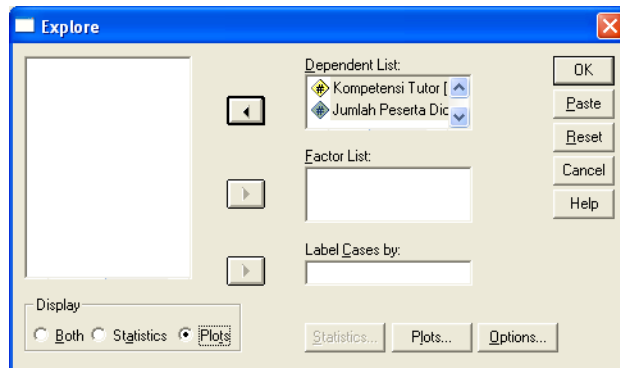
Prosedur untuk menguji distribusi data adalah sebagai berikut:

1. Setelah data dimasukkan klik **Analyze ► Descriptive Statistics ► Explore**, sebagaimana gambar berikut ini:

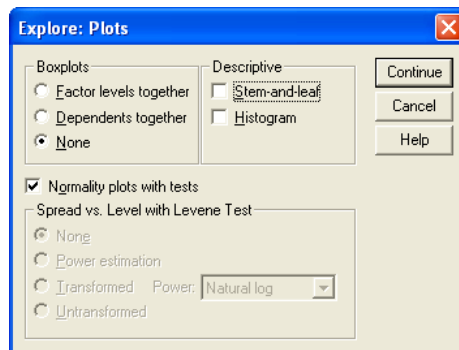


2. Setelah itu akan keluar gambar berikut ini. Destinaskan kedua variabel itu ke dalam kotak **Dependent List**, dengan memblok kedua variabel itu dan mengeklik panah yang ada di sebelah kiri

kotak itu. Pada bagian **Display**, klik kotak **Plots**. Kemudian klik kotak **Plots**.



3. Setelah itu akan keluar gambar seperti berikut ini. Pilih **None** pada bagian **Boxplot**. Nonaktifkan pilihan **stem and leaf**. Aktifkan kotak **Normality Plots with tests**. Setelah itu klik **Ok**



4. Berikut ini adalah **output** uji normalitas data.

Case Processing Summary

	Cases					
	Valid		Missing		Total	
	N	Percent	N	Percent	N	Percent
Kompetensi Tutor	35	100,0%	0	,0%	35	100,0%
Jumlah Peserta Didik Baru	35	100,0%	0	,0%	35	100,0%

Tests of Normality

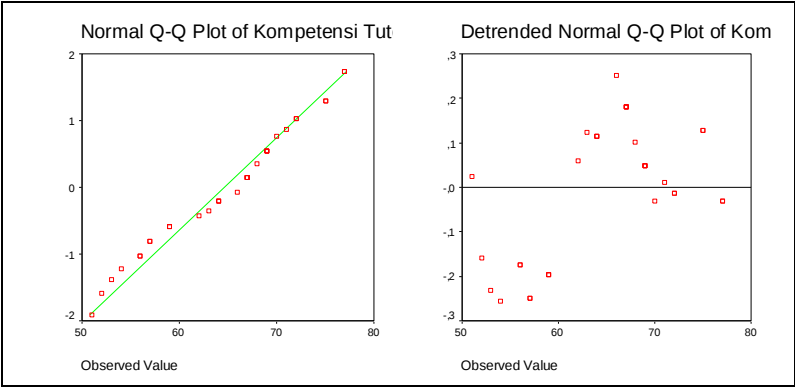
	Kolmogorov-Smirnov ^a			Shapiro-Wilk		
	Statistic	df	Sig.	Statistic	df	Sig.
Kompetensi Tutor	,140	35	,082	,961	35	,241
Jumlah Peserta Didik Baru	,118	35	,200*	,945	35	,079

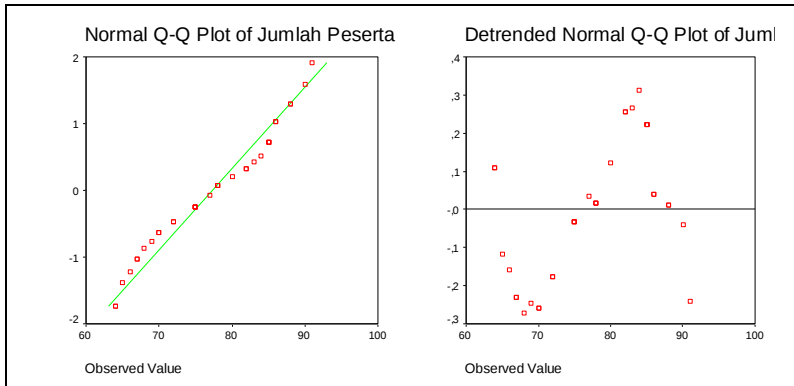
*. This is a lower bound of the true significance.

a. Lilliefors Significance Correction

Untuk mengetahui normalitas dapat digunakan skor **Sig.** yang ada pada hasil penghitungan **Kolmogorov-Smirnov**. Bila angka **Sig.** Lebih besar atau sama dengan 0,05, maka berdistribusi normal, tetapi apabila kurang, maka data tidak berdistribusi normal. Karena **Sig.** Untuk variabel x (0,082), dan variabel y (0,200) itu lebih besar dari 0,05, maka kedua data variabel itu berdistribusi normal.

Hal ini juga dapat dilihat pada grafik Normal Q-Q Plot maupun Detrended Normal Q-Q Plot. Untuk Normal Q-Q Plot itu apabila sebaran data dari variabel itu bergerombol di sekitar garis uji yang mengarah ke kanan atas dan tidak ada yang terletak jauh dari sebaran data, maka data tersebut berdistribusi normal. Sementara untuk Detrended Normal Q-Q Plot, apabila datanya tidak membentuk suatu pola tertentu atau menyebar secara acak, maka data itu berarti berdistribusi normal. Sesuai dengan nilai **Sig.** di atas, maka hasil uji dengan dua model grafik di bawah ini juga menunjukkan kedua data dari variabel itu berdistribusi normal.



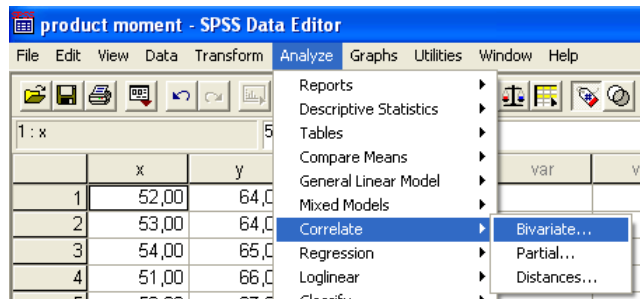


Dikarenakan kedua variabel itu distribusi datanya normal, maka dapat dilakukan analisis dengan product moment. Seandainya distribusi datanya tidak normal, maka harus menggunakan analisis non-parametrik, seperti kendall's tau atau spearman rank.

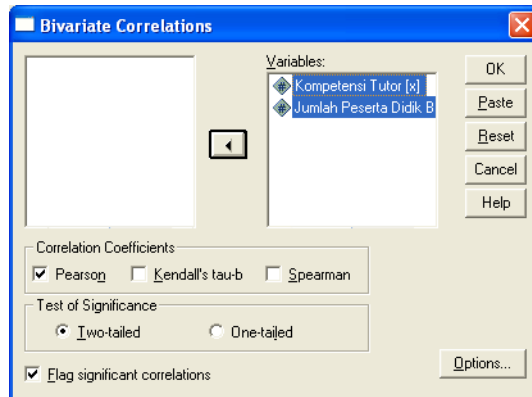
a. Aplikasi dengan SPSS

Prosedur untuk menguji hipotesis dengan product moment adalah sebagai berikut:

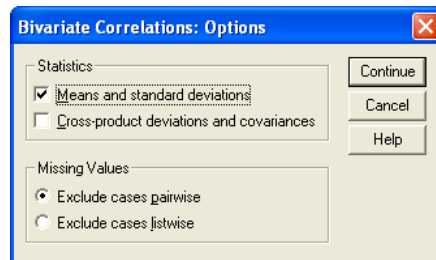
1. Ketika data yang diinput pada data view, klik **Analyze** ➤ **Correlate** ➤ **Bivariate**, sebagaimana gambar di bawah ini:



2. Ketika sudah keluar seperti gambar berikut ini, destinasikan kedua variabel itu ke dalam kotak **Variables**. Kalau berkeinginan untuk menampilkan rata-rata dan deviasi standar, klik **Options**, maka akan keluar gambar di bawah ini.:



3. Ketika sudah keluar seperti gambar berikut ini, klik pada kotak yang ada di depan **Means and** Setelah itu klik **Continue**. Abaikan yang lain dan klik **Ok**.



4. Berikut ini adalah **Output** proses analisis di atas.

Descriptive Statistics

	Mean	Std. Deviation	N
Kompetensi Tutor	64,6857	7,23867	35
Jumlah Peserta Didik Baru	77,2857	8,19100	35

Correlations

		Kompetensi Tutor	Jumlah Peserta Didik Baru
Kompetensi Tutor	Pearson Correlation	1	,949**
	Sig. (2-tailed)	.	,000
	N	35	35
Jumlah Peserta Didik Baru	Pearson Correlation	,949**	1
	Sig. (2-tailed)	,000	.
	N	35	35

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

Dari output di atas, diketahui bahwa jumlah sampel dari penelitian ini adalah **35**. Skor rata-rata variabel x adalah **64,6857** dengan deviasi standar **7,23867**. Sementara untuk

skor rata-rata variabel y adalah **77,2857** dengan deviasi standar **8,19100**.

Adapun nilai koefisien korelasi antara variabel x dengan variabel y adalah **0,949**, dengan nilai **Sig**nya 0,000 dalam arti kesalahan menolak H_0 hanyalah 0% atau mendekati 0%. Hasil pada **Sig.** itu dapat dicek ulang dengan membanding r_{hitung} **0,949** dengan r_{tabel} untuk dk.: 35 (jumlah sampel) dikurangi 2 (jumlah variabel) = 33. Nilai tabel untuk 33 dengan kesalahan 5%: **0,344** dan 1%: **0,442**. Karena r_{hitung} lebih besar dari r_{tabel} , maka berarti H_0 ditolak dan H_a diterima. Ini mengandung pengertian bahwa kesimpulan dari sampel ini dapat digeneralisasikan untuk populasi. Dan karena r_{hitung} tidak bertanda negatif, maka menunjukkan arah korelasi ini positif. Jadi, semakin tinggi kompetensi tutor semakin tinggi pula jumlah siswa baru kursus tersebut.

b. Aplikasi dengan Microsoft Excel

Setidaknya ada dua rumus product moment yaitu:

$$r_{xy} = \frac{\sum xy}{\sqrt{(\sum x^2)(\sum y^2)}}$$

Dimana :

r_{xy} = Korelasi antara variabel x dan y

x = $(X_i - \bar{X})$

y = $(Y_i - \bar{Y})$

$$r_{xy} = \frac{n \sum x_i y_i - (\sum x_i)(\sum y_i)}{\sqrt{\left\{n \sum x_i^2 - (\sum x_i)^2\right\} \left\{n \sum y_i^2 - (\sum y_i)^2\right\}}}$$

Rumus terakhir digunakan bila sekaligus akan menghitung persamaan regresi.

Aplikasi sesuai dengan rumus di atas dilakukan dengan langkah-langkah sebagai berikut:

1. Setelah data dimasukkan, baik untuk variabel x maupun variabel y, buatlah rata-rata dari masing-masing variabel itu.
2. Kurangi masing-masing skor responden dengan nilai rata-ratanya.

3. Kuadratkan hasil pengurangan itu.
4. Jumlahkan hasil pengkuadratkan itu.
5. Untuk mengetahui varians untuk sampel, hasil penjumlahan no.4 dibagi dengan jumlah sampel dikurangi 1.
6. untuk mengetahui deviasi standar, Carilah akar varians di atas.
7. Kalikan hasil pengurangan dari masing-masing reponden antara variabel x dan y.
8. Jumlahkan skor hasil dari no. 7 di atas.
9. Setelah itu masukkan ke dalam rumus.

$$r_{xy} = \frac{\sum xy}{\sqrt{(\sum x^2 y^2)}}$$

Prosedur di atas dapat dilihat pada aplikasi Microsoft Excel berikut ini.

.....											
A	B	C	D	E	F	G	H	I	J	K	L
1	KOMPETENSI TUTOR	JML MURID BARU	x	x ²	y	y ²	xy				
2	52	64	-12,686	160,927	-13,286	176,510	168,539				
3	53	64	-11,686	136,556	-13,286	176,510	155,253				
4	54	65	-10,686	114,184	-12,286	150,939	131,282				
5	51	66	-13,686	187,299	-11,286	127,367	154,453				
6	56	67	-8,686	75,442	-10,286	105,796	89,339				
7	56	67	-8,686	75,442	-10,286	105,796	89,339				
8	57	68	-7,686	59,070	-9,286	86,224	71,367				
9	62	69	-2,686	7,213	-8,286	68,653	22,253				
10	59	70	-5,686	32,327	-7,286	53,082	41,424				
11	59	70	-5,686	32,327	-7,286	53,082	41,424				
12	57	72	-7,686	59,070	-5,286	27,939	40,524				
13	59	72	-5,686	32,327	-5,286	27,939	30,053				
14	64	75	-0,686	0,470	-2,286	5,224	1,567				
15	64	75	-0,686	0,470	-2,286	5,224	1,567				
16	63	75	-1,686	2,842	-2,286	5,224	3,853				
17	67	75	2,314	5,356	-2,286	5,224	-5,290				
18	66	77	1,314	1,727	-0,286	0,082	-0,376				
19	67	78	2,314	5,356	0,714	0,510	1,653				
20	67	78	2,314	5,356	0,714	0,510	1,653				
21	64	78	-0,686	0,470	0,714	0,510	-0,490				
22	72	80	7,314	53,499	2,714	7,367	19,853				
23	67	82	2,314	5,356	4,714	22,224	10,910				
24	67	82	2,314	5,356	4,714	22,224	10,910				
25	68	83	3,314	10,984	5,714	32,653	18,939				
26	69	84	4,314	18,613	6,714	45,082	28,967				
27	70	85	5,314	28,242	7,714	59,510	40,996				
28	75	85	10,314	106,384	7,714	59,510	79,567				
29	69	85	4,314	18,613	7,714	59,510	33,282				
30	69	85	4,314	18,613	7,714	59,510	33,282				
31	71	86	6,314	39,870	8,714	75,939	55,024				
32	72	86	7,314	53,499	8,714	75,939	63,739				
33	77	88	12,314	151,642	10,714	114,796	131,939				
34	69	88	4,314	18,613	10,714	114,796	46,224				
35	75	90	10,314	106,384	12,714	161,853	131,139				
36	77	91	12,314	151,642	13,714	188,082	168,882				
37	64,68671429	77,28571429		1781,543		2281,143	1913,143				
38				52,398		67,092					
39				7,239		8,191					
40											

Ternyata hasilnya sama persis dengan penghitungan dengan SPSS.

Apabila kita menggunakan rumus 2, maka aplikasinya sebagaimana berikut:

	A	B	C	D	E	F	G	H	I	J	K	L
	x	y	x ²	y ²	xy		APLIKASI RUMUS					
1	52	64	2.704	4.096	3.328							
2	53	64	2.809	4.096	3.392							
3	54	65	2.916	4.225	3.510							
4	51	66	2.601	4.356	3.366							
5	56	67	3.136	4.489	3.752							
6	56	67	3.136	4.489	3.752							
7	56	67	3.136	4.489	3.752							
8	57	68	3.249	4.624	3.876							
9	62	69	3.844	4.761	4.278							
10	59	70	3.481	4.900	4.130							
11	59	70	3.481	4.900	4.130							
12	57	72	3.249	5.184	4.104							
13	59	72	3.481	5.184	4.248							
14	64	75	4.096	5.625	4.800							
15	64	75	4.096	5.625	4.800							
16	63	75	3.969	5.625	4.725							
17	67	75	4.489	5.625	5.025							
18	66	77	4.356	5.929	5.082							
19	67	78	4.489	6.084	5.226							
20	67	78	4.489	6.084	5.226							
21	64	78	4.096	6.084	4.992							
22	72	80	5.184	6.400	5.760							
23	67	82	4.489	6.724	5.494							
24	67	82	4.489	6.724	5.494							
25	68	83	4.624	6.889	5.644							
26	69	84	4.761	7.056	5.796							
27	70	85	4.900	7.225	5.950							
28	75	85	5.625	7.225	6.375							
29	69	85	4.761	7.225	5.865							
30	69	85	4.761	7.225	5.865							
31	71	86	5.041	7.396	6.106							
32	72	86	5.184	7.396	6.192							
33	77	88	5.929	7.744	6.776							
34	69	88	4.761	7.744	6.072							
35	75	90	5.625	8.100	6.750							
36	77	91	5.929	8.281	7.007							
37	2.264	2.705	148.230	211.339	176.888							
38												

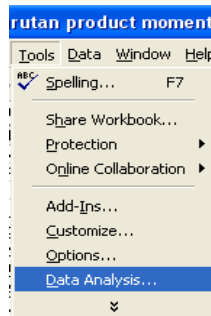
Caranya:

1. Setelah data dimasukkan, baik untuk variabel x maupun variabel y, jumlahkan masing-masing variabel itu.
2. Kuadratkan masing-masing skor yang ada di variabel x.
3. Jumlahkan hasil pengkuadratkan itu.
4. Kuadratkan masing-masing skor yang ada di variabel y.

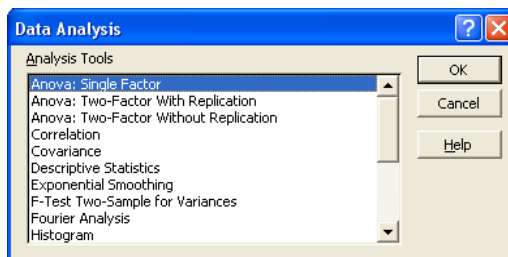
5. Jumlahkan hasil pengkuadratkan itu.
6. Kalikan masing-masing skor antara variabel x dan y.
7. Jumlahkan skor hasil pengkalian di atas.
8. Setelah itu masukkan ke dalam rumus.

$$r_{xy} = \frac{n \sum x_i y_i - (\sum x_i)(\sum y_i)}{\sqrt{\{n \sum x_i^2 - (\sum x_i)^2\} \{n \sum y_i^2 - (\sum y_i)^2\}}}$$

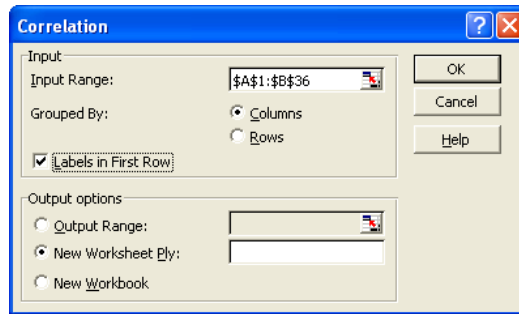
Sedangkan cara yang paling cepat untuk menguji hipotesis asosiatif yang datanya interval/rasio dan berdistribusi normal adalah menggunakan **Data Analysis** dari menu **Tools**, sebagaimana ada pada gambar di bawah ini:



1. Ketika keluar gambar sebagaimana di bawah ini klik **Correlation** dan klik **Ok.**



2. Selah keluar gambar di bawah ini, maka bloklah seluruh skor variabel x dan variabel y. Klik **Label in First Row** kalau label dari variabel akan ditampilkan. Lalu klik **New Worksheet Ply** agar hasil analisis nanti ditampilkan pada worksheet tersendiri. Selanjutnya klik **Ok.**



3. Inilah hasil dari penghitungan itu, ternyata hasilnya, manakala di belakang koma dijadikan tiga angka maka skornya sama persis dengan penghitungan SPSS dan Microsoft Excel di atas.

	A	B	C
1		x	y
2	x	1	
3	y	0,949	1
4			

2. Korelasi Ganda

Korelasi ganda merupakan angka yang menunjukkan arah dan kuatnya hubungan antara dua variabel atau lebih secara bersama-sama dengan variabel yang lain.

Contoh:

Setelah mengetahui korelasi antara kompetensi tutor dengan jumlah siswanya, Seorang pemilik tersebut ingin meneliti korelasi antara kualitas pelayanan kepada siswa dan kompetensi tutor secara bersama-sama dengan jumlah peserta didik.

Proses Pengambilan Keputusan.

Hipotesis:

- Ho Tidak ada korelasi antara kualitas pelayanan kepada siswa dan kompetensi tutor secara bersama-sama dengan jumlah peserta didik
- Ha Terdapat korelasi antara kualitas pelayanan kepada siswa dan kompetensi tutor secara bersama-sama dengan jumlah peserta didik

Dasar Pengambilan Keputusan:

Dengan membandingkan r_{hitung} dengan r_{tabel} dengan ketentuan:

Ho diterima $r_{hitung} < r_{tabel}$

Ho ditolak $r_{hitung} \geq r_{tabel}$

Dengan menggunakan angka probabilitas, dengan ketentuan:

Ho diterima Probabilitas > taraf nyata (α)

Ho ditolak Probabilitas $\leq$ taraf nyata (α)

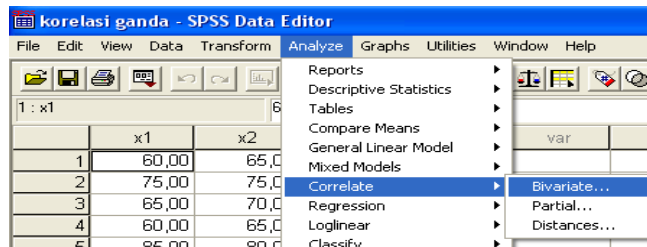
a. Aplikasi dengan SPSS

Data ketiga variabel itu dimasukkan pada data view SPSS terlihat sebagai berikut:

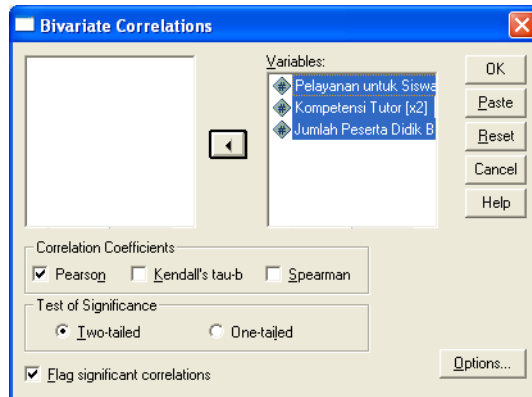
	x1	x2	y
1	52,00	52,00	64,00
2	62,00	53,00	64,00
3	56,00	54,00	65,00
4	76,00	51,00	66,00
5	62,00	56,00	67,00
6	66,00	56,00	67,00
7	56,00	57,00	68,00
8	59,00	62,00	69,00
9	58,00	59,00	70,00
10	65,00	59,00	70,00
11	79,00	57,00	72,00
12	82,00	59,00	72,00
13	66,00	64,00	75,00
14	67,00	64,00	75,00
15	76,00	63,00	75,00
16	79,00	67,00	75,00
17	97,00	66,00	77,00
18	68,00	67,00	78,00
19	68,00	67,00	78,00
20	69,00	64,00	78,00
21	76,00	72,00	80,00
22	58,00	67,00	82,00
23	72,00	67,00	82,00
24	76,00	68,00	83,00
25	89,00	69,00	84,00
26	74,00	70,00	85,00
27	78,00	75,00	85,00
28	85,00	69,00	85,00
29	92,00	69,00	85,00
30	76,00	71,00	86,00
31	99,00	72,00	86,00
32	85,00	77,00	88,00
33	86,00	69,00	88,00
34	99,00	75,00	90,00
35	89,00	77,00	91,00

Prosedur untuk menguji hipotesis dengan korelasi ganda adalah sama dengan aplikasi product moment:

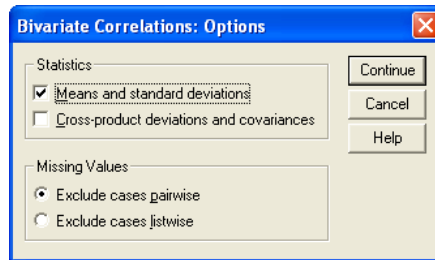
1. Ketika data sudah diinput pada data view klik **Analyze** ➤ **Correlate** ➤ **Bivariate**, sebagaimana gambar di bawah ini:



2. Setelah itu keluar gambar seperti di bawah ini. Destinaskan ketiga variabel itu ke dalam kotak **Variables**. Kalau berkeinginan untuk menampilkan rata-rata dan deviasi standar, klik **Options**.



3. Setelah itu klik pada kotak yang ada di depan **Means and**
Setelah itu klik **Continue**. Abaikan yang lain dan klik **Ok**.



4. Berikut ini adalah hasil analisis di atas.

Descriptive Statistics

	Mean	Std. Deviation	N
Pelayanan untuk Siswa	74,2000	12,83790	35
Kompetensi Tutor	64,6857	7,23867	35
Jumlah Peserta Didik Baru	77,2857	8,19100	35

Correlations

		Pelayanan untuk Siswa	Kompetensi Tutor	Jumlah Peserta Didik Baru
Pelayanan untuk Siswa	Pearson Correlation	1	,650**	,714**
	Sig. (2-tailed)	.	,000	,000
	N	35	35	35
Kompetensi Tutor	Pearson Correlation	,650**	1	,949**
	Sig. (2-tailed)	,000	.	,000
	N	35	35	35
Jumlah Peserta Didik Baru	Pearson Correlation	,714**	,949**	1
	Sig. (2-tailed)	,000	,000	.
	N	35	35	35

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

Dari output di atas, dapat diketahui bahwa jumlah sampel dari penelitian ini adalah 35. Skor rata-rata variabel x_1 adalah **74,2000** dengan deviasi standar **12,83790**, Skor rata-rata variabel x_2 adalah **64,6857** dengan deviasi standar **7,23867**. Sementara untuk skor rata-rata variabel y adalah **77,2857** dengan deviasi standar **8,19100**.

Adapun nilai koefisien korelasi antara variabel x_1 dengan variabel y adalah **0,714**, nilai koefisien korelasi antara variabel x_2 dengan variabel y adalah **0,949**, dan nilai koefisien korelasi antara variabel x_1 dengan variabel x_2 adalah **0,650**. Nilai **Sig** untuk ketiga korelasi itu adalah 0,000 dalam arti kesalahan menolak H_0 hanyalah 0% atau mendekati 0%. Hasil pada **Sig.** itu dapat dicek ulang dengan membanding r_{hitung} dengan r_{tabel} untuk dk.: 35 (jumlah sampel) dikurangi 2 (jumlah variabel) = 33. Nilai tabel untuk 33 dengan kesalahan 5%: **0,344** dan 1%: **0,442**. Karena r_{hitung} lebih besar dari r_{tabel} , maka berarti H_0 ditolak dan H_a diterima. Ini mengandung pengertian bahwa kesimpulan dari sampel ini dapat digeneralisasikan untuk populasi. Dan karena r_{hitung} tidak bertanda negatif, maka menunjukkan arah korelasi ini positif. Jadi, semakin tinggi pelayanan untuk siswa maka semakin tinggi pula jumlah siswanya, demikian juga semakin tinggi kompetensi tutor maka semakin tinggi pula jumlah siswanya. Sayangnya output itu tidak memperlihatkan skor korelasi kedua variabel x secara bersama-sama dengan variabel y . Untuk dapat mengetahui korelasi kedua variabel x secara bersama-sama dengan variabel y dapat dilakukan dengan Microsoft Excel.

b. Aplikasi dengan Microsoft Excel

Rumus Korelasi Ganda adalah sebagai berikut:

$$r_{y \cdot x_1 x_2} = \sqrt{\frac{r_{yx_1}^2 + r_{yx_2}^2 - 2r_{yx_1} r_{yx_2} r_{x_1 x_2}}{1 - r_{x_1 x_2}^2}}$$

Untuk dapat mengaplikasikan rumus di atas terlebih dahulu harus dicari koefisien korelasi antar variabel yang caranya sama dengan pencarian koefisien korelasi product moment sebagaimana telah dipraktikkan di atas. Setelah masing-masing koefisien korelasi antar variabel

diketemukan, maka dapat dimasukkan ke dalam rumus korelasi ganda di atas. prosedurnya aplikasi Microsoft Excel pada halaman berikutnya.

Hasil penghitungan dengan Microsoft Excel menghasilkan skor yang sama dengan skor koefisien korelasi hasil penghitungan SPSS. Kelebihannya ia juga bisa digunakan mencari skor korelasi kedua variabel x secara bersama-sama dengan variabel y , yaitu **0,958**. Kalau kita membandingkan r_{hitung} dengan r_{tabel} untuk dk.: 35 (jumlah sampel) dikurangi 3 (jumlah variabel) = 32. Nilai tabel untuk 32 dengan kesalahan 5%: **0,349** dan 1%: **0,449**. Karena r_{hitung} lebih besar dari r_{tabel} , maka berarti H_0 ditolak dan H_a diterima. Ini mengandung pengertian bahwa kesimpulan dari sampel ini dapat digeneralisasikan untuk populasi.

Sedangkan cara yang paling cepat adalah menggunakan **Data Analysis** dari menu **Tools** ternyata mempunyai hasil yang sama.

	A	B	C	D	
1		x^2	x^2	y	
2	x^2	1			
3	x^2	0,650	1		
4	y	0,714	0,949	1	
5					

	A	B	C	D	E	F	G	H	I	J	K	L	M	N	O
1	x ₁	x ₂	y	x ₁ ²	x ₂ ²	y ²	x ₁ y	x ₂ y	x ₁ x ₂	APLIKASI RUMUS					
2	52	52	64	2.704	2.704	4.096	3.328	3.328	2.704	r _{x1y}					
3	62	53	64	3.844	2.809	4.096	3.968	3.392	3.266	= (35*G37)-(A37*C37))					
4	56	54	65	3.136	2.916	4.225	3.640	3.510	3.024	= (35*D37)-(A37*A37))*((35*F37)-(C37*C37))					
5	76	51	66	5.776	2.601	4.356	5.016	3.366	3.876	= SORT(J4)					
6	62	56	67	3.844	3.136	4.489	4.154	3.752	3.472	=J3,J5					
7	66	56	67	4.356	3.136	4.489	4.422	3.752	3.696	r _{x2y}					
8	56	57	68	3.136	3.249	4.624	3.808	3.876	3.192	= (35*H37)-(B37*C37))					
9	59	62	69	3.481	3.844	4.761	4.071	4.278	3.658	= (35*F37)-(C37*C37))*((35*E37)-(B37*B37))					
10	58	59	70	3.364	3.481	4.900	4.060	4.130	3.422	= SORT(J10)					
11	65	59	70	4.225	3.481	4.900	4.550	4.130	3.896	=J9,J11					
12	79	57	72	6.241	3.249	5.184	5.688	4.104	4.503	r _{x1x2}					
13	82	59	72	6.724	3.481	5.184	5.904	4.248	4.898	= (35*I37)-(A37*B37))					
14	66	64	75	4.356	4.096	5.625	4.950	4.800	4.224	= (35*D37)-(A37*A37))*((35*E37)-(B37*B37))					
15	67	64	75	4.489	4.096	5.625	5.025	4.800	4.288	= SORT(J16)					
16	76	63	75	5.776	3.969	5.625	5.700	4.725	4.788	=J15,J17					
17	79	67	75	6.241	4.489	5.625	5.925	5.025	5.293	r _{x1x2y}					
18	97	66	77	9.409	4.356	5.929	7.469	5.082	6.402	= (J6*J6)+(J12*J12)-2*(J6*J12*J18)					
19	68	67	78	4.624	4.489	6.084	5.304	5.226	4.556	=1-(J18*J18)					
20	68	67	78	4.624	4.489	6.084	5.304	5.226	4.556	=J21,J22					
21	69	64	78	4.761	4.096	6.084	5.362	4.992	4.416	= SORT(J23)					
22	76	72	80	5.776	5.184	6.400	6.080	5.760	5.472	r _{yx1} + r _{yx2} - 2r _{yx1} r _{yx2} r _{x1x2}					
23	58	67	82	3.364	4.489	6.724	4.756	5.494	3.866	1 - r _{x1x2}					
24	72	67	82	5.184	4.489	6.724	5.904	4.824	4.824	r _{yx1} + r _{yx2} - 2r _{yx1} r _{yx2} r _{x1x2}					
25	76	68	83	5.776	4.624	6.889	6.308	5.644	5.168	1 - r _{x1x2}					
26	69	69	84	7.921	4.761	7.056	7.476	5.796	6.141	r _{yx1} + r _{yx2} - 2r _{yx1} r _{yx2} r _{x1x2}					
27	74	70	85	5.476	4.900	7.225	6.290	5.950	5.180	1 - r _{x1x2}					
28	78	75	85	6.084	5.625	7.225	6.630	6.375	5.850	r _{yx1} + r _{yx2} - 2r _{yx1} r _{yx2} r _{x1x2}					
29	85	69	85	7.225	4.761	7.225	7.225	5.865	5.865	1 - r _{x1x2}					
30	92	69	85	8.464	4.761	7.225	7.820	5.865	6.348	r _{yx1} + r _{yx2} - 2r _{yx1} r _{yx2} r _{x1x2}					
31	76	71	86	5.776	5.041	7.396	6.536	6.106	5.396	1 - r _{x1x2}					
32	99	72	86	9.801	5.184	7.396	8.514	6.192	7.128	r _{yx1} + r _{yx2} - 2r _{yx1} r _{yx2} r _{x1x2}					
33	85	77	88	7.225	5.929	7.744	7.480	6.776	6.545	1 - r _{x1x2}					
34	86	69	88	7.396	4.761	7.744	7.568	6.072	5.934	r _{yx1} + r _{yx2} - 2r _{yx1} r _{yx2} r _{x1x2}					
35	99	75	90	9.801	5.625	8.100	9.910	6.750	7.425	1 - r _{x1x2}					
36	89	77	91	7.921	5.929	8.281	8.099	7.007	6.853	r _{yx1} + r _{yx2} - 2r _{yx1} r _{yx2} r _{x1x2}					
37	2.597	2.264	2.705	198.301	148.230	211.339	203.264	176.888	170.044	r _{yx1} + r _{yx2} - 2r _{yx1} r _{yx2} r _{x1x2}					
38										r _{yx1} + r _{yx2} - 2r _{yx1} r _{yx2} r _{x1x2}					

3. Korelasi Parsial

Korelasi parsial digunakan untuk menganalisis hubungan atau pengaruh variabel independen terhadap variabel dependen,

di mana salah satu variabel independennya dibuat tetap atau dikendalikan.

Contoh:

Setelah mengetahui korelasi antara kualitas pelayanan kepada siswa dan kompetensi tutor secara bersama-sama dengan jumlah peserta didik, Seorang pemilik kursus tersebut ingin meneliti korelasi antara kualitas pelayanan kepada siswa dengan jumlah peserta didik, manakala kompetensi tutornya dibuat sama (diparsialkan), dan korelasi antara kompetensi tutor dengan jumlah peserta didik, manakala kualitas pelayanan kepada siswa dibuat sama (diparsialkan).

Proses Pengambilan Keputusan.

Hipotesis:

- | | |
|----|---|
| Ho | Tidak ada korelasi antara kualitas pelayanan kepada siswa dengan jumlah peserta didik, manakala kompetensi tutornya dibuat sama (diparsialkan). |
| Ha | Terdapat korelasi antara kualitas pelayanan kepada siswa dengan jumlah peserta didik, manakala kompetensi tutornya dibuat sama (diparsialkan). |
| Ho | Tidak terdapat korelasi antara kompetensi tutor dengan jumlah peserta didik, manakala kualitas pelayanan kepada siswa dibuat sama (diparsialkan). |
| Ha | Terdapat korelasi antara kompetensi tutor dengan jumlah peserta didik, manakala kualitas pelayanan kepada siswa dibuat sama (diparsialkan). |

Dasar Pengambilan Keputusan:

Dengan membandingkan t_{hitung} dengan t_{tabel} dengan ketentuan:

Ho diterima $t_{hitung} < t_{tabel}$

Ho ditolak $t_{hitung} \geq t_{tabel}$

Dengan menggunakan angka probabilitas, dengan ketentuan:

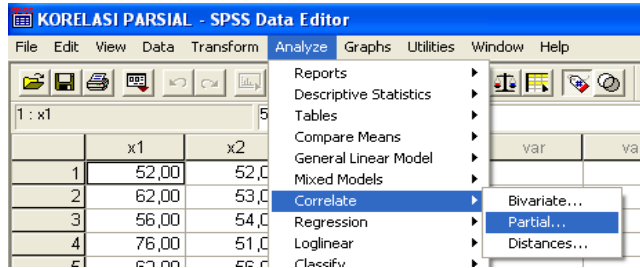
Ho diterima Probabilitas $>$ taraf nyata (α)

Ho ditolak Probabilitas $\leq$ taraf nyata (α)

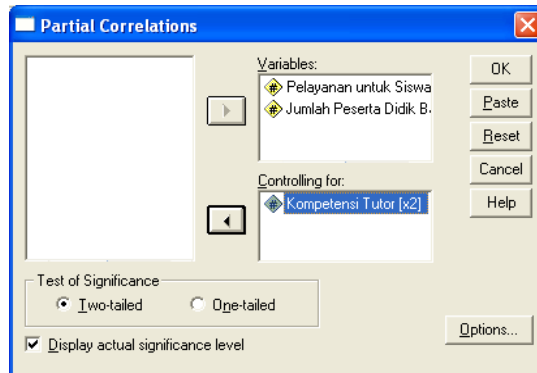
a. Aplikasi dengan SPSS

Untuk menganalisis korelasi parsial adalah sebagai berikut:

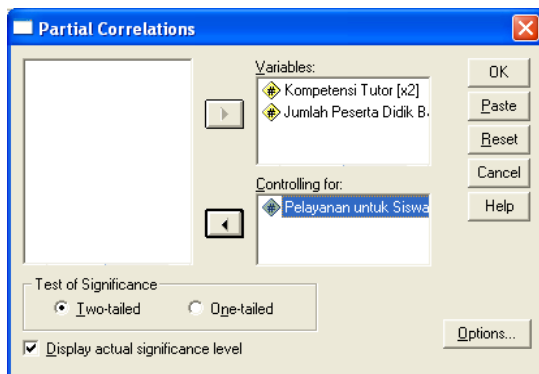
- 1) Setelah data diinput, lalu klik **Analyze ► Correlate ► Partial...**, sebagaimana gambar berikut ini:



- 2) Setelah keluar gambar seperti berikut ini, masukkan variabel yang akan dicari korelasinya pada kolom **Variables:** sedangkan variabel yang dibuat tetap atau diparsialkan masukkan ke kolom **Controlling for:** Setelah itu klik **Ok**. Misalkan: yang dibuat tetap itu waktu belajarnya, maka yang dimasukkan kolom **Controlling for:** adalah variabel itu.



akan tetapi kalau yang dibuat tetap itu kualitas layanan kepada siswa, maka yang dimasukkan kolom **Controlling for:** adalah variabel itu.



3) Berikut ini adalah output analisis di atas:

- - - P A R T I A L C O R R E L A T I O N C O
E F F I C I E N T S - - -

Controlling for.. X2

	X1	Y
X1	1,0000	,4042
	(0)	(32)
	P= .	P= ,018
Y	,4042	1,0000
	(32)	(0)
	P= ,018	P= .

(Coefficient / (D.F.) / 2-tailed Significance)

" , " is printed if a coefficient cannot be
computed

Output di atas menunjukkan bahwa koefisien korelasi antara kualitas layanan kepada siswa dengan jumlah siswa yang dimiliki apabila kompetensi tutor dibuat tetap (diparsialkan) adalah **0,4042**. Skor koefisien korelasi parsial ini ternyata lebih rendah dibanding skor koefisien antara 2 (dua) variabel di atas manakala tidak ada yang diparsialkan, yaitu **0,714**.

- - - P A R T I A L C O R R E L A T I O N C O
E F F I C I E N T S - - -

Controlling for.. X1

	Y	X2
Y	1,0000	,9112

```

(      0)      (      32)
P= .          P= ,000

X2            ,9112      1,0000
(      32)      (      0)
P= ,000      P= .

(Coefficient / (D.F.) / 2-tailed Significance)
" , " is printed if a coefficient cannot be
computed

```

Output ini menunjukkan bahwa koefisien korelasi antara kompetensi guru dengan jumlah peserta yang dimiliki apabila kualitas layanan kepada siswa dibuat tetap (diparsialkan) adalah **0,9112**. Skor koefisien korelasi parsial ini ternyata lebih rendah dibanding skor koefisien antara 2 (dua) variabel di atas manakala tidak ada yang diparsialkan, yaitu **0,949**.

b. Aplikasi dengan Microsoft Excel

Rumus korelasi parsial apabila X_2 yang dibuat sama atau diparsialkan adalah sebagai berikut:

$$r_{x_2(x_1y)} = \frac{r_{yx_1} - r_{yx_2} \cdot r_{x_1x_2}}{\sqrt{(1 - r_{x_2y}^2)(1 - r_{x_1x_2}^2)}}$$

Rumus korelasi parsial apabila X_1 yang dibuat sama atau diparsialkan adalah sebagai berikut:

$$r_{x_1(x_2y)} = \frac{r_{yx_2} - r_{yx_1} \cdot r_{x_1x_2}}{\sqrt{(1 - r_{x_1y}^2)(1 - r_{x_1x_2}^2)}}$$

Sedangkan uji korelasi parsial dapat dihitung dengan rumus berikut:

$$t = \frac{r_p \sqrt{n-3}}{\sqrt{1 - r_p^2}}$$

Untuk dapat mengaplikasikan rumus di atas terlebih dahulu harus dicari koefisien korelasi antar variabel yang caranya sama dengan pencarian koefisien korelasi product moment sebagaimana telah dipraktikkan di atas. Setelah masing-masing koefisien korelasi antar variabel

diketemukan, maka dapat dimasukkan ke dalam rumus korelasi parsial di atas.

Sedangkan prosedurnya dapat dilihat pada aplikasi Microsoft Excel sebagai berikut.

	A	B	C	D	E	F	G	H	I	J	K	L	M
1	x_1	x_2	y	x_1^2	x_2^2	y^2	x_1y	x_2y	x_1x_2				
2	52	52,00	64,0	2704	2704,000	4096	3328	3328,000	2704,000		r_{xy}		r_{xy}
3	62	62,00	64,0	3844	2809,000	4096	3968	3968,000	2704,000		89355,000		$=((35^*G37)-(A37^*C37))$
4	56	54,00	65,0	3136	2916,000	4225	3640	3510,000	3024,000		15658699840,000		$=((35^*D37)-(A37^*2))^{1/2}((35^*F37)-(C37^*2))$
5	76	51,00	66,0	5776	2601,000	4356	5016	3366,000	3876,000		125134,727		$=SQRT(K3)$
6	62	56,00	67,0	3844	3136,000	4489	4154	3752,000	3472,000		0,714		$=K2/K4$
7	66	56,00	67,0	4356	3136,000	4489	4422	3752,000	3696,000				
8	56	57,00	68,0	3136	3249,000	4624	3808	3876,000	3192,000		r_{yz}		r_{yz}
9	59	62,00	69,0	3481	3844,000	4761	4071	4278,000	3698,000		66980,000		$=((35^*H37)-(B37^*C37))$
10	58	59,00	70,0	3364	3481,000	4900	4060	4130,000	3422,000		4978343360,000		$=((35^*E37)-(B37^*2))^{1/2}((35^*F37)-(C37^*2))$
11	65	59,00	70,0	4225	3481,000	4900	4550	4130,000	3835,000		70557,376		$=SQRT(K9)$
12	79	57,00	72,0	6241	3249,000	5184	5688	4104,000	4503,000		0,949		$=K8/K10$
13	82	59,00	72,0	6724	3481,000	5184	5904	4248,000	4838,000		r_{xz}		r_{xz}
14	66	64,00	75,0	4356	4096,000	5625	4950	4800,000	4224,000		71932,000		$=((35^*I37)-(A37^*B37))$
15	67	64,00	75,0	4489	4096,000	5625	5025	4800,000	4288,000		12229240604,000		$=((35^*D37)-(A37^*2))^{1/2}((35^*E37)-(B37^*2))$
16	76	63,00	75,0	5776	3969,000	5625	5700	4725,000	4788,000		110585,897		$=SQRT(K15)$
17	79	67,00	75,0	6241	4489,000	5625	5925	5025,000	5293,000		0,650		$=K14/K16$
18	97	66,00	77,0	9409	4356,000	5929	7468	5082,000	6402,000				
19	68	67,00	78,0	4624	4489,000	6084	5304	5226,000	4556,000		x_2 dikonstantkan		x_2 dikonstantkan
20	68	67,00	78,0	4624	4489,000	6084	5304	5226,000	4556,000		0,097		$=K5-K11^*K17$
21	69	64,00	78,0	4761	4096,000	6084	5382	4992,000	4416,000		0,057		$=((1-(K1^*2))^{1/2}((1-(K1^*2))$
22	76	72,00	80,0	5776	5184,000	6400	6080	5760,000	5472,000		0,239		$=SQRT(K21)$
23	58	67,00	82,0	3364	4489,000	6724	4756	5494,000	3886,000		0,404		=K20/K22
24	72	67,00	82,0	5184	4489,000	6724	5904	5494,000	4824,000		2,500		$=K23^*(SQRT(35-3)/(1-(K23^*2)))$
25	76	68,00	83,0	5776	4624,000	6889	6308	5644,000	5168,000		2,042		$dk = n - m $ variabel
26	89	69,00	84,0	7921	4761,000	7056	7476	5796,000	6141,000				
27	74	70,00	85,0	5476	4900,000	7225	6290	5950,000	5180,000		x_1 dikonstantkan		x_1 dikonstantkan
28	78	75,00	85,0	6084	5625,000	7225	6630	6375,000	5850,000		0,485		$=K11-K5^*K17$
29	85	69,00	85,0	7225	4761,000	7225	7225	5885,000	5885,000		0,283		$=((1-(K5^*2))^{1/2}((1-(K1^*2))$
30	92	69,00	85,0	8464	4761,000	7225	7820	5885,000	6348,000		0,532		$=SQRT(K29)$
31	76	71,00	86,0	5776	5041,000	7396	6536	6106,000	5386,000		0,911		=K28/K30
32	99	72,00	86,0	9801	5184,000	7396	8514	6192,000	7128,000		12,516		$=K31^*(SQRT(35-3)/(1-(K31^*2)))$
33	85	77,00	88,0	7225	5929,000	7744	7480	6776,000	6545,000		2,042		
34	86	69,00	88,0	7396	4761,000	7744	7568	6072,000	5934,000				
35	99	75,00	90,0	9801	5625,000	8100	8910	6750,000	7425,000				
36	89	77,00	91,0	7921	5929,000	8281	8099	7007,000	6853,000				
37	2597	2264,000	2705	198301	148230,000	211339	203264	176888,000	170044,000				

Hasil penghitungan dengan Microsoft Excel ternyata sama dengan skor koefisien korelasi hasil penghitungan SPSS. Selanjutnya skor koefisien korelasi tersebut dimasukkan ke dalam rumus uji korelasi parsial. Koefisien korelasi antara kualitas layanan kepada siswa dengan jumlah siswa yang dimiliki apabila kompetensi tutor dibuat tetap (diparsialkan) adalah **0,4042**. Untuk menguji hipotesis, skor koefisien tersebut dimasukkan dalam rumus uji t di atas. Hasil hitung uji t didapatkan skor **2,5000**. Sedangkan koefisien korelasi antara kompetensi guru dengan jumlah peserta yang dimiliki apabila kualitas layanan kepada siswa dibuat tetap (diparsialkan) adalah **0,9112**. Ketika skor koefisien tersebut dimasukkan dalam rumus uji t , didapatkan skor **12,516**. Masing-masing skor dari t_{hitung} tersebut dibandingkan dengan t tabel dengan dk jumlah sampel (35) dikurangi variabel (3), yaitu 32, sebesar 2,036932. Karena t_{hitung} lebih tinggi dibanding t_{tabel} , sehingga H_0 ditolak dan H_a diterima.

F. Statistik Non-parametrik

1. Koefisien Kontingency

Teknik ini digunakan untuk menghitung hubungan antarvariabel bila datanya berbentuk nominal. Teknik ini mempunyai kaitan erat dengan Chi Kuadrat yang digunakan untuk menguji komparatif k -sampel independent. Oleh karena itu rumus yang digunakan mengandung nilai Chi Kuadrat.

Contoh:

Dilakukan penelitian untuk mengetahui ada tidaknya hubungan antara penggunaan alat (teknologi) hitung dengan benar tidaknya hasil hitung dalam penghitungan analisis statistik. Sampel diambil 45 orang, 15 menggunakan kalkulator, 15 menggunakan komputer dengan program microsoft Excel, dan 15 orang menggunakan SPSS.

Proses Pengambilan Keputusan.

Hipotesis:

- | | |
|-------|--|
| H_0 | Tidak ada hubungan antara penggunaan alat (teknologi) hitung dengan benar tidaknya hasil hitung dalam penghitungan analisis statistik. |
| H_a | Terdapat hubungan antara penggunaan alat (teknologi) hitung dengan benar tidaknya hasil |

hitung dalam penghitungan analisis statistik.

Dasar Pengambilan Keputusan:

Dengan membandingkan χ^2_{hitung} dengan χ^2_{tabel} dengan ketentuan:

Ho diterima $\chi^2_{\text{hitung}} < \chi^2_{\text{tabel}}$

Ho ditolak $\chi^2_{\text{hitung}} \geq \chi^2_{\text{tabel}}$

Dengan menggunakan angka probabilitas, dengan ketentuan:

Ho diterima Probabilitas > taraf nyata (α)

Ho ditolak Probabilitas $\leq$ taraf nyata (α)

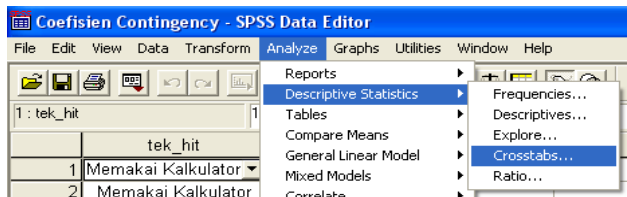
a. Aplikasi dengan SPSS

Data dari penelitian ini adalah data kategori, maka cara untuk memasukkan data ke SPSS juga harus sesuai dengan aturan data kategori. Bila sub menu **value labels** dalam menu **view** diaktifkan maka akan keluar seperti gambar di bawah ini.

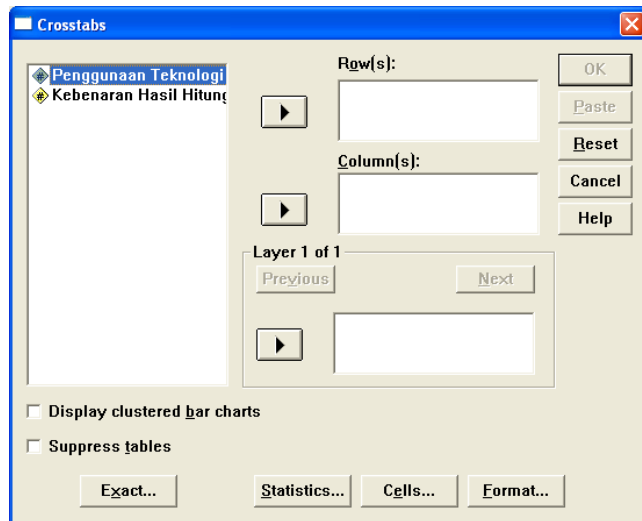
	tek_hit	valid
1	Memakai Kalkulator	Benar
2	Memakai Kalkulator	Salah
3	Memakai Kalkulator	Salah
4	Memakai Kalkulator	Benar
5	Memakai Kalkulator	Benar
6	Memakai Kalkulator	Benar
7	Memakai Kalkulator	Salah
8	Memakai Kalkulator	Salah
9	Memakai Kalkulator	Salah
10	Memakai Kalkulator	Benar
11	Memakai Kalkulator	Salah
12	Memakai Kalkulator	Salah
13	Memakai Kalkulator	Salah
14	Memakai Kalkulator	Salah
15	Memakai Kalkulator	Benar
16	Microsoft Excel	Benar
17	Microsoft Excel	Benar
18	Microsoft Excel	Benar
19	Microsoft Excel	Benar
20	Microsoft Excel	Benar
21	Microsoft Excel	Salah
22	Microsoft Excel	Benar
23	Microsoft Excel	Benar
24	Microsoft Excel	Benar
25	Microsoft Excel	Benar

	tek_hit	valid
26	Microsoft Excel	Benar
27	Microsoft Excel	Benar
28	Microsoft Excel	Salah
29	Microsoft Excel	Benar
30	Microsoft Excel	Benar
31	SPSS	Benar
32	SPSS	Benar
33	SPSS	Benar
34	SPSS	Benar
35	SPSS	Benar
36	SPSS	Benar
37	SPSS	Benar
38	SPSS	Benar
39	SPSS	Benar
40	SPSS	Benar
41	SPSS	Benar
42	SPSS	Benar
43	SPSS	Benar
44	SPSS	Benar
45	SPSS	Benar

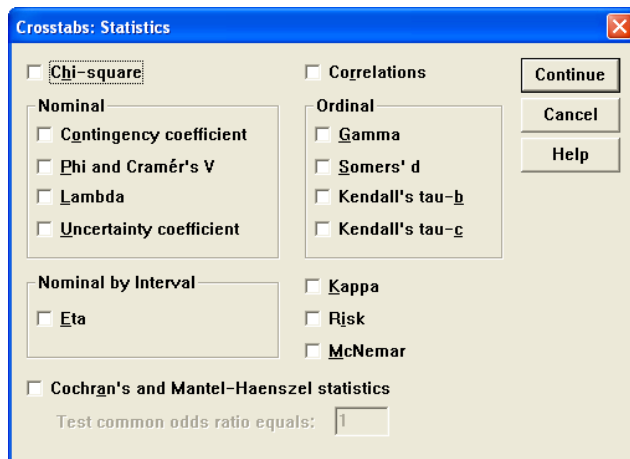
1. Setelah data diinput, klik **Analyze** ➤ **Descriptive Statistics** ➤ **Crosstabs** seperti di bawah ini:



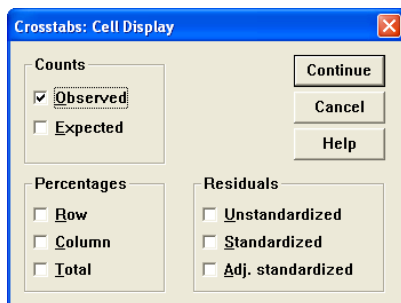
- Setelah itu keluar gambar seperti di bawah ini, destinasikan salah satu variabel itu ke kotak **Row(s)** dan variabel satunya ke **Column(s)**. Lalu klik **Statistics**.



- Setelah itu Aktifkan **Chi-square** dan **Contingency Coefficient** dalam **Nominal**, selanjutnya klik **Continue**. Setelah itu klik **cells**.



- Setelah itu aktifkan **expected** lalu klik **Continue**. Abaikan yang lain dan klik **Ok**.



5. Berikut ini adalah output analisis di atas.

Case Processing Summary

Cases		Penggunaan Teknologi Hitung * Kebenaran Hasil Hitung
Valid	N	45
	Percent	100.0%
Missing	N	0
	Percent	.0%
Total	N	45
	Percent	100.0%

Output ini menunjukkan bahwa jumlah sampel penelitian ini adalah 45 orang.

Penggunaan Teknologi Hitung * Kebenaran Hasil Hitung Crosstabulation

			Kebenaran Hasil Hitung		Total
			Benar	Salah	
Penggunaan Teknologi Hitung	Memakai Kalkulator	Count	6	9	15
		Expected Count	11.3	3.7	15.0
	Microsoft Excel	Count	13	2	15
		Expected Count	11.3	3.7	15.0
	SPSS	Count	15	0	15
		Expected Count	11.3	3.7	15.0
Total	Count		34	11	45
	Expected Count		34.0	11.0	45.0

Output ini menunjukkan bahwa ada 6 orang yang dapat menyelesaikan analisis kuantitatif secara benar dengan menggunakan kalkulator, sedangkan yang menggunakan Microsoft Excel ada 13 orang dan 15 orang yang menggunakan SPSS seluruhnya benar. Sedangkan yang salah terdiri 9 orang yang menggunakan kalkulator dan 2 orang yang menggunakan Microsoft Excel.

Chi-Square Tests

	Value	df	Asymp. Sig. (2-sided)
Pearson Chi-Square	16.123 ^a	2	.000
Likelihood Ratio	18.083	2	.000
Linear-by-Linear Association	14.294	1	.000
N of Valid Cases	45		

a. 3 cells (50.0%) have expected count less than 5. The minimum expected count is 3.67.

Symmetric Measures

	Value	Approx. Sig.
Nominal by Nominal Contingency Coefficient	.514	.000
N of Valid Cases	45	

a. Not assuming the null hypothesis.

b. Using the asymptotic standard error assuming the null hypothesis.

Dari Output di atas dapat diketahui bahwa nilai hitung koefisien kontingensi adalah 0,514 sementara nilai Chi kuadratnya adalah 16,123. Karenanya Approx. Sig nya 0,000, yang lebih kecil dibanding 0,05 (kesalahan yang ditoleransi), maka berarti kesalahan untuk menolak H_0 mendekati 0 persen, oleh karena itu, H_0 ditolak dan H_a diterima.

b. Aplikasi dengan Microsoft Excel

Sementara penghitungan dengan Excel sebagaimana contoh berikut ini. Dari kedua penghitungan ini ternyata menghasilkan nilai yang sama baik untuk nilai **Chi kuadrat** maupun **Korelasi Koefisien Kontingency**.

Rumus itu adalah sebagai berikut:

$$C = \sqrt{\frac{\chi^2}{N + \chi^2}}$$

Nilai Chi Kuadrat dihitung dengan rumus:

$$\chi^2 = \sum \frac{(f_o - f_h)^2}{f_h}$$

	A	B	C	D	E	F
1		ALAT HITUNG				
2	KEBENARAN	KALKULATOR	EXCEL	SPSS	JUMLAH	%
3	BENAR	6	13	15	34	0,756
4	HARAPAN	11,333	11,333	11,333		
5	SALAH	9	2	0	11	0,244
6	HARAPAN	3,667	3,667	3,667		
7	JUMLAH	15	15	15	45	
8						
9		2,510	0,245	1,186		
10		7,758	0,758	3,667		
11		16,123				
12		0,264				
13		0,514				
14						
15						
16		FORMULANYA				
17	Menjumlah di sel E3		=SUM(B3:D3)			
18	Prosentase di Sel F3		=(E3/E7)			
19	Harapan Benar		=(F3*B7)			
20	Harapan Salah		=(F5*B7)			
21						
22	$\chi^2 = \sum \frac{(f_o - f_h)^2}{f_h}$		=((B3-B4)*(B3-B4)/B4)			
23						
24			=(B9+C9+D9+B10+C10+D10)			
25						
26	$C = \sqrt{\frac{\chi^2}{N + \chi^2}}$		=(B11/(E7+B11))			
27			=SQRT(B12)			
28						

Cara Mencari Expexted adalah dengan mengkalikan persentase jumlah yang benar dan yang salah dengan jumlah masing-masing sampel:

Benar : $0.755556 \times 15 = 11.3333$

Salah : $0.244444 \times 15 = 3.6666$

Untuk menguji Nilai Koefisien Kontingency menggunakan nilai Chi Kuadrat dengan derajat kebebasan: kategori variabel x dikurangi 1 dan kategori variabel y dikurangi 1. Jadi $3 - 1 = 2$ ditambah $2 - 1 = 1$ sama dengan 3. $16.123 > 7,815$ (kesalahan 5%) maupun $11,341$ (kesalahan 1%). Kesimpulannya H_0 ditolak H_a diterima. Jadi, jenis teknologi hitung mempunyai korelasi yang signifikan dengan kebenaran hitung sebesar 0.514. Dan kesimpulan dari data sampel berlaku untuk populasi.

2. Spearman Rank

Sebagaimana dipaparkan pada tabel teknik analisis korelasi, maka spearman rank adalah teknik analisis manakala datanya berbentuk ordinal atau interval yang diordinalkan.

Contoh:

Untuk mendapatkan data tentang karakteristik pemimpin yang berpengaruh, pada tahun 1995 dan 2002 telah diadakan penelitian di Afrika, Asia, Eropa, Amerika Utara, Amerika Selatan, dan Australia diperoleh data dengan peringkat sebagai berikut:

	karakter	th2002	th1995
1	Jujur	1.00	1.00
2	Berpikiran Maju	2.00	2.00
3	Kompeten	3.00	4.00
4	Dapat Memberi Inspirasi	4.00	3.00
5	Cerdas	5.00	7.50
6	Adil	6.00	5.00
7	Berpandangan Luas	7.00	7.50
8	Suka Mendukung	8.00	6.00
9	Terus Terang	9.00	9.00
10	Dapat Diandalkan	10.00	10.00
11	Bekerjasama	11.00	12.50
12	Tegas	12.00	15.00
13	Berdaya Imajinasi	13.00	12.50
14	Berambisi	14.00	16.50
15	Berani	15.50	11.00
16	Perhatian	15.50	14.00
17	Dewasa Berfikir dan Bertindak	17.00	16.50
18	Setia	18.00	18.00
19	Penguasaan Diri	19.00	19.50
20	Mandiri	20.00	19.50

Proses Pengambilan Keputusan.**Hipotesis:**

Ho Tidak terdapat korelasi tentang peringkat karakteristik pemimpin yang berpengaruh antara tahun 1995 dan 2002.

Ha Terdapat korelasi tentang peringkat karakteristik pemimpin yang berpengaruh antara tahun 1995 dan 2002.

Dasar Pengambilan Keputusan:

Dengan membandingkan ρ/t_{hitung} dengan ρ/t_{tabel} dengan ketentuan:

Ho diterima $\rho/t_{hitung} < \rho/t_{tabel}$

Ho ditolak $\rho/t_{hitung} \geq \rho/t_{tabel}$

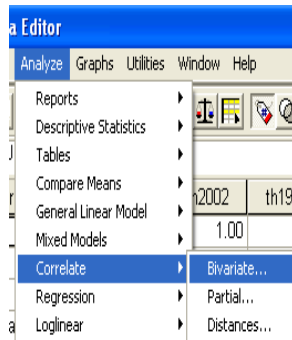
Dengan menggunakan angka probabilitas, dengan ketentuan:

Ho diterima Probabilitas $>$ taraf nyata (α)

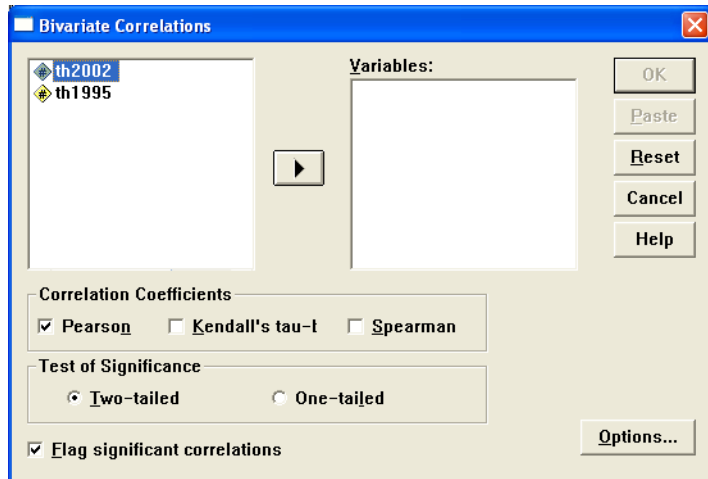
Ho ditolak Probabilitas $\leq$ taraf nyata (α)

a. Aplikasi dengan SPSS

Setelah data diinput sebagaimana di atas, maka klik **Analyze** ► **Correlate** ► **Bivariate** sebagaimana gambar berikut ini:



Setelah keluar gambar di bawah ini, masukkan dua variabel tersebut (**th2002** dan **th1995**) ke dalam kolak Variables. Klik **Pearson** untuk menghilangkan contengannya sebagai gantinya klik **Spearman**, lalu klik **Ok**.



Output di bawah ini adalah hasil penghitungan **SPSS**, yaitu korelasi antara hasil penelitian th 1995 dengan th 2002 yang mempunyai skor_{hitung}: 0,959.

Nonparametric Correlations

Correlations			TH2002	TH1995
Spearman's rho	TH2002	Correlation Coefficient	1.000	.959**
		Sig. (2-tailed)	.	.000
		N	20	20
	TH1995	Correlation Coefficient	.959**	1.000
		Sig. (2-tailed)	.000	.
		N	20	20

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

Bila dibandingkan ρ_{tabel} untuk $n = 20$ dengan taraf kesalahan 5%: 0,450 dan taraf kesalahan 1%: 0,591. Dikarenakan ρ_{hitung} (0,959) adalah lebih besar dari ρ_{tabel} baik untuk kesalahan 5% maupun 1%, maka korelasi antara penilaian tahun 1995 dan 2002 adalah positif dan signifikan dan berlaku juga untuk populasi.

b. Aplikasi dengan Microsoft Excel

Sedangkan aplikasi dengan Excel adalah sebagai berikut:

1. masing-masing skor pada variabel x dikurangi dengan masing-masing skor pada variabel y.
2. masing-masing hasil pengurangan itu dikuadratkan.
3. Seluruh hasil pengkuadratan itu dijumlahkan.
4. Hasil penjumlahan itu dimasukkan ke dalam rumus di bawah ini

$$\rho = 1 - \frac{6 \sum b_i^2}{n(n^2 - 1)}$$

Untuk lebih jelas dapat dilihat pada aplikasikan Microsoft Excel pada halaman berikutnya. Hasilnya ternyata sama persis dengan hasil penghitungan dengan SPSS.

	A	B	C	D	E
1	Katakteristik	(Xi)	(Yi)	bi	bi ²
2	Jujur	1	1	0	0
3	Berpikiran Maju	2	2	0	0
4	Kompeten	3	4	-1	1
5	Dapat Memberi Inspirasi	4	3	1	1
6	Cerdas	5	7,5	-2,5	6,25
7	Adil	6	5	1	1
8	Berpandangan Luas	7	7,5	-0,5	0,25
9	Suka Mendukung	8	6	2	4
10	Terus Terang	9	9	0	0
11	Dapat Diandalkan	10	10	0	0
12	Bekerjasama	11	12,5	-1,5	2,25
13	Tegas	12	15	-3	9
14	Berdaya Imajinasi	13	12,5	0,5	0,25
15	Berambisi	14	16,5	-2,5	6,25
16	Berani	15,5	11	4,5	20,25
17	Perhatian	15,5	14	1,5	2,25
18	Dewasa Berfikir dan Bertindak	17	16,5	0,5	0,25
19	Setia	18	18	0	0
20	Penguasaan Diri	19	19,5	-0,5	0,25
21	Mandiri	20	19,5	0,5	0,25
22				0	54,5
23					
24	6 x jumlah bi ²	327,000		= (6*E22)	
25	n x (n ² - 1)	7980,000		= 20*((20*20)-1)	
26	hasil 6 x jumlah bi ² dibagi hasil n x (n ² - 1)	0,041		= (B24/B25)	
27	1 dikurangi hasil pembagian di atas	0,959		= 1-B26	

Apabila jumlah sampel lebih dari 30, maka ρ_{tabel} tidak tersedia. Oleh karena itu, untuk menguji hipotesis digunakan rumus sebagai berikut:

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

Untuk sekedar contoh aplikasi rumus di atas digunakan menguji hipotesis di atas

$$14,3606 = 0,959 \sqrt{\frac{20-2}{1-(0,959^2)}}$$

Sedangkan t_{tabel} untuk taraf nyata (α) 5% dan derajat kebebasan 18 adalah 2,10092. Karena thitung lebih besar dibanding $t_{\text{tabel}:0,05;18}$, maka H_0 ditolak dan H_a diterima.

3. Kendall's tau

Fungsi dari Kendall's tau semestinya sama dengan spearman, bedanya hanyalah kalau spearman biasanya digunakan

untuk sampel kecil, tetapi kendall's tau dapat digunakan untuk sampel besar.

Kendall's tau juga sering digunakan untuk menganalisis data yang semula direncanakan dianalisis dengan product moment. Setelah diuji distribusi datanya ternyata tidak normal atau sampelnya kurang dari 30, maka akhirnya dianalisis dengan Kendall's tau. Untuk kasus seperti itu, bila dianalisis dengan menggunakan excel, maka data yang semula interval atau rasio harus dirangking.

Contoh:

Seorang peneliti ingin mengetahui korelasi antara IQ dengan prestasi belajar siswa. Untuk penelitian ini digunakan sampel 25 orang. Karena jumlahnya yang kurang dari 30, maka digunakan analisis non-parametrik, yaitu kendall's tau.

Proses Pengambilan Keputusan.

Hipotesis:

Ho Tidak terdapat korelasi antara IQ dengan prestasi belajar siswa

Ha Terdapat korelasi antara IQ dengan prestasi belajar siswa

Dasar Pengambilan Keputusan:

Dengan membandingkan z_{hitung} dengan z_{tabel} dengan ketentuan:

Ho diterima $z_{hitung} < z_{tabel}$

Ho ditolak $z_{hitung} \geq z_{tabel}$

Dengan menggunakan angka probabilitas, dengan ketentuan:

Ho diterima Probabilitas > taraf nyata (α)

Ho ditolak Probabilitas $\leq$ taraf nyata (α)

a. Aplikasi dengan SPSS

Sedangkan prosedur untuk analisisnya sama dengan spearman, bedanya kalau yang diaktifkan untuk kendall's tau adalah kendall.

	iq	prestasi			
1	132	68	13	117	47
2	131	70	14	116	50
3	130	65	15	113	46
4	129	67	16	111	38
5	125	61	17	110	43
6	124	60	18	107	44
7	123	59	19	105	42
8	122	58	20	103	41
9	121	45	21	97	49
10	120	64	22	96	35
11	119	62	23	93	39
12	118	51	24	89	40
			25	87	37

Setelah seluruh prosedur sebagaimana di spearman dilakukan maka akan keluar output sebagaimana di bawah ini:

Correlations

			IQ	PRESTASI
Kendall's tau_b	IQ	Correlation Coefficient	1.000	.760**
		Sig. (2-tailed)	.	.000
		N	25	25
	PRESTASI	Correlation Coefficient	.760**	1.000
		Sig. (2-tailed)	.000	.
		N	25	25

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

Berdasarkan output di atas, maka diketahui bahwa korelasi antara variabel x (IQ) dan y (prestasi) adalah 0,760. Selanjutnya hasil itu dimasukkan dalam rumus z, yang akan penulis jelaskan berbarengan dengan penjelasan aplikasi dengan excel.

Dari perhitungan ini diketahui bahwa terdapat hubungan positif antara IQ dan prestasi sebesar 0,76. Hal ini berarti semakin tinggi IQ seseorang semakin tinggi prestasinya.

c. Aplikasi dengan Microsoft Excel

Sedangkan aplikasi analisis kendall's dengan excel adalah sebagai berikut:

1. Skor dari variabel x dirangking dari rangking 1 dan seterusnya.
2. Skor dari variabel y juga dirangking.
3. Berdasarkan rangking dari variabel y itu dicari rangking atas (Ra), yaitu menghitung skor yang lebih tinggi bila dibandingkan dengan masing-masing skor pada variabel y.

4. Berdasarkan rangking dari variabel y itu dicari rangking bawah (Rb), yaitu menghitung skor yang lebih rendah bila dibandingkan dengan masing-masing skor pada variabel y.
5. Selanjutnya masing-masing skor Ra dan Rb dijumlahkan.
6. Hasilnya penjumlahan itu dimasukkan dalam rumus.

$$\tau = \frac{\Sigma A - \Sigma B}{\frac{N(N-1)}{2}}$$

Untuk mendapatkan kejelasan prosedur tersebut dapat dilihat pada aplikasi Microsoft Excel pada halaman berikutnya.

Hasil penghitungan dengan Excel ini ternyata mempunyai hasil yang sama dengan SPSS. Selanjutnya τ_{hitung} 0,760 itu dimasukkan ke dalam rumus di bawah ini.

$$z = \frac{\tau}{\sqrt{\frac{2(2N+5)}{9N(N-1)}}} \quad 5,3249 = \frac{0,760}{\sqrt{\frac{2(2 \times 25 + 5)}{9 \times 25(25-1)}}$$

Z_{hitung} : 5,3249, bila dibandingkan dengan z_{tabel} pada kesalahan 1% dibagi dua ($0,01/2 = 0,005$) atau kesalahan 5% dibagi dua ($0,05/2 = 0,025$). Demikian juga 100% dibagi dua ($0,100/2 = 0,50$) setelah dikurangi dengan $0,005 = 0,495$ skor tabelnya adalah 2,58 atau dikurangi $0,025 = 0,475$ skor tabelnya 1,96. Ternyata Z_{hitung} lebih besar dibanding z_{tabel} . Oleh karena itu H_0 ditolak dan H_a diterima.

RANKING										
A	B	C	D	E	F	G	H	I	J	K
1	IQ	Prestasi	R1	R2	Jumlah Ra	Jumlah Rb	FORMULA			
2	132	68	1	2	23	1	Cara Mencari Ranking R1 (dari IQ)			
3	131	70	2	1	23	0	Cara Mencari Ranking R2 (dari Prestasi)			
4	130	65	3	4	21	1	Cara Mencari Ra			
5	129	67	4	3	21	0	Cara Mencari Rb			
6	125	61	5	7	18	2	Keterangan:			
7	124	60	6	8	17	2	>2 (untuk penghitungan di bawahnya 2 nya diganti dengan masing-masing skor yang ada di R2 secara manual.			
8	123	59	7	9	16	2				
9	122	58	8	10	15	2				
10	121	45	9	16	9	7				
11	120	64	10	5	15	0				
12	119	62	11	6	14	0				
13	118	51	12	11	13	0				
14	117	47	13	14	10	2				
15	116	50	14	12	11	0				
16	113	46	15	15	9	1	Aplikasi Rumus Kendall			
17	111	38	16	23	2	7				
18	110	43	17	18	6	2				
19	107	44	18	17	6	1				
20	105	42	19	19	5	1	Aplikasi Rumus Z			
21	103	41	20	20	4	1				
22	97	49	21	13	4	0				
23	96	35	22	25	0	3				
24	93	39	23	22	1	1				
25	89	40	24	21	1	0				
26	87	37	25	24	0	0				
27				264	36					
28										
29				228	110					
30				300	5400					
31				0,76	0,02037					
32					0,14272					
33					5,32493					
34										

G. Koefisien Penentu

Koefisien penentu adalah digunakan untuk menjawab berapa persen variabel x mempengaruhi variabel y. Rumus untuk mencari Koefisien penentu adalah: $(\text{Koefisien Korelasi})^2 \times 100$.

Dari analisis korelasi yang sudah dilakukan diatas didapatkan skor korelasi untuk:

1. Product Moment **0,949**. Skor penentunya sebesar **90,060%**. Ini artinya pengaruh variabel x terhadap variabel y sebesar 90,060%, sedangkan yang 9,940% ditentukan oleh variabel lainnya.
2. Korelasi Ganda juga sebesar **0,958**. Skor penentunya sebesar **91,776%**. Ini artinya pengaruh kedua variabel x_1 x_2 secara bersama-sama terhadap variabel y sebesar 91,776%, sedangkan yang 8,224% ditentukan oleh variabel lainnya.
3. Korelasi Parsial ketika Variabel X_2 yang dibuat tetap (diparsialkan) sebesar **0,404**. Skor penentunya sebesar **16,322%**. Ini artinya pengaruh variabel x_1 terhadap y manakala X_2 diparsialkan sebesar 16,322%, sedangkan yang 83,678% ditentukan oleh variabel lainnya.
4. Korelasi Parsial ketika Variabel X_1 yang dibuat tetap (diparsialkan) sebesar **0,911**. Skor penentunya sebesar **82,922%**. Ini artinya pengaruh variabel x_2 terhadap y manakala X_1 diparsialkan sebesar 82,922%, sedangkan yang 17,008% ditentukan oleh variabel lainnya.
5. Koefisien Kontingensi **0,514**. Skor penentunya sebesar **26,420%**. Ini artinya pengaruh variabel x terhadap variabel y sebesar 26,420%, sedangkan yang 73,580% ditentukan oleh variabel lainnya.
6. Spearman Rank sebesar **0,959**. Skor penentunya sebesar **91,968%**. Ini artinya pengaruh variabel x terhadap variabel y sebesar 91,968%, sedangkan yang 8,032% ditentukan oleh variabel lainnya.
7. Kendall's tau **0,760**. Skor penentunya sebesar **57,760%**. Ini artinya pengaruh variabel x terhadap variabel y sebesar 57,760%, sedangkan yang 42,240% ditentukan oleh variabel lainnya.

Aplikasi untuk mencari Koefisien Penentu dengan Microsoft Excel seperti terlihat di bawah ini:

	A	B	C	D	E	F
1	CARA MENCARI KOEFISIEN PENENTU					
2						
3	NO	KETERANGAN	KK	KP		Rumus
4	01	Product Moment	0,949	90,060		=(C4*2)*100
5	02	Korelasi Ganda	0,958	91,776		=(C5*2)*100
6	03	Korelasi Parsial (X_2)	0,404	16,322		=(C6*2)*100
7	04	Korelasi Parsial (X_1)	0,911	82,992		=(C7*2)*100
8	05	Koefisien kontingensi	0,514	26,420		=(C8*2)*100
9	06	Spearman Rank	0,959	91,968		=(C9*2)*100
10	07	Kendall's tau	0,760	57,760		=(C10*2)*100
11						

BAB V

ANALISIS REGRESI

A. Pendahuluan

Korelasi dan regresi mempunyai hubungan yang sangat erat. Setiap regresi selalu ada korelasinya, tetapi belum tentu korelasi dilanjutkan dengan regresi. Korelasi yang dapat dilanjutkan dengan regresi adalah korelasi antara dua variabel atau lebih yang secara teori atau konsep mempunyai hubungan kausal (sebab akibat) atau hubungan fungsional.

Regresi digunakan manakala ingin diketahui bagaimana variabel y dapat diprediksikan melalui variabel x . Hasil analisis regresi dapat digunakan untuk memutuskan apakah naik dan turunnya skor variabel y dapat dilakukan melalui menaikkan dan menurunkan skor variabel x .

Analisis regresi adalah alat analisis yang termasuk dalam statistik parametrik. Dengan demikian, untuk menggunakan regresi, seorang peneliti harus melakukan pengujian asumsi terlebih dahulu. Asumsi yang harus diuji adalah, normalitas distribusi data, linieritas (jika kita hendak mempergunakan regresi linier), tiadanya heteroskedastisitas yang sering juga disebut heteroginitas, yaitu terjadinya ketidaksamaan varians dari residual satu data ke data yang lain, tiadanya multikolinearitas untuk regresi ganda, yaitu korelasi yang tinggi (di atas 0,5) antar variabel independen, serta tiadanya autokorelasi, yaitu korelasi antara kesalahan pengganggu pada periode t dengan kesalahan pada periode sebelumnya.

Contoh:

Seorang pemilik Kursus yang memiliki cabang di seluruh Indonesia berkeinginan untuk memprediksikan perubahan

jumlah peserta didik bila Tutornya ditingkatkan atau diturunkan kompetensinya.

Proses Pengambilan Keputusan.

Hipotesis:

Ho Kompetensi Tutor tidak berpengaruh terhadap jumlah peserta didik kursus.

Ha Kompetensi Tutor berpengaruh terhadap jumlah peserta didik kursus.

Dasar Pengambilan Keputusan:

Dengan membandingkan $r/t/F_{hitung}$ dengan $r/t/F_{tabel}$ dengan ketentuan:

Ho diterima $r/t/F_{hitung} < r/t/F_{tabel}$

Ho ditolak $r/t/F_{hitung} \geq r/t/F_{tabel}$

Dengan menggunakan angka probabilitas, dengan ketentuan:

Ho diterima Probabilitas $>$ taraf nyata (α)

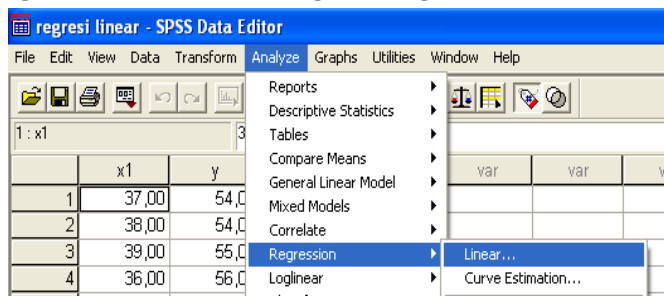
Ho ditolak Probabilitas $\leq$ taraf nyata (α)

B. Regresi Linear Sederhana

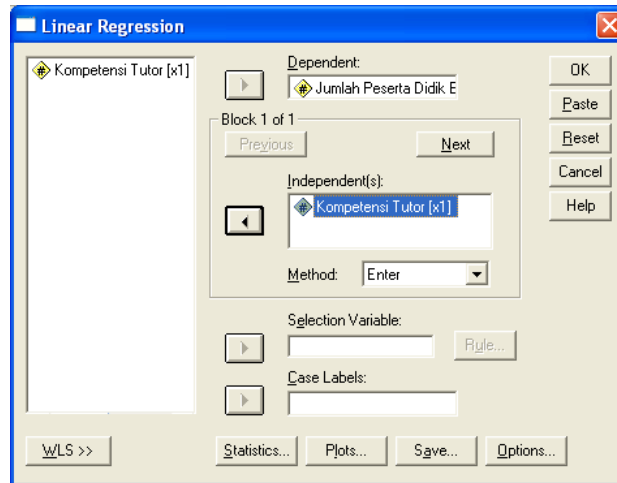
Di dalam regresi linear, variabel yang terlibat di dalamnya hanya dua.

1. Aplikasi dengan SPSS

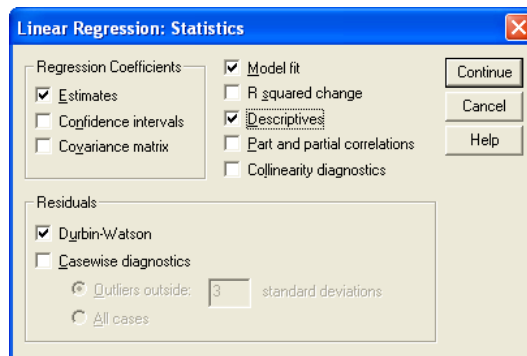
- a. Data diinput pada worksheet data view, lalu klik **Analyze** ➤ **Regression** ➤ **Linear**, sebagaimana gambar di bawah ini:



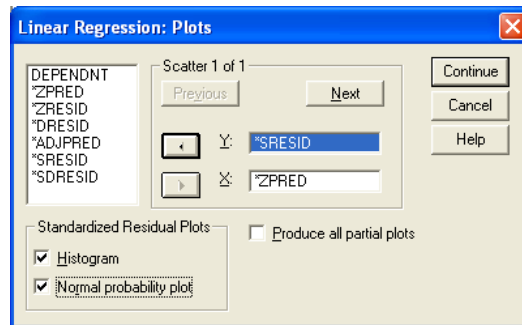
- b. Destinaskan variabel dependen (y) pada kolom yang ada di bawah **Dependent**, demikian juga variabel independen ke kolom **Independent(s)**, sebagaimana gambar di bawah ini. Selanjutnya klik **Statistics...**



- c. Setelah keluar gambar seperti berikut ini aktifkan juga **Descriptive** untuk mendeskripsikan data, Aktifkan **Dubin-Watson** pada **Residuals** untuk mengetahui ada tidaknya autokorelasi. Selanjutnya klik **Continue** untuk kembali pada tampilan sebelumnya, lalu klik **Plots...**



- d. Destinaskan **ZPRED** pada kolom X dan **SRESID** pada kolom Y untuk mengetahui terpenuhi tidaknya Homoskedastisitas atau sering juga disebut homogenitas. Selanjutnya aktifkan juga **Histogram** dan **Normal Probability Plot** untuk mengetahui normalitas distribusi data, lalu klik **Continue** dan selanjutnya klik **Ok**.



d. Berikut ini adalah **Output** analisis regresi linear di atas..

Descriptive Statistics

	Mean	Std. Deviation	N
Jumlah Peserta Didik Baru	67,2857	8,19100	35
Kompetensi Tutor	49,6857	7,23867	35

Dikarenakan jumlah Peserta Didik adalah termasuk tipe deskrit, maka angka desimal pada mean (rata-rata) dapat ditiadakan, sehingga dapat dikatakan bahwa rata-rata peserta didik di seluruh kursus adalah 67 dengan standard deviasi 8; sedangkan rata-rata kompetensi tutornya adalah 49,6857 dengan standard deviasinya 7,23867. Sampel institusi kursus yang diteliti sebanyak 35.

Correlations

		Jumlah Peserta Didik Baru	Kompetensi Tutor
Pearson Correlation	Jumlah Peserta Didik Baru	1,000	,949
	Kompetensi Tutor	,949	1,000
Sig. (1-tailed)	Jumlah Peserta Didik Baru	.	,000
	Kompetensi Tutor	,000	.
N	Jumlah Peserta Didik Baru	35	35
	Kompetensi Tutor	35	35

Korelasi kompetensi tutor dengan jumlah peserta kursus adalah 0,949. Bila digunakan skor koefisien determinasi sebesar 0,9060, maka dapat dikatakan bahwa 90,60% jumlah peserta kursus dipengaruhi oleh kompetensi tutor. Data serupa juga ditampilkan pada Output **Model Summary** di bawah ini.

Variables Entered/Removed^a

Model	Variables Entered	Variables Removed	Method
1	Kompetensi Tutor	.	Enter

a. All requested variables entered.

b. Dependent Variable: Jumlah Peserta Didik Baru

Model Summary^b

Model	R	R Square	Adjusted R Square	Std. Error of the Estimate	Durbin-Watson
1	,949 ^a	,901	,898	2,62089	1,847

a. Predictors: (Constant), Kompetensi Tutor

b. Dependent Variable: Jumlah Peserta Didik Baru

Berdasarkan skor Durbin-Watson sebesar 1,847 berarti tidak terjadi korelasi antara kesalahan pengganggu pada periode t dengan kesalahan pada periode sebelumnya. Model regresi yang baik adalah regresi yang bebas dari autokorelasi. Secara garis besar tolok ukur untuk menyimpulkan adanya autokorelasi atau tidak adalah sebagai berikut:

- Angka D-W di bawah -2 berarti terjadi autokorelasi positif.
- Angka D-W di antara -2 sampai dengan +2 berarti tidak terjadi autokorelasi.
- Angka D-W di atas +2 berarti terjadi autokorelasi negatif.

ANOVA^a

Model		Sum of Squares	df	Mean Square	F	Sig.
1	Regression	2054,464	1	2054,464	299,090	,000 ^a
	Residual	226,679	33	6,869		
	Total	2281,143	34			

a. Predictors: (Constant), Kompetensi Tutor

b. Dependent Variable: Jumlah Peserta Didik Baru

Berdasarkan F_{hitung} sebesar 299,090 yang lebih tinggi dibandingkan dengan $F_{tabel: 0,05;1;33}$ dengan $dk\ v_1 = 1$ dan $dk\ v_2 = 33$, skornya sebesar 4,139252 atau dengan tingkat signifikansi sebesar 0,000 yang jauh lebih rendah dari alpha sebesar 0,05, maka H_0 ditolak dan H_a diterima karena kesalahan untuk menolak H_0 mendekati 0%, oleh karena itu dapat disimpulkan bahwa model regresi ini, yaitu kompetensi tutor dapat digunakan untuk memprediksi jumlah siswa. Tidak hanya itu, hasil ini juga menunjukkan bahwa antara variabel kompetensi tutor dengan variabel jumlah siswa ada hubungan linear.

Coefficients^a

Model	Unstandardized Coefficients		Standardized Coefficients	t	Sig.
	B	Std. Error	Beta		
1 (Constant)	13,930	3,117		4,469	,000
Kompetensi Tutor	1,074	,062	,949	17,294	,000

a. Dependent Variable: Jumlah Peserta Didik Baru

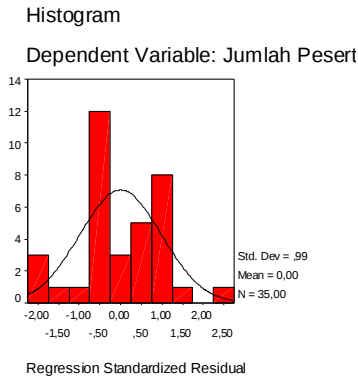
Output di atas menyajikan koefisien regresi. Persamaan regresinya adalah $\text{Jumlah peserta} = 13,930 + 1,074 \times \text{kompetensi tutor}$. Apabila pada sebuah kursus skor kompetensi tutornya sebesar 35, maka dapat diramalkan peserta didik akan berjumlah 51,52 atau dapat dibulatkan menjadi 53 orang. Output di atas memperlihatkan T_{hitung} yang berada pada garis Kompetensi Tutor sebesar 17,294 ternyata lebih besar bila dibandingkan dengan $T_{\text{tabel: } 0,025;33}$ sebesar 2,348334, dengan signifikansi 0,0000 maka H_0 dapat ditolak. Hal ini mengandung pengertian Perubahan jumlah peserta didik juga ditentukan oleh perubahan kompetensi tutor.

Perlu dijelaskan bahwa di dalam analisis regresi ganda, T-test digunakan untuk menguji pengaruh masing-masing variabel x terhadap y , sedangkan F-test yang ditampilkan di atasnya adalah untuk menguji pengaruh seluruh variabel x secara bersama-sama terhadap variabel y .

Residuals Statistics^a

	Minimum	Maximum	Mean	Std. Deviation	N
Predicted Value	52,5891	80,5096	67,2857	7,77338	35
Std. Predicted Value	-1,891	1,701	,000	1,000	35
Standard Error of Predicted Value	,44505	,95834	,60708	,15709	35
Adjusted Predicted Value	52,0626	80,8316	67,2843	7,80719	35
Residual	-5,4016	6,0813	,0000	2,58206	35
Std. Residual	-2,061	2,320	,000	,985	35
Stud. Residual	-2,095	2,367	,000	1,012	35
Deleted Residual	-5,5838	6,3282	,0014	2,72393	35
Stud. Deleted Residual	-2,216	2,558	-,002	1,043	35
Mahal. Distance	,009	3,575	,971	1,045	35
Cook's Distance	,001	,151	,028	,039	35
Centered Leverage Value	,000	,105	,029	,031	35

a. Dependent Variable: Jumlah Peserta Didik Baru

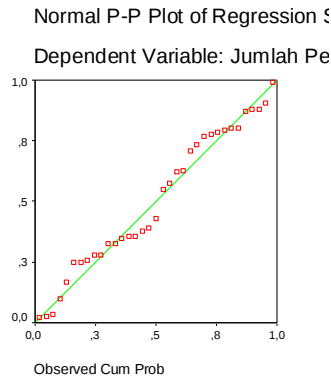


Berdasarkan output histogram di atas, terlihat bahwa sebaran data yang ada menyebar merata ke hampir semua daerah kurva normal. Dengan demikian dapat disimpulkan bahwa data yang kita miliki mempunyai distribusi normal. Demikian juga dengan normal P-P Plot yang terlihat pada halaman berikutnya memperlihatkan hasil yang sama. Untuk lebih meyakinkan normalitas data tersebut dapat digunakan uji kolmogorov-semirnov dengan hasil berikut:

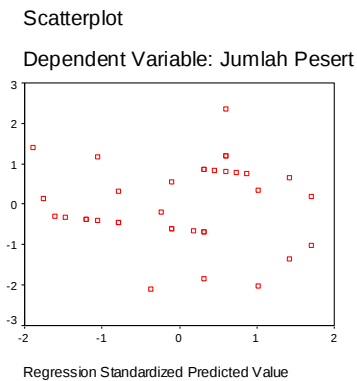
	Kolmogorov-Smirnov ^a			Shapiro-Wilk		
	Statistic	df	Sig.	Statistic	df	Sig.
Kompetensi Tutor	,140	35	,082	,961	35	,241
Jumlah Peserta Didik Baru	,118	35	,200*	,945	35	,079

*. This is a lower bound of the true significance.
a. Lilliefors Significance Correction

Output ini memperlihatkan bahwa angka Sig. dari Kolmogorov Semirnov untuk variabel x (0,082), dan variabel y (0,200) yang lebih besar dari 0,05 sebagai toleransi kesalahan, maka kedua data variabel itu berdistribusi normal.



Gambar di atas yang menunjukkan nilai data ada di sekitar baris diagonal ini menguatkan kesimpulan dari hasil uji F, yang tertera dalam Output ANOVA di atas, sehingga menunjukkan bahwa antara variabel kompetensi tutor dengan variabel jumlah siswa ada hubungan linear. Di samping itu, gambar tersebut juga menunjukkan bahwa distribusi data tersebut adalah normal karena titik-titik berada di sekitar garis diagonal.



Output di atas memperlihatkan bahwa tidak terjadi heteroskedastisitas dikarenakan penyebaran nilai-nilai residual terhadap harga-harga prediksi tidak membentuk suatu pola sehingga varians data bersifat homogen.

Untuk memastikan hasil uji homoskedastisitas dapat digunakan metode levene test. Dan hasilnya juga sama karena signifikasinya sebesar 0,287 yang berarti lebih tinggi dari 0,05, artinya peneliti tidak dapat menolak H_0 . Oleh karena H_0 yang diterima, maka dapat disimpulkan bahwa asumsi homoskedastisitas terpenuhi.

Test of Homogeneity of Variance

		Levene Statistic	df1	df2	Sig.
skor	Based on Mean	1,153	1	68	,287
	Based on Median	1,149	1	68	,287
	Based on Median and with adjusted df	1,149	1	66,760	,288
	Based on trimmed mean	1,161	1	68	,285

Berdasarkan penjelasan di atas, maka seluruh asumsi yang dibutuhkan terpenuhi sehingga model yang **BLUE (Best Linear Unbiased Estimator)** terpenuhi.

2. Aplikasi dengan Microsoft Excel

Setidaknya ada tiga cara yang dapat digunakan untuk melakukan aplikasi analisis regresi dengan Microsoft Excel, yaitu menjabarkan rumus, aplikasi dari function, maupun dari Data Analysis. Hanya saja ketiganya lebih mengkonsentrasinya diri pada pencarian persamaan regresi.

Rumus persamaan regresi linear sederhana adalah sebagai berikut:

$$Y = a + bX$$

Dimana :

- Y = Subyek dalam variabel dependen yang diprediksikan
- a = Harga Y bila X = 0 (harga konstan)
- b = Angka arah atau koefisien regresi, yang menunjukkan angka peningkatan ataupun penurunan variabel dependen yang didasarkan pada variabel independen. Bila b (+) maka naik, dan bila (-) maka terjadi penurunan
- X = Subyek pada variabel independen yang mempunyai nilai tertentu

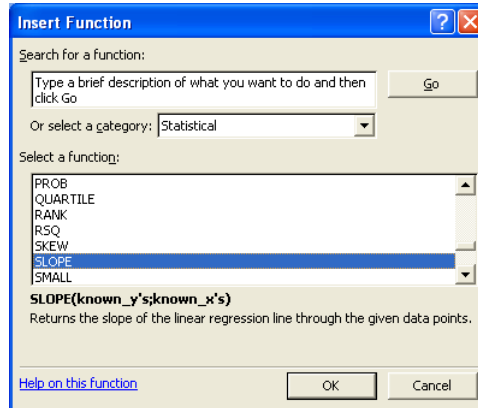
Sementara rumus untuk mencari a dan b adalah sebagai berikut:

$$a = \frac{\sum XY - n(\bar{X})(\bar{Y})}{\sum X^2 - n(\bar{X}^2)}$$

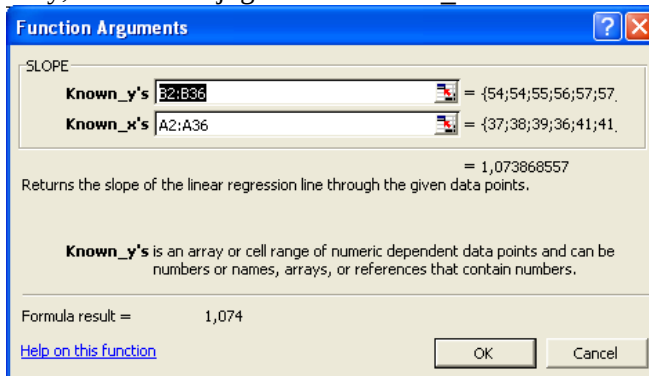
$$b = \bar{Y} - b(\bar{X})$$

Sedangkan aplikasi dari function dapat diaplikasikan dengan fungsi slope dan intercept. Adapaun cara untuk mencari

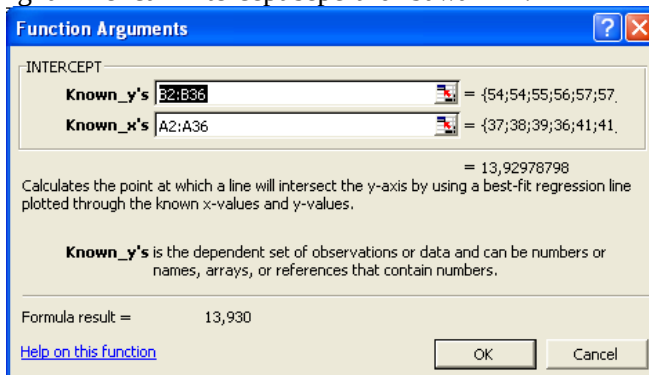
slope adalah dengan mencari menu tersebut pada function seperti di bawah ini.



Setelah klik **Ok**, maka akan keluar seperti gambar berikut ini, maka klik di kolom **Known_y's** lalu bloklah seluruh data variabel y, sedemikian juga untuk **Known_x's** untuk variabel x.



Sedangkan mencari intercept seperti di bawah ini:



Aplikasi dari function ternyata mempunyai hasil yang sama dibandingkan pencarian persamaan regresi dengan menjabarkan rumus, seperti terlihat pada tampilan berikut ini.

A	B	C	D	E	F	G	H	I	J	K
1	x	y	x ²	y ²	xy					
2	37	54	1.369	2.916	1.998					
3	38	54	1.444	2.916	2.052					
4	39	55	1.521	3.025	2.145					
5	36	56	1.296	3.136	2.016					
6	41	57	1.681	3.249	2.337					
7	41	57	1.681	3.249	2.337					
8	42	58	1.764	3.364	2.436					
9	47	59	2.209	3.481	2.773					
10	44	60	1.936	3.600	2.640					
11	44	60	1.936	3.600	2.640					
12	42	62	1.764	3.844	2.604					
13	44	62	1.936	3.844	2.728					
14	49	65	2.401	4.225	3.185					
15	49	65	2.401	4.225	3.185					
16	48	65	2.304	4.225	3.120					
17	52	65	2.704	4.225	3.380					
18	51	67	2.601	4.489	3.417					
19	52	68	2.704	4.624	3.536					
20	52	68	2.704	4.624	3.536					
21	49	68	2.401	4.624	3.332					
22	57	70	3.249	4.900	3.990					
23	52	72	2.704	5.184	3.744					
24	52	72	2.704	5.184	3.744					
25	53	73	2.809	5.329	3.869					
26	54	74	2.916	5.476	3.996					
27	55	75	3.025	5.625	4.125					
28	60	75	3.600	5.625	4.500					
29	54	75	2.916	5.625	4.050					
30	54	75	2.916	5.625	4.050					
31	56	76	3.136	5.776	4.256					
32	57	76	3.249	5.776	4.332					
33	62	78	3.844	6.084	4.836					
34	54	78	2.916	6.084	4.212					
35	60	80	3.600	6.400	4.800					
36	62	81	3.844	6.561	5.022					
37	1.739	2.365	88.185	160.739	118.923					
38	49.686	67.286								

$$b = \frac{66.960.000}{62.354.000}$$

$$b = 1,074$$

$$a = 13,930$$

$$a = \frac{\sum XY - n(\bar{X})(\bar{Y})}{\sum X^2 - n(\bar{X}^2)}$$

$$b = \bar{Y} - b(\bar{X})$$

$$=SLOPE(B2:B36;A2:A36)$$

$$=INTERCEPT(B2:B36;A2:A36)$$

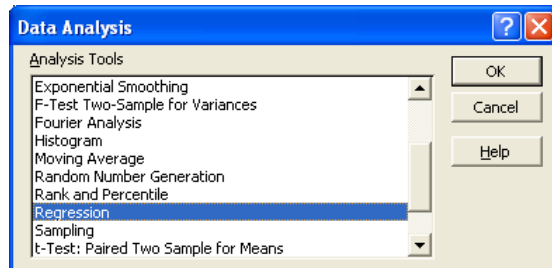
$$=35^*E37-A37^*B37$$

$$=35^*C37-A37^*2$$

$$=G4/G5$$

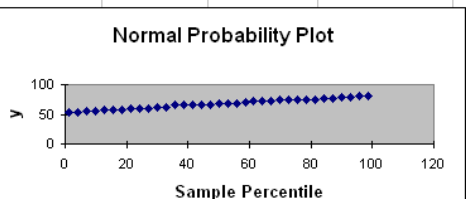
$$=B38-G4^*A38$$

Di samping cara di atas, dapat juga dilakukan dengan mengaplikasi Microsoft Excel melalui Data Analysis seperti di bawah ini:



Setelah regression diklik, dan variabel diisi sesuai dengan tempatnya, dan menu-menu lain yang dibutuhkan diklik, maka akan keluar output yang sebagiannya ditampilkan berikut ini.

	A	B	C	D	E	F
1	SUMMARY OUTPUT					
2						
3	<i>Regression Statistics</i>					
4	Multiple R	0,949014879				
5	R Square	0,900629241				
6	Adjusted R Square	0,897618006				
7	Standard Error	2,620888688				
8	Observations	35				
9						
10	ANOVA					
11		<i>df</i>	<i>SS</i>	<i>MS</i>	<i>F</i>	<i>Significance F</i>
12	Regression	1	2054,463959	2054,463959	299,089643	0,000
13	Residual	33	226,6788979	6,869057512		
14	Total	34	2281,142857			
15						
16		<i>Coefficients</i>	<i>Standard Error</i>	<i>t Stat</i>	<i>P-value</i>	<i>Lower 95%</i>
17	Intercept	13,92978798	3,116834745	4,46920967	8,7308E-05	7,588534988
18	x	1,073868557	0,062094115	17,29420837	4,12599E-18	0,947537031
19						
20	PROBABILITY OUTPUT					
21						
22	<i>Percentile</i>	<i>y</i>				
23	1,428571429	54				
24	4,285714286	54				
25	7,142857143	55				
26	10	56				
27	12,85714286	57				
28	15,71428571	57				
29	18,57142857	58				
30	21,42857143	59				
31	24,28571429	60				



Dari tiga cara di atas ternyata menghasilkan skor yang sama, yaitu slopenya = 1,074 sementara intercept (konstanta)nya sebesar = 13,930.

C. Regresi Ganda Dua Prediktor

Regresi ganda dua prediktor adalah regresi di mana ada tiga variabel yang terlibat di dalamnya. Dua di antara tiga variabel tersebut menjadi variabel independen dan satu menjadi variabel dependen.

Contoh:

Peneliti berkeinginan untuk memprediksikan naik dan turunnya jumlah peserta kursus apabila terjadi perubahan/perbedaan kompetensi tutor dan dana promosi.

Proses Pengambilan Keputusan.

Hipotesis:

Ho Kompetensi tutor dan dana promosi tidak berpengaruh terhadap jumlah peserta didik kursus.

Ha Kompetensi tutor dan dana promosi berpengaruh terhadap jumlah peserta didik kursus.

Dasar Pengambilan Keputusan:

Dengan membandingkan $r/t/F_{hitung}$ dengan $r/t/F_{tabel}$ dengan ketentuan:

Ho diterima $r/t/F_{hitung} < r/t/F_{tabel}$

Ho ditolak $r/t/F_{hitung} \geq r/t/F_{tabel}$

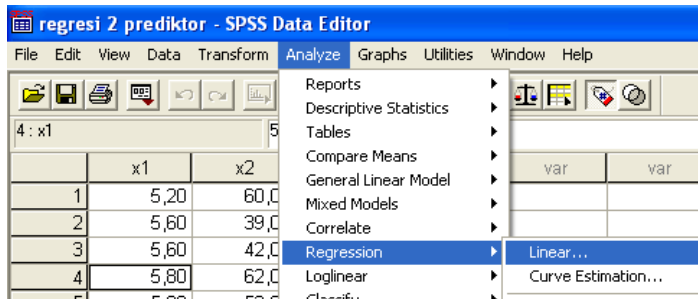
Dengan menggunakan angka probabilitas, dengan ketentuan:

Ho diterima Probabilitas $>$ taraf nyata (α)

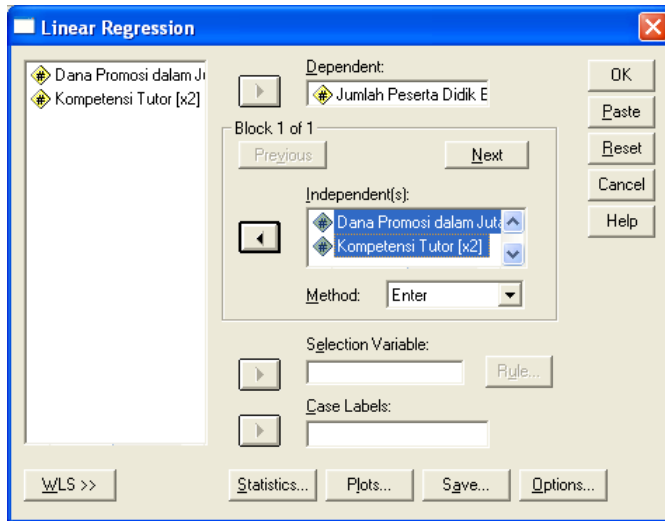
Ho ditolak Probabilitas $\leq$ taraf nyata (α)

1. Aplikasi dengan SPSS

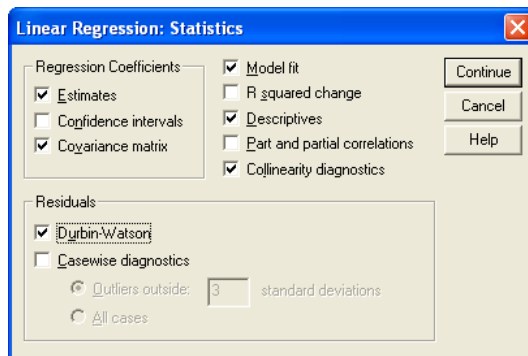
- a. Data diinput pada worksheet data view, lalu klik **Analyze** ➤ **Regression** ➤ **Linear**, sebagaimana gambar di bawah ini:



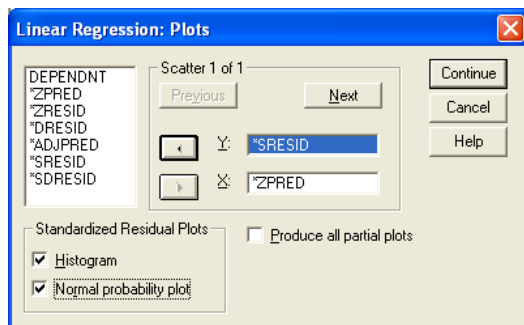
- b. Destinaskan variabel dependen (y) pada kolom yang ada di bawah **Dependent**, demikian juga kedua variabel independen ke kolom **Independent(s)**, sebagaimana gambar di bawah ini. Selanjutnya klik **Statistics...**



- c. Setelah keluar gambar seperti berikut ini aktifkan juga **Descriptive** untuk mendeskripsikan data, aktifkan **Collinearity diagnostic** untuk mengetahui apakah terjadi korelasi antar variabel indenpenden atau tidak, aktifkan **Dubin-Watson** pada **Residuals** untuk mengetahui ada tidaknya autokorelasi. Selanjutnya klik **Continue** untuk kembali pada tampilan sebelumnya, lalu klik **Plots...**



- d. Destinasikan **ZPRED** pada kolom **X** dan **SRESID** pada kolom **Y** untuk mengetahui terpenuhi tidaknya Homoskedastisitas atau sering juga disebut homogenitas. Selanjutnya aktifkan juga **Histogram** dan **Normal Probability Plot** untuk mengetahui normalitas distribusi data, lalu klik **Continue** dan selanjutnya klik **Ok**.



d. Berikut ini adalah **Output** analisis regresi di atas.

Descriptive Statistics

	Mean	Std. Deviation	N
Jumlah Peserta Didik Baru	67,2857	8,19100	35
Dana Promosi dalam Jutaan	7,4200	1,28379	35
Kompetensi Tutor	51,9429	7,74575	35

Dikarenakan jumlah Peserta Didik adalah termasuk tipe deskrit, maka angka desimal pada mean (rata-rata) dapat ditiadakan, sehingga dapat dikatakan bahwa rata-rata peserta didik di seluruh kursus adalah 67 dengan standard deviasi 8; sedangkan rata-rata dana promosi sebesar 7.420.000,- dengan standard deviasi sebesar 1.283.790,-; serta kompetensi tutornya adalah 51,9429 dengan standard deviasinya 7,74575. Sampel institusi kursus yang diteliti sebanyak 35.

Correlations

		Jumlah Peserta Didik Baru	Dana Promosi dalam Jutaan	Kompetensi Tutor
Pearson Correlation	Jumlah Peserta Didik Baru	1,000	,714	,588
	Dana Promosi dalam Jutaan	,714	1,000	,319
	Kompetensi Tutor	,588	,319	1,000
Sig. (1-tailed)	Jumlah Peserta Didik Baru	.	,000	,000
	Dana Promosi dalam Jutaan	,000	.	,031
	Kompetensi Tutor	,000	,031	.
N	Jumlah Peserta Didik Baru	35	35	35
	Dana Promosi dalam Jutaan	35	35	35
	Kompetensi Tutor	35	35	35

Korelasi kompetensi tutor dengan jumlah peserta kursus adalah 0,588, antara Dana promosi dan jumlah peserta didik

sebesar 0,714, sedangkan korelasi antar dua variabel independen, yaitu dana promosi dan kompetensi tutor sebesar 0,031.

Variables Entered/Removed^a

Model	Variables Entered	Variables Removed	Method
1	Kompetensi Tutor, Dana Promosi dalam Jutaan ^a	.	Enter

a. All requested variables entered.

b. Dependent Variable: Jumlah Peserta Didik Baru

Model Summary^b

Model	R	R Square	Adjusted R Square	Std. Error of the Estimate	Durbin-Watson
1	,809 ^a	,654	,633	4,96497	1,487

a. Predictors: (Constant), Kompetensi Tutor, Dana Promosi dalam Jutaan

b. Dependent Variable: Jumlah Peserta Didik Baru

Berdasarkan skor Durbin-Watson sebesar 1,487 berarti tidak terjadi korelasi antara kesalahan pengganggu pada periode t dengan kesalahan pada periode sebelumnya. Model regresi yang baik adalah regresi yang bebas dari autokorelasi. Secara garis besar tolok ukur untuk menyimpulkan adanya autokorelasi atau tidak adalah sebagai berikut:

- d. Angka D-W di bawah -2 berarti terjadi autokorelasi positif.
- e. Angka D-W di antara -2 sampai dengan +2 berarti tidak terjadi autokorelasi.
- f. Angka D-W di bawah +2 berarti terjadi autokorelasi negatif.

ANOVA^a

Model		Sum of Squares	df	Mean Square	F	Sig.
1	Regression	1492,313	2	746,157	30,269	,000 ^a
	Residual	788,829	32	24,651		
	Total	2281,143	34			

a. Predictors: (Constant), Kompetensi Tutor, Dana Promosi dalam Jutaan

b. Dependent Variable: Jumlah Peserta Didik Baru

Berdasarkan F_{hitung} sebesar 30,269 yang lebih tinggi dibandingkan dengan F_{tabel} : 0,05;2;32. dengan dk $v_1 = 2$ dan dk $v_2 = 32$, skornya sebesar 3,294531 atau dengan tingkat signifikansi sebesar 0,000 yang jauh lebih rendah dari alpha sebesar 0,05, maka H_0 ditolak dan H_a diterima karena kesalahan untuk menolak H_0 mendekati 0%, oleh karena itu dapat disimpulkan

bahwa model regresi ini, yaitu variabel dana promosi dan kompetensi tutor secara bersama-sama dapat digunakan untuk memprediksi jumlah siswa. Tidak hanya itu, hasil ini juga menunjukkan bahwa antara variabel dana promosi, kompetensi tutor dengan variabel jumlah siswa ada hubungan linear.

Coefficients ^a				
		Model		
		1		
		(Constant)	Dana Promosi dalam Jutaan	Kompetensi Tutor
Unstandardized Coefficients	B	17,513	3,741	,424
	Std. Error	6,635	,700	,116
Standardized Coefficients	Beta		,586	,401
t		2,640	5,347	3,654
Sig.		,013	,000	,001
Collinearity Statistics	Tolerance		,898	,898
	VIF		1,113	1,113

a. Dependent Variable: Jumlah Peserta Didik Baru

Output di atas menyajikan tiga hal penting, yaitu persamaan regresi, t_{hitung} dan uji multikolinearitas. Persamaan regresinya adalah Jumlah peserta = $17,513 + (0,424 \times \text{kompetensi tutor}) + (3,741 \times \text{Dana Promosi})$. Di samping itu, output di atas memperlihatkan T_{hitung} yang berada pada kolom Dana promosi sebesar 5,347 dan pada kolom Kompetensi Tutor sebesar 3,654 ternyata lebih besar bila dibandingkan dengan $T_{tabel: 0,025;33}$ sebesar 2,348334, dengan signifikansi 0,0000 maka H_0 dapat ditolak. Hal ini mengandung pengertian Perubahan jumlah peserta didik juga ditentukan oleh perubahan Dana Promosi dan juga oleh kompetensi tutor. Sedangkan uji multikolinearitas dapat diketahui dari skor Tolenace dan VIF. Apabila skor **VIF** (Variance Inflation Factor) di sekitar angka 1 dan mempunyai angka **Tolerance** mendekati membuktikan bahwa model tersebut terbebas dari multikolinearitas. Artinya, tidak diketemukan korelasi yang tinggi antar variabel independent. VIF Dana Promosi sebesar 1,113 dan angka tolenace sebesar 0,898 sedangkan VIF Kompetensi Tutor sebesar 1,113 dengan angka tolerance sebesar 0,898.

Coefficient Correlations^a

Model			Kompetensi Tutor	Dana Promosi dalam Jutaan
1	Correlations	Kompetensi Tutor	1,000	-,319
		Dana Promosi dalam Jutaan	-,319	1,000
	Covariances	Kompetensi Tutor	,013	-,026
		Dana Promosi dalam Jutaan	-,026	,490

a. Dependent Variable: Jumlah Peserta Didik Baru

Output di atas juga digunakan untuk mendeteksi adanya masalah multikolinearitas. Apabila koefisien korelasi antar variabel independen di bawah 0,5 membuktikan bahwa model tersebut terbebas dari multikolinearitas. Artinya, tidak ditemukan korelasi yang tinggi antar variabel independent. Berdasarkan output di atas diketahui bahwa koefisien antara Kompetensi Tutor dengan Dana promosi adalah -0,319. Skor ini lebih kecil dibanding 0,5.

Collinearity Diagnostics^a

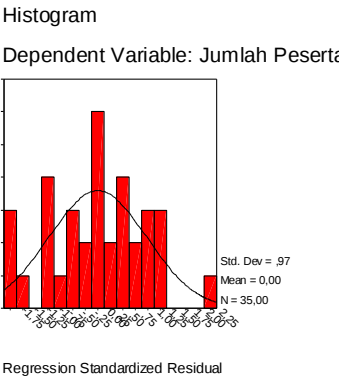
Model	Dimension	Eigenvalue	Condition Index	Variance Proportions		
				(Constant)	Dana Promosi dalam Jutaan	Kompetensi Tutor
1	1	2,972	1,000	,00	,00	,00
	2	,018	13,007	,05	,92	,33
	3	,010	17,003	,94	,08	,67

a. Dependent Variable: Jumlah Peserta Didik Baru

Residuals Statistics^a

	Minimum	Maximum	Mean	Std. Deviation	N
Predicted Value	54,9916	80,8260	67,2857	6,62507	35
Std. Predicted Value	-1,856	2,044	,000	1,000	35
Standard Error of Predicted Value	,85168	2,21714	1,40011	,39639	35
Adjusted Predicted Value	54,9903	81,7017	67,4676	6,64853	35
Residual	-8,4160	10,7508	,0000	4,81673	35
Std. Residual	-1,695	2,165	,000	,970	35
Stud. Residual	-1,890	2,259	-,017	1,022	35
Deleted Residual	-10,4860	11,6970	-,1819	5,35378	35
Stud. Deleted Residual	-1,973	2,425	-,019	1,047	35
Mahal. Distance	,029	5,809	1,943	1,601	35
Cook's Distance	,000	,296	,039	,064	35
Centered Leverage Value	,001	,171	,057	,047	35

a. Dependent Variable: Jumlah Peserta Didik Baru



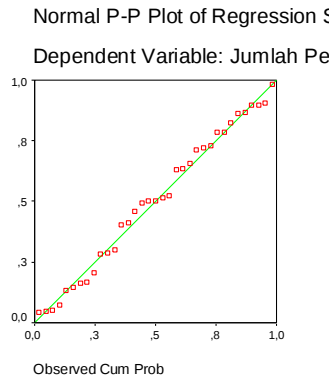
Berdasarkan output histogram di atas, terlihat bahwa sebaran data yang ada menyebar merata ke hampir semua daerah kurva normal. Dengan demikian dapat disimpulkan bahwa data yang kita miliki mempunyai distribusi normal. Demikian juga dengan normal P-P Plot yang akan terlihat nanti memperlihatkan hasil yang sama. Untuk lebih meyakinkan normalitas data tersebut dapat digunakan uji kolmogorov-semirnov dengan hasil berikut:

	Kolmogorov-Smirnov ^a			Shapiro-Wilk		
	Statistic	df	Sig.	Statistic	df	Sig.
Dana Promosi dalam Jutaan	,086	35	,200*	,968	35	,384
Kompetensi Tutor	,132	35	,132	,929	35	,026
Jumlah Peserta Didik Baru	,118	35	,200*	,945	35	,079

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

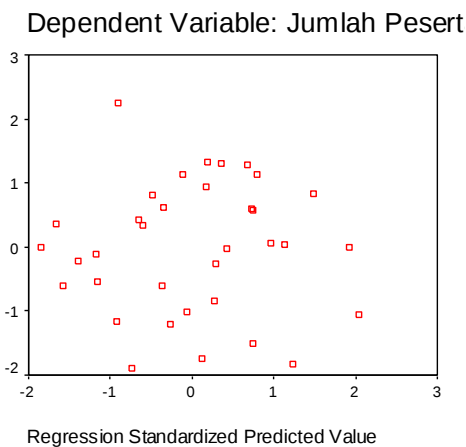
a. Lilliefors Significance Correction

Output ini memperlihatkan bahwa angka Sig. dari Kolmogorov Seminrnov untuk variabel x_1 (0,200), variabel x_2 (0,132), dan variabel y (0,200) yang lebih besar dari 0,05 sebagai toleransi kesalahan, maka ketiga data variabel itu berdistribusi normal.



Gambar di atas menguatkan kesimpulan hasil uji F di atas sehingga menunjukkan bahwa antara variabel kompetensi tutor dan dana promosi dengan variabel jumlah siswa ada hubungan linear. Di samping itu, gambar tersebut juga menunjukkan normalitas distribusi data karena titik-titik berada di sekitar garis diagonal.

Scatterplot



Output di atas memperlihatkan bahwa tidak terjadi heteroskedastisitas dikarenakan penyebaran nilai-nilai residual terhadap harga-harga prediksi tidak membentuk suatu pola sehingga varians data bersifat homogen.

Berdasarkan penjelasan di atas, maka seluruh asumsi yang dibutuhkan terpenuhi sehingga model yang **BLUE (Best Linear Unbiased Estimator)** terpenuhi.

2. Aplikasi dengan Microsoft Excel

Untuk mencari permasalahan regresi ganda dapat digunakan rumus berikut ini:

$$Y = a + b_1x_1 + b_2x_2.$$

Keterangan

Y = adalah skor yang diprediksikan

a = *intercept* atau Konstanta

X_1 dan X_2 = variabel bebas I dan II

b_1 dan b_2 = koefisien regresi

Sedangkan cara untuk menghitung harga a, b_1 , dan b_2 menggunakan persamaan rumus sebagai berikut:

$$b_1 = \frac{(\sum x_2^2)(\sum x_1 y) - (\sum x_2 y)(\sum x_1 x_2)}{(\sum x_1^2)(\sum x_2^2) - (\sum x_1 x_2)^2}$$

$$b_2 = \frac{(\sum x_1^2)(\sum x_2 y) - (\sum x_1 y)(\sum x_1 x_2)}{(\sum x_1^2)(\sum x_2^2) - (\sum x_1 x_2)^2}$$

$$a = \bar{Y} - b_1(\bar{X}_1) - b_2(\bar{X}_2)$$

Keterangan:

$$\sum x_1^2 = \sum x_1^2 - \frac{(\sum x_1)^2}{n}$$

$$\sum x_2^2 = \sum x_2^2 - \frac{(\sum x_2)^2}{n}$$

$$\sum x_1 x_2 = \sum x_1 x_2 - \frac{(\sum x_1)(\sum x_2)}{n}$$

$$\sum x_1 y = \sum x_1 y - \frac{(\sum x_1)(\sum y)}{n}$$

$$\sum x_2 y = \sum x_2 y - \frac{(\sum x_2)(\sum y)}{n}$$

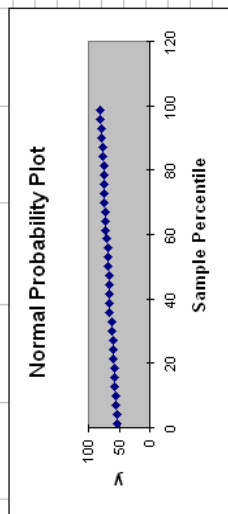
$$\sum y^2 = \sum y^2 - \frac{(\sum y)^2}{n}$$

Sedangkan aplikasinya sebagai berikut:

	A	B	C	D	E	F	G	H	I	J	K	L	M	N	O
1	x ₁	x ₂	y	x ₁ ²	x ₂ ²	y ²	x ₁ y	x ₂ y	x ₁ x ₂						
2	5,21	60	54	27	3.600	2.916	281	3.240	312		56,04				
3	5,6	39	55	31	1.521	3.025	308	2.145	218		2039,89				
4	5,6	42	58	31	1.764	3.364	325	2.436	235		107,74				
5	5,8	62	60	34	3.844	3.600	348	3.720	360		255,30				
6	5,8	52	72	34	2.704	5.184	418	3.744	302		1267,57				
7	5,9	47	59	35	2.209	3.481	348	2.773	277		2281,14				
8	6,2	36	54	38	1.444	2.916	335	2.052	236						
9	6,2	41	57	38	1.681	3.249	353	2.337	254		384214,68				
10	6,5	62	60	42	3.844	3.600	390	3.720	403		102699,13				
11	6,6	41	57	44	1.681	3.249	376	2.337	271		3.741				
12	6,6	49	65	44	2.401	4.225	429	3.185	323						
13	6,7	49	65	45	2.401	4.225	436	3.185	328		43523,611				
14	6,8	52	68	46	2.704	4.624	462	3.536	354		102699,128				
15	6,8	62	68	46	3.844	4.624	462	4.216	422		0,424				
16	6,9	49	68	48	2.401	4.624	469	3.332	338						
17	7,2	52	72	52	2.704	5.184	518	3.744	374		17,513				
18	7,4	55	75	55	3.025	5.625	555	4.125	407						
19	7,6	36	56	58	1.296	3.136	426	2.016	274						
20	7,6	62	65	58	3.844	4.225	494	4.030	471						
21	7,6	57	70	58	3.249	4.900	532	3.990	433						
22	7,6	53	73	58	2.809	5.329	555	3.869	403						
23	7,6	56	76	58	3.136	5.776	578	4.256	426						
24	7,8	80	75	61	3.600	5.625	585	4.500	468						
25	7,9	42	62	62	1.764	3.844	490	2.604	332						
26	7,9	52	65	62	2.704	4.225	514	3.360	411						
27	8,2	44	62	67	1.936	3.844	508	2.728	361						
28	8,5	54	75	72	2.916	5.625	636	4.050	459						
29	8,5	53	78	72	2.809	6.084	663	4.134	451						
30	8,6	54	78	74	2.916	6.084	671	4.212	464						
31	8,9	54	74	79	2.916	5.476	659	3.996	481						
32	8,9	62	81	79	3.844	6.561	721	5.022	552						
33	9,2	54	75	85	2.916	5.625	690	4.050	497						
34	9,7	51	67	94	2.601	4.489	650	3.417	495						
35	9,9	62	76	98	3.844	5.776	752	4.712	614						
36	9,9	60	80	98	3.600	6.400	792	4.800	594						
37	260	1.818	2.355	1.963	96.472	160.739	17.729	123.593	13.597						
38	7,4	51,9	67,3												

Di samping cara di atas, dapat juga dilakukan dengan mengaplikasi Microsoft Excel melalui Data Analysis. Setelah regression diklik, dan variabel diisi sesuai dengan tempatnya, dan menu-menu lain yang dibutuhkan diklik, maka akan keluar output yang sebagiannya ditampilkan berikut ini.

	A	B	C	D	E	F	G	H	I
1	SUMMARY OUTPUT								
2									
3	Regression Statistics								
4	Multiple R	0,808823549							
5	R Square	0,654195533							
6	Adjusted R Square	0,632582764							
7	Standard Error	4,964969125							
8	Observations	35							
9									
10	ANOVA								
11		df	SS	MS	F	Significance F			
12	Regression	2	1492,313468	746,1567341	30,26892231	4,18109E-08			
13	Residual	32	788,823389	24,56091841					
14	Total	34	2281,142857						
15									
16	Coefficients	Standard Error	t Stat	P-value	Lower 95%	Upper 95%	Lower 95,0%	Upper 95,0%	
17	Intercept	17,51300768	6,634552605	2,6396667	0,012717986	3,998877605	31,02713766	3,998877605	31,02713766
18	x1	3,741167862	0,689739019	5,346518846	7,24902E-06	2,315847319	5,166488384	2,315847319	5,166488384
19	x2	0,423797274	0,1163975629	3,654192498	0,000915135	0,187562248	0,660031699	0,187562248	0,660031699
20									
21									
22									
23	PROBABILITY OUTPUT								
24		y							
25	Percentile								
26	1,428571429	54							
27	4,285714286	54							
28	7,142857143	55							
29	10	56							
30	12,85714286	57							
31	15,71428571	57							
32	18,57142857	58							
33	21,42857143	59							
34	24,28571429	60							
35	27,14285714	60							
36	30	62							



Dari tiga cara di atas ternyata menghasilkan skor persamaan regresinya adalah Jumlah peserta = $17,513 + (0,424 \times \text{kompetensi tutor}) + (3,741 \times \text{Dana Promosi})$.

BAB VI

ANALISIS KOMPARASI

A. Pendahuluan

Analisis komparasi berarti menguji parameter populasi yang berbentuk perbandingan melalui ukuran sampel yang juga berbentuk perbandingan. Hal ini juga berarti menguji hipotesis mengenai ada tidaknya perbedaan antarvariabel yang sedang diteliti. Jika ada perbedaan, apakah perbedaan itu signifikan atau hanya terjadi secara kebetulan.

B. Macam-macam Teknik Analisis Komparasi

Terdapat dua model komparasi, yaitu komparasi antara dua sampel dan komparasi antara lebih dari dua sampel. Masing-masing juga dibagi menjadi dua jenis, yaitu komparasi antar sampel yang berkorelasi dan komparasi antara sampel yang tidak berkorelasi. Untuk lebih jelasnya dapat dilihat pada gambar berikut ini:

MACAM DATA	BENTUK HIPOTESIS			
	Komparatif (Dua Sampel)		Komparatif (Lebih dari dua sampel)	
	Related	Independen	Related	Independen
NOMINAL	Mc Nemar	Fisher Exact Probability χ^2 Two Sample	χ^2 for k Sample Cochran Q	χ^2 for k Sample
ORDINAL	Sign Test Wilcoxon Matched Pairs	Median Test Mann-Whitney U Test Kolmogorov- Smirnov Wald- Woldfowitz	Friedman Two-way Anova	Median Extention Kruskal- Wallis One Way Anova

INTERVAL RASIO	T-test of Related	T-test of independent	One-way Anova Two-way Anova	One-way Anova Two-way Anova
-------------------	----------------------	--------------------------	--------------------------------------	--------------------------------------

C. Komparasi Dua Sampel Berkorelasi

1. Statistik Parametrik: T-test of Related

T-test of Related adalah statistik parametrik yang digunakan untuk menguji hipotesis komparatif rata-rata dua sampel bila datanya berbentuk interval atau rasio yang menggunakan rancangan perbandingan sebelum dan sesudah adanya perlakuan/treatment.

Contoh:

Dilakukan penelitian untuk mengetahui ada tidaknya perbedaan prestasi belajar siswa sebelum dan sesudah diberi kursus. Berdasarkan data dari 35 siswa yang dipilih secara random diketahui bahwa prestasi belajar siswa sebelum dan sesudah mengikuti kursus adalah sebagai berikut:

	sebelum	sesudah			
1	75,00	85,00	18	70,00	80,00
2	80,00	90,00	19	80,00	90,00
3	65,00	75,00	20	65,00	60,00
4	70,00	75,00	21	75,00	75,00
5	75,00	75,00	22	80,00	85,00
6	80,00	90,00	23	70,00	80,00
7	65,00	70,00	24	90,00	95,00
8	80,00	85,00	25	70,00	75,00
9	90,00	95,00	26	75,00	79,00
10	75,00	70,00	27	87,00	85,00
11	60,00	65,00	28	78,00	79,00
12	70,00	75,00	29	85,00	83,00
13	75,00	85,00	30	87,00	85,00
14	70,00	65,00	31	79,00	80,00
15	80,00	95,00	32	80,00	82,00
16	65,00	65,00	33	86,00	84,00
17	75,00	80,00	34	82,00	85,00
			35	84,00	85,00

Proses Pengambilan Keputusan.

Hipotesis:

Ho Prestasi belajar siswa antara sebelum dan sesudah kursus adalah identik/sama

Ha Prestasi belajar siswa antara sebelum dan sesudah kursus adalah berbeda

Dasar Pengambilan Keputusan:

Dengan membandingkan t_{hitung} dengan t_{tabel} dengan ketentuan:

Ho diterima $t_{hitung} < t_{tabel}$

Ho ditolak $t_{hitung} \geq t_{tabel}$

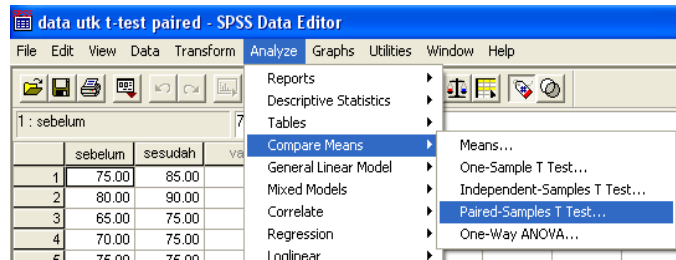
Dengan menggunakan angka probabilitas, dengan ketentuan:

Ho diterima Probabilitas > taraf nyata (α)

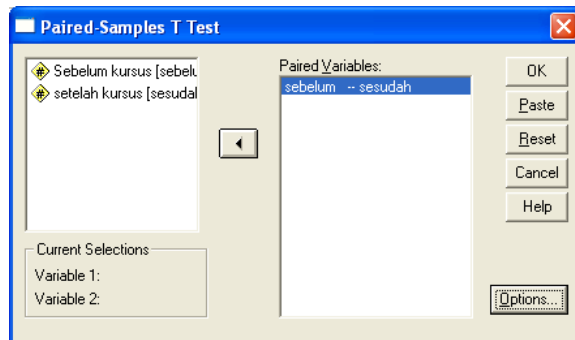
Ho ditolak Probabilitas $\leq$ taraf nyata (α)

a. Aplikasi dengan SPSS

1. Setelah data diinput, lalu klik **Analyze** ➤ **Compare Means** ➤ **Paired Samples T Test** sebagaimana gambar di bawah ini:



2. Setelah keluar gambar sebagai berikut: arahkan variabel x (sebelum) dan y (sesudah) ke dalam **Paired Variables**, lalu klik **Ok**.



3. Tampilan di bawah ini adalah hasil penghitungan itu.

Paired Samples Statistics

		Mean	N	Std. Deviation	Std. Error Mean
Pair 1	Prestasi Belajar Sebelum Kursus	76,3714	35	7,66614	1,29581
	Prestasi Belajar Sesudah Kursus	80,2000	35	8,79438	1,48652

Output ini menunjukkan bahwa sampel penelitian ini adalah 35, rata-rata prestasi belajar sebelum mengikuti kursus adalah 76,3714 dengan standard deviasi sebesar 7,66614 dan sesudahnya adalah 80,2 dengan standard deviasi sebesar 8,79438. Jadi rentang data sebelum kursus adalah antara 53,37298 sampai dengan 99,36982. Sedangkan rentang data setelah kursus adalah antara 53,81686 sampai dengan 106,5831. Skor ini merupakan penjabaran dari $\text{Mean} \pm 3 \text{ Standard Deviasi}$. Sedangkan rentang rata-rata sebelum kursus adalah antara 75,07559 sampai dengan 77,66721. Sedangkan rentang rata-rata setelah kursus adalah antara 78,71348 sampai dengan 81,68652. Skor interval ini adalah dari $\text{Mean} \pm 1 \text{ Standard Error of Mean}$.

Paired Samples Correlations

		N	Correlation	Sig.
Pair 1	Prestasi Belajar Sebelum Kursus & Prestasi Belajar Sesudah Kursus	35	,816	,000

Output ini menunjukkan bahwa korelasi antara variabel sebelum dan sesudah adalah 0,816. Hal ini menunjukkan bahwa korelasi data dari masing-masing responden adalah sangat erat.

Paired Samples Test

		Pair 1	
		Prestasi Belajar Sebelum Kursus - Prestasi Belajar Sesudah Kursus	
Paired Differences	Mean	-3,8286	
	Std. Deviation	5,10182	
	Std. Error Mean	,86237	
	95% Confidence Interval of the Difference	Lower	-5,5811
		Upper	-2,0760
t		-4,440	
df		34	
Sig. (2-tailed)		,000	

Output ini menunjukkan bahwa perbedaan rata-rata prestasi belajar sebelum dan sesudah kursus adalah 3,8286, dengan standard deviasi 5,10182 dan standard error of mean sebesar 0,86237. Hal yang sangat penting dari output di atas adalah t_{hitung} : -4,440. Bila t_{hitung} ini dimutlakkan akan menjadi: 4,440. Skor ini ternyata lebih tinggi dari $t_{tabel} (0,5;34)$: 2,032243. Dengan demikian, maka H_0 ditolak dan H_a diterima. Kesimpulan ini sama apabila digunakan skor sig untuk 2 sisi, yaitu 0,000 yang jauh lebih kecil bila dibandingkan dengan kesalahan yang ditoleransi yaitu 0,05 (5%). Berangkat dari hasil analisis ini dapat disimpulkan bahwa prestasi belajar siswa antara sebelum dan sesudah kursus adalah berbeda.

b. Aplikasi dengan Microsoft Excel

Langkah-langkah untuk mengaplikasikan analisis t-test of related dengan Microsoft Excel adalah sebagai berikut:

1. Carilah selisih masing-masing skor sebelum dan sesudah mendapatkan perlakuan.
2. Carilah rata-rata skor dari selisih di atas.
3. Kurangilah masing-masing skor hasil selisih di atas dengan rata-rata selisihnya.
4. Kuadratkan masing-masing hasil pengurangan di nomor 3.
5. Jumlahkan seluruh skor dari hasil di nomor 4.
6. Hasil dari nomor 5 dimasukkan ke rumus

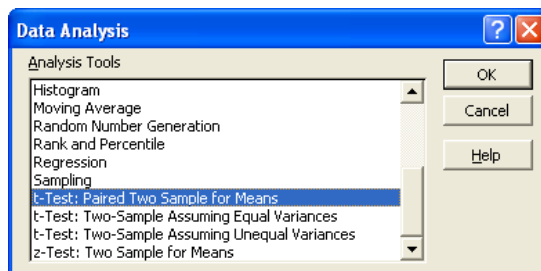
$$t = \frac{Md}{\sqrt{\frac{\sum x^2 d}{N(N-1)}}}$$

Untuk lebih jelasnya dapat dilihat pada aplikasi Microsoft Excel berikut ini.

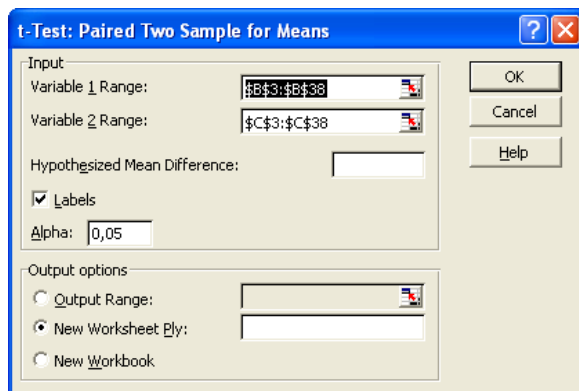
A4B										
	A	B	C	D	E	F	G	H	I	J
1	URUTAN ANALISIS T-TEST PAIRED									
2										
3	NO	Pre	Post	d	d-Md	xd ²				
4	01	75	85	-10	-6,17	38,09		0,744 =F39/(35*34)		
5	02	80	90	-10	-6,17	38,09		0,862 =SQRT(H4)		
6	03	65	75	-10	-6,17	38,09				
7	04	70	75	-5	-1,17	1,37		-4,440 =D39/H5		
8	05	75	75	0	3,83	14,66				
9	06	80	90	-10	-6,17	38,09				
10	07	65	70	-5	-1,17	1,37				
11	08	80	85	-5	-1,17	1,37				
12	09	90	95	-5	-1,17	1,37				
13	10	75	70	5	8,83	77,94				
14	11	60	65	-5	-1,17	1,37				
15	12	70	75	-5	-1,17	1,37				
16	13	75	85	-10	-6,17	38,09				
17	14	70	65	5	8,83	77,94				
18	15	80	95	-15	-11,17	124,80				
19	16	65	65	0	3,83	14,66				
20	17	75	80	-5	-1,17	1,37				
21	18	70	80	-10	-6,17	38,09				
22	19	80	90	-10	-6,17	38,09				
23	20	65	60	5	8,83	77,94				
24	21	75	75	0	3,83	14,66				
25	22	80	85	-5	-1,17	1,37				
26	23	70	80	-10	-6,17	38,09				
27	24	90	95	-5	-1,17	1,37				
28	25	70	75	-5	-1,17	1,37				
29	26	75	79	-4	-0,17	0,03				
30	27	87	85	2	5,83	33,97				
31	28	78	79	-1	2,83	8,00				
32	29	85	83	2	5,83	33,97				
33	30	87	85	2	5,83	33,97				
34	31	79	80	-1	2,83	8,00				
35	32	80	82	-2	1,83	3,34				
36	33	86	84	2	5,83	33,97				
37	34	82	85	-3	0,83	0,69				
38	35	84	85	-1	2,83	8,00				
39		76,371	80,200	-3,829		884,97				

Di samping cara di atas, t-test juga dapat diaplikasikan melalui menu **Data Analysis** sebagai berikut:

1. Setelah dipilih **Data Analysis** dari **Tools**, maka pilihlah **t-test: Paired Two Sample for Means** sebagaimana gambar berikut ini.



- Setelah keluar gambar seperti berikut ini, klik di kolom **Variable 1 Range**, lalu bloklah seluruh data sebelum kursus. Selanjutnya klik pada kolom **Variable 2 Range**, lalu bloklah seluruh data setelah kursus. Aktifkan **Labels** kalau di atas data yang kita ketik diberi label dan juga termasuk diblok untuk ditampilkan. Setelah itu klik **Ok**.



- Berikut ini adalah output dari aplikasi di atas. Terlihat bahwa hasilnya sama dengan hasil penghitungan dengan SPSS dan penjabaran rumus dengan Microsoft Excel. Kelebihan dari aplikasi **Data Analysis** adalah adanya tampilan t_{tabel} dengan alpha 5% baik untuk uji satu sisi maupun dua sisi.

	A	B	C
1	t-Test: Paired Two Sample for Means		
2			
3		<i>Pre</i>	<i>Post</i>
4	Mean	76,37142857	80,2
5	Variance	58,7697479	77,34117647
6	Observations	35	35
7	Pearson Correlation	0,816404533	
8	Hypothesized Mean Difference	0	
9	df	34	
10	t Stat	-4,439618065	
11	P(T<=t) one-tail	4,51832E-05	
12	t Critical one-tail	1,690923455	
13	P(T<=t) two-tail	9,03663E-05	
14	t Critical two-tail	2,032243174	

Dengan derajat kebebasan $35-1 = 34$, untuk alpha 5% t_{tabel} nya adalah 2,032243 untuk uji dua sisi atau 1,690923 untuk uji satu sisi dan untuk alpha 1% = 2,728393 untuk uji dua sisi atau 2,441147 untuk uji satu sisi. Karena $t_{\text{hitung}} -4,439618065$ lebih besar dari t_{tabel} baik untuk alpha 1% maupun 5% maka H_0 ditolak dan H_a diterima, artinya kesimpulan dari sampel bahwa

prestasi belajar siswa berbeda antara sebelum dan sesudah kursus, dalam kasus ini perubahan itu menuju kepada peningkatan, berlaku untuk populasi, baik dalam kesalahan 1% maupun 5%.

2. Statistik Non-Parametrik

a. McNemar

McNemar digunakan untuk menguji hipotesis komparative apabila datanya bertipe nominal yang berkorelasi. Rancangan penelitiannya biasanya berbentuk perbandingan skor sebelum dan sesudah adanya perlakuan/treatment. Jadi desain penelitiannya sering disebut *"before after"*.

Contoh:

Seorang peneliti ingin mengetahui pengaruh perilaku seorang Ketua yang bernama Ikuelaji selama satu periode menjabat terhadap jumlah suara yang diperoleh dalam pemilihan Ketua pada suatu Sekolah Tinggi pada periode berikutnya. Selama ia menjabat, mayoritas Tenaga Pendidik menganggap ketua tersebut sering melakukan korupsi, mengelola keuangan secara tidak transparan, dan gemar menekan sivitas akademika yang kritis. Pada pemilihan tahun 2005 mendapat suara 180 dari 200 Tenaga Pendidik, Sedangkan pada pemilihan 2009 dia mendapat suara 173 suara. Dari 20 Tenaga Pendidik yang pada pemilihan 2005 tidak memilihnya 10 di antaranya menjadi memilihnya pada tahun 2009, sedangkan 10 lainnya tetap tidak memilihnya. Dari 180 Tenaga Pendidik pendukungnya pada tahun 2005 ada sebanyak 16 Tenaga Pendidik akhirnya tidak memilihnya, sedangkan 164 lainnya tetap setia memilihnya. Tabulasi Perbandingan hasil suara dapat dilihat pada tabel berikut ini.

Pemilihan 2005 * Pemilihan 2009 Crosstabulation

Count		Pemilihan 2009		Total
		Tidak Memilih	Memilih	
Pemilihan 2005	Tidak Memilih	10	10	20
	Memilih	16	164	180
Total		26	174	200

Proses Pengambilan Keputusan.

Hipotesis:

Ho Tidak terdapat perubahan jumlah pemilih antara pemilihan tahun 2005 dan 2009

Ha Terdapat perubahan jumlah pemilih antara pemilihan tahun 2005 dan 2009

Dasar Pengambilan Keputusan:

Dengan membandingkan χ^2_{hitung} dengan χ^2_{tabel} dengan ketentuan:

Ho diterima $\chi^2_{hitung} < \chi^2_{tabel}$

Ho ditolak $\chi^2_{hitung} \geq \chi^2_{tabel}$

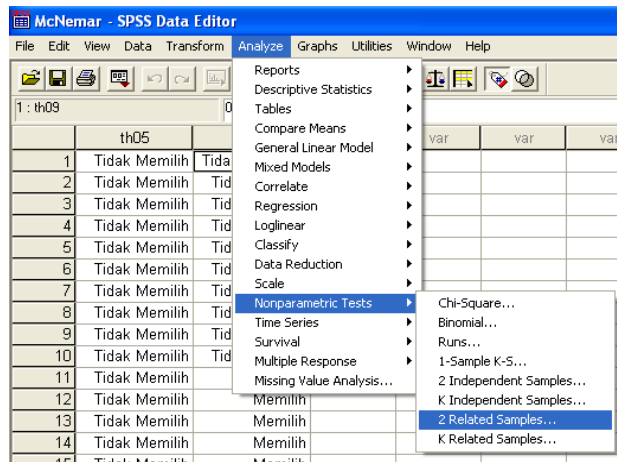
Dengan menggunakan angka probabilitas, dengan ketentuan:

Ho diterima Probabilitas > taraf nyata (α)

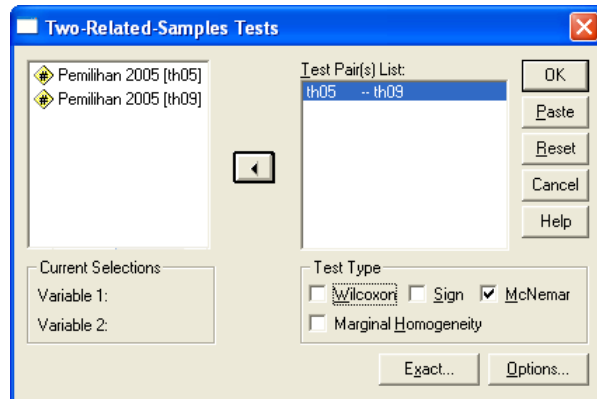
Ho ditolak Probabilitas $\leq$ taraf nyata (α)

1) Aplikasi dengan SPSS

- a) Setelah data diinput, maka klik **Analyze** ➤ **Nonparametric Test** ➤ **2 Related Samples...**, sebagaimana gambar berikut ini.



- b) Setelah keluar seperti gambar berikut ini, destinasikan kedua kumpulan data sampel tersebut ke dalam kotak **Test Pais(s) List**, lalu non-aktifkan **Wilcoxon** dan aktifkan **McNemar** lalu klik **Ok**.



c) Berikut ini adalah Output hasil analisis di atas.

Pemilihan 2005 & Pemilihan 2009

Pemilihan 2005	Pemilihan 2009	
	0	1
0	10	10
1	16	164

Output di atas memperlihatkan bahwa:

1. Jumlah Tenaga Pendidik yang pada pemilihan 2005 tidak memilih Ikuelaji dan pada pemilihan 2009 berubah memilihnya sebanyak 10 orang (+/A).
2. Jumlah Tenaga Pendidik yang pada pemilihan 2005 tidak memilih Ikuelaji dan pada pemilihan 2009 juga tetap tidak memilih memilihnya sebanyak 10 orang.
3. Jumlah Tenaga Pendidik yang pada pemilihan 2005 memilih Ikuelaji dan pada pemilihan 2009 juga tetap memilihnya sebanyak 164 orang.
4. Jumlah Tenaga Pendidik yang pada pemilihan 2005 memilih Ikuelaji dan pada pemilihan 2009 berubah tidak memilihnya sebanyak 16 orang (-/D).

Test Statistics^b

	Pemilihan 2005 & Pemilihan 2009
N	200
Chi-Square ^a	,962
Asymp. Sig.	,327

a. Continuity Corrected

b. McNemar Test

Output di atas memperlihatkan bahwa jumlah data sebanyak 200. Skor $\chi^2 = 0,962$ apabila hal ini dibandingkan χ^2_{tabel} dengan derajat kebebasan 1 = 3,841455 untuk alpha 5% dan 6,634891 untuk alpha 1%, maka H_0 diterima dan H_a ditolak, artinya kesimpulan dari sampel bahwa perilaku seorang ketua tidak mempunyai pengaruh terhadap perolehan suara dalam pemilihan pada pemilihan periode berikutnya. Kesimpulan ini juga sama ketika digunakan skor signifikansi output SPSS 0,327 yang lebih tinggi dibandingkan alpha yang ditoleransi yaitu 0,05.

2) Aplikasi dengan Microsoft Excel

Untuk keperluan analisis McNemar dengan Microsoft Excel, maka data perlu dikelompokkan ke dalam 4 bagian. Kelompok A adalah data yang berubah dari negatif menjadi positif, contoh untuk kasus ini adalah dari tidak memilih menjadi memilih. Kelompok B adalah data yang tetap negatif, kelompok C adalah data yang telah positif, sedangkan kelompok D adalah data yang ada perubahan dari positif menjadi negatif, untuk kasus ini adalah dari memilih menjadi tidak memilih. Setelah itu dihitung dengan rumus sebagai berikut:

$$\chi^2 = \frac{(|A - D| - 1)^2}{A + D}$$

Sedangkan aplikasi dengan Microsoft Excel dapat dilihat pada tampilan berikut ini.

	A	B	C	D
1			Pemilihan Tahun 2009	
2			Memilih	Tidak Memilih
3	Pemilihan	Tidak Memilih	10 (A)	10 (B)
4	Tahun 2005	Memilih	164 (C)	16 (D)
5				
6			1,884615385	
7			=((C3-D4-1)*2)/(C3+D4)	
8			0,942307692	
9			=C6/2	
10				

Hasil penghitungan Microsoft Excel ternyata tidak menghasilkan Skor χ^2 yang sama persis dengan hasil hitung SPSS. Hanya saja, dari tiga rumus McNemar yang penulis temukan dari berbagai buku, rumus di atas yang hasilnya mendekati hasil hitung SPSS. Karena selisihnya yang sangat sedikit, hasil SPSS: 0,962 dan Microsoft Excel: 0,942307692, maka kesimpulan akhirnya tetap

sama, yaitu menerima H_0 karena χ^2_{hitung} lebih kecil dibanding χ^2_{tabel} baik untuk alpha 1% maupun 5%.

b. Sign Test

Sign test digunakan untuk menguji hipotesis komparatif dua sampel yang berkorelasi apabila datanya bertipe ordinal atau interval/rasio yang ditransform menjadi ordinal. Analisis ini disebut sign test (uji tanda) karena analisisnya didasarkan pada tanda + (positif) dan – (negatif). Sebuah data akan diberi tanda + manakala skor data setelah ada perlakuan lebih tinggi dibanding sebelumnya. Sedangkan data yang skornya lebih rendah bila dibandingkan sebelum adanya perlakuan, maka akan diberi tanda –.

Contoh:

Peneliti ingin mengetahui pengaruh Tunjangan Sertifikasi Tenaga Pendidik terhadap pengeluaran untuk makan keluarga.

Proses Pengambilan Keputusan.

Hipotesis:

- | | |
|----|---|
| Ho | Jumlah pengeluaran untuk makan keluarga antara sebelum dan sesudah mendapatkan tunjangan sertifikasi tenaga pendidik adalah sama |
| Ha | Jumlah pengeluaran untuk makan keluarga antara sebelum dan sesudah mendapatkan tunjangan sertifikasi tenaga pendidik adalah berbeda |

Dasar Pengambilan Keputusan:

Dengan membandingkan χ^2_{hitung} dengan χ^2_{tabel} dengan ketentuan:

Ho diterima $\chi^2_{hitung} < \chi^2_{tabel}$

Ho ditolak $\chi^2_{hitung} \geq \chi^2_{tabel}$

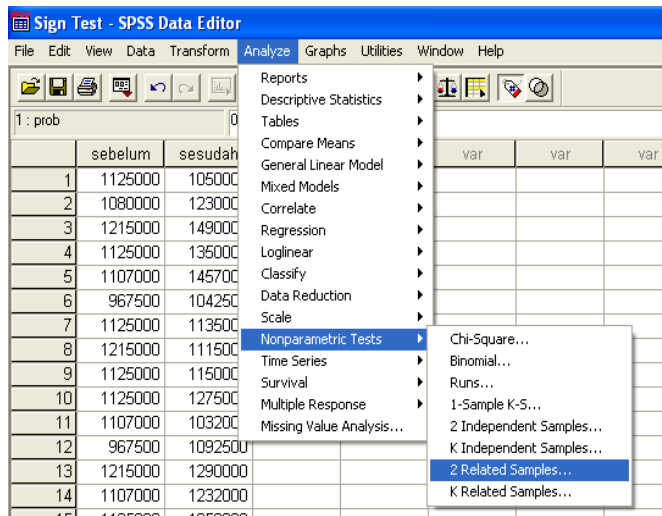
Dengan menggunakan angka probabilitas, dengan ketentuan:

Ho diterima Probabilitas > taraf nyata (α)

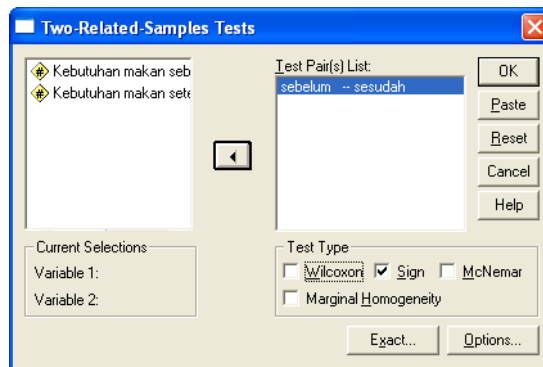
Ho ditolak Probabilitas $\leq$ taraf nyata (α)

1) Aplikasi dengan SPSS

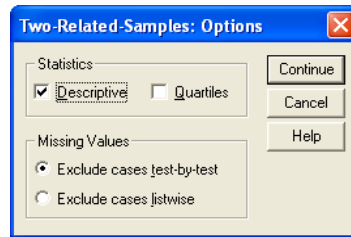
- a) Setelah data diinput, maka klik **Analyze** ➤ **Nonparametric Test** ➤ **2 Related Samples...**, sebagaimana gambar berikut ini.



- b) Setelah keluar seperti gambar berikut ini, destinasikan kedua kumpulan data sampel tersebut ke dalam kotak **Test Pair(s) List**, lalu non-aktifkan **Wilcoxon** dan aktifkan **Sign**. Apabila berkeinginan mengaktifkan deskripsi data, maka klik **Option**.



- c) Setelah keluar gambar berikut ini klik kotak yang ada di depan **Descriptive** lalu klik **Continue**, lalu klik **Ok**.



d) Berikut ini adalah Output hasil analisis di atas.

Descriptive Statistics

	N	Mean	Std. Deviation	Minimum	Maximum
Kebutuhan makan sebelum Sertifikasi	24	1115437	72260,145	967500	1215000
Kebutuhan makan setelah Sertifikasi	24	1206896	132295,534	1032000	1490000

Output di atas memperlihatkan bahwa rata-rata pengeluaran untuk makan keluarga tenaga pendidik sebelum mendapatkan tunjangan sertifikasi adalah Rp.1.115.437, sedangkan sesudahnya sebanyak Rp.1.206.896.

Frequencies

	N
Kebutuhan makan setelah Sertifikasi - Kebutuhan makan sebelum Sertifikasi	6
Negative Difference ^a	18
Positive Difference ^b	0
Ties ^c	24
Total	

a. Kebutuhan makan setelah Sertifikasi <
Kebutuhan makan sebelum Sertifikasi

b. Kebutuhan makan setelah Sertifikasi >
Kebutuhan makan sebelum Sertifikasi

c. Kebutuhan makan setelah Sertifikasi =
Kebutuhan makan sebelum Sertifikasi

Output ini memperlihatkan terdapat 6 keluarga yang pengeluaran untuk makan keluarga setelah mendapat tunjangan sertifikasi lebih kecil dibanding sebelumnya, sementara 18 keluarga menunjukkan bahwa pengeluaran untuk makan keluarga setelah sertifikasi lebih tinggi dibanding sebelumnya.

Test Statistics^b

	Kebutuhan makan setelah Sertifikasi - Kebutuhan makan sebelum Sertifikasi
Exact Sig. (2-tailed)	,023 ^a

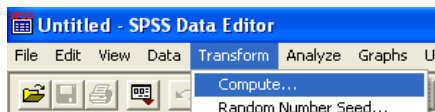
a. Binomial distribution used.

b. Sign Test

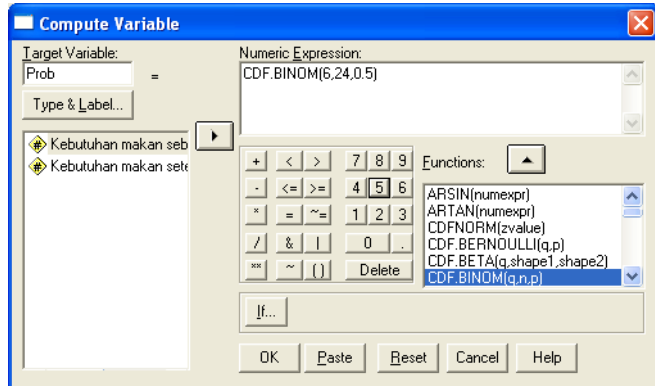
Output ini menghasilkan bahwa skor **Exact Sig untuk dua sisi** sebesar 0,023 yang hal ini ternyata lebih kecil dibanding alpha yang ditoleransi yaitu 0,05. Oleh karena itu, H_0 ditolak dan H_a diterima. Dengan demikian, hipotesis yang menyatakan bahwa pengeluaran untuk makan keluarga antara sebelum dan setelah mendapat tunjangan sertifikasi adalah berbeda dapat diberlakukan untuk populasi.

Untuk mencari skor signifikansi dari uji tanda semestinya menggunakan harga-harga dalam test binomial. Sedangkan proses untuk mengaplikasikan dapat menggunakan langkah sebagai berikut:

- a) Bukalah file **Sign Test**.
- b) Klik **Transform ► Compute**, sebagaimana gambar berikut ini.



- c) Berilah nama variabel pada kotak **Target Variable**, lalu pilihlah **CDF.BINOM(q.n.p)** pada **Function** lalu klik tanda segitiga. Setelah itu pada **q** isikan angka terkecil, untuk kasus ini 6, karena 6 lebih kecil dibanding 18, pada **n** isikan **jumlah data**, untuk kasus ini **24**, dan pada **p** isikan probabilitasnya, karena kita berkeinginan tanda + dan – sama besarnya, maka isikan **0,5**, sebagaimana gambar berikut ini. Setelah itu klik **Ok**.



- d) Berikut ini adalah hasilnya, yaitu akan ada satu tambahan variabel baru yang berisi angka probabilitasnya untuk uji satu sisi. Ketika skor tersebut dikalikan 2 maka hasilnya akan sama dengan skor **Exact Sig untuk dua sisi**, yaitu sebesar 0,023.

Sign Test - SPSS Data Editor

File Edit View Data Transform Analyze Graphs Utilities Window

1 : sebelum 1125000

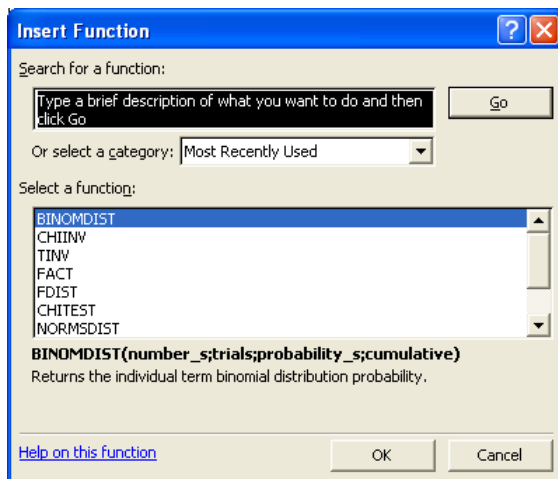
	sebelum	sesudah	prob
1	1125000	1050000	,011327922344208
2	1080000	1230000	,011327922344208

2) Aplikasi dengan Microsoft Excel

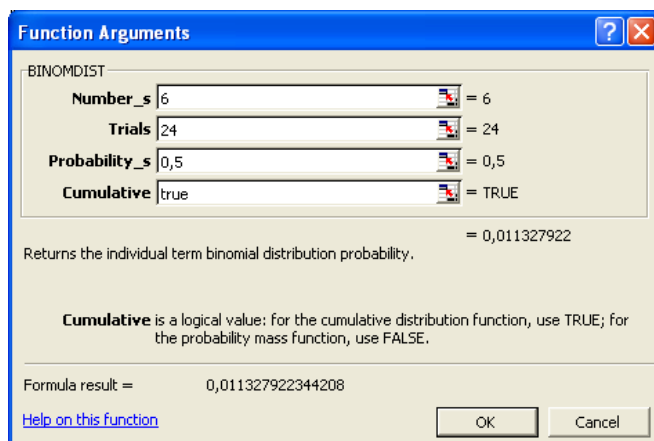
Analisis tanda dapat dilakukan dengan langkah-langkah sebagai berikut:

- Setelah data diinput, baik untuk sebelum maupun sesudahnya, maka masing-masing skor sesudah adanya berlakuan dikurangi dengan masing-masing skor sebelum adanya perlakuan.
- Setelah itu, buatlah satu kolom untuk digunakan mencatat tanda + manakala hasil pengurangan tadi positif (artinya tidak ada tanda positif), dan tanda – apabila hasil pengurangan tadi negatif.
- Jumlahkan masing-masing tanda tersebut.
- Berdasarkan jumlah terkecil dan jumlah sampel dicari skor alpha-nya dengan menggunakan menu binomial. Sedangkan untuk aplikasi Microsoft Excel dapat

dilakukan dengan prosedur berikut. Carilah menu **BINOMDIST** lalu klik dua kali.



- e) Setelah itu isilah pada kolom **Number's** dengan jumlah terkecil, lalu pada kolom **Trials** dengan jumlah seluruh data, pada kolom **Probability_s** dengan **0,5** kalau dikehendaki perubahan data sebelum dan sesudah tadi seimbang, dan yang terakhir pada kolom **Cumulative** diisi dengan **True**. Selanjutnya klik **Ok**. Sebagaimana gambar di bawah ini.



Aplikasi di atas dapat dilakukan dengan cara pintas sebagaimana tersebut di bawah. Dan ternyata hasil penghitungan kedua cara tersebut juga sama.

=BINOMDIST(6;24;0,5;TRUE)
0,011327922344208

- e) Sedangkan aplikasi prosedur sign test dengan menjabarkan rumus dapat dilihat pada gambar berikut ini.

	A	B	C	D	E	F	G	H	I	J	K
1	SEBELUM										
2	1.125.000	1.050.000	-Rp75.000	-							
3	1.080.000	1.230.000	Rp150.000	+							
4	1.215.000	1.490.000	Rp275.000	+							
5	1.125.000	1.350.000	Rp225.000	+							
6	1.107.000	1.457.000	Rp350.000	+							
7	967.500	1.042.500	Rp75.000	+							
8	1.125.000	1.135.000	Rp10.000	+							
9	1.215.000	1.115.000	-Rp100.000	-							
10	1.125.000	1.150.000	Rp25.000	+							
11	1.125.000	1.275.000	Rp150.000	+							
12	1.107.000	1.032.000	-Rp75.000	-							
13	967.500	1.092.500	Rp125.000	+							
14	1.215.000	1.290.000	Rp75.000	+							
15	1.107.000	1.232.000	Rp125.000	+							
16	1.125.000	1.050.000	-Rp75.000	-							
17	1.215.000	1.340.000	Rp125.000	+							
18	1.125.000	1.110.000	-Rp15.000	-							
19	1.080.000	1.330.000	Rp250.000	+							
20	1.125.000	1.250.000	Rp125.000	+							
21	1.107.000	1.182.000	Rp75.000	+							
22	967.500	1.217.500	Rp250.000	+							
23	1.215.000	1.340.000	Rp125.000	+							
24	1.080.000	1.155.000	Rp75.000	+							
25	1.125.000	1.050.000	-Rp75.000	-							
26	Bertanda "+" berjumlah			18							
27	Bertanda "-" berjumlah			6							
28	Bertanda "0" berjumlah			0							
29											

formula C2	=B2-A2
formula D2	=IF(A2>B2,"-";IF(A2<B2,"+",IF(A2=B2,"0")))
formula D26	=COUNTIF(D2:D25,"+")
formula D27	=COUNTIF(D2:D25,"-")
formula D28	=COUNTIF(D2:D25,"0")

N	24
jml tanda terkecil	6
skor dalam tabel	0,011
0,011 dari	0,011327922344208
kalau dikalikan 2	0,023
Kesimpulan	0,023 < 0,05 Ho ditolak, Ha diterima

Dari seluruh prosedur ternyata menghasilkan skor yang sama, yaitu skor alpha 0,023 yang lebih kecil dari alpha maksimal yang ditoleransi yaitu 0,05. Oleh karena itu, Ho ditolak dan Ha diterima.

Apabila jumlah sampel lebih besar dari 25, maka uji statistiknya menggunakan Chi Kuadrat yang rumusnya sebagai berikut:

$$\chi^2 = \frac{[(n_1 - n_2) - 1]^2}{n_1 + n_2}$$

Di mana:

n_1 = Banyak data positif

n_2 = Banyak data negatif.

Chi Kuadrat tersebut dibandingkan dengan Chi Kuadrat tabel dengan $dk = 1$.

c. Wilcoxon Matched Pairs

Analisis Wilcoxon Matched Pairs pada dasarnya merupakan penyempurnaan dari sign test. Sebagaimana dijelaskan di atas bahwa dalam sign test tidak memperhitungkan besarnya selisih nilai angka positif dan negatif. Dalam Wilcoxon, besarnya selisih ikut diperhitungkan.

Contoh:

Seorang peneliti berkeinginan untuk mengetahui perbedaan angka kredit penelitian sebelum dan sesudah pelatihan. Atau dengan kata lain, peneliti berkeinginan untuk mengetahui pengaruh pelatihan terhadap produktifitas penelitian dosen.

Proses Pengambilan Keputusan.

Hipotesis:

H_0 Jumlah angka kredit dosen dari unsur penelitian antara sebelum dan sesudah mengikuti pelatihan adalah sama

H_a Jumlah angka kredit dosen dari unsur penelitian antara sebelum dan sesudah mengikuti pelatihan adalah berbeda/tidak sama

Dasar Pengambilan Keputusan:

Dengan membandingkan z_{hitung} dengan z_{tabel} dengan ketentuan:

H_0 diterima $z_{hitung} < z_{tabel}$

H_0 ditolak $z_{hitung} \geq z_{tabel}$

Dengan menggunakan angka probabilitas, dengan ketentuan:

Ho diterima Probabilitas > taraf nyata (α)

Ho ditolak Probabilitas $\leq$ taraf nyata (α)

Dikarenakan tipe data penelitian ini adalah data rasio, maka direncanakan digunakan analisis t-test apabila asumsinya terpenuhi.

1) Aplikasi dengan SPSS

Salah satu asumsi yang harus terpenuhi untuk menggunakan t-test sebagai analisis parametrik adalah distribusi datanya harus normal. Prosedur untuk menguji distribusi data adalah sebagai berikut.

Setelah data diinput, lalu klik **Analyze** ➤ **Descriptive Statistics** ➤ **Explore**. Setelah itu destinaskan variabel itu ke dalam kotak **Dependent List**, lalu pada bagian **Display**, klik kotak **Plots**, Kemudian klik kotak **Plots**. Setelah itu aktifkan kotak **Normality Plots with tests**, lalu nonaktifkan pilihan **stem and leaf**, kemudian pilih **None** pada bagian **Boxplot**, Setelah itu klik **Ok**. Berikut ini adalah hasil analisis normalitas data tersebut.

Case Processing Summary

Tests of Normality

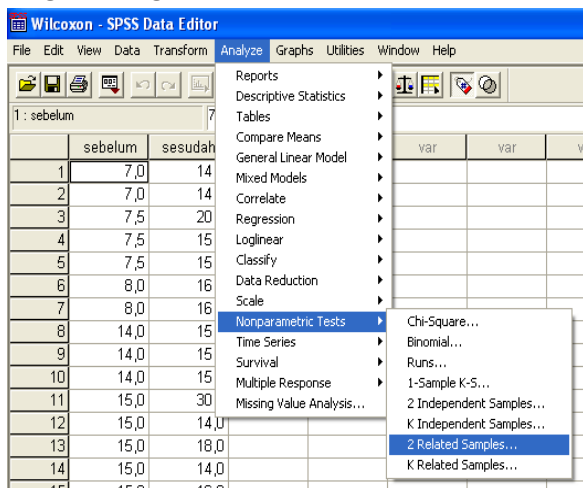
	Kolmogorov-Smirnov ^a			Shapiro-Wilk		
	Statistic	df	Sig.	Statistic	df	Sig.
Angka Kredit Penelitian Dosen sebelum pelatihan	,179	30	,015	,896	30	,007
Angka Kredit Penelitian Dosen sesudah pelatihan	,275	30	,000	,756	30	,000

a. Lilliefors Significance Correction

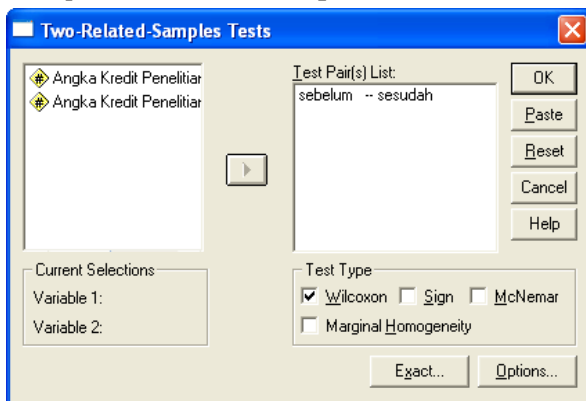
Untuk membuat kesimpulan tentang normalitas data dapat digunakan skor **Sig.** Yang ada pada hasil penghitungan **Kolmogorov-Smirnov**. Bila angka **Sig.** Lebih besar atau sama dengan 0,05, maka data tersebut berdistribusi normal, tetapi apabila kurang, maka data itu tidak berdistribusi normal. Karena **Sig.** Untuk data sebelum pelatihan skor sig-nya = 0,015 dan yang sesudah pelatihan sebesar 0,000 yang lebih kecil dibanding 0,05, maka data kedua variabel itu tidak berdistribusi normal.

Karena distribusi datanya tidak normal, maka tidak dapat digunakan analisis t-test. Sebagai gantinya akan digunakan analisis Wilcoxon. Sedangkan langkah-langkah untuk analisis tersebut adalah sebagai berikut.

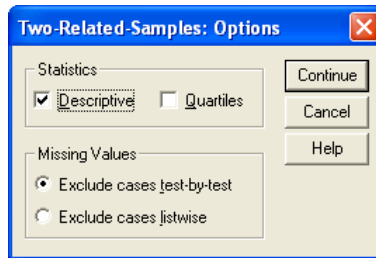
- a) Setelah data diinput, maka klik **Analyze** ➤ **Nonparametric Test** ➤ **2 Related Samples...**, sebagaimana gambar berikut ini.



- b) Setelah keluar seperti gambar berikut ini, destinasikan kedua kumpulan data sampel tersebut ke dalam kotak **Test Pair(s) List**. Apabila berkeinginan mengaktifkan deskripsi data, maka klik **Option**.



- c) Setelah keluar gambar berikut ini klik kotak yang ada di depan **Descriptive** lalu klik **Continue**, lalu klik **Ok**.



d) Berikut ini adalah Output hasil analisis di atas.

Descriptive Statistics

	N	Mean	Std. Deviation	Minimum	Maximum
Angka Kredit Penelitian Dosen sebelum pelatihan	30	15,117	5,0355	7,0	22,5
Angka Kredit Penelitian Dosen sesudah pelatihan	30	18,350	5,1947	14,0	30,0

Output di atas memperlihatkan bahwa sebelum mengikuti pelatihan, rata-rata dosen memiliki 15,117 angka kredit untuk penelitian dalam satu tahun; sedangkan setelah pelatihan rata-rata angka kreditnya adalah 18,350.

Ranks

		N	Mean Rank	Sum of Ranks
Angka Kredit Penelitian Dosen sesudah pelatihan - Angka Kredit Penelitian Dosen sebelum pelatihan	Negative Ranks	8 ^a	8,94	71,50
	Positive Ranks	21 ^b	17,31	363,50
	Ties	1 ^c		
	Total	30		

- a. Angka Kredit Penelitian Dosen sesudah pelatihan < Angka Kredit Penelitian Dosen sebelum pelatihan
- b. Angka Kredit Penelitian Dosen sesudah pelatihan > Angka Kredit Penelitian Dosen sebelum pelatihan
- c. Angka Kredit Penelitian Dosen sesudah pelatihan = Angka Kredit Penelitian Dosen sebelum pelatihan

Output di atas memberitahukan bahwa terdapat 8 orang yang angka kredit penelitian menurun setelah pelatihan, 21 orang meningkat, dan 1 di antara 30 orang mempunyai angka kredit penelitian tetap.

Test Statistics^b

	Angka Kredit Penelitian Dosen sesudah pelatihan - Angka Kredit Penelitian Dosen sebelum pelatihan
Z	-3,186 ^a
Asymp. Sig. (2-tailed)	,001

a. Based on negative ranks.

b. Wilcoxon Signed Ranks Test

Output terakhir ini adalah skor z_{hitung} : -3,186. Untuk alpha 5% maka z_{tabel} nya adalah 1,96. Dikarenakan skor z_{hitung} : -3,186 merupakan skor mutlak, maka ia lebih besar dibandingkan z_{tabel} sehingga H_0 ditolak dan H_a diterima. Hal ini juga sesuai dengan skor Asymp. Sig. (2-tailed) sebesar 0,001 yang jauh lebih kecil dibanding alpha 0,05.

2) Aplikasi dengan Microsoft Excel

Prosedur untuk menganalisis dengan menggunakan Wilcoxon adalah sebagai berikut:

- Input data dari kedua sampel tersebut ke dalam dua kolom.
- Masing-masing skor data sesudah pelatihan dikurangi dengan masing-masing data sebelum pelatihan.
- Berdasarkan skor hasil pengurangan tersebut, buatlah ranking dari skor terkecil menuju kepada skor terbesar tanpa memperhatikan tanda negatif. Apabila ada skor yang sama, maka skor ranking tersebut dijumlahkan lalu dibagi dengan jumlahnya. Misalkan untuk kasus data ini yaitu ranking 2 (berupa angka 1) ternyata ada 12. Cara untuk membuat ranking yang skornya sama adalah sebagai berikut:

$$\frac{2+3+4+5+6+7+8+9+10+11+12+13}{12} = \frac{90}{12} = 7,5$$

- Memisahkan antara ranking yang positif dan negatif. Maksud positif dan negatif di sini adalah ranking tersebut berasal dari hasil pengurangan yang bertanda negatif atau tidak.

e) Selanjutnya hasil ranking tersebut dijumlahkan.

Jumlah terkecil dari hasil pengurangan tersebut dibandingkan dengan skor tabel yang berdasarkan pada jumlah sampel. Seandainya jumlah penelitian itu 21, maka skor tabelnya adalah 59. Tabel yang tersedia untuk Wilcoxon adalah jumlah sampel maksimal 25. Apabila jumlah sampel lebih besar dari 25, maka distribusi akan mendekati distribusi normal. Oleh karena itu digunakan rumus z dalam pengujiannya. Sedangkan rumusnya adalah sebagai berikut:

$$z = \frac{T - \mu_T}{\sigma_T}$$

Di mana:

T merupakan jumlah jenjang/ranking yang kecil. Untuk kasus contoh di atas adalah 79,5. Sedangkan penjabaran μ_T dan σ_T sebagaimana dijabarkan di bawah ini.

$$\mu_T = \frac{n(n+1)}{1}$$

$$\sigma_T = \sqrt{\frac{n(n+1)(2n+1)}{24}}$$

Untuk lebih jelasnya dapat dilihat pada aplikasi berikut ini.

	A	B	C	D	E	F	G	H	I
1	sblm	ssdh	beda	ranking	+	-			
2	7	14	7	20,5	20,5	0		$232,5 = (30 \cdot 31) / 4$	
3	7	14	7	20,5	20,5	0		$2325,042 = (30 \cdot 31 \cdot 2 \cdot 30 + 1) / 24$	
4	7,5	20	12,5	29	29	0		$48,21869 = \text{SQRT}(H3)$	
5	7,5	15	7,5	23,5	23,5	0			
6	7,5	15	7,5	23,5	23,5	0		$-3,173 = (F32 - H2) / H4$	
7	8	16	8	26,5	26,5	0			
8	8	16	8	26,5	26,5	0			
9	14	15	1	7,5	7,5	0		$z = \frac{T - \mu_T}{\sigma_T}$	
10	14	15	1	7,5	7,5	0			
11	14	15	1	7,5	7,5	0		$\mu_T = \frac{n(n+1)}{1}$	
12	15	30	15	30	30	0			
13	15	14	-1	7,5	0	7,5		$\sigma_T = \sqrt{\frac{n(n+1)(2n+1)}{24}}$	
14	15	18	3	18	18	0			
15	15	14	-1	7,5	0	7,5			
16	15	16	1	7,5	7,5	0			
17	16	15	-1	7,5	0	7,5			
18	16	18	2	15,5	15,5	0			
19	16	17	1	7,5	7,5	0			
20	16	14	-2	15,5	0	15,5			
21	16	15	-1	7,5	0	7,5			
22	17	16	-1	7,5	0	7,5			
23	17	16	-1	7,5	0	7,5			
24	20	22	2	15,5	15,5	0			
25	20	21	1	7,5	7,5	0			
26	20	22	2	15,5	15,5	0			
27	20	29	9	28	28	0			
28	22,5	30	7,5	23,5	23,5	0			
29	22,5	16	-6,5	19	0	19			
30	22,5	22,5	0	1	1	0			
31	22,5	30	7,5	23,5	23,5	0			
32					385,5	79,5			

Skor z_{hitung} adalah -3,173. Apabila skor ini dibandingkan dengan $z_{\text{tabel};0,05}$ sebesar 1,96, maka H_0 ditolak dan H_a diterima, karena ketika z_{hitung} dimutlukkan maka menjadi 3,173 itu lebih besar dibanding 1,96. Berdasarkan hasil analisis ini dapat disimpulkan bahwa hipotesis yang berbunyi terdapat perbedaan angka kredit dosen dari penelitian antara sebelum dan sesudah pelatihan dapat diterima dan berlaku untuk populasi.

D. Komparasi Dua Sampel Independen

1. Statistik Parametrik: T-test of Independent

T-test of Independent digunakan untuk menguji hipotesis komparatif dua sampel independent bila tipe datanya adalah interval atau rasio.

Contoh:

Dilakukan penelitian untuk mengetahui kecepatan memasuki dunia kerja antara lulusan SMA dan STM.

Proses Pengambilan Keputusan.

Hipotesis:

Ho Masa menunggu untuk mendapatkan pekerjaan antara alumni STM dan SMA adalah sama.

Ha Masa menunggu untuk mendapatkan pekerjaan antara alumni STM dan SMA adalah tidak sama.

Dasar Pengambilan Keputusan:

Dengan membandingkan t_{hitung} dengan t_{tabel} dengan ketentuan:

Ho diterima $t_{hitung} < t_{tabel}$

Ho ditolak $t_{hitung} \geq t_{tabel}$

Dengan menggunakan angka probabilitas, dengan ketentuan:

Ho diterima Probabilitas > taraf nyata (α)

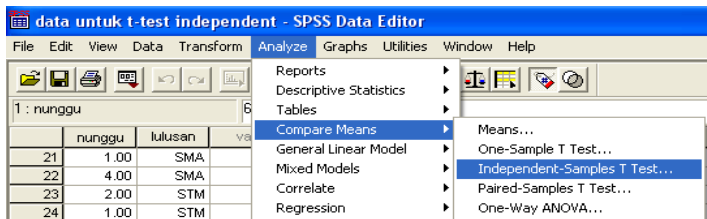
Ho ditolak Probabilitas $\leq$ taraf nyata (α)

a. Aplikasi dengan SPSS

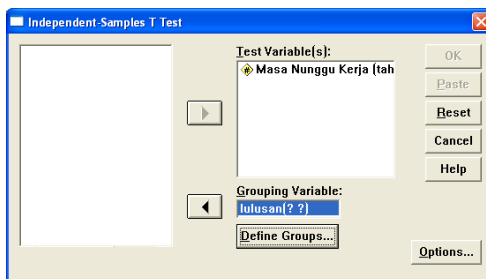
Sebelum menjelaskan prosedur analisis t-test independent dengan SPSS, akan dijelaskan data hasil penelitian tentang perbandingan waktu mendapatkan pekerjaan antara lulusan SMA dan STM sebagai berikut:

	nunggu	lulusan
1	6.00	SMA
2	3.00	SMA
3	5.00	SMA
4	2.00	SMA
5	5.00	SMA
6	1.00	SMA
7	2.00	SMA
8	3.00	SMA
9	1.00	SMA
10	3.00	SMA
11	2.00	SMA
12	4.00	SMA
13	3.00	SMA
14	4.00	SMA
15	2.00	SMA
16	3.00	SMA
17	1.00	SMA
18	5.00	SMA
19	1.00	SMA
20	3.00	SMA
21	1.00	SMA
22	4.00	SMA
23	2.00	STM
24	1.00	STM
25	3.00	STM
26	1.00	STM
27	3.00	STM
28	2.00	STM
29	2.00	STM
30	1.00	STM
31	3.00	STM
32	1.00	STM
33	1.00	STM
34	1.00	STM
35	3.00	STM
36	2.00	STM
37	1.00	STM
38	2.00	STM
39	2.00	STM
40	1.00	STM

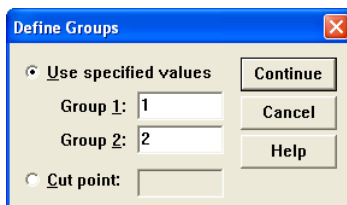
- 1) Setelah data diinput, lalu klik **Analyze ► Compare Means ► Independent Samples T Test** sebagaimana gambar di bawah ini:



- 2) Setelah keluar gambar seperti di bawah ini destinasikan variabel masa nunggu di **Test variable(s)** dan lulusan pada **Grouping variable**. Setelah itu klik **Define Groups**.



- 3) Setelah keluar gambar seperti di bawah ini ketiklah 1 di group 1, dan 2 di kolom Group 2. lalu klik **Continue**, lalu klik **Ok**.



- 4) Tampilan di bawah ini adalah out-put hasil penghitungan t-test independent dengan SPSS.

Group Statistics		
	Masa Nunggu Kerja (tahun)	
	LULUSAN	
	SMA	STM
N	22	18
Mean	2.9091	1.7778
Std. Deviation	1.50899	.80845
Std. Error Mean	.32172	.19055

Tampilan di atas menunjukkan bahwa sampel lulusan SMA sejumlah 22 dan lulusan STM 18 orang. Rata-rata lulusan

SMA memasuki dunia kerja setelah menunggu selama 2,9091 tahun sementara lulusan STM menunggu 1,7778 tahun.

Independent Samples Test			Masa Nunggu Kerja (tahun)	
			Equal variances assumed	Equal variances not assumed
Levene's Test for Equality of Variances	F		5.162	
	Sig.		.029	
t-test for Equality of Means	t		2.858	3.026
	df		38	33.262
	Sig. (2-tailed)		.007	.005
	Mean Difference		1.1313	1.1313
	Std. Error Difference		.39578	.37392
95% Confidence Interval of the Difference	Lower		.33009	.37080
	Upper		1.93253	1.89182

Hasil analisis t-test independent dengan SPSS senantiasa ada dua kolom, yaitu kolom untuk hasil analisis manakala variansnya homogen dan satunya untuk yang hiterogin. Penggunaan kolom ditentukan oleh hasil F-test atau signifikansinya. Apabila signifikansinya paling tinggi 0,05, maka hasil yang ada di kolom hiterogin (equal variances not assumed) yang digunakan. Sebaliknya kalau signifikansi untuk F-test lebih tinggi dari 0,05 maka yang digunakan adalah kolom yang homogen (Equal variances assumed). Karena signifikansi dari F-test 0,025 yang lebih kecil dari 0,05, maka yang digunakan adalah hasil hitung dalam kolom hiterogin.

T_{hitung} untuk analisis ini adalah 3,026. Apabila dibandingkan dengan $t_{tabel0,05;33}$ sebesar 2,034517, maka H_0 ditolak dan H_a diterima karena t_{hitung} lebih besar dibandingkan $t_{tabel0,05;33}$. Kesimpulan yang sama kita dapatkan ketika digunakan skor signifikansi sebesar 0,005 yang jauh lebih rendah dibandingkan dengan skor alpha 0,05. Berangkat dari hasil analisis ini, maka H_a yang berbunyi masa menunggu untuk mendapatkan kerja antara lulusan SMA dan STM itu berbeda dapat diterima dan berlaku untuk populasi.

b. Aplikasi dengan Microsoft Excel

Terdapat 2 rumus t-test yang dapat digunakan untuk menguji hipotesis komparatif dua sampel independent bila tipe datanya adalah interval atau rasio, yaitu.

$$t = \frac{\bar{X}_1 - \bar{X}_2}{\sqrt{\frac{S_1^2}{n_1} + \frac{S_2^2}{n_2}}}$$

Rumus 1

$$t = \frac{\bar{X}_1 - \bar{X}_2}{\sqrt{\frac{(n_1 - 1)S_1^2 + (n_2 - 1)S_2^2}{n_1 + n_2 - 2} \left(\frac{1}{n_1} + \frac{1}{n_2} \right)}}$$

Rumus 2

1. Bila jumlah anggota sampel 1 dan 2 sama dan varians homogens, maka dapat digunakan rumus 1 dan 2. Untuk mengetahui t tabel digunakan dk yang besarnya = $n_1 + n_2 - 2$.
2. Bila jumlah anggota sampel 1 dan 2 tidak sama dan varians homogen, maka dapat menggunakan rumus 2. Besarnya dk adalah $n_1 - n_2 - 2$.
3. Bila jumlah anggota sampel 1 dan 2 sama dan varians tidak homogens, maka dapat digunakan rumus 1 dan 2. Untuk mengetahui t tabel digunakan dk yang besarnya = $n_1 - 1$ atau $n_2 - 1$.
4. Bila jumlah anggota sampel 1 dan 2 tidak sama dan varians tidak homogens, maka dapat digunakan rumus 1. Untuk mengetahui t tabel digunakan dk yang besarnya = $n_1 - 1$ dan $n_2 - 1$, dibagi dua dan kemudian ditambah dengan harga t yang terkecil. Sebagai contoh $n_1 = 25$, berarti $dk = 24$, maka harga t tabel = 2,797. $n_2 = 13$, $dk = 12$, harga t tabel = 3,005 (untuk kesalahan 1%, uji dua fihak. Jadi harga t tabel yang digunakan adalah $(3,005 - 2,797) : 2 = 0,104$. Selanjutnya harga ini ditambah dengan t yang terkecil. Jadi $0,104 + 2,797 = 2,901$.

Untuk menguji homogenitas varians adalah dengan menggunakan rumus:

$$F = \frac{\text{Varians Terbesar}}{\text{Varians terkecil}}$$

Bila F hitung lebih kecil atau sama dengan F tabel, maka varians homogens.

Manakala dihitung dengan excel, maka terlihat aplikasinya sebagai berikut:

	A	B	C	D	E	F	G	H	I	J
1	URUTAN ANALISIS T-TEST INDEPENDENT									
2										
3	SMA	STM								
4	6	2	3,09	9,563719008		0,22	0,049382716			= (A4-\$A\$26)
5	3	1	0,09	0,008264463		-0,78	0,604938272			= (C4*C4)
6	5	3	2,09	4,371900826		1,22	1,49382716			= (B4*\$B\$26)
7	2	1	-0,91	0,826446281		-0,78	0,604938272			= (F4*F4)
8	5	3	2,09	4,371900826		1,22	1,49382716			= AVERAGE(A4:A25)
9	1	2	-1,91	3,644628099		0,22	0,049382716			= AVERAGE(B4:B25)
10	2	2	-0,91	0,826446281		0,22	0,049382716			
11	3	1	0,09	0,008264463		-0,78	0,604938272			D26
12	1	3	-1,91	3,644628099		1,22	1,49382716			= SUM(D4:D25)
13	3	1	0,09	0,008264463		-0,78	0,604938272			= (D26/21)
14	2	1	-0,91	0,826446281		-0,78	0,604938272			= SORT(D27)
15	4	1	1,09	1,190082645		-0,78	0,604938272			
16	3	3	0,09	0,008264463		1,22	1,49382716			Aplikasi Rumus
17	4	2	1,09	1,190082645		0,22	0,049382716			= (A26-B26)
18	2	1	-0,91	0,826446281		-0,78	0,604938272			= (D27/22)
19	3	2	0,09	0,008264463		0,22	0,049382716			= (G27/18)
20	1	2	-1,91	3,644628099		0,22	0,049382716			= (O31+D32)
21	5	1	2,09	4,371900826		0,22	0,049382716			= SORT(O33)
22	1	-1,91	3,644628099			-0,78	0,604938272			= (O30/O34)
23	3	0,09	0,008264463							
24	1	-1,91	3,644628099							
25	4	1,09	1,190082645							
26	2,91	1,78	47,81818182				11,11111111			
27			2,277056277				0,653594771			
28			1,508991808				0,808452083			
29										
30										
31										
32										
33										
34										
35										
36										

Cara mencari t tabel	2,080
22-1 = 21	2,110
18-1 = 17	0,030
	0,015
Inilah t-tabelnya	2,095

$\bar{X}_1 - \bar{X}_2$	1,131
$\sqrt{\frac{S_1^2}{n_1} + \frac{S_2^2}{n_2}}$	0,104
	0,036
	0,140
	0,374
	3,026

F hitung > F tabel (2,23)	3,483886104
	2,277056277
	0,653594771
	1,508991808
	0,808452083

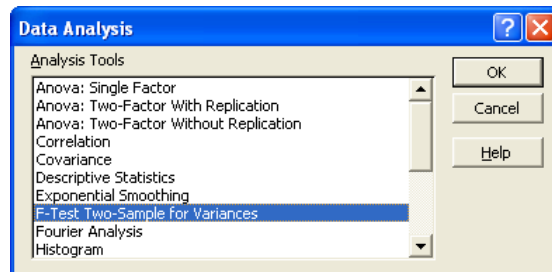
- Carilah mean dari sampel SMA maupun STM.
- Kurangi masing-masing skor dengan mean masing-masing
- Kuadratkan masing-masing skor hasil pengurangan nomor 2.
- Jumlahkan seluruh skor hasil pengkuadratan pada sampel masing-masing
- Hasil nomor 4 dibagi dengan jumlah sampel dikurangi 1 (Inilah variansnya).

6. Carilah homogenitas varians dengan cara varians terbesar dibagi dengan varians terkecil.
7. Bandingkan hasilnya itu dengan F_{tabel} , manakala F_{hitung} lebih kecil atau sama dengan F_{tabel} , maka berarti varians homogen.

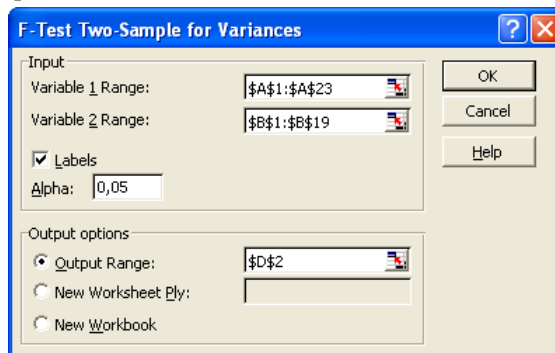
Ternyata hasilnya: 3,483896104 lebih besar dari tabel: 2,218897, sehingga variannys hiterogin. Karena jumlah sampel tidak sama dan variansnya hiterogin, maka berlaku ketentuan nomor 4 (pada halaman 194).

Di samping cara di atas, Untuk menguji homogenitas varians juga dapat diaplikasikan melalui menu **Data Analysis** sebagai berikut:

1. Setelah dipilih **Data Analysis** dari **Tools**, maka pilihlah **F-test Two Samples for Variances** sebagaimana gambar berikut ini.



2. Setelah keluar gambar seperti berikut ini, klik di kolom **Variable 1 Range**, lalu bloklah seluruh data SMA. Selanjutnya klik pada kolom **Variable 2 Range**, lalu bloklah seluruh data STM. Aktifkan **Labels** kalau di atas data yang kita ketik diberi label dan tadi juga termasuk diblok untuk ditampilkan. Setelah itu klik **Ok**.



3. Berikut ini adalah output dari aplikasi di atas.

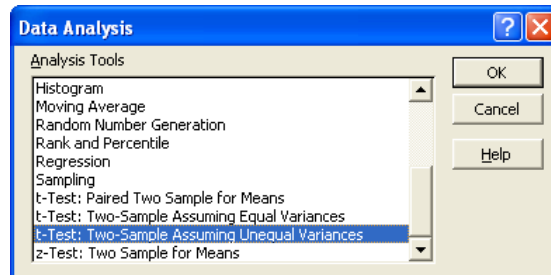
	A	B	C	D	E	F
1	SMA	STM				
2	6	2	F-Test Two-Sample for Variances			
3	3	1				
4	5	3				
5	2	1				
6	5	3				
7	1	2				
8	2	2				
9	3	1				
10	1	3				
11	3	1				

	SMA	STM
Mean	2,909090909	1,777777778
Variance	2,277056277	0,653594771
Observations	22	18
df	21	17
F	3,483896104	
P(F<=f) one-tail	0,005793262	
F Critical one-tail	2,218897066	

Skor F_{hitung} sebesar 3,483896104 dan nilai kritis sebesar 2,218897 ternyata juga sama dengan hasil penjabaran rumus di atas.

Di samping cara menjabarkan rumus di atas, t-test independen yang variansnya heterogin juga dapat diaplikasikan melalui menu **Data Analysis** sebagai berikut:

1. Setelah dipilih **Data Analysis** dari **Tools**, maka pilihlah **t-test: Two Sample Assuming Unequal Variances** sebagaimana gambar berikut ini.



2. Setelah keluar gambar seperti berikut ini, klik di kolom **Variable 1 Range**, lalu bloklah seluruh data SMA. Selanjutnya klik pada kolom **Variable 2 Range**, lalu bloklah seluruh data STM. Apabila peneliti berhipotesis bahwa tidak terjadi perbedaan masa menunggu untuk mendapatkan kerja antara lulusan SMA dan STM, maka pada kolom **Hypothesized Mean Different** diisi dengan **0**. Selanjutnya aktifkan **Labels** kalau di atas data yang kita ketik diberi label dan tadi juga termasuk diblok untuk ditampilkan. Setelah itu klik **Ok**.

t-Test: Two-Sample Assuming Unequal Variances

Input

Variable 1 Range:

Variable 2 Range:

Hypothesized Mean Difference:

☒ Labels

Alpha:

Output options

☒ Output Range:

☐ New Worksheet Ply:

☐ New Workbook

OK Cancel Help

3. Berikut ini adalah output dari aplikasi di atas.

	A	B	C	D	E	F
1	SMA	STM				
2	6	2		t-Test: Two-Sample Assuming Unequal Variances		
3	3	1				
4	5	3				
5	2	1				
6	5	3				
7	1	2				
8	2	2				
9	3	1				
10	1	3				
11	3	1				
12	2	1				
13	4	1				
14	3	3				

	SMA	STM
Mean	2,909090909	1,777777778
Variance	2,277056277	0,653594771
Observations	22	18
Hypothesized Mean Difference	0	
df	33	
t Stat	3,02557876	
P(T<=t) one-tail	0,002390792	
t Critical one-tail	1,692360456	
P(T<=t) two-tail	0,004781585	
t Critical two-tail	2,03451691	

Dari seluruh penghitungan di atas ternyata hasilnya sama. Karena t_{hitung} lebih besar dari t_{tabel} , maka H_0 ditolak dan H_a diterima. Oleh karena itu, H_a yang berbunyi terjadi perbedaan masa menunggu untuk mendapatkan pekerjaan antara lulusan SMA dan STM dapat diterima dan berlaku untuk populasi.

Contoh 2:

Peneliti berkeinginan mengetahui perbedaan masa menunggu untuk mendapatkan pekerjaan antara lulusan Perguruan Tinggi Negeri (PTN) dan Perguruan Tinggi Swasta (PTS). Peneliti mengambil sampel 22 alumni PTN dan 22 alumni PTS.

Proses Pengambilan Keputusan.

Hipotesis:

H_0 Masa menunggu untuk mendapatkan pekerjaan antara alumni PTN dan PTS adalah sama.

Ha Masa menunggu untuk mendapatkan pekerjaan antara alumni PTN dan PTS adalah tidak sama.

Dasar Pengambilan Keputusan:

Dengan membandingkan t_{hitung} dengan t_{tabel} dengan ketentuan:

Ho diterima $t_{hitung} < t_{tabel}$

Ho ditolak $t_{hitung} \geq t_{tabel}$

Dengan menggunakan angka probabilitas, dengan ketentuan:

Ho diterima Probabilitas > taraf nyata (α)

Ho ditolak Probabilitas $\leq$ taraf nyata (α)

Data penelitian tersebut adalah sebagai berikut:

	ptn	alumni		ptn	alumni
1	6.00	PTN	23	3.00	PTS
2	3.00	PTN	24	5.00	PTS
3	5.00	PTN	25	4.00	PTS
4	2.00	PTN	26	3.00	PTS
5	5.00	PTN	27	4.00	PTS
6	1.00	PTN	28	2.00	PTS
7	2.00	PTN	29	5.00	PTS
8	3.00	PTN	30	2.00	PTS
9	1.00	PTN	31	3.00	PTS
10	3.00	PTN	32	2.00	PTS
11	2.00	PTN	33	1.00	PTS
12	4.00	PTN	34	2.00	PTS
13	3.00	PTN	35	2.00	PTS
14	4.00	PTN	36	3.00	PTS
15	2.00	PTN	37	4.00	PTS
16	3.00	PTN	38	3.00	PTS
17	1.00	PTN	39	2.00	PTS
18	5.00	PTN	40	4.00	PTS
19	1.00	PTN	41	2.00	PTS
20	3.00	PTN	42	2.00	PTS
21	1.00	PTN	43	2.00	PTS
22	4.00	PTN	44	3.00	PTS

a. Aplikasi dengan SPSS

Adapun caranya sama dengan yang tertera pada halaman 190-192. Tampilan di bawah ini adalah output hasil penghitungan t-test independent dengan SPSS.

Group Statistics		
	Masa Nunggu Iuluan PT	
	ALUMNI	
	PTN	PTS
N	22	22
Mean	2.9091	2.8636
Std. Deviation	1.50899	1.08213
Std. Error		
Mean	.32172	.23071

Jumlah sampel lulusan PTN dan PTS adalah sama 22. Lulusan PTN rata-rata menunggu untuk mendapatkan pekerjaan selama 2,9 tahun, sedangkan lulusan PTS rata-rata menunggu untuk mendapatkan pekerjaan selama 2,86 tahun.

Independent Samples Test			
		Masa Nunggu lulusan PT	
		Equal variances assumed	Equal variances not assumed
Levene's Test for Equality of Variances	F	2.005	
	Sig.	.164	
t-test for Equality of Means	t	.115	.115
	df	42	38.081
	Sig. (2-tailed)	.909	.909
	Mean Difference	.0455	.0455
	Std. Error Difference	.39589	.39589
95% Confidence Interval of the Difference	Lower	-.75349	-.75593
	Upper	.84439	.84684

Berdasarkan signifikansi dari F-test sebesar 0,164 yang lebih tinggi dibanding dengan alpha 0,05, maka varians data penelitian ini adalah homogen. Oleh karena itu, seluruh analisis didasarkan pada hasil yang ada pada kolom Equal Varians Assumed. T_{hitung} : 0,115 dan Sig (2-tailed): 0,909. Karena skor signifikansi lebih besar dibanding alpha 0,05, maka H_0 diterima dan H_a ditolak. Kesimpulan penelitian ini adalah tidak terjadi perbedaan masa menunggu untuk mendapatkan pekerjaan antara lulusan PTN dan PTS.

b. Aplikasi dengan Microsoft Excel

Cara untuk mengaplikasikan analisis t-test independen adalah sama dengan yang dipaparkan pada halaman 195-196. Yang membedakan dengan analisis tersebut adalah jumlah sampelnya sama dan variansnya ternyata homogen. Oleh karena itu, sebagaimana dijelaskan pada halaman 193, manakala sampel 1 dan 2 jumlahnya sama dan variannya homogen, maka dapat digunakan rumus 1 atau rumus 2.

Selain dijelaskan dengan aplikasi pada halaman berikut ini pada aplikasi Microsoft Excel, homogenitas varians juga akan dideteksi dengan aplikasi **Data Analysis**, dengan hasil berikut ini:

	A	B	C
1	F-Test Two-Sample for Variances		
2			
3		<i>PTN</i>	<i>PTS</i>
4	Mean	2,909090909	2,863636364
5	Variance	2,277056277	1,170995671
6	Observations	22	22
7	df	21	21
8	F	1,944547135	
9	P(F<=f) one-tail	0,067800709	
10	F Critical one-tail	2,084188822	

Karena F_{hitung} 1,944547135 lebih kecil dibanding F_{tabel} : 2,084188822, maka H_0 diterima, artinya variansnya homogen.

Sedangkan untuk aplikasi Microsoft dengan menjabarkan rumus dapat diperhatikan pada contoh berikut.

	A	B	C	D	F	G	H	J	K	L	M	N	O	P	Q	R
1	TAN ANALISIS T-TEST INDEPEND															
2																
3	PTN	PTS														
4	6	3	3,09	9,553719	0,14	0,018595										
5	3	5	0,09	0,008264	2,14	4,56405										
6	5	4	2,09	4,371901	1,14	1,291322										
7	2	3	-0,91	0,826446	0,14	0,018595										
8	5	4	2,09	4,371901	1,14	1,291322										
9	1	2	-1,91	3,644628	-0,86	0,745868										
10	2	5	-0,91	0,826446	2,14	4,56405										
11	3	2	0,09	0,008264	-0,86	0,745868										
12	1	3	-1,91	3,644628	0,14	0,018595										
13	3	2	0,09	0,008264	-0,86	0,745868										
14	2	1	-0,91	0,826446	-1,86	3,47314										
15	4	2	1,09	1,190083	-0,86	0,745868										
16	3	2	0,09	0,008264	-0,86	0,745868										
17	4	3	1,09	1,190083	0,14	0,018595										
18	2	4	-0,91	0,826446	1,14	1,291322										
19	3	3	0,09	0,008264	0,14	0,018595										
20	1	2	-1,91	3,644628	-0,86	0,745868										
21	5	4	2,09	4,371901	1,14	1,291322										
22	1	2	-1,91	3,644628	-0,86	0,745868										
23	3	2	0,09	0,008264	-0,86	0,745868										
24	1	2	-1,91	3,644628	-0,86	0,745868										
25	4	3	1,09	1,190083	0,14	0,018595										
26	64,00	63,00		47,81818		24,59091										
27	2,91	2,86		2,277056		1,170996										
28				1,506992		1,062126										

1,944547

2,09

Maka varians homogens

$$\begin{aligned}
 \text{Rumus 2} &= (21 \cdot D27) + (21 \cdot G27) & \text{Rumus 1} &= (D27/22) + (G27/22) \\
 &= (1/22 + 1/22) & &= 0,395891 \\
 &= (L9 \cdot L10)/42 & &= 0,15673 \\
 &= \text{SQRT}(L11) & &= 0,395891 \\
 &= (A27 - B27)/L12 & &= 0,115 \\
 & & &= 0,115
 \end{aligned}$$

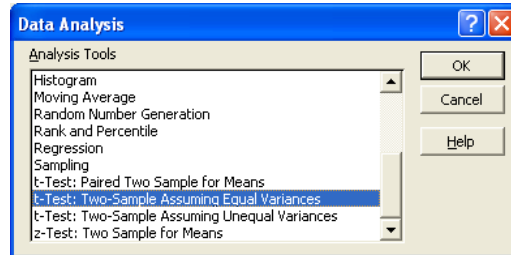
t tabel untuk 22 + 22 - 2 = 42 adalah 2,018082

$$t = \frac{\bar{X}_1 - \bar{X}_2}{\sqrt{\frac{S_1^2}{n_1} + \frac{S_2^2}{n_2}}} = \frac{\bar{X}_1 - \bar{X}_2}{\sqrt{\frac{(n_1 - 1)S_1^2 + (n_2 - 1)S_2^2}{n_1 + n_2 - 2} \left(\frac{1}{n_1} + \frac{1}{n_2} \right)}}$$

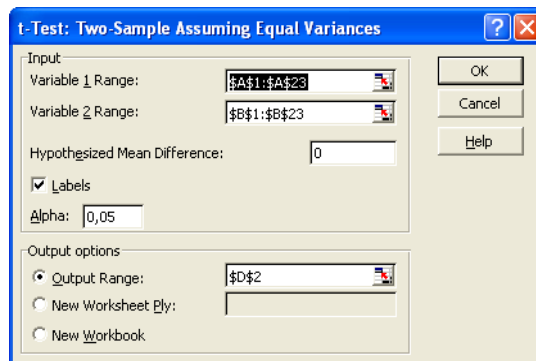
Untuk menguji hipotesis dengan membandingkan antara t_{hitung} dengan t_{tabel} . T_{tabel} untuk derajat kebebasan $22 + 22 - 2 = 42$ untuk kesalahan 5%, dengan uji 2 fihak adalah 2,018082341, ternyata t_{hitung} 0,115 lebih kecil dari pada t_{tabel} 2,018082341, maka H_0 diterima dan H_a ditolak, berarti tidak terdapat perbedaan masa menunggu untuk mendapatkan pekerjaan antara alumni PTN dan alumni PTS.

Di samping cara menjabarkan rumus di atas, t-test independen yang variansnya homogen juga dapat diaplikasikan melalui menu **Data Analysis** sebagai berikut:

1. Setelah dipilih **Data Analysis** dari **Tools**, maka pilihlah **t-test: Two Sample Assuming Equal Variances** sebagaimana gambar berikut ini.



2. Setelah keluar gambar seperti berikut ini, klik di kolom **Variable 1 Range**, lalu bloklah seluruh data PTN. Selanjutnya klik pada kolom **Variable 2 Range**, lalu bloklah seluruh data PTS. Apabila peneliti berhipotesis bahwa tidak terjadi perbedaan masa menunggu untuk mendapatkan kerja antara lulusan PTN dan PTS, maka pada kolom **Hypothesized Mean Different** diisi dengan **0**. Selanjutnya aktifkan **Labels** kalau di atas data yang kita ketik diberi label dan tadi juga termasuk diblok untuk ditampilkan. Setelah itu klik **Ok**.



3. Berikut ini adalah output dari aplikasi di atas. Terlihat bahwa hasilnya sama dengan hasil penghitungan dengan SPSS dan penjabaran rumus dengan Microsoft Excel. Kelebihan dari aplikasi **Data Analysis** adalah adanya tampilan t_{tabel} dengan alfa 5% baik untuk uji satu sisi maupun dua sisi.

	A	B	C	D	E	F
1	PTN	PTS				
2	6	3		t-Test: Two-Sample Assuming Equal Variances		
3	3	5				
4	5	4				
5	2	3				
6	5	4				
7	1	2				
8	2	5				
9	3	2				
10	1	3				
11	3	2				
12	2	1				
13	4	2				
14	3	2				
15	4	3				
16	2	4				

	PTN	PTS
Mean	2,909090909	2,863636364
Variance	2,277056277	1,170995671
Observations	22	22
Pooled Variance	1,724025974	
Hypothesized Mean Difference	0	
df	42	
t Stat	0,114815827	
P(T<=t) one-tail	0,454569131	
t Critical one-tail	1,681951289	
P(T<=t) two-tail	0,909138263	
t Critical two-tail	2,018082341	

Dari seluruh penghitungan di atas ternyata hasilnya sama. Karena t_{hitung} : 0,114815827 itu lebih kecil dibandingkan t_{tabel} : 2,018082341 untuk alpha 5% dan uji dua sisi, maka H_0 diterima dan H_a ditolak. Oleh karena itu, H_0 yang berbunyi tidak terjadi perbedaan masa menunggu untuk mendapatkan pekerjaan antara lulusan PTN dan PTS dapat diterima dan berlaku untuk populasi.

2. Statistik Non-Parametrik

a. Fisher Exact Probability Test

Fisher Exact Probability Test digunakan untuk menguji hipotesis komparatif apabila datanya bertipe nominal yang jumlah sampelnya kecil. Apabila jumlah sampelnya besar digunakan Chi Kuadrat (χ^2). Untuk lebih jelasnya dapat diperhatikan ketentuan berikut ini:

1. Jika $n_1 + n_2 > 40$, dapat dipakai test Chi Kuadrat dengan koreksi kontinuitas dari Yates.
2. Jika $n_1 + n_2$ antara 20 – 40 dan jika tidak satu selpun memiliki frekuensi yang diharapkan ≥ 5 , dapat digunakan Chi Kuadrat dengan koreksi kontinuitas dari Yates. Bila frekuensinya < 5 maka dipakai test Fisher.
3. Jika $n_1 + n_2 < 20$ maka digunakan test Fisher.¹

Contoh:

Peneliti berkeinginan untuk membandingkan frekuensi dosen menulis karya ilmiah antara yang mempunyai golongan kepangkatan III dan IV. Bagi yang menulis 4

¹ Sugiyono, *Statistika untuk Penelitian*, (Bandung: CV Alfabeta, 2003), hlm. 249.

artikel atau lebih dalam waktu dua tahun, maka dikategorikan sering; sedangkan yang kurang dari 4 dikategorikan jarang. Dalam penelitian ini diambil sampel 7 dosen untuk masing-masing golongan.

Proses Pengambilan Keputusan.

Hipotesis:

- Ho Tidak terdapat perbedaan frekuensi menulis artikel karya ilmiah antara dosen yang mempunyai golongan kepangkatan III dan IV.
- Ha Terdapat perbedaan frekuensi menulis artikel karya ilmiah antara dosen yang mempunyai golongan kepangkatan III dan IV.

Dasar Pengambilan Keputusan:

Dengan membandingkan χ^2_{hitung} dengan χ^2_{tabel} dengan ketentuan:

Ho diterima $\chi^2_{hitung} < \chi^2_{tabel}$

Ho ditolak $\chi^2_{hitung} \geq \chi^2_{tabel}$

Dengan menggunakan angka probabilitas, dengan ketentuan:

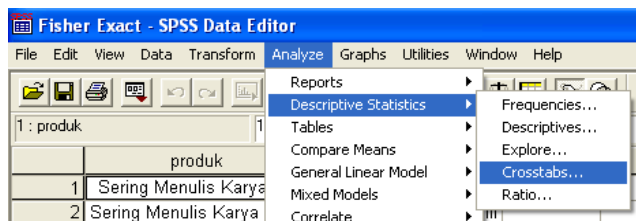
Ho diterima Probabilitas > taraf nyata (α)

Ho ditolak Probabilitas $\leq$ taraf nyata (α)

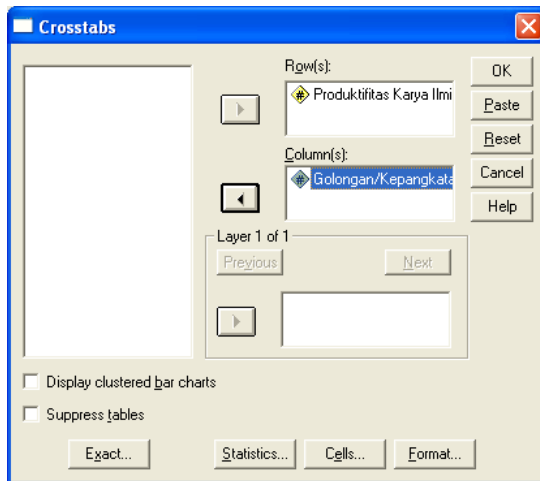
1. Aplikasi dengan SPSS

Langkah-langkah analisis Fisher Exact Probability dengan SPSS adalah sebagai berikut:

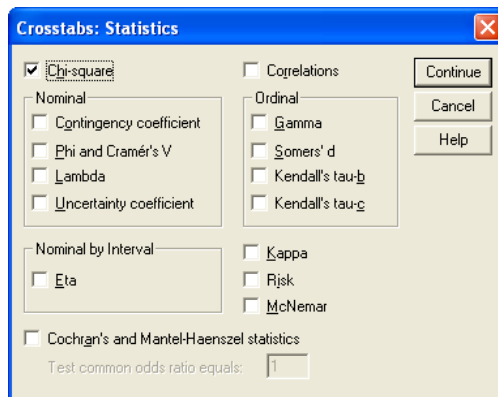
- a) Setelah data diinput, klik **Analyze** ➤ **Descriptive Statistics** ➤ **Crosstabs**, sebagaimana gambar berikut ini.



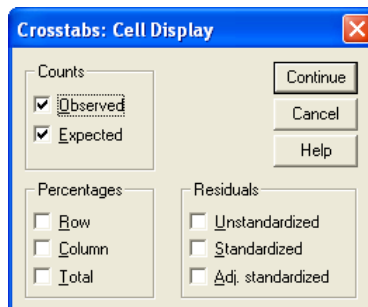
- b) Setelah keluar gambar seperti berikut ini destinasikan salah satu variabel ke kotak **Row(s)** dan variabel satunya pada kotak **Column(s)**. Selanjutnya klik **Statistics**.



- c) Setelah keluar gambar seperti berikut ini klik kotak yang ada di depan **Chi-square**. Selanjutnya klik **Continue**, setelah itu klik **Cells**.



- d) Setelah keluar gambar seperti berikut ini klik kotak yang ada di depan **Expected**. Selanjutnya klik **Continue**, setelah itu klik **Ok**.



e) Berikut ini adalah Output analisis di atas.

Case Processing Summary		
Cases		Produktifitas Karya Ilmiah * Golongan/Kepangkatan
Valid	N	14
	Percent	100,0%
Missing	N	0
	Percent	,0%
Total	N	14
	Percent	100,0%

Output ini memperlihatkan bahwa jumlah sampel adalah 14.

Produktifitas Karya Ilmiah * Golongan/Kepangkatan Crosstabulation					
			Golongan/Kepangkatan		Total
			Dosen Golongan III	Dosen Golongan IV	
Produktifitas Karya Ilmiah	Sering Menulis Karya Ilmiah	Count	5	3	8
		Expected Count	4,0	4,0	8,0
	Jarang Menulis Karya Ilmiah	Count	2	4	6
		Expected Count	3,0	3,0	6,0
Total	Count		7	7	14
	Expected Count		7,0	7,0	14,0

Output ini menyajikan bahwa ada 5 orang dosen dari golongan III yang sering menulis karya ilmiah, 2 lainnya termasuk jarang. Sedangkan dosen dari golongan IV yang sering menulis karya ilmiah sebanyak 3 sedangkan 4 lainnya jarang menulis karya ilmiah

Chi-Square Tests					
	Value	df	Asymp. Sig. (2-sided)	Exact Sig. (2-sided)	Exact Sig. (1-sided)
Pearson Chi-Square	1,167 ^b	1	,280	,592	,296
Continuity Correction ^a	,292	1	,589		
Likelihood Ratio	1,185	1	,276		
Fisher's Exact Test					
Linear-by-Linear Association	1,083	1	,298		
N of Valid Cases	14				

- a. Computed only for a 2x2 table
- b. 4 cells (100,0%) have expected count less than 5. The minimum expected count is 3,00.

Dikarenakan jumlah sampel hanya 14, dan frekuensi yang diharapkan (**expected count**) seluruhnya lebih kecil dari lima, maka hasil analisis yang digunakan adalah dari Fisher's Exact Test. Skor Exact Sig untuk dua sisi sebesar 0,592 yang lebih tinggi dibandingkan alpha: 0,05, maka

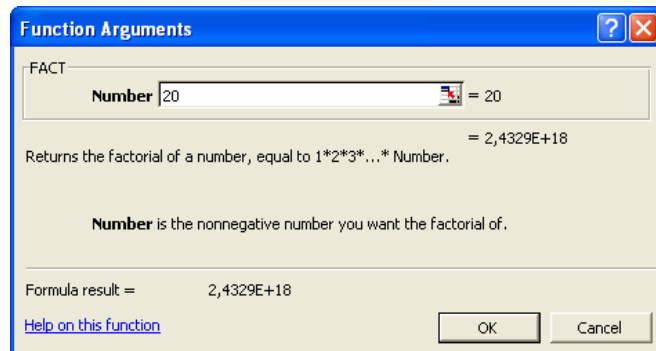
H_0 diterima dan H_a ditolak. Dari hasil ini dapat disimpulkan bahwa tidak terdapat perbedaan frekuensi penulis karya ilmiah antara dosen yang mempunyai golongan III dan IV.

2. Aplikasi dengan Microsoft Excel

Rumus yang digunakan untuk analisis Fisher's Exact Test adalah sebagai berikut:

$$p = \frac{(A+B)! (C+D)! (A+C)! (B+D)!}{N! A! B! C! D!}$$

Nilai Faktorial dapat dilihat pada tabel V (lampiran) atau dapat dicari dari menu function Microsoft Excel; misalnya $20! = 2.432.902.008.176.640.000$. Skor ini merupakan hasil dari aplikasi menu fungsi FACT dan pada kolom number diketik 20, lalu **Ok**.



Atau langsung diketik pada sembarang sel pada Microsoft Excel seperti di bawah ini:

`=FACT(20)`

Untuk mempu dahkan aplikasi rumus Fisher's Exact Test di atas, maka perlu dibuat tabel kontingensi 2 x 2 seperti berikut ini:

Kelompok	I	II	Jumlah
I	A	B	A + B
II	C	D	C + D
Jumlah	A + C	B + D	A+B+C+D

Sedangkan aplikasi analisis Fisher's exact Test dengan Microsoft excel seperti tercantum di bawah ini:

	A	B	C	D	E	F
1		Kepangkatan/Golongan Dosen				
2		Golongan III	Golongan IV	JUMLAH		
3	Sering Menulis	5	3	8		
4	Jarang Menulis	2	4	6		
5	Jumlah	7	7	14		
6						
7		737.418.608.640.000 =(FACT(8))*(FACT(6))*(FACT(7))*(FACT(7))				
8		3.012.881.743.872.000 =(FACT(14))*(FACT(5))*(FACT(3))*(FACT(2))*(FACT(4))				
9		0,245 =B10/B11				
10		0,490 =A9*2				
11						
12		$p = \frac{(A+B)! (C+D)! (A+C)! (B+D)!}{N! A! B! C! D!}$				
13						
14						

Hasil penghitungan dengan Microsoft Excel ini ternyata tidak sama persis dengan hasil hitung dengan SPSS. Hanya saja, karena selisihnya sedikit --skor dari SPSS: 0,592 sedangkan penjabaran rumus skornya: 0,490-- maka kesimpulan akhirnya masih tetap sama, yaitu menerima H_0 dan menolak H_a .

b. χ^2 Two Sample Independent

χ^2 Two Sample Independent digunakan untuk menguji hipotesis komparatif apabila datanya bertipe nominal yang jumlah sampelnya besar.

Contoh:

Peneliti berkeinginan untuk membandingkan frekuensi dosen menulis karya ilmiah antara yang mempunyai golongan kepangkatan III dan IV. Bagi yang menulis 4 artikel atau lebih dalam waktu dua tahun, maka dikategorikan sering; sedangkan yang kurang dari 4 artikel dikategorikan jarang. Dalam penelitian ini diambil sampel 25 dosen untuk masing-masing golongan.

Proses Pengambilan Keputusan.

Hipotesis:

- | | |
|-------|--|
| H_0 | Tidak terdapat perbedaan frekuensi menulis artikel karya ilmiah antara dosen yang mempunyai golongan kepangkatan III dan IV. |
| H_a | Terdapat perbedaan frekuensi menulis artikel karya ilmiah antara dosen yang mempunyai golongan kepangkatan III dan IV. |

Dasar Pengambilan Keputusan:

Dengan membandingkan χ^2_{hitung} dengan χ^2_{tabel} dengan ketentuan:

Ho diterima $\chi^2_{\text{hitung}} < \chi^2_{\text{tabel}}$

Ho ditolak $\chi^2_{\text{hitung}} \geq \chi^2_{\text{tabel}}$

Dengan menggunakan angka probabilitas, dengan ketentuan:

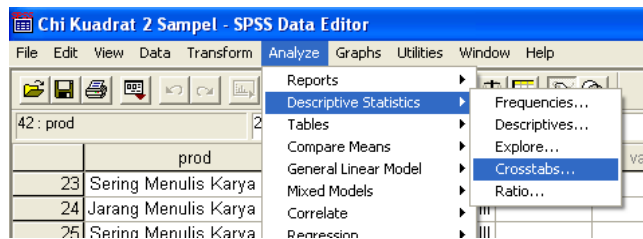
Ho diterima Probabilitas > taraf nyata (α)

Ho ditolak Probabilitas $\leq$ taraf nyata (α)

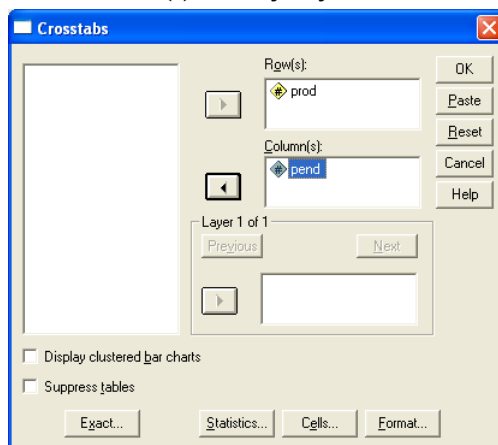
1. Aplikasi dengan SPSS

Langkah-langkah analisis χ^2 Two Sample Independent dengan SPSS adalah sebagai berikut:

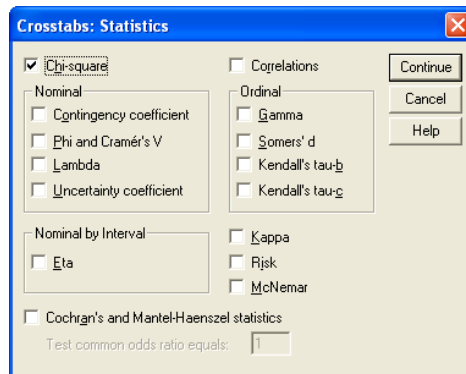
- a) Setelah data diinput, klik **Analyze** ➤ **Descriptive Statistics** ➤ **Crosstabs**, sebagaimana gambar berikut ini.



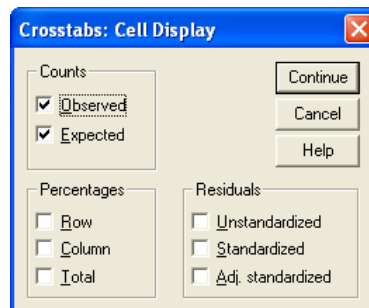
- b) Setelah keluar gambar seperti berikut ini destinasikan salah satu variabel ke kotak **Row(s)** dan variabel satunya pada kotak **Column(s)**. Selanjutnya klik **Statistics**.



- c) Setelah keluar gambar seperti berikut ini klik kotak yang ada di depan **Chi-square**. Selanjutnya klik **Continue**, setelah itu klik **Cells**.



- d) Setelah keluar gambar seperti berikut ini klik kotak yang ada di depan **Expected**. Selanjutnya klik **Continue**, setelah itu klik **Ok**.



- e) Berikut ini adalah Output analisis di atas.

Case Processing Summary		
Cases		PROD * PEND
Valid	N	50
	Percent	100,0%
Missing	N	0
	Percent	,0%
Total	N	50
	Percent	100,0%

Output ini memperlihatkan bahwa jumlah sampel adalah 50.

PROD * PEND Crosstabulation

			PEND		Total
			Dosen Golongan III	Dosen Golongan IV	
PROD	Sering Menulis Karya Ilmiah	Count	19	8	27
		Expected Count	13,5	13,5	27,0
	Jarang Menulis Karya Ilmiah	Count	6	17	23
		Expected Count	11,5	11,5	23,0
Total		Count	25	25	50
		Expected Count	25,0	25,0	50,0

Output ini menyajikan bahwa ada 19 orang dosen dari golongan III yang sering menulis karya ilmiah, 6 lainnya termasuk jarang. Sedangkan dosen dari golongan IV yang sering menulis karya ilmiah sebanyak 8 sedangkan 17 lainnya jarang menulis karya ilmiah

Chi-Square Tests

	Value	df	Asymp. Sig. (2-sided)	Exact Sig. (2-sided)	Exact Sig. (1-sided)
Pearson Chi-Square	9,742 ^b	1	,002		
Continuity Correction ^a	8,052	1	,005		
Likelihood Ratio	10,097	1	,001		
Fisher's Exact Test				,004	,002
Linear-by-Linear Association	9,548	1	,002		
N of Valid Cases	50				

- a. Computed only for a 2x2 table
b. 0 cells (.0%) have expected count less than 5. The minimum expected count is 11,50.

Dikarenakan jumlah sampelnya lebih dari 40, maka digunakan skor Chi Kuadrat dengan koreksi kontinuitas dari Yates (**Continuity Correction**). Skor χ^2_{hitung} sebesar 8,052 yang lebih besar dibanding $\chi^2_{tabel: 0,05;1} = 3,841455$. Demikian juga skor Asymp. Sig untuk dua sisi sebesar 0,005 yang jauh lebih kecil dari alpha, yaitu 0,05, maka H_0 ditolak dan H_a diterima. Dari hasil ini dapat disimpulkan bahwa terdapat perbedaan frekuensi penulis karya ilmiah antara dosen yang mempunyai golongan III dan IV.

2. Aplikasi dengan Microsoft Excel

Rumus yang digunakan untuk analisis Chi Kuadrat dengan koreksi kontinuitas dari Yates adalah sebagai berikut:

$$\chi^2 = \frac{n \left(|ad - bc| - \frac{1}{2}n \right)^2}{(a + b)(a + c)(b + d)(c + d)}$$

Sebagaimana untuk aplikasi Fisher's Exact Test, dalam aplikasi Chi Kuadrat dengan koreksi kontinuitas dari Yates juga perlu dibuatkan tabel kontingensi sebagai berikut:

Sampel	Frekuensi Pada		Jumlah Sampel
	Obyek I	Obyek I	
Sampel I	a	b	a + b
Sampel II	c	d	c + d
Jumlah	a + c	b + d	n

Aplikasi dari rumus di atas dapat dilihat pada gambar berikut ini:

	A	B	C	D	E	F	G	H
1		Dosen						
2		Golongan III	Golongan IV	JUMLAH	%			
3	Sering Menulis	19	8	27	0,54			
4	Harapan	13,5	13,5	27				
5	Jarang Menulis	6	17	23	0,46			
6	Harapan	11,5	11,5	23				
7	Jumlah	25	25	50				
8								
9								
10	2,241	=((B3-B4)^2)/B4						
11	2,241	=((C3-C4)^2)/C4						
12	2,630	=((B5-B6)^2)/B6						
13	2,630	=((C5-C6)^2)/C6						
14	9,742	=SUM(A10:A13)						
15								
16	3125000	=D7*(((B3*C5)-(C3*B5))-(0,5*D7))^2						
17	388125	=(B3+C3)*(B3+B5)*(C3+C5)*(B5+C5)						
18	8,052	=A16/A17						
19								

$$\chi^2 = \sum_{i=1}^k \frac{(f_o - f_h)^2}{f_h}$$

$$\chi^2 = \frac{n \left(|ad - bc| - \frac{1}{2}n \right)^2}{(a+b)(a+c)(b+d)(c+d)}$$

Hasil hitung dengan Microsoft Excel ternyata sama persis dengan SPSS, baik untuk skor χ^2_{hitung} : 9,742 maupun yang dikoreksi kontinuitasnya dari Yates: 8,052. Skor ini lebih besar dibanding χ^2_{tabel} : 0,05;1: 3,841455; oleh karena itu, H_0 ditolak dan H_a diterima.

c. Median Test

Penggunaan median test pada dasarnya sama dengan dua teknik analisis yang digunakan di atas, yaitu Fisher's Exact Probability Test dan χ^2 Two Sample Independent. Bedanya, kalau Fisher Exact Probability Test digunakan untuk menganalisis sampel kecil, χ^2 Two Sample Independent untuk sampel besar, maka Median Test digunakan untuk menganalisis jumlah sampel di antara keduanya. Teknik analisis ini disebut Median Test karena pengujian didasarkan atas median dari sampel yang dianalisis.

Contoh:

Peneliti berkeinginan untuk membandingkan prestasi belajar antara mahasiswa yang jarak tempat tinggalnya dengan kampus kurang dari 5 km dan yang jaraknya 5 km atau lebih. Penelitian ini menggunakan sampel 11 mahasiswa yang jarak tempat tinggalnya dengan kampus kurang dari 5 km dan 13 mahasiswa yang jarak tempat tinggalnya dengan kampus 5 km atau lebih. Data penelitiannya adalah sebagai berikut:

		JARAK	
		< 5 km	>= 5 km
IPK	2,75	0	1
	2,80	0	1
	2,85	1	0
	2,90	1	0
	2,95	1	1
	3,00	0	1
	3,10	0	1
	3,20	0	2
	3,23	2	2
	3,30	1	2
	3,43	1	0
	3,45	0	1
	3,52	1	0
	3,54	1	0
	3,56	1	1
	3,60	1	0
Total		11	13

Proses Pengambilan Keputusan.

Hipotesis:

- Ho

Tidak terdapat perbedaan prestasi belajar antara mahasiswa yang rumahnya berjarak kurang 5 km dengan yang 5 km atau lebih berdasarkan mediannya.
- Ha

Terdapat perbedaan prestasi belajar antara mahasiswa yang rumahnya berjarak kurang 5 km dengan yang 5 km atau lebih berdasarkan mediannya.

Dasar Pengambilan Keputusan:

Dengan membandingkan χ^2_{hitung} dengan χ^2_{tabel} dengan ketentuan:

Ho diterima $\chi^2_{hitung} < \chi^2_{tabel}$

Ho ditolak $\chi^2_{hitung} \geq \chi^2_{tabel}$

Dengan menggunakan angka probabilitas, dengan ketentuan:

Ho diterima Probabilitas > taraf nyata (α)

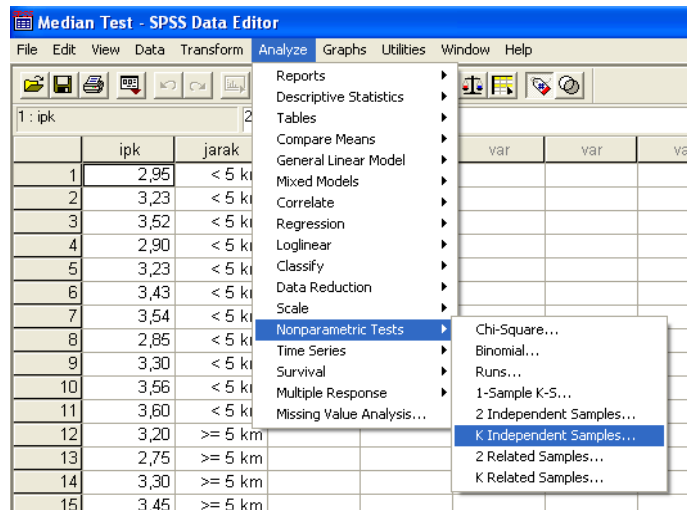
Ho ditolak

Probabilitas $\leq$ taraf nyata (α)

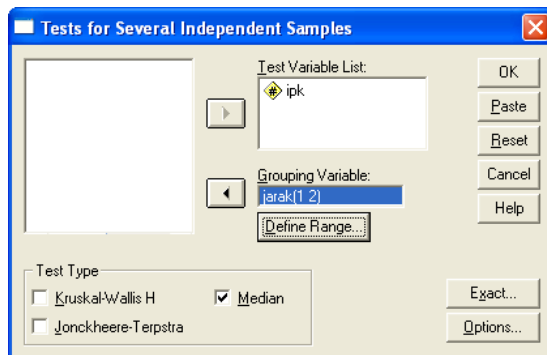
1. Aplikasi dengan SPSS

Langkah-langkah yang dapat ditempuh untuk menganalisis dengan Median test adalah sebagai berikut:

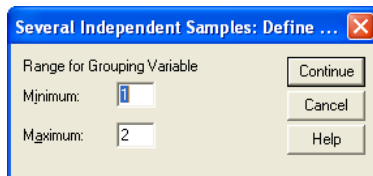
- a) Setelah data diinput, klik **Analyze** ➤ **Nonparametric Tests** ➤ **K Independent Samples**, sebagaimana gambar berikut ini.



- b) Setelah keluar gambar seperti berikut ini, destinasikan data tentang **IPK** ke **Test Variable List**, variable kategori, yaitu **jarak** ke **Grouping Variable**, Non aktifkan **Kruskal-Wallis H**, dan aktifkan **Median**, lalu klik **Define Range..**.



- c) Setelah keluar gambar seperti berikut ini, Isilah angka **1** pada kotak **Minimum**, dan angka **2** pada kotak **Maximum**, lalu klik **Continue**, lalu klik **Ok**.



- d) Berikut ini adalah outputnya.

Frequencies			
		JARAK	
		< 5 km	>= 5 km
IPK	> Median	6	4
	<= Median	5	9

Output ini memperlihatkan bahwa terdapat 6 mahasiswa yang jarak tempat tinggalnya dengan kampus kurang dari 5 km yang mempunyai IPK lebih dari 3,23; 5 mahasiswa lainnya mempunyai IPK 3,23 atau lebih rendah. Sedangkan mahasiswa yang jarak tempat tinggalnya dengan kampus 5 km atau lebih yang mempunyai IPK lebih tinggi dibanding 3,23 sebanyak 4 orang, sedangkan 9 lainnya mempunyai IPK 3,23 atau lebih rendah.

Test Statistics ^a	
	IPK
N	24
Median	3,2300
Exact Sig.	,408

a. Grouping Variable: JARAK

Output di atas memperlihatkan bahwa sampel penelitian ini adalah 24 mahasiswa. Median IPK seluruh mahasiswa yang dijadikan sampel adalah 3,23. Skor Exact Signifikansi adalah 0,408 yang lebih tinggi dibanding $\alpha=5\%$, oleh karena itu H_0 diterima dan H_a ditolak. Dari hasil ini dapat disimpulkan bahwa tidak terdapat perbedaan IPK mahasiswa yang jarak tempat tinggalnya dengan kampus kurang dari 5 km dan mahasiswa yang jarak tempat tinggalnya dengan kampus 5 km atau lebih.

2. Aplikasi dengan Microsoft Excel

Untuk mengaplikasikan test median, maka seluruh data harus diurutkan dan dicari mediannya. Skor data yang lebih tinggi dari median dipisahkan dengan data yang sama dengan median atau lebih rendah. Selanjutnya data tersebut dikelompokkan seperti berikut ini.

Kelompok	I	II	Jumlah
> Median	A	B	A + B
<= Median	C	D	C + D
Jumlah	A + C = n_1	B + D = n_2	N = $n_1 + n_2$

Untuk menguji hipotesis dapat digunakan rumus χ^2 sebagai berikut:

$$\chi^2 = \frac{N \left[(AD - BC) - \frac{N}{2} \right]^2}{(A+B)(C+D)(A+C)(B+D)}$$

Untuk lebih jelasnya dapat dilihat pada aplikasi Microsoft Excel berikut ini.

	A	B	C	D
1		JARAK		
2		< 5 km	>= 5 km	JUMLAH
3	IPK: > Median	6	4	10
4	IPK: <= Median	5	9	14
5	JUMLAH	11	13	24
6				
7		11616 = D5*(((B3*C4)-(C3*B4))-(D5/2))^2		
8		20020 = D3*D4*B5*C5		
9		0,58021978 = A7/A8		
10				
11		$\chi^2 = \frac{N \left[(AD - BC) - \frac{N}{2} \right]^2}{(A+B)(C+D)(A+C)(B+D)}$		
12				
13				
14				

Analisis dengan Microsoft Excel menghasilkan skor χ^2_{hitung} sebesar 0,58021978. Apabila dibandingkan dengan $\chi^2_{tabel;0,05;1}$ yaitu 3,841455338, maka H_0 diterima dan H_a ditolak karena χ^2_{hitung} lebih kecil dibandingkan dengan χ^2_{tabel} . Kesimpulan ini juga sama dengan skor signifikansi hasil aplikasi SPSS.

d. Mann-Whitney U-Test

U-test, begitu biasanya teknik analisis ini disebut, digunakan untuk menguji hipotesis komparatif dua sampel

independen bila datanya berbentuk ordinal. Teknik ini sering juga digunakan untuk menganalisis data penelitian yang direncanakan menggunakan t-test of independent tetapi ternyata sebagian asumsi untuk menggunakan t-test tidak terpenuhi. Untuk kasus terakhir ini, maka data yang semula bertipe interval atau rasio harus diubah menjadi ordinal.

Contoh:

Peneliti ingin membandingkan motivasi berprestasi antara alumni pesantren dan alumni non-pesantren. Dalam penelitian itu diambil 30 sampel dari alumni pesantren dan 40 sampel dari alumni non-pesantren. Dikarenakan data dalam penelitian itu bertipe interval, maka direncanakan untuk menggunakan t-test of independent. Setelah data diuji distribusinya dan ternyata tidak normal, sebagaimana output berikut ini, maka teknik analisis yang digunakan akhirnya adalah Mann-Whitney U-Test.

Tests of Normality			
		Motivasi Berprestasi	
		LULUSAN	
		alumni pesantren	alumni non-pesantren
Kolmogorov-Smirnov	Statistic	,247	,250
	df	30	40
	Sig.	,000	,000
Shapiro-Wilk	Statistic	,814	,820
	df	30	40
	Sig.	,000	,000

a. Lilliefors Significance Correction

Kedua skor signifikansi Kolmogorov Smirnov, baik untuk data alumni pesantren maupun non-pesantren ternyata 0,000 yang berarti di bawah 0,05. Oleh karena itu, distribusi kedua data tersebut tidak normal.

Proses Pengambilan Keputusan.

Hipotesis:

- Ho

Tidak terdapat perbedaan tentang motivasi berprestasi antara alumni pesantren dan non-pesantren
- Ha

Terdapat perbedaan tentang motivasi berprestasi antara alumni pesantren dan non-pesantren

Dasar Pengambilan Keputusan:

Dengan membandingkan z_{hitung} dengan z_{tabel} dengan ketentuan:

Ho diterima $z_{hitung} < z_{tabel}$

Ho ditolak $z_{hitung} \geq z_{tabel}$

Dengan menggunakan angka probabilitas, dengan ketentuan:

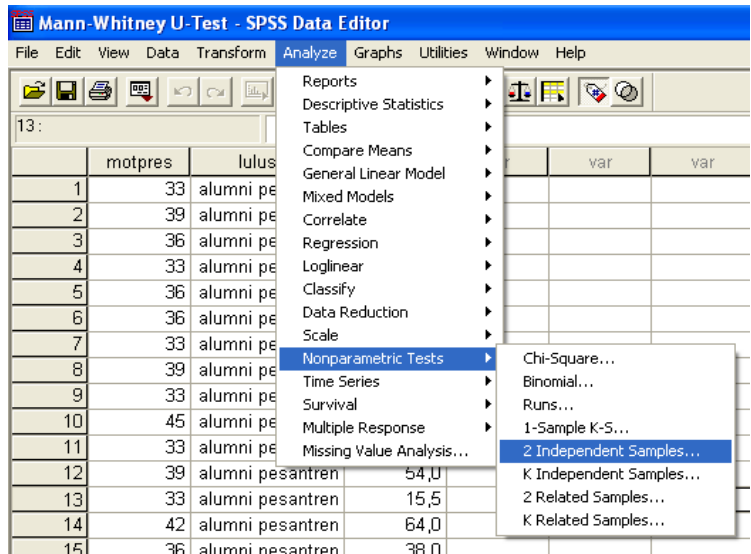
Ho diterima Probabilitas > taraf nyata (α)

Ho ditolak Probabilitas $\leq$ taraf nyata (α)

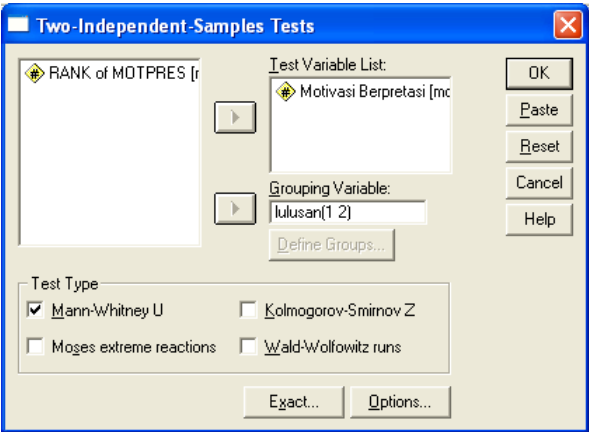
1. Aplikasi dengan SPSS

Langkah-langkah yang dapat ditempuh untuk menganalisis dengan Mann-Whitney U-Test adalah sebagai berikut:

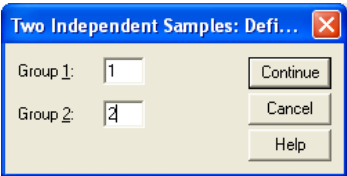
- a) Setelah data diinput, klik **Analyze** ➤ **Nonparametric Tests** ➤ **2 Independent Samples**, sebagaimana gambar berikut ini.



- b) Setelah keluar gambar seperti berikut ini, destinasikan variabel **Motivasi Berprestasi** ke kotak **Test Variable List**, dan **lulusan** ke **Grouping Variable**. Setelah itu klik **Define Group**.



c) Setelah keluar gambar seperti berikut ini, Isilah **1** untuk **Group 1** dan **2** untuk **Group 2**, selanjutnya klik **Continue**, lalu klik **Ok**.



d) Berikut ini adalah Output analisis Mann-Whitney U-Test.

Ranks				
LULUSAN		N	Mean Rank	Sum of Ranks
Motivasi Berpretasi	alumni pesantren	30	33,32	999,50
	alumni non-pesantren	40	37,14	1485,50
	Total	70		

Output ini memperlihatkan bahwa jumlah sampel dari alumni pesantren sebanyak 30, dengan rata-rata ranking: 33,32, dan jumlah ranking: 999,50, sedangkan jumlah sampel alumni non-pesantren sebanyak 40 dengan rata-rata ranking: 37,17 dan jumlah ranking: 1485,50.

Test Statistics^a

	Motivasi Berpretasi
Mann-Whitney U	534,500
Wilcoxon W	999,500
Z	-,821
Asymp. Sig. (2-tailed)	,412

a. Grouping Variable: LULUSAN

Output terakhir ini menghasilkan skor Mann-Whitney: 534,500 (menggunakan rumus 2), skor z_{hitung} : -0,821. Untuk tingkat kepercayaan 95% dan uji dua sisi didapat skor $z_{tabel} \pm 1,96$. Karena z_{hitung} lebih kecil dibandingkan z_{tabel} maka H_0 diterima dan H_a ditolak. Kesimpulan ini juga sama apabila digunakan skor **Asymp. Sig.** Dua sisi: 0,412 yang lebih besar dibanding α : 0,05.

2. Aplikasi dengan Microsoft Excel

Setidaknya ada dua rumus untuk mencari skor Mann-Whitey U-Test. Kedua rumus tersebut harus digunakan semuanya dan yang nantinya digunakan untuk menguji hipotesis penelitian adalah hasil yang lebih kecil. Sedangkan rumus yang dimaksud sebagaimana tercantum berikut ini:

Rumus 1

$$U_1 = n_1 n_2 + \frac{n_1 (n_1 + 1)}{2} - R_1$$

Rumus 2

$$U_2 = n_1 n_2 + \frac{n_2 (n_2 + 1)}{2} - R_2$$

Hasil penghitungan dengan rumus di atas langsung dapat dibandingkan Mann-Whitney tabel manakala jumlah masing-masing n_1 dan n_2 paling banyak 20. Manakala salah satu dari n_1 atau n_2 ada yang lebih banyak dibandingkan 20, maka digunakan dengan pendekatan kurva normal rumus z sebagai berikut:

$$z = \frac{\sum Rn_1 - n_1 \left(\frac{N+1}{2} \right)}{\sqrt{\frac{n_1 \cdot n_2}{N \cdot (N-1)} \cdot \left[\sum Rn_1^2 + \sum Rn_2^2 \right] - \frac{n_1 \cdot n_2 \cdot ((N+1)^2)}{4 \cdot (N-1)}}}$$

Untuk dapat mengaplikasikan seluruh rumus tersebut, langkah-langkahnya adalah sebagai berikut::

- a) Seluruh data, baik n_1 maupun n_2 , secara bersama dicari rankingnya.
- b) Setelah dibuat ranking secara bersama, data dan ranking dari n_1 dipisahkan dari data n_1 .
- c) Masing-masing kelompok ranking, baik dari n_1 maupun n_2 dijumlahkan.
- d) Setelah itu dapat diaplikasikan rumus 1 maupun rumus 2 Mann-Whitney U-Test.
- e) Sedangkan untuk mengaplikasikan rumus z, masing-masing ranking dikuadratkan.
- f) Setelah itu, hasil pengkuadratan tersebut dijumlahkan untuk setiap kelompok sampel.

Untuk lebih dijelasnya dapat dilihat aplikasinya pada Microsoft Excel berikut ini.

I	A	B	C	D	E	F	G	H	I	J	K
1	pesautan										
2	Mot-Pres	Ranking	R2	Mot-Pres	Ranking	R2					
3	33	15,5	240,25	39	54	2916					
4	39	54	2916	33	15,5	240,25					
5	36	38	1444	45	68	4624					
6	33	15,5	240,25	33	15,5	240,25					
7	36	38	1444	39	54	2916					
8	36	38	1444	39	54	2916					
9	33	15,5	240,25	33	15,5	240,25					
10	39	54	2916	45	68	4624					
11	33	15,5	240,25	33	15,5	240,25					
12	45	68	4624	39	54	2916					
13	33	15,5	240,25	33	15,5	240,25					
14	39	54	2916	40	64	4096					
15	33	15,5	240,25	36	38	1444					
16	42	64	4096	33	15,5	240,25					
17	36	38	1444	39	54	2916					
18	33	15,5	240,25	36	38	1444					
19	39	54	2916	36	38	1444					
20	36	38	1444	33	15,5	240,25					
21	36	38	1444	39	54	2916					
22	33	15,5	240,25	33	15,5	240,25					
23	39	54	2916	33	15,5	240,25					
24	33	15,5	240,25	33	15,5	240,25					
25	33	15,5	240,25	39	54	2916					
26	33	15,5	240,25	36	38	1444					
27	39	54	2916	33	15,5	240,25					
28	36	38	1444	36	38	1444					
29	33	15,5	240,25	36	38	1444					
30	36	38	1444	33	15,5	240,25					
31	36	38	1444	39	54	2916					
32	33	15,5	240,25	33	15,5	240,25					
33				45	68	4624					
34				33	15,5	240,25					
35				39	54	2916					
36				39	54	2916					
37				33	15,5	240,25					
38				45	68	4624					
39				33	15,5	240,25					
40				39	54	2916					
41				33	15,5	240,25					
42				42	64	4096					
43		999,5	42335,25		1485,5	71512,25					

$$U_1 = n_1 n_2 + \frac{n_1 (n_1 + 1)}{2} - R_1$$

$$U_2 = n_1 n_2 + \frac{n_2 (n_2 + 1)}{2} - R_2$$

$$665,5 = 30 \cdot 40 + \frac{30 \cdot 31}{2} - 5345$$

$$5345 = 30 \cdot 40 + \frac{40 \cdot 41}{2} - 5345$$

$$-65,5 = 5345 - 30 \cdot 40 + \frac{30 \cdot 31}{2}$$

$$0,2484472 = \frac{30 \cdot 40 \cdot (70 + 69)}{113847,5} = \frac{C43 \cdot H43}{21917,391} = \frac{30 \cdot 40 \cdot (71^2 - 2)}{6367,7019} = \frac{H15 \cdot H16 - H17}{79,797881} = \frac{SQRT(H18)}{-0,8211} = \frac{H14}{H19}$$

$$999,5 - 30 \left(\frac{70 + 1}{2} \right) = \frac{30 \cdot 40}{\sqrt{70 \cdot (70 - 1)}} \cdot [42335,25 + 71512,25] - \frac{30 \cdot 40 \cdot ((70 + 1)^2)}{4 \cdot (70 - 1)}$$

$$z = \frac{\sum R_{n_1} - n_1 \left(\frac{N + 1}{2} \right)}{\sqrt{\frac{n_1 \cdot n_2}{N \cdot (N - 1)} \left[\sum R_{n_1}^2 + \sum R_{n_2}^2 \right] - \frac{n_1 \cdot n_2 \cdot (N + 1)^2}{4 \cdot (N - 1)}}}$$

Dari aplikasi di atas diketahui bahwa hasil penghitungan rumus 1 adalah 665,5, dan rumus 2 sebesar 5345. Karena hasil dari rumus 2 yang lebih kecil, maka yang digunakan adalah hasil yang kecil tersebut. Hanya saja karena jumlah sampel baik untuk n_1 maupun n_2 lebih

banyak dibandingkan 20, maka digunakan hasil rumus z , yaitu $-0,8208238$. Untuk tingkat kepercayaan 95% dan uji dua sisi didapat skor z_{tabel} sebesar $\pm 1,96$. Karena z_{hitung} lebih kecil dibandingkan z_{tabel} maka H_0 diterima dan H_a ditolak. Hasil ini sama persis yang Output SPSS di atas.

e. Kolmogorov Semirnov

Kolmogorov Semirnov termasuk statistik non-parametrik yang digunakan untuk menguji hipotesis komparatif dua sampel independen bila datanya bertipe ordinal yang tersusun pada tabel distribusi frekuensi kumulatif.

Contoh:

Dilakukan penelitian untuk membandingkan kedisiplinan kerja antara alumni Madrasah Aliyah dan SMA. Masing-masing kelompok sampel diambil 19 orang. Dengan menggunakan checklist dengan berisi instrumen tentang kedisiplinan kerja, peneliti melakukan observasi selama 3 bulan.

Proses Pengambilan Keputusan.

Hipotesis:

H_0 Kedisiplinan kerja antara alumni Madrasah Aliyah dan SMA adalah sama

H_a Kedisiplinan kerja antara alumni Madrasah Aliyah dan SMA adalah tidak sama

Dasar Pengambilan Keputusan:

Dengan membandingkan x_{hitung} dengan x_{tabel} dengan ketentuan:

H_0 diterima $x_{\text{hitung}} < x_{\text{tabel}}$

H_0 ditolak $x_{\text{hitung}} \geq x_{\text{tabel}}$

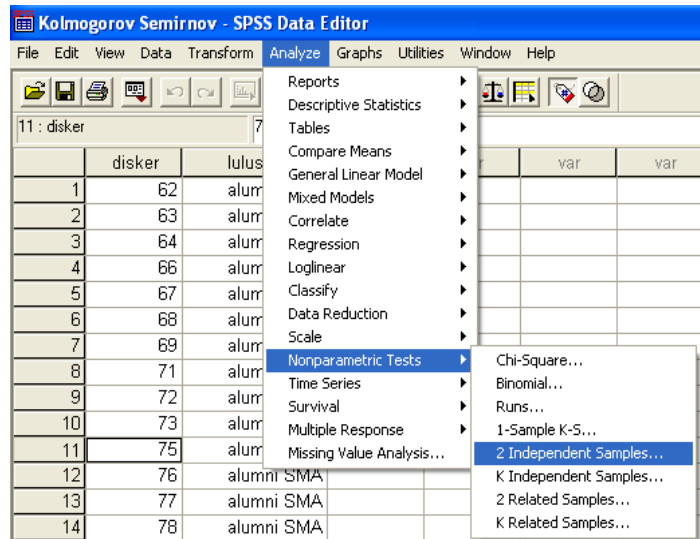
Dengan menggunakan angka probabilitas, dengan ketentuan:

H_0 diterima Probabilitas $>$ taraf nyata (α)

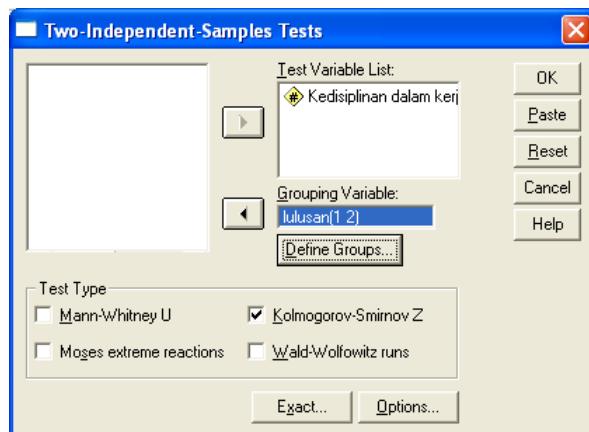
H_0 ditolak Probabilitas $\leq$ taraf nyata (α)

1. Aplikasi dengan SPSS

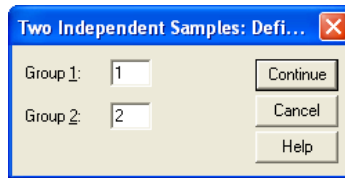
- a) Setelah data diinput, klik **Analyze ► Nonparametric Tests ► 2 Independent Samples**, sebagaimana gambar berikut ini.



- b) Setelah keluar gambar seperti berikut ini, destinasikan variabel **kedisiplinan dalam kerja** ke kotak **Test Variable List**, sedangkan variabel **lulusan** ke kotak **Grouping Variable**, lalu non-aktifkan **Mann-Whitney U** dalam Test Type, lalu aktifkan **Kolmogorov-Smirnov**, lalu klik **Define Groups**.



- c) Setelah keluar gambar seperti berikut ini, ketiklah **1** pada kotak **Group 1** dan **2** pada kotak **Group 2**, lalu klik **Continue**, lalu klik **Ok**.



d) Berikut ini adalah outputnya.

Frequencies

	LULUSAN	N
Kedisiplinan dalam kerja	alumni SMA	19
	alumni MA	19
	Total	38

Output ini memperlihatkan bahwa sampel untuk masing-masing alumni SMA dan MA adalah 19 orang.

Test Statistics^a

		Kedisiplinan dalam kerja
Most Extreme Differences	Absolute	,158
	Positive	,158
	Negative	-,158
Kolmogorov-Smirnov Z		,487
Asymp. Sig. (2-tailed)		,972

a. Grouping Variable: LULUSAN

Output ini menunjukkan skor hasil hitung Kolmogorov Smirnov yang dalam keadaan positif adalah 0,158 dan dalam keadaan negatif juga sama yaitu -0,158. Skor tersebut kemudian dibandingkan dengan tabel untuk n_1 19 dan n_2 19 diketahui skor tabelnya adalah 0,44124. Karena kolmogorov smirnov hitung lebih kecil dibanding kolmogorov smirnov tabel, maka H_0 diterima dan H_a ditolak. Kesimpulan yang sama akan diperoleh manakala digunakan skor Asymp. Sig. (2 sisi) sebesar 0,972 yang jauh lebih tinggi dibandingkan taraf nyata (α), yaitu 0,05.

2. Aplikasi dengan Microsoft Excel

Langkah-langkah yang perlu ditempuh ketika menganalisis data dengan kolmogorov smirnov adalah:

- Masing-masing data dibuat ranking dengan nomor urut dari yang terkecil sampai kepada yang terbesar.
- Dicocokkan skor yang sama antara n_1 dan n_2 dengan menempatkan skor yang sama pada baris yang sama.

- c) Membuat nomor urut (kumulatif) yang dijadikan pembilang ketika diaplikasikan pengurangan $Sn_1(x)$ dengan $Sn_2(x)$. Sedangkan penyebutnya adalah jumlah masing-masing data n_1 dan n_2 .
- d) Dari hasil pengurangan tersebut dicari skor terbesar baik dalam keadaan positif maupun negatifnya. Skor terbesar dari yang positif itulah skor Kolmogorov Smirnov.
- e) Sedangkan untuk mencari tabel digunakan rumus di bawah ini.

$$D = 1,36 \sqrt{\frac{x_1 + x_2}{x_1 \cdot x_2}}$$

Untuk lebih jelasnya dapat diamati dari aplikasi Microsoft Excel berikut ini.

	A	B	C	D	E	F	G	H
1	x_1	x_2				Formula di C2		
2	62		0,053	1	0	=D2/19-E2/19		
3	63		0,105	2	0			
4	64		0,158	3	0	Formula di C29		
5	66	66	0,158	4	1	=MAX(C2:C26)		
6	67	67	0,158	5	2	Formula di C30		
7	68	68	0,158	6	3	=MIN(C2:C26)		
8	69	69	0,158	7	4			
9		70	0,105	7	5	Formula di E30		
10	71	71	0,105	8	6	=(19+19)/(19*19)		
11	72	72	0,105	9	7	Formula di E31		
12	73	73	0,105	10	8	=1,36*(SQRT(E33))		
13		73	0,053	10	9			
14		74	0,000	10	10			
15		74	-0,053	10	11			
16		75	-0,105	10	12			
17	75	75	-0,105	11	13			
18	76	76	-0,105	12	14			
19		76	-0,158	12	15			
20	77	77	-0,158	13	16			
21	78		-0,105	14	16			
22	80	80	-0,105	15	17			
23	81	81	-0,105	16	18			
24	82		-0,053	17	18			
25	83		0,000	18	18			
26	84	84	0,000	19	19			
27								
28								
29			0,158					
30			-0,158		0,105263	$D = 1,36 \sqrt{\frac{x_1 + x_2}{x_1 \cdot x_2}}$		
31					0,44124			
32								

Skor hasil hitung Kolmogorov Smirnov yang dalam keadaan positif adalah 0,158 dan dalam keadaan negatif

juga sama yaitu -0,158. Skor ini sama dengan hasil hitung dengan SPSS. Skor tersebut kemudian dibandingkan dengan tabel untuk n_1 19 dan n_2 19 diketahui skor tabelnya adalah 0,44124. Karena kolmogorov smirnov hitung lebih kecil dibanding kolmogorov smirnov tabel, maka H_0 diterima dan H_a ditolak.

f. Wald Woldfowitz Runs

Teknik analisis ini digunakan untuk menganalisis hipotesis komparatif dua sampel independen bila datanya bertipe ordinal. Untuk aplikasi dengan Microsoft Excel, kedua datanya disatukan dan disusun dalam bentuk run.

Contoh:

Dalam rangka mengetahui performance sebuah pendidikan tinggi, diadakan penelitian tentang keterlambatan kehadiran pimpinan dan staf lembaga tersebut dalam seminggu.

Proses Pengambilan Keputusan.

Hipotesis:

H_0 Lamanya keterlambatan datang antara pimpinan dan staf adalah sama

H_a Lamanya keterlambatan datang antara pimpinan dan staf adalah sama

Dasar Pengambilan Keputusan:

Dengan membandingkan run_{hitung} dengan run_{tabel} dengan ketentuan:

H_0 diterima $run_{hitung} > run_{tabel}$

H_0 ditolak $run_{hitung} \leq run_{tabel}$

Apabila salah satu atau kedua-duanya n_1 atau $n_2 > 20$, maka pengambailan keputusan dengan membandingkan z_{tabel} dengan taraf nyata (α) dengan ketentuan:

H_0 diterima $z_{tabel} > \text{taraf nyata } (\alpha)$

H_0 ditolak $z_{tabel} \leq \text{taraf nyata } (\alpha)$

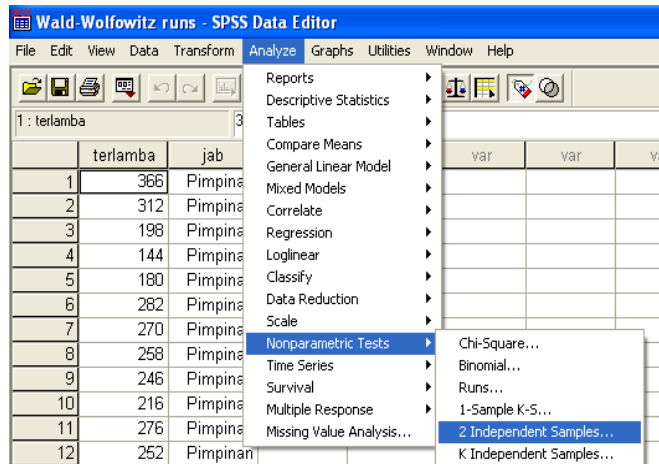
Dengan menggunakan angka probabilitas, dengan ketentuan:

H_0 diterima Probabilitas $>$ taraf nyata (α)

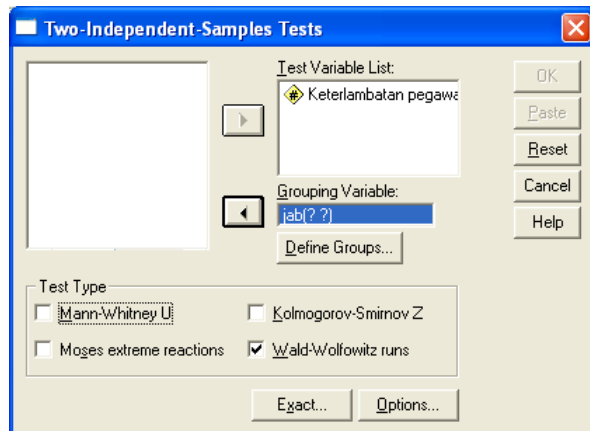
H_0 ditolak Probabilitas $\leq$ taraf nyata (α)

1. Aplikasi dengan SPSS

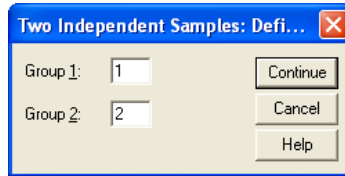
- a) Setelah data diinput, klik **Analyze ► Nonparametric Tests ► 2 Independent Samples**, sebagaimana gambar berikut ini.



- b) Setelah keluar gambar seperti berikut ini destinasikan variabel **Keterlambatan pegawai** pada kolom **Test Variable List**, dan variabel **Jab** pada **Grouping Variable**. Selanjutnya non-aktifkan **Mann-Whitney U** dan aktifkan **Wald-Wolfowitz runs**, lalu klik **Define Groups**



- c) Setelah keluar gambar seperti berikut ini, ketiklah **1** pada kotak **Group 1** dan **2** pada kotak **Group 2**, lalu klik **Continue**, lalu klik **Ok**.



d) Berikut ini adalah outputnya.

Frequencies

	Jabatan	N
Keterlambatan pegawai dalam menit	Pimpinan	15
	Staf	15
	Total	30

Output ini memperlihatkan bahwa sampel untuk masing-masing pimpinan dan staf adalah 15 orang.

Test Statistics^{b,c}

		Number of Runs	Z	Exact Sig. (1-tailed)
Keterlambatan pegawai dalam menit	Minimum Possible	12 ^a	-1,301	,097
	Maximum Possible	15 ^a	-,186	,424

a. There are 2 inter-group ties involving 4 cases.

b. Wald-Wolfowitz Test

c. Grouping Variable: Jabatan

Output ini memperlihatkan bahwa jumlah run dari kedua kelompok data tersebut adalah minimal 12 manaka skor yang sama antara data keterlambatan pimpinan dan staf tidak dihitung, akan tetapi kalau data yang sama itu dihitung, maka maksimal jumlah run adalah 15. Apabila jumlah run ini dibandingkan dengan run tabel (Tabel VII a) dengan n_1 dan n_2 masing-masing 15 adalah 10. Karena run_{hitung} lebih tinggi dibanding run_{tabel} , maka H_0 diterima dan H_a ditolak. Kesimpulan sama akan didapatkan ketika digunakan skor Exact Sig sebesar 0,424 yang jauh lebih tinggi dibanding taraf nyata (α): 0,05. Oleh karena itu, H_0 yang berbunyi Lamanya keterlambatan datang antara pimpinan dan staf adalah sama diterima.

2. Aplikasi dengan Microsoft Excel

Langkah-langkah yang perlu ditempuh ketika menganalisis data dengan Wald-Wolfowitz runs adalah:

- Data dari kedua sampel itu dijadikan satu dan diurutkan.
- Dihitung run-nya dengan cara setiap pergantian antara data dari satu sampel ke sampel lainnya dihitung satu.

- c) Jumlah run tersebut dibandingkan dengan run tabel.
 d) Apabila jumlah n_1 atau n_2 ada yang lebih banyak dari 20, maka uji statistiknya menggunakan rumus berikut ini.

$$z = \frac{r - \mu_r}{\sigma_r} = \frac{r - \left(\frac{2n_1n_2}{n_1 + n_2} + 1 \right) - 0,5}{\sqrt{\frac{2n_1n_2(2n_1n_2 - n_1 - n_2)}{(n_1 + n_2)^2 (n_1 + n_2 - 1)}}}$$

Untuk lebih jelasnya dapat diperhatikan aplikasi dengan Microsoft Excel berikut ini.

	A	B	C	D	E	F	G	H	I
1	48	Pimpinan	1						
2	48	Staf	2						
3	54	Pimpinan	3						
4	66	Staf	4						
5	96	Staf							
6	108	Staf							
7	108	Staf							
8	138	Staf							
9	138	Staf							
10	138	Staf							
11	144	Pimpinan	5						
12	156	Staf	6						
13	162	Staf							
14	180	Pimpinan	7						
15	192	Staf	8						
16	198	Pimpinan	9						
17	204	Staf	10						
18	210	Staf							
19	216	Pimpinan	11						
20	216	Staf	12						
21	246	Pimpinan	13						
22	252	Pimpinan							
23	258	Pimpinan							
24	270	Pimpinan							
25	276	Pimpinan							
26	282	Pimpinan							
27	288	Staf	14						
28	312	Pimpinan	15						
29	330	Pimpinan							
30	366	Pimpinan							

$z = \frac{r - \mu_r}{\sigma_r} = \frac{r - \left(\frac{2n_1n_2}{n_1 + n_2} + 1 \right) - 0,5}{\sqrt{\frac{2n_1n_2(2n_1n_2 - n_1 - n_2)}{(n_1 + n_2)^2 (n_1 + n_2 - 1)}}}$
$15 - \left(\frac{2.15.15}{15 + 15} + 1 \right) - 0,5$
$\sqrt{\frac{2.15.15 (2.15.15 - 15 - 15)}{(15 + 15)^2 (15 + 15 - 1)}}$

-1,5	=15-(((2*15*15)/(15+15))+1)-0,5
189000	=2*15*15*(2*15*15-15-15)
26100	=((15+15)^2)*(15+15-1)
7,241379	=E14/E15
2,690981	=SQRT(E16)
-0,55742	=E13/E17
0,2886	=NORMSDIST(-0,55742)

Hasil aplikasi Microsoft Excel ini menghasilkan jumlah run sebanyak 15. Hal ini sama dengan jumlah maksimal run hasil aplikasi SPSS.

E. Komparasi k Sampel Berkorelasi

1. Statistik Parametrik: One-way Anova

One-way Anova digunakan untuk menguji hipotesis komparatif 3 sampel atau lebih bila datanya berbentuk interval atau rasio, berdistribusi normal, dan variannya homogen. Untuk sub bab ini menguji sampel berpasangan dalam arti sebelum dan sesudah perlakuan/treatment. Oleh karena itu, maksud dari 3 sampel atau lebih, artinya setiap sampel mendapatkan perlakuan tiga kali atau lebih.

Contoh:

Dilakukan penelitian tentang perbedaan prestasi belajar siswa sebelum kursus, setelah kursus berjalan 3 bulan, dan setelah mendapatkan kursus 6 bulan.

Proses Pengambilan Keputusan.

Hipotesis:

- | | |
|----|---|
| Ho | Tidak terdapat perbedaan prestasi belajar dengan adanya kursus (kursus tidak berpengaruh terhadap prestasi belajar) |
| Ha | Terdapat perbedaan prestasi belajar dengan adanya kursus (kursus berpengaruh terhadap prestasi belajar) |

Dasar Pengambilan Keputusan:

Dengan membandingkan F_{hitung} dengan F_{tabel} dengan ketentuan:

Ho diterima $F_{hitung} < F_{tabel}$

Ho ditolak $F_{hitung} \geq F_{tabel}$

Dengan menggunakan angka probabilitas, dengan ketentuan:

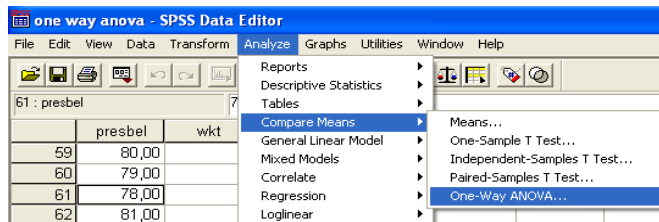
Ho diterima Probabilitas > taraf nyata (α)

Ho ditolak Probabilitas $\leq$ taraf nyata (α)

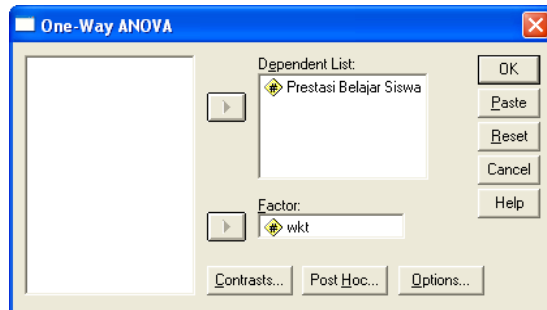
a. Aplikasi dengan SPSS

Langkah-langkah dalam menggunakan analisis ini adalah sebagai berikut:

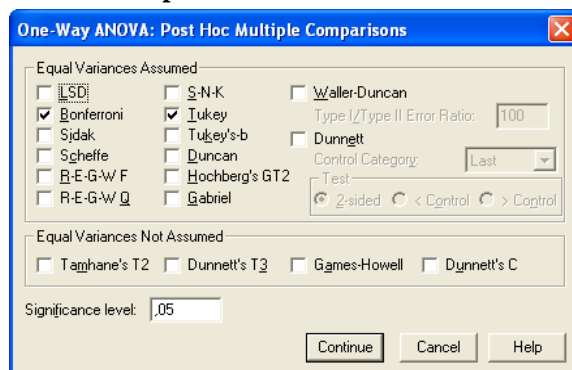
1. Setelah data diinput, lalu klik **Analyze** ➤ **Compare Means** ➤ **One-way ANOVA** sebagaimana gambar di bawah ini:



2. Setelah keluar gambar seperti berikut ini, destinasikan **Prestasi Belajar Siswa** ke kotak **Dependent List** dan **wkt** ke kotak **Factor**. Selanjutnya klik **Post Hoc...**

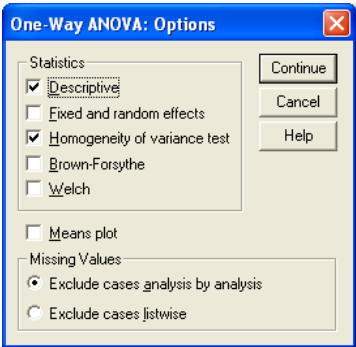


3. Setelah keluar gambar seperti berikut ini, aktifkan **Bonferroni** dan **Tukey** untuk mengetahui perbedaan antara sampel satu dengan lainnya. Selanjutnya klik **Continue** dan **Options**.



3. Setelah keluar gambar seperti berikut ini, aktifkan **Descriptive** untuk mendiskripsikan data ketiga sampel tersebut dan **Homogeneity of variance test** untuk

mengetahui varians ketiga sampel tersebut homogen atau tidak.



3. Berikut ini adalah Output one-way ANOVA..

Descriptives

Prestasi Belajar Siswa				
	Prestasi Belajar sebelum Kursus	Prestasi Belajar setelah 3 bulan Kursus	Prestasi Belajar setelah 6 bulan Kursus	Total
N	30	30	30	90
Mean	74,7667	76,4333	78,6667	76,6222
Std. Deviation	5,98091	5,85858	5,68321	5,99546
Std. Error	1,09196	1,06963	1,03761	,63198
95% Confidence Interval for Mean	Lower Bound	72,5334	74,2457	76,5445
	Upper Bound	77,0000	78,6210	80,7888
Minimum	65,00	67,00	70,00	65,00
Maximum	87,00	89,00	92,00	92,00

Output pertama ini mendiskripsikan data, misalnya rata-rata prestasi belajar siswa sebelum mengikuti kursus adalah 74,7667, standard deviasinya 5,98091, dan standard error of mean-nya 1,09196.

Test of Homogeneity of Variances

Prestasi Belajar Siswa			
Levene Statistic	df1	df2	Sig.
,128	2	87	,880

Output kedua ini terkait uji asumsi untuk ANOVA yaitu variannya harus homogen. Berdasarkan skor signifikansi sebesar 0,880 yang lebih tinggi dibanding taraf nyata (α): 0,05, maka varians ketiga sampel adalah homogen.

ANOVA

Prestasi Belajar Siswa					
	Sum of Squares	df	Mean Square	F	Sig.
Between Groups	229,756	2	114,878	3,366	,039
Within Groups	2969,400	87	34,131		
Total	3199,156	89			

Output ketiga ini terkait uji hipotesis. Skor F_{hitung} nya sebesar 3,366. Bila dibandingkan dengan F_{tabel} untuk pembilang 2 (m-1) dan penyebut 87 (N-m) maka diketahui sebesar:

3,101292

=FINV(0,05;2;87)

karena F_{hitung} lebih besar dibandingkan F_{tabel} , maka H_a diterima dan H_o ditolak. Oleh karena itu, kesimpulan dari penelitian ini adalah terdapat perbedaan prestasi belajar dengan adanya kursus (kursus berpengaruh terhadap prestasi belajar). Kesimpulan yang sama akan didapatkan ketika digunakan skor sig sebesar 0,039 yang lebih kecil dibanding taraf nyata (α): 0,05.

Multiple Comparisons

Dependent Variable: Prestasi Belajar Siswa							
		Mean Difference (I-J)	Std. Error	Sig.	95% Confidence Interval		
(I) WKT	(J) WKT				Lower Bound	Upper Bound	
Tukey HSD	Prestasi Belajar sebelum Kursus	Prestasi Belajar setelah 3 bulan Kursus	-1,6667	1,50844	,514	-5,2635	1,9302
		Prestasi Belajar setelah 6 bulan Kursus	-3,9000*	1,50844	,030	-7,4969	-,3031
	Prestasi Belajar setelah 3 bulan Kursus	Prestasi Belajar sebelum Kursus	1,6667	1,50844	,514	-1,9302	5,2635
		Prestasi Belajar setelah 6 bulan Kursus	-2,2333	1,50844	,305	-5,8302	1,3635
	Prestasi Belajar setelah 6 bulan Kursus	Prestasi Belajar sebelum Kursus	3,9000*	1,50844	,030	,3031	7,4969
		Prestasi Belajar setelah 3 bulan Kursus	2,2333	1,50844	,305	-1,3635	5,8302
Bonferroni	Prestasi Belajar sebelum Kursus	Prestasi Belajar setelah 3 bulan Kursus	-1,6667	1,50844	,817	-5,3490	2,0157
		Prestasi Belajar setelah 6 bulan Kursus	-3,9000*	1,50844	,034	-7,5823	-,2177
	Prestasi Belajar setelah 3 bulan Kursus	Prestasi Belajar sebelum Kursus	1,6667	1,50844	,817	-2,0157	5,3490
		Prestasi Belajar setelah 6 bulan Kursus	-2,2333	1,50844	,427	-5,9157	1,4490
	Prestasi Belajar setelah 6 bulan Kursus	Prestasi Belajar sebelum Kursus	3,9000*	1,50844	,034	,2177	7,5823
		Prestasi Belajar setelah 3 bulan Kursus	2,2333	1,50844	,427	-1,4490	5,9157

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

Output keempat ini mendeteksi perbedaan antar sampel. Untuk menentukan antar sampel mana yang berbeda secara signifikan digunakan skor signifikansi dibandingkan dengan taraf nyata (α). Apabila skor signifikansi $\leq$ taraf nyata (α), maka antar sampel tersebut berbeda secara signifikan. Output di atas memperlihatkan bahwa yang berbeda hanyalah prestasi belajar sebelum kursus dengan prestasi belajar setelah 6 bulan kursus sebesar 3,9000, karena skor

signifikansinya sebesar 0,034 yang $<$ taraf nyata (α): 0,05. Sedangkan antara prestasi belajar sebelum kursus dengan setelah kursus 3 bulan dan prestasi belajar setelah kurus 3 bulan dan setelah kurus 6 bulan ternyata tidak berbeda secara signifikan karena skor significansinya yang lebih tinggi dibanding taraf nyata (α): 0,05.

Prestasi Belajar Siswa				
WKT		N	Subset for alpha = .05	
			1	2
Tukey HSD	Prestasi Belajar sebelum Kursus	30	74,7667	
	Prestasi Belajar setelah 3 bulan Kursus	30	76,4333	76,4333
	Prestasi Belajar setelah 6 bulan Kursus	30		78,6667
	Sig.		,514	,305

Means for groups in homogeneous subsets are displayed.

a. Uses Harmonic Mean Sample Size = 30,000.

Output kelima ini menunjukkan bahwa ketiga kelompok data tersebut dapat dikelompokkan menjadi 2, yaitu kelompok pertama berisi data sampel sebelum kursus dan setelah kursus 3 bulan; dan kelompok kedua data sampel setelah kurus 3 bulan dan data setelah kursus 6 bulan.

b. Aplikasi dengan Microsoft Excel

Sedangkan prosedur analisis one-way ANOVA dengan Microsoft Excel adalah sebagai berikut:

1. Data diinput ke dalam kelompok sampel masing-masing.
2. Mencari varians dari masing-masing sampel.
3. Menguji homogenitas varians dengan menggunakan uji F, yaitu dengan cara varians terbesar dibagi dengan varians terkecil. Apabila $F_{hitung} < F_{tabel}$, maka variansnya homegen, akan tetapi kalau $F_{hitung} \geq F_{tabel}$, maka variansnya hiterogen. Apabila variansnya homogen, maka analisis one-way ANOVA dapat diteruskan.
4. Masing-masing skor pada masing-masing sampel dikuadratkan.
5. Masing-masing skor pada masing-masing sampel dijumlahkan.
6. Masing-masing skor hasil pengkuadratan dari masing-masing skor sampel dijumlahkan.
7. Hasil penjuamlahan nomor 5 dijumlahkan.

8. Hasil penjumlahan nomor 6 dijumlahkan.
9. Selanjutnya dimasukkan ke dalam rumus berikut ini.
 - a. Menghitung jumlah kuadrat total (JK_{Tot}) dengan rumus:

$$JK_{tot} = \sum \sum X_{tot}^2 - \frac{(\sum x_{tot})^2}{N}$$
 - b. Menghitung jumlah kuadrat antar kelompok (JK_{antar}) dengan rumus:

$$JK_{ant} = \sum \frac{(\sum_k)^2}{n_k} - \frac{(\sum x_{tot})^2}{N}$$
 - c. Menghitung jumlah kuadrat dalam kelompok (JK_{dalam}) dengan rumus:

$$JK_{dalam} = JK_{tot} - JK_{antar}$$
 - d. Menghitung mean kuadrat antar kelompok (Mk_{antar}) dengan rumus:

$$MK_{antar} = \frac{JK_{antar}}{m - 1}$$
 - e. Menghitung mean kuadrat dalam kelompok (MK_{dalam}) dengan rumus:

$$MK_{dalam} = \frac{JK_{dalam}}{N - m}$$
 - f. Menghitung F_{hitung} dengan rumus:

$$F_{hitung} = \frac{MK_{antar}}{MK_{dalam}}$$

Untuk lebih jelasnya tentang prosedur analisis one-way ANOVA dapat dilihat pada aplikasi Microsoft Excel berikut ini.

	A		B		C		D		E		F		G		H		I		J		K		L		M		N		O		P	
1	x ₁	x ₂	x ₃	x ₄	x ₅	x ₆	x ₇	x ₈	x ₉	x ₁₀	x ₁₁	x ₁₂	x ₁₃	x ₁₄	x ₁₅	x ₁₆	x ₁₇	x ₁₈	x ₁₉	x ₂₀	x ₂₁	x ₂₂	x ₂₃	x ₂₄	x ₂₅	x ₂₆	x ₂₇	x ₂₈	x ₂₉	x ₃₀		
2	75	76	78	78	0,233	0,054	-0,433	0,188	-0,867	0,444																						
3	76	77	81	81	1,233	1,521	0,567	0,321	-2,333	5,444																						
4	67	69	70	72	-7,767	60,321	-7,433	55,254	-8,667	75,111																						
5	65	67	70	70	-9,767	95,368	-9,433	88,968	-8,667	75,111																						
6	67	69	71	72	-7,767	60,321	-7,433	55,254	-7,667	58,778																						
7	78	80	83	83	3,233	10,454	3,567	12,721	4,333	18,778																						
8	85	87	90	90	10,233	104,721	10,567	111,854	11,333	128,444																						
9	67	69	73	73	-7,767	60,321	-7,433	55,254	-5,667	32,111																						
10	87	89	92	92	12,233	149,654	12,567	157,921	13,333	177,778																						
11	78	80	80	80	3,233	10,454	3,567	12,721	1,333	1,778																						
12	77	79	79	79	2,233	4,968	2,567	6,568	0,333	0,111																						
13	75	76	80	80	0,233	0,054	-0,433	0,188	1,333	1,778																						
14	76	78	81	81	1,233	1,521	1,567	2,454	2,333	5,444																						
15	76	78	81	81	1,233	1,521	1,567	2,454	2,333	5,444																						
16	78	79	80	80	3,233	10,454	2,567	6,568	1,333	1,778																						
17	76	76	81	81	1,233	1,521	1,567	2,454	2,333	5,444																						
18	79	81	82	82	4,233	17,921	4,567	20,854	3,333	11,111																						
19	81	83	86	86	6,233	36,854	6,567	43,121	7,333	53,778																						
20	82	82	83	83	7,233	52,321	5,567	30,968	4,333	18,778																						
21	65	67	70	70	-9,767	95,368	-9,433	88,968	-8,667	75,111																						
22	66	68	71	71	-8,767	76,854	-8,433	71,121	-7,667	58,778																						
23	67	69	72	72	-7,767	60,321	-7,433	55,254	-6,667	44,444																						
24	76	76	79	79	1,233	1,521	-0,433	0,188	0,333	0,111																						
25	70	72	75	75	-4,767	22,721	-4,433	19,654	-3,667	13,444																						
26	67	69	72	72	-7,767	60,321	-7,433	55,254	-6,667	44,444																						
27	78	80	83	83	3,233	10,454	3,567	12,721	4,333	18,778																						
28	76	77	80	80	1,233	1,521	0,567	0,321	1,333	1,778																						
29	77	79	79	79	2,233	4,968	2,567	6,568	0,333	0,111																						
30	78	80	78	78	3,233	10,454	3,567	12,721	-0,667	0,444																						
31	78	79	80	80	3,233	10,454	2,567	6,568	1,333	1,778																						
32	74,767	76,433	78,667		1037,367		995,367		936,667																							
33					35,771		34,323		32,289																							
34					5,981		5,859		5,683																							
35	2243	2293	2360						168739	176257	186590																					
37	1	3199,156	P=35-((O35*2)/90)																													
38	2	229,756	=(A35*2)/30)+((B35*2)/30)+((C35*2)/90)																													
39	3	2969,400	=B37-B38																													
40	4	114,878	=B38/(3-1)																													
41	5	34,131	=B39/(90-3)																													
42	6	3,366	=B40/B41																													

2. Statistik Non-Parametrik

a. χ^2 k Sample Related

χ^2 k Sample Related digunakan untuk menguji hipotesis komparatif tiga sampel atau lebih apabila datanya bertipe nominal.

Contoh:

Peneliti berkeinginan untuk membandingkan angka kredit yang dapat dikumpulkan dosen di 4 fakultas antara yang dapat mengumpulkan lebih dari 100 dan yang mengumpulkan 100 atau kurang dalam 1 tahun.

Proses Pengambilan Keputusan.

Hipotesis:

- | | |
|----|--|
| Ho | Tidak terdapat perbedaan angka kredit yang dapat dikumpulkan dosen dalam 1 tahun di 4 fakultas tersebut. |
| Ha | Terdapat perbedaan angka kredit yang dapat dikumpulkan dosen dalam 1 tahun di 4 fakultas tersebut. |

Dasar Pengambilan Keputusan:

Dengan membandingkan χ^2_{hitung} dengan χ^2_{tabel} dengan ketentuan:

Ho diterima $\chi^2_{hitung} < \chi^2_{tabel}$

Ho ditolak $\chi^2_{hitung} \geq \chi^2_{tabel}$

Dengan menggunakan angka probabilitas, dengan ketentuan:

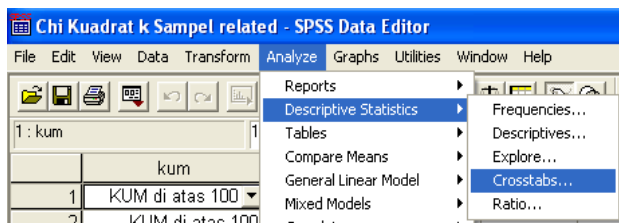
Ho diterima Probabilitas > taraf nyata (α)

Ho ditolak Probabilitas $\leq$ taraf nyata (α)

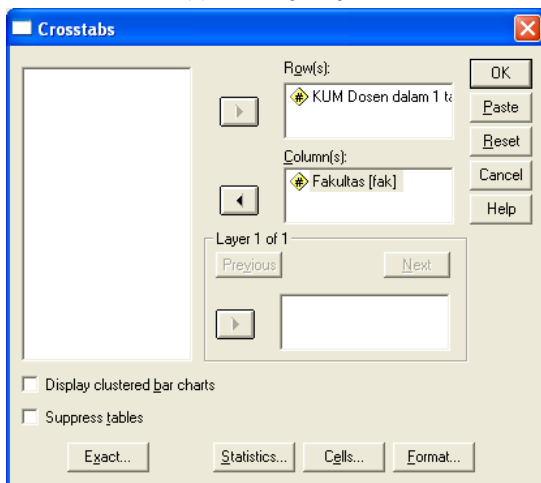
1. Aplikasi dengan SPSS

Langkah-langkah analisis χ^2 k Sample Related dengan SPSS adalah sebagai berikut:

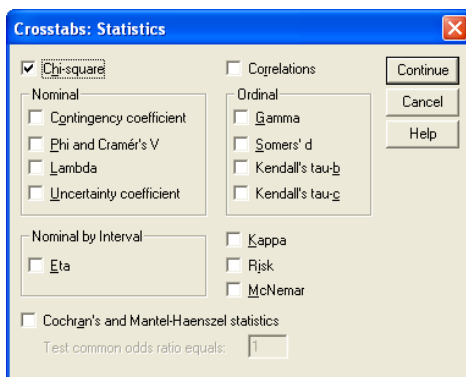
- a) Setelah data diinput, klik **Analyze ► Descriptive Statistics ► Crosstabs**, sebagaimana gambar berikut ini.



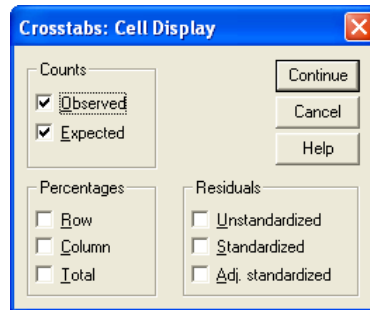
- b) Setelah keluar gambar seperti berikut ini destinasikan salah satu variabel ke kotak **Row(s)** dan variabel satunya pada kotak **Column(s)**. Selanjutnya klik **Statistics**.



- c) Setelah keluar gambar seperti berikut ini klik kotak yang ada di depan **Chi-square**. Selanjutnya klik **Continue**, setelah itu klik **Cells**.



- d) Setelah keluar gambar seperti berikut ini klik kotak yang ada di depan **Expected**. Selanjutnya klik **Continue**, setelah itu klik **Ok**.



e) Berikut ini adalah Output analisis di atas.

Case Processing Summary

	Cases					
	Valid		Missing		Total	
	N	Percent	N	Percent	N	Percent
KUM Dosen dalam 1 tahun * Fakultas	70	100,0%	0	,0%	70	100,0%

Output ini memperlihatkan bahwa jumlah sampel adalah 70.

KUM Dosen dalam 1 tahun * Fakultas Crosstabulation

		KUM Dosen dalam 1 tahun				Total	
		KUM di atas 100		KUM 100 atau kurang			
		Count	Expected Count	Count	Expected Count	Count	Expected Count
Fakultas	Tarbiyah	14	7,4	6	12,6	20	20,0
	Syari'ah	7	7,4	13	12,6	20	20,0
	Dakwah	3	5,6	12	9,4	15	15,0
	Ushuluddin	2	5,6	13	9,4	15	15,0
Total		26	26,0	44	44,0	70	70,0

Output ini menyajikan bahwa ada 14 orang dosen dari Fakultas Tarbiyah yang mempunyai angka kredit di atas 100 dalam 1 tahun dan 6 lainnya mempunyai angka kredit 100 atau kurang dalam 1 tahun. Di Fakultas Syari'ah ada 7 orang dosen yang mempunyai angka kredit di atas 100 dalam 1 tahun dan 13 lainnya mempunyai angka kredit 100 atau kurang dalam 1 tahun. Di Fakultas Dakwah ada 3 orang dosen yang mempunyai angka kredit di atas 100 dalam 1 tahun dan 12 lainnya mempunyai angka kredit 100 atau kurang dalam 1 tahun. Di Fakultas Ushuluddin ada 2 orang dosen yang mempunyai angka kredit di atas 100 dalam 1 tahun dan 13 lainnya mempunyai angka kredit 100 atau kurang dalam 1 tahun.

Chi-Square Tests

	Value	df	Asymp. Sig. (2-sided)
Pearson Chi-Square	14,818 ^a	3	,002
Likelihood Ratio	15,235	3	,002
Linear-by-Linear Association	13,010	1	,000
N of Valid Cases	70		

a. 0 cells (0,0%) have expected count less than 5. The minimum expected count is 5,57.

Skor χ^2_{hitung} sebesar 14,818 yang lebih besar dibanding $\chi^2_{\text{tabel: } 0,05;3}$: 7,8147247. Demikian juga skor Asymp. Sig untuk dua sisi sebesar 0,002 yang jauh lebih kecil dari alpha, yaitu 0,05, maka H_0 ditolak dan H_a diterima. Hal ini berarti terjadi perbedaan angka kredit yang dapat dikumpulkan dosen dalam satu tahun di 4 fakultas tersebut.

2. Aplikasi dengan Microsoft Excel

Rumus yang digunakan adalah sebagai berikut:

$$\chi^2 = \sum \frac{(f_o - f_h)^2}{f_h}$$

Untuk dapat mengaplikasikan rumus tersebut, maka perlu dibuat tabel kotingensi. Untuk lebih jelasnya dapat dilihat pada aplikasi Microsoft Excel berikut ini.

	A	B	C	D	E
1		KUM	Jumlah	Harapan	
2	Fakultas	> 100	14	7,43	5,81
3	Tarbiyah	<= 100	6	12,57	3,44
4	Fakultas	> 100	7	7,43	0,02
5	Syariah	<= 100	13	12,57	0,01
6	Fakultas	> 100	3	5,57	1,19
7	Dakwah	<= 100	12	9,43	0,70
8	Fakultas	> 100	2	5,57	2,29
9	Ushuluddin	<= 100	13	9,43	1,35
10			26		14,818
11			44		
12			70		
13					

Hasil penghitungan ini juga menghasilkan skor χ^2 sebesar 14,818. Apabila ini dibandingkan dengan χ^2_{tabel} dengan dk 2:

=CHIINV(0,05;3)
7,814725

maka kita menerima H_a dan menolak H_0 karena χ^2_{hitung} lebih besar di banding χ^2_{tabel} .

Sedangkan prosedur pembuatan tabel kontingensi dan aplikasi rumus dapat dilihat pada beberapa formula di bawah ini.

14	Formula C10	=C2+C4+C6+C8		
15	Formula C11	=C3+C5+C7+C9		
16	Formula C12	=C10+C11		
17				
18	Formula D2	=(C\$10/C\$12)*20		
19	Formula D3	=(C\$11/C\$12)*20		
20	Formula D4	=(C\$10/C\$12)*20		
21	Formula D5	=(C\$11/C\$12)*20		
22	Formula D6	=(C\$10/C\$12)*15		
23	Formula D7	=(C\$11/C\$12)*15		
24	Formula D8	=(C\$10/C\$12)*15		
25	Formula D9	=(C\$11/C\$12)*15		
26				
27	Formula E2	=((C2-D2)*2/D2)		
28				
29	Formula E10	=SUM(E2:E9)		

b. Cochran's Q

Teknik analisis ini digunakan untuk menguji hipotesis komparatif tiga sampel atau lebih apabila datanya bertipe nominal dikhotomis. Misalnya: ya-tidak, sukses-gagal, disiplin-tidak disiplin, dan terjual-tidak terjual.

Contoh:

Peneliti berkeinginan untuk mengetahui keberhasilan aplikasi analisis statistik antara yang menggunakan kalkulator, Microsoft excel, dan SPSS. Ada 25 mahasiswa yang dilatih analisis statistik menggunakan tiga media penghitungan yang berbeda. Setelah dirasa memahami dan dapat mengaplikasikan, maka 25 mahasiswa tersebut diuji dengan diberikan data yang hampir sama untuk dianalisis dengan menggunakan 3 media penghitungan tersebut.

Proses Pengambilan Keputusan.

Hipotesis:

- Ho Tidak terdapat perbedaan keberhasilan disebabkan perbedaan media analisis statistik.
- Ha Terdapat perbedaan keberhasilan disebabkan perbedaan media analisis statistik.

Dasar Pengambilan Keputusan:

Dengan membandingkan Q_{hitung} dengan χ^2_{tabel} dengan ketentuan:

Ho diterima $Q_{hitung} < \chi^2_{tabel}$

Ho ditolak $Q_{hitung} \geq \chi^2_{tabel}$

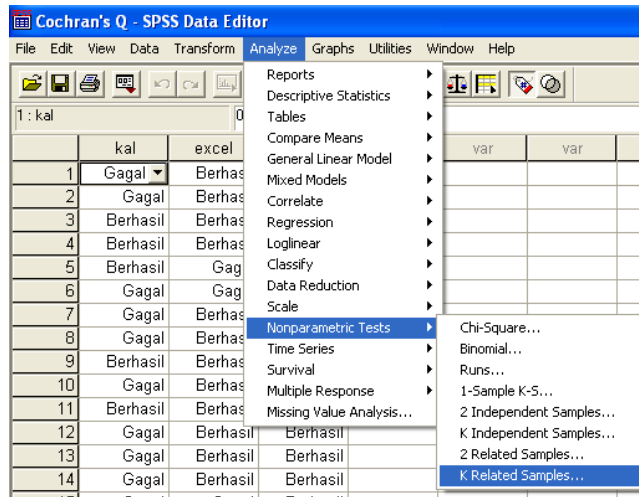
Dengan menggunakan angka probabilitas, dengan ketentuan:

Ho diterima Probabilitas > taraf nyata (α)

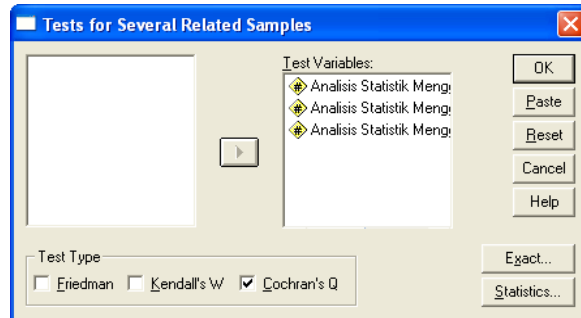
Ho ditolak Probabilitas $\leq$ taraf nyata (α)

1. Aplikasi dengan SPSS

- a) Setelah data diinput, klik **Analyze** ➤ **Nonparametric Test** ➤ **K Related Samples**, sebagaimana gambar berikut ini.



- b) Setelah keluar gambar seperti berikut ini, destinasikan ketiga variabel yang akan dianalisis ke **Test Variables**. Selanjutnya non-aktifkan **Friedman** dan aktifkan **Cochran's**, lalu klik **ok**.



- c) Berikut ini output analisis di atas.

Frequencies

	Value	
	0	1
Analisis Statistik Menggunakan Kalkulator	14	11
Analisis Statistik Menggunakan Microsoft Excel	4	21
Analisis Statistik Menggunakan SPSS	0	25

Output ini memperlihatkan ada 14 mahasiswa ketika menggunakan kalkulator untuk analisis statistik yang gagal, 11 lainnya mengalami keberhasilan. Terdapat 4 mahasiswa ketika menggunakan Microsoft Excel untuk analisis statistik yang gagal, 21 lainnya mengalami keberhasilan. Sedangkan ketika 25 mahasiswa tersebut menggunakan SPSS seluruhnya berhasil.

Test Statistics

N	25
Cochran's Q	19,500 ^a
df	2
Asymp. Sig.	,000

a. 0 is treated as a success.

Output ini menunjukkan bahwa jumlah sampel yang diteliti sebanyak 25 mahasiswa. Skor Cochran's sebesar 19,5. Skor tersebut dibandingkan dengan tabel χ^2 dengan dk. 2:

=CHIIINV(0,05;2)
5,991476

maka H_a diterima dan H_o ditolak karena Q_{hitung} lebih besar dibanding χ^2_{tabel} . Hal ini berarti terdapat perbedaan keberhasilan disebabkan perbedaan media analisis statistik.

Penggunaan χ^2_{tabel} ini disebabkan distribusi sampling Q mendekati distribusi χ^2 .

2. Aplikasi dengan Microsoft Excel

Rumus Cochran's Q adalah sebagai berikut:

$$Q = \frac{(k-1) \left[k \sum_{j=1}^k G_j^2 - \left(\sum_{j=1}^k G_j \right)^2 \right]}{k \sum_{i=1}^N L_i - \sum_{i=1}^N L_i^2}$$

Untuk dapat mengaplikasikan rumus di atas, maka jumlah yang berhasil dijumlahkan untuk setiap kelompok, lalu hasil penjumlahan tersebut dijumlahkan. Masing keberhasilan individu sampel dijumlahkan, lalu hasil penjumlahan tersebut dikuadratkan, dan hasil pengkuadratan itu juga dikuadratkan. Selanjutnya diaplikasikan dalam rumus.

	A	B	C	D	E	F
1	Kalkulator	Excel	SPSS	jml berhasil	jml ²	
2	Gagal	Berhasil	Berhasil	2	4	
3	Gagal	Berhasil	Berhasil	2	4	
4	Berhasil	Berhasil	Berhasil	3	9	
5	Berhasil	Berhasil	Berhasil	3	9	
6	Berhasil	Gagal	Berhasil	2	4	
7	Gagal	Gagal	Berhasil	1	1	
8	Gagal	Berhasil	Berhasil	2	4	
9	Gagal	Berhasil	Berhasil	2	4	
10	Berhasil	Berhasil	Berhasil	3	9	
11	Gagal	Berhasil	Berhasil	2	4	
12	Berhasil	Berhasil	Berhasil	3	9	
13	Gagal	Berhasil	Berhasil	2	4	
14	Gagal	Berhasil	Berhasil	2	4	
15	Gagal	Berhasil	Berhasil	2	4	
16	Gagal	Gagal	Berhasil	1	1	
17	Berhasil	Berhasil	Berhasil	3	9	
18	Berhasil	Gagal	Berhasil	2	4	
19	Berhasil	Berhasil	Berhasil	3	9	
20	Berhasil	Berhasil	Berhasil	3	9	
21	Gagal	Berhasil	Berhasil	2	4	
22	Gagal	Berhasil	Berhasil	2	4	
23	Berhasil	Berhasil	Berhasil	3	9	
24	Berhasil	Berhasil	Berhasil	3	9	
25	Gagal	Berhasil	Berhasil	2	4	
26	Gagal	Berhasil	Berhasil	2	4	
27	11	21	25	57	139	
28						
29	formula A27 =COUNTIF(A2:A26;"Berhasil")					
30	formula D2 =COUNTIF(A2:C2;"Berhasil")					
31						
32	624,000 =((3-1)*((3*(A27*2+B27*2+C27*2)-(D27*2))))					
33	32,000 =(3*D27)-E27					
34	19,500 =A32/A33					
35						

Hasil analisis di atas ternyata juga sama dengan hasil analisis SPSS, di mana Q_{hitung} sebesar: 19,5.

c. Friedman Two-way anova

Teknik analisis ini digunakan untuk menguji hipotesis komparatif 3 sampel atau lebih yang berpasangan apabila datanya berbentuk ordinal.

Contoh:

Seorang pimpinan berkeinginan untuk mengetahui efektifitas kerja stafnya apabila gaya kepemimpinan berbeda. Dalam tahun pertama dia menggunakan gaya

kepemimpinan partisipatif, tahun kedua menggunakan gaya kepemimpinan direktif, dan pada tahun ketiga dia menggunakan gaya kepemimpinan supportif.

Proses Pengambilan Keputusan.

Hipotesis:

Ho Tidak terdapat perbedaan efektivitas kerja staf disebabkan perbedaan gaya kepemimpinan atasannya.

Ha Terdapat perbedaan efektivitas kerja staf disebabkan perbedaan gaya kepemimpinan atasannya.

Dasar Pengambilan Keputusan:

Dengan membandingkan χ^2_{hitung} dengan χ^2_{tabel} dengan ketentuan:

Ho diterima $\chi^2_{hitung} < \chi^2_{tabel}$

Ho ditolak $\chi^2_{hitung} \geq \chi^2_{tabel}$

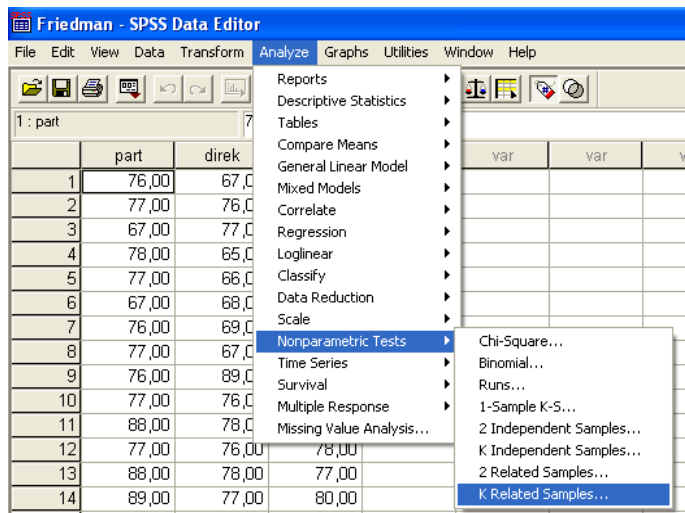
Dengan menggunakan angka probabilitas, dengan ketentuan:

Ho diterima Probabilitas > taraf nyata (α)

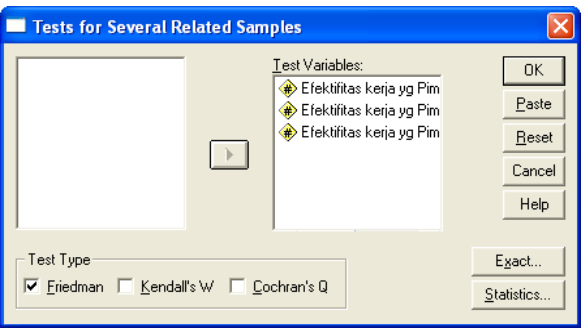
Ho ditolak Probabilitas $\leq$ taraf nyata (α)

1. Aplikasi dengan SPSS

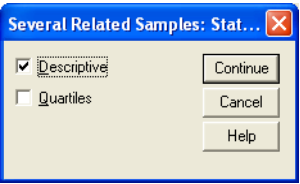
- a) Setelah data diinput, klik **Analyze** ➤ **Nonparametric Test** ➤ **K Related Samples**, sebagaimana gambar berikut ini.



- b) Setelah keluar gambar seperti berikut ini, destinasikan ketiga variabel yang akan dianalisis ke **Test Variables**. Selanjutnya klik **Statistics**.



- c) Setelah keluar gambar berikut ini, aktifkan **Descriptive**, lalu klik **Continue**, lalu klik **Ok**.



- d) Berikut ini output analisis di atas.

Descriptive Statistics					
	N	Mean	Std. Deviation	Minimum	Maximum
Efektifitas Kerja yg pimpinannya Partisipatif	30	80,2000	7,03391	66,00	90,00
Efektifitas Kerja yg pimpinannya Direktif	30	78,0667	7,49682	65,00	90,00
Efektifitas Kerja yg pimpinannya Supportif	30	75,7000	5,85544	65,00	89,00

Output ini menunjukkan bahwa rata-rata skor efektifitas kerja staf yang pimpinannya partisipatif adalah 80,20 dengan standard deviasi sebesar 7,03391. Hal ini menunjukkan bahwa rentang efektifitas kerja staf yang pimpinannya partisipatif berkisar antara 59,09827 sampai dengan 101,3017. Rentang data ini hasil dari $\text{mean} \pm 3$ standard deviasi. Demikian juga untuk data lainnya.

Ranks

	Mean Rank
Efektifitas Kerja yg dipimpinnya Partisipatif	2,40
Efektifitas Kerja yg dipimpinnya Direktif	1,90
Efektifitas Kerja yg dipimpinnya Supportif	1,70

Output ini memperlihatkan rata-rata ranking efektifitas kerja staf yang pimpinan partisipatif sebesar 2,4, yang dipimpinnya direktif sebesar 1,9, dan yang dipimpinnya supportif sebesar 1,7.

Test Statistics^a

N	30
Chi-Square	7,800
df	2
Asymp. Sig.	,020

a. Friedman Test

Output ini menunjukkan bahwa jumlah sampel yang diteliti sebanyak 30 staf. Skor χ^2_{hitung} sebesar 7,8. Skor tersebut dibandingkan dengan tabel χ^2 dengan dk. 2:

=CHIINV(0,05;2)
5,991476

maka H_a diterima dan H_o ditolak karena χ^2_{hitung} lebih besar dibanding χ^2_{tabel} . Kesimpulan yang sama akan didapatkan kalau digunakan skor signifikansi sebesar 0,02 yang jauh lebih kecil dibanding dengan taraf nyata (α): 0,05. Hal ini berarti efektifitas kerja staf berbeda dikarenakan perbedaan gaya kepemimpinan atasannya diterima dan berlaku untuk populasi.

2. Aplikasi dengan Microsoft Excel

Rumus teknik analisis ini sebagai berikut.

$$\chi^2 = \frac{12}{Nk(k-1)} \sum_{j=1}^k (R_j)^2 - 3N(k+1)$$

Untuk dapat mengaplikasikan rumus di atas, setiap skor efektifitas kerja dari masing-masing individu dibuat ranking. Selanjutnya ranking tentang efektifitas kerja staf dari setiap gaya kepemimpinan dijumlahkan. Selanjutnya dimasukkan dalam rumus di atas.

	A	B	C	D	E	F	G	H	I	J	K
1	part	direct	supp	rank part	rank direct	rank supp					
2	76	67	65	3	2	1		formula D2			
3	77	76	66	3	2	1		=RANK(A2;A2:C2;1)			
4	67	77	68	1	3	2		formula E2			
5	78	65	69	3	1	2		=RANK(B2;A2:C2;1)			
6	77	66	78	2	1	3		formula F2			
7	67	68	65	2	3	1		=RANK(C2;A2:C2;1)			
8	76	69	77	2	1	3					
9	77	67	76	3	1	2					
10	76	89	67	2	3	1					
11	77	76	78	2	1	3		0,033333333=12/(30*3*(3+1))			
12	88	78	77	3	2	1		11034=D32*2+E32*2+F32*2			
13	77	76	78	2	1	3		360=3*30*(3+1)			
14	88	78	77	3	2	1		7,800=H11*H12-H13			
15	89	77	80	3	1	2					
16	66	76	88	1	2	3					
17	78	88	77	2	3	1					
18	88	78	77	3	2	1					
19	78	88	77	2	3	1					
20	80	90	78	2	3	1					
21	77	76	78	2	1	3					
22	78	88	77	2	3	1					
23	87	78	89	2	1	3					
24	88	87	76	3	2	1					
25	90	88	77	3	2	1					
26	76	89	67	2	3	1					
27	77	76	78	2	1	3					
28	88	78	77	3	2	1					
29	88	78	77	3	2	1					
30	88	78	77	3	2	1					
31	89	77	80	3	1	2					
32				72	57	51					
33	80,20	78,07	75,70	2,40	1,90	1,70					

Aplikasi rumus di atas ternyata juga menghasilkan skor yang sama dengan aplikasi SPSS.

F. Komparasi k Sampel Independent

1. Statistik Parametrik: One-way Anova

One-way Anova digunakan untuk menguji hipotesis komparatif 3 sampel atau lebih bila datanya berbentuk interval atau rasio, berdistribusi normal, dan variannya homogen. Untuk sub bab ini menguji sampel independen.

Contoh:

Dilakukan penelitian tentang perbedaan produktivitas kerja dosen antara dosen alumni UIN Jakarta, UIN Yogyakarta, dan UIN Malang.

Proses Pengambilan Keputusan.

Hipotesis:

- Ho Tidak terdapat perbedaan produktivitas kerja dosen berdasarkan perbedaan latar belakang pendidikan (latar belakang pendidikan tidak berpengaruh terhadap produktivitas kerja dosen)
- Ha Terdapat perbedaan produktivitas kerja dosen berdasarkan perbedaan latar belakang pendidikan (latar belakang pendidikan berpengaruh terhadap produktivitas kerja dosen)

Dasar Pengambilan Keputusan:

Dengan membandingkan F_{hitung} dengan F_{tabel} dengan ketentuan:

Ho diterima $F_{hitung} < F_{tabel}$

Ho ditolak $F_{hitung} \geq F_{tabel}$

Dengan menggunakan angka probabilitas, dengan ketentuan:

Ho diterima Probabilitas > taraf nyata (α)

Ho ditolak Probabilitas $\leq$ taraf nyata (α)

Salah satu asumsi yang harus terpenuhi ketika akan menggunakan analisis Varians adalah datanya berdistribusi normal. Dari ketiga data tersebut diketahui bahwa ketiganya berdistribusi normal karena signifikansi > taraf nyata (α): 0,05, di mana sig untuk data alumni UIN Jakarta: 0,200, alumni UIN Yogyakarta: 0,200, dan alumni UIN Malang: 0,198. Karena datanya berdistribusi normal, maka analisis one-way ANOVA dapat diteruskan.

Tests of Normality

		Produktivitas Kerja Dosen		
		Latar Belakang Perguruan Tinggi		
		Alumni UIN Jakarta	Alumni UIN Yogyakarta	Alumni UIN Malang
Kolmogorov-Smirnov	Statistic	,126	,113	,131
	df	30	30	30
	Sig.	,200*	,200*	,198
Shapiro-Wilk	Statistic	,933	,933	,939
	df	30	30	30
	Sig.	,059	,059	,088

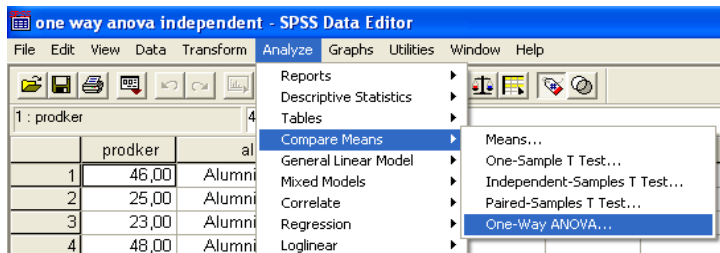
*. This is a lower bound of the true significance.

a. Lilliefors Significance Correction

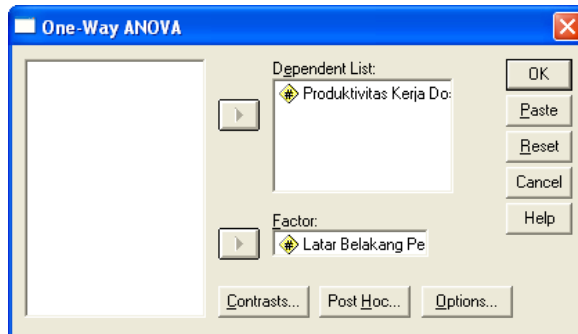
a. Aplikasi dengan SPSS

Langkah-langkah dalam menggunakan analisis ini adalah sebagai berikut:

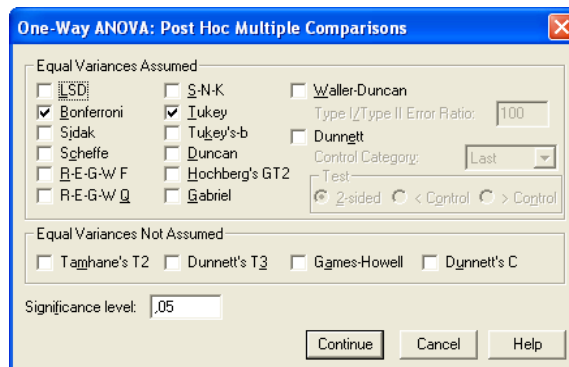
1. Setelah data diinput, lalu klik **Analyze** ➤ **Compare Means** ➤ **One-way ANOVA** sebagaimana gambar di bawah ini:



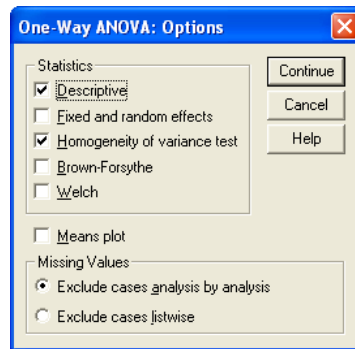
2. Setelah keluar gambar seperti berikut ini, destinasikan **Produktifitas Kerja Dosen** ke kotak **Dependent List** dan **Latar Belakang Pendidikan** ke kotak **Factor**. Selanjutnya klik **Post Hoc...**



3. Setelah keluar gambar seperti berikut ini, aktifkan **Bonferroni** dan **Tukey** untuk mengetahui perbedaan antara sampel satu dengan lainnya. Selanjutnya klik **Continue** dan **Options**.



3. Setelah keluar gambar seperti berikut ini, aktifkan **Descriptive** untuk mendiskripsikan data ketiga sampel tersebut dan **Homogeneity of variance test** untuk mengetahui varians ketiga sampel tersebut homogen atau tidak.



3. Berikut ini adalah Output one-way ANOVA..

Descriptives

Produktivitas Kerja Dosen

	Alumni UIN Jakarta	Alumni UIN Yogyakarta	Alumni UIN Malang	Total
N	30	30	30	90
Mean	38,0333	36,0333	33,0333	35,7000
Std. Deviation	9,76441	9,75381	9,55017	9,80077
Std. Error	1,78273	1,78079	1,74361	1,03309
95% Confidence Interval for Mean				
Lower Bound	34,3872	32,3912	29,4672	33,6473
Upper Bound	41,6794	39,6755	36,5994	37,7527
Minimum	23,00	21,00	18,00	18,00
Maximum	53,00	51,00	49,00	53,00

Output pertama ini mendiskripsikan data, misalnya rata-rata produktivitas kerja dosen alumni UIN Jakarta adalah 38,0333, standard deviasinya 9,76441, dan standard error of mean-nya 1,78273.

Test of Homogeneity of Variances

Produktivitas Kerja Dosen

Levene Statistic	df1	df2	Sig.
,020	2	87	,980

Output kedua ini terkait uji asumsi untuk ANOVA yaitu variannya harus homogen. Berdasarkan skor signifikansi sebesar 0,980 yang lebih tinggi dibanding taraf nyata (α): 0,05, maka varians ketiga sampel adalah homogen.

ANOVA

Produktivitas Kerja Dosen					
	Sum of Squares	df	Mean Square	F	Sig.
Between Groups	380,000	2	190,000	2,024	,138
Within Groups	8168,900	87	93,895		
Total	8548,900	89			

Output ketiga ini terkait uji hipotesis. Skor F_{hitung} nya sebesar 2,024. Bila dibandingkan dengan F_{tabel} untuk pembilang 2 (m-1) dan penyebut 87 (N-m) maka diketahui sebesar:

3,101292

=FINV(0,05;2;87)

maka H_0 diterima dan H_a ditolak karena F_{hitung} lebih kecil dibandingkan F_{tabel} . Oleh karena itu, kesimpulan dari penelitian ini adalah Tidak terdapat perbedaan produktivitas kerja dosen berdasarkan perbedaan latar belakang pendidikan (latar belakang pendidikan tidak berpengaruh terhadap produktivitas kerja dosen). Kesimpulan yang sama akan didapatkan ketika digunakan skor sig sebesar 0,138 yang lebih besar dibanding taraf nyata (α): 0,05.

Multiple Comparisons

Dependent Variable: Produktivitas Kerja Dosen

	(I) Latar Belakang Perguruan Tinggi	(J) Latar Belakang Perguruan Tinggi	Mean Difference (I-J)	Std. Error	Sig.	95% Confidence Interval	
						Lower Bound	Upper Bound
Tukey HSD	Alumni UIN Jakarta	Alumni UIN Yogyakarta	2,0000	2,50194	,704	-3,9658	7,9658
		Alumni UIN Malang	5,0000	2,50194	,119	-,9658	10,9658
	Alumni UIN Yogyakarta	Alumni UIN Jakarta	-2,0000	2,50194	,704	-7,9658	3,9658
		Alumni UIN Malang	3,0000	2,50194	,457	-2,9658	8,9658
	Alumni UIN Malang	Alumni UIN Jakarta	-5,0000	2,50194	,119	-10,9658	,9658
		Alumni UIN Yogyakarta	-3,0000	2,50194	,457	-8,9658	2,9658
Bonferroni	Alumni UIN Jakarta	Alumni UIN Yogyakarta	2,0000	2,50194	1,000	-4,1076	8,1076
		Alumni UIN Malang	5,0000	2,50194	,146	-1,1076	11,1076
	Alumni UIN Yogyakarta	Alumni UIN Jakarta	-2,0000	2,50194	1,000	-8,1076	4,1076
		Alumni UIN Malang	3,0000	2,50194	,701	-3,1076	9,1076
	Alumni UIN Malang	Alumni UIN Jakarta	-5,0000	2,50194	,146	-11,1076	1,1076
		Alumni UIN Yogyakarta	-3,0000	2,50194	,701	-9,1076	3,1076

Oleh karena H_0 yang diterima, maka Output keempat menjadi tidak ada artinya, karena output keempat ini digunakan mendeteksi perbedaan antar sampel.

Produktivitas Kerja Dosen

		N	Subset for alpha = .05
Latar Belakang Perguruan Tinggi			1
Tukey HSD	Alumni UIN Malang	30	33,0333
	Alumni UIN Yogyakarta	30	36,0333
	Alumni UIN Jakarta	30	38,0333
	Sig.		,119

Means for groups in homogeneous subsets are displayed.

a. Uses Harmonic Mean Sample Size = 30,000.

Demikian juga Output kelima ini, karena H_0 diterima, maka wajar apabila ketiga kelompok data tersebut hanya dikelompokkan menjadi satu.

b. Aplikasi dengan Microsoft Excel

Sedangkan prosedur analisis one-way ANOVA dengan Microsoft Excel adalah sebagai berikut:

1. Data diinput ke dalam kelompok sampel masing-masing.
2. Mencari varians dari masing-masing sampel.
3. Menguji homogenitas varians dengan menggunakan uji F, yaitu dengan cara varians terbesar dibagi dengan varians terkecil. Apabila $F_{hitung} < F_{tabel}$, maka variansnya homogen, akan tetapi kalau $F_{hitung} \geq F_{tabel}$, maka variansnya heterogen. Apabila variansnya homogen, maka analisis one-way ANOVA dapat diteruskan.
4. Masing-masing skor pada masing-masing sampel dikuadratkan.
5. Masing-masing skor pada masing-masing sampel dijumlahkan.
6. Masing-masing skor hasil pengkuadratan dari masing-masing skor sampel dijumlahkan.
7. Hasil penjumlahan nomor 5 dijumlahkan.
8. Hasil penjumlahan nomor 6 dijumlahkan.
9. Selanjutnya dimasukkan ke dalam rumus berikut ini.

a. Menghitung jumlah kuadrat total (JK_{Tot}) dengan rumus:

$$JK_{tot} = \sum \sum X_{tot}^2 - \frac{(\sum X_{tot})^2}{N}$$

b. Menghitung jumlah kuadrat antar kelompok (JK_{antar}) dengan rumus:

$$JK_{\text{ant}} = \sum \frac{(\sum_k)^2}{n_k} - \frac{(\sum x_{\text{tot}})^2}{N}$$

- c. Menghitung jumlah kuadrat dalam kelompok (JK_{dalam}) dengan rumus:

$$JK_{\text{dalam}} = JK_{\text{tot}} - JK_{\text{antar}}$$

- d. Menghitung mean kuadrat antar kelompok (MK_{antar}) dengan rumus:

$$MK_{\text{antar}} = \frac{JK_{\text{antar}}}{m - 1}$$

- e. Menghitung mean kuadrat dalam kelompok (MK_{dalam}) dengan rumus:

$$MK_{\text{dalam}} = \frac{JK_{\text{dalam}}}{N - m}$$

- f. Menghitung F_{hitung} dengan rumus:

$$F_{\text{hitung}} = \frac{MK_{\text{antar}}}{MK_{\text{dalam}}}$$

Untuk lebih jelasnya tentang prosedur analisis one-way ANOVA dapat dilihat pada aplikasi Microsoft Excel berikut ini.

A	B	C	D	E	F	G	H	I	J	K	L	M	N	O	P	Q
1	x_1		x_1^2	x_2			x_2^2	x_3			x_3^2	JUMLAH				
2	46	7,97	63,47	43	6,97	48,53	18,49	29	-4,03	16,27	841,00	118,00	4806,00			
3	25	-13,03	169,87	625,00	32	-4,03	16,27	1024,00	42	8,97	80,40	1764,00	99,00	3413,00	F_{hitung}	$1,05 = C33/K33$
4	23	-15,03	226,00	529,00	34	-2,03	4,13	1156,00	21	-12,03	144,80	441,00	78,00	2126,00	F_{tabel}	$1,86 = FINV(0,05;29;29)$
5	48	9,97	99,33	2304,00	47	10,97	120,27	2209,00	19	-14,03	196,93	361,00	114,00	4874,00	Homogen	
6	37	-1,03	1,07	1369,00	26	-10,03	100,67	676,00	44	10,97	120,27	1936,00	107,00	3981,00		
7	39	0,97	0,93	1521,00	24	-12,03	144,80	576,00	33	-0,03	0,00	1089,00	96,00	3186,00	1	$8,548,9 = N32 - ((M32^2)/90)$
8	52	13,97	195,07	2704,00	49	12,97	168,13	2401,00	35	1,97	3,87	1225,00	136,00	6330,00	2	$380 = (((A33^2)/30) + ((E33^2)/30) + ((I33^2)/30)) - ((M32^2)/90)$
9	31	-7,03	49,47	961,00	38	1,97	3,87	1444,00	48	14,97	224,00	2304,00	117,00	4709,00		$= P7 - P8$
10	29	-9,03	81,60	841,00	40	3,97	15,73	1600,00	27	-6,03	36,40	729,00	96,00	3170,00	3	$8168,9 = P8/(3-1)$
11	52	13,97	195,07	2704,00	51	14,97	224,00	2601,00	23	-10,03	100,67	529,00	126,00	5834,00	4	$190 = P10/(90-3)$
12	39	0,97	0,93	1521,00	28	-8,03	64,53	784,00	25	-8,03	64,53	625,00	92,00	2930,00	5	$93,90 = P10/(90-3)$
13	39	0,97	0,93	1521,00	24	-12,03	144,80	576,00	21	-12,03	144,80	441,00	84,00	2538,00	6	$2024 = P11/P12$
14	50	11,97	143,20	2500,00	47	10,97	120,27	2209,00	44	10,97	120,27	1936,00	141,00	6645,00		
15	27	-11,03	121,73	729,00	34	-2,03	4,13	1156,00	31	-2,03	4,13	961,00	92,00	2846,00		
16	23	-15,03	226,00	529,00	34	-2,03	4,13	1156,00	31	-2,03	4,13	961,00	88,00	2646,00		
17	46	7,97	63,47	2116,00	44	7,97	63,47	1936,00	41	7,97	63,47	1681,00	131,00	5733,00		
18	33	-5,03	25,33	1089,00	23	-13,03	169,87	529,00	20	-13,03	169,87	400,00	76,00	2018,00		
19	34	-4,03	16,27	1156,00	21	-15,03	226,00	441,00	18	-15,03	226,00	324,00	73,00	1921,00		
20	47	8,97	80,40	2209,00	46	9,97	99,33	2116,00	43	9,97	99,33	1849,00	136,00	6174,00		
21	26	-12,03	144,80	676,00	35	-1,03	1,07	1225,00	32	-1,03	1,07	1024,00	93,00	2925,00		
22	24	-14,03	196,93	576,00	37	0,97	0,93	1369,00	45	11,97	143,20	2025,00	106,00	3970,00		
23	49	10,97	120,27	2401,00	50	13,97	195,07	2500,00	34	0,97	0,93	1156,00	133,00	6057,00		
24	38	-0,03	0,00	1444,00	29	-7,03	49,47	841,00	36	2,97	8,80	1296,00	103,00	3581,00		
25	40	1,97	3,87	1600,00	27	-9,03	81,60	729,00	49	15,97	254,93	2401,00	116,00	4730,00		
26	53	14,97	224,00	2809,00	50	13,97	195,07	2500,00	26	-7,03	49,47	676,00	129,00	5985,00		
27	32	-6,03	36,40	1024,00	37	0,97	0,93	1369,00	22	-11,03	121,73	484,00	91,00	2877,00		
28	29	-9,03	81,60	841,00	37	0,97	0,93	1369,00	45	11,97	143,20	2025,00	111,00	4235,00		
29	52	13,97	195,07	2704,00	48	11,97	143,20	2304,00	32	-1,03	1,07	1024,00	132,00	6032,00		
30	39	0,97	0,93	1521,00	25	-11,03	121,73	625,00	32	-1,03	1,07	1024,00	96,00	3170,00		
31	39	0,97	0,93	1521,00	21	-15,03	226,00	441,00	43	9,97	99,33	1349,00	103,00	3811,00		
32	38,03		2764,97	46161,00	36,03		2738,97	41711,00	33,03		2644,97	35381,00	3213,00	123253,00		
33	1141		95,34	1081			95,14	991			91,21					

Hasil aplikasi Microsoft Excel tentang F_{hitung} sebesar 2,024 juga sama dengan hasil hitung SPSS.

2. Statistik Non-Parametrik

a. χ^2 k Sample Independent

χ^2 k Sample Independent digunakan untuk menguji hipotesis komparatif tiga sampel atau lebih apabila datanya bertipe nominal.

Contoh:

Peneliti berkeinginan untuk membandingkan perbedaan alasan orang tua memilihkan anaknya sekolah.

Proses Pengambilan Keputusan.

Hipotesis:

Ho Tidak terdapat perbedaan alasan memilihkan anaknya tipe sekolah berdasarkan perbedaan profesi orang tua.

Ha Terdapat perbedaan alasan memilihkan anaknya tipe sekolah berdasarkan perbedaan profesi orang tua.

Dasar Pengambilan Keputusan:

Dengan membandingkan χ^2_{hitung} dengan χ^2_{tabel} dengan ketentuan:

Ho diterima $\chi^2_{hitung} < \chi^2_{tabel}$

Ho ditolak $\chi^2_{hitung} \geq \chi^2_{tabel}$

Dengan menggunakan angka probabilitas, dengan ketentuan:

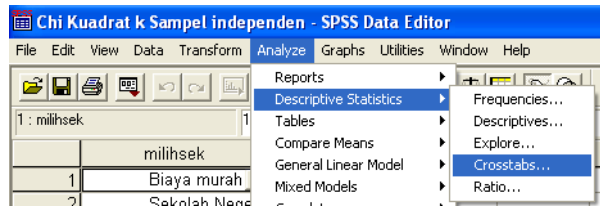
Ho diterima Probabilitas > taraf nyata (α)

Ho ditolak Probabilitas $\leq$ taraf nyata (α)

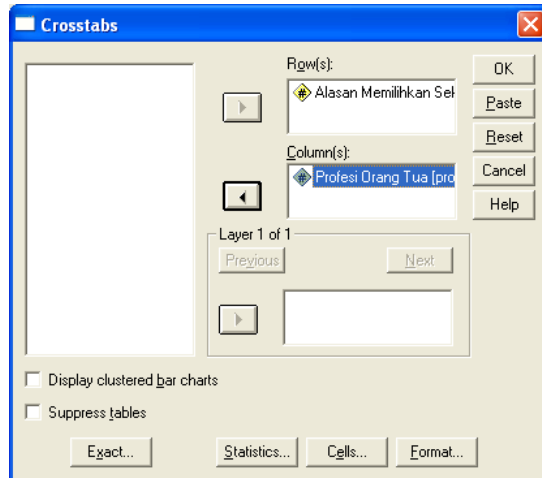
1. Aplikasi dengan SPSS

Langkah-langkah analisis χ^2 k Sample Independent dengan SPSS adalah sebagai berikut:

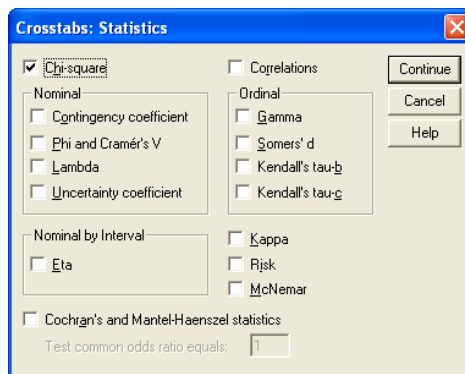
- a) Setelah data diinput, klik **Analyze ► Descriptive Statistics ► Crosstabs**, sebagaimana gambar berikut ini.



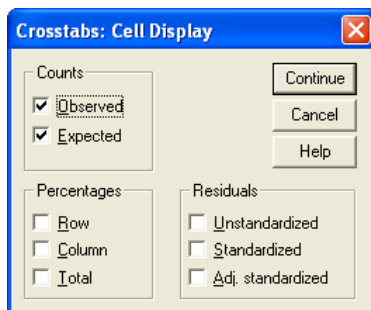
- b) Setelah keluar gambar seperti berikut ini destinasikan salah satu variabel ke kotak **Row(s)** dan variabel satunya pada kotak **Column(s)**. Selanjutnya klik **Statistics**.



- c) Setelah keluar gambar seperti berikut ini klik kotak yang ada di depan **Chi-square**. Selanjutnya klik **Continue**, setelah itu klik **Cells**.



- d) Setelah keluar gambar seperti berikut ini klik kotak yang ada di depan **Expected**. Selanjutnya klik **Continue**, setelah itu klik **Ok**.



e) Berikut ini adalah Output analisis di atas.

Case Processing Summary

	Cases					
	Valid		Missing		Total	
	N	Percent	N	Percent	N	Percent
Alasan Memilih Sekolah Anak * Profesi Orang Tua	90	100,0%	0	,0%	90	100,0%

Output ini memperlihatkan bahwa jumlah sampel adalah 90.

Alasan Memilih Sekolah Anak * Profesi Orang Tua Crosstabulation

			Profesi Orang Tua			Total
			Pegawai Negeri	Pegawai swasta	Wiraswasta	
Alasan Memilih Sekolah Anak	Biaya murah	Count	7	4	1	12
		Expected Count	4,0	4,0	4,0	12,0
	Kwalitas Tinggi Walau Mahal	Count	0	17	24	41
		Expected Count	13,7	13,7	13,7	41,0
	Sekolah Negeri	Count	23	9	5	37
		Expected Count	12,3	12,3	12,3	37,0
	Total	Count	30	30	30	90
		Expected Count	30,0	30,0	30,0	90,0

Output ini menyajikan bahwa ada 7 orang orang tua yang berprofesi sebagai pegawai negeri, 4 orang orang tua yang berprofesi sebagai pegawai swasta, dan 1 orang orang tua yang berprofesi sebagai wiraswasta yang memilihkan sekolah anaknya dengan alasan biaya murah. Terdapat 17 orang orang tua yang berprofesi sebagai pegawai swasta dan 24 orang orang tua yang berprofesi sebagai wiraswasta yang memilihkan sekolah anaknya dengan alasan kwalitas tinggi walau biaya mahal. Terdapat 23 orang orang tua yang berprofesi sebagai pegawai negeri, 9 orang orang tua yang berprofesi sebagai pegawai swasta, dan 5 orang orang tua yang berprofesi sebagai wiraswasta yang memilihkan sekolah anaknya dengan alasan sekolah negeri.

Chi-Square Tests

	Value	df	Asymp. Sig. (2-sided)
Pearson Chi-Square	41,279 ^a	4	,000
Likelihood Ratio	53,478	4	,000
Linear-by-Linear Association	5,079	1	,024
N of Valid Cases	90		

a. 3 cells (33,3%) have expected count less than 5. The minimum expected count is 4,00.

Skor χ^2_{hitung} sebesar 41,279 yang lebih besar dibanding $\chi^2_{tabel: 0,05;4}$: 9,487728. Demikian juga skor Asymp. Sig untuk dua sisi sebesar 0,000 yang jauh lebih kecil dari alpha, yaitu 0,05, maka H_0 ditolak dan H_a diterima. Hal ini berarti terdapat perbedaan alasan memilihkan anaknya tipe sekolah berdasarkan perbedaan profesi orang tua.

2. Aplikasi dengan Microsoft Excel

Rumus yang digunakan adalah sebagai berikut:

$$\chi^2 = \sum \frac{(f_o - f_h)^2}{f_h}$$

Untuk dapat mengaplikasikan rumus tersebut, maka perlu dibuat tabel kotingensi. Untuk lebih jelasnya dapat dilihat pada aplikasi Microsoft Excel berikut ini.

	A	B	C	D	E	F
1			Profesi Orang Tua			
2			Peg. Negeri	Peg. Swasta	Wiraswasta	Jumlah
3		Biaya Murah	7	4	1	12
4		Harapan	4	4	4	
5	Alasan	Kwalitas Tinggi Walau Mahal	0	17	24	41
6	Memilihkan	Harapan	13,67	13,67	13,67	
7	Sekolah	Sekolah Negeri	23	9	5	37
8	Anak	Harapan	12,33	12,33	12,33	
9			30	30	30	90
10	2,250	=((C3-C4)*2)/C4				
11	0,000	=((D3-D4)*2)/D4				
12	2,250	=((E3-E4)*2)/E4				
13	13,667	=((C5-C6)*2)/C6				
14	0,813	=((D5-D6)*2)/D6				
15	7,813	=((E5-E6)*2)/E6				
16	9,225	=((C7-C8)*2)/C8				
17	0,901	=((D7-D8)*2)/D8				
18	4,360	=((E7-E8)*2)/E8				
19	41,279	=SUM(H3:H11)				

Hasil penghitungan ini juga menghasilkan skor χ^2 sebesar 41,279. Apabila ini dibandingkan dengan χ^2 tabel dengan dk 4:

=CHIINV(0,05;4)
9,487728

maka H_a diterima dan H_o ditolak karena χ^2_{hitung} lebih besar di banding χ^2_{tabel} .

b. Median Extention

Teknik analisis ini digunakan untuk menguji komparatif tiga sampel independen atau lebih yang datanya bertipe ordinal. Sesuai dengan namanya, pengujian didasarkan atas median dari sampel yang dianalisis.

Contoh:

Peneliti berkeinginan untuk membandingkan nilai MK Statistika Pendidikan antara mahasiswa yang dibimbing oleh Dosen yang berbeda. Sampel penelitian ini adalah 20 mahasiswa dibimbing oleh Dosen A, 19 mahasiswa dibimbing oleh Dosen B, dan 20 mahasiswa dibimbing oleh Dosen C.

Proses Pengambilan Keputusan.

Hipotesis:

H_o Tidak terdapat perbedaan nilai Mata Kuliah Statistika Pendidikan mahasiswa berdasarkan perbedaan Dosen Pembimbingnya.

H_a Terdapat perbedaan nilai Mata Kuliah Statistika Pendidikan mahasiswa berdasarkan perbedaan Dosen Pembimbingnya.

Dasar Pengambilan Keputusan:

Dengan membandingkan χ^2_{hitung} dengan χ^2_{tabel} dengan ketentuan:

H_o diterima $\chi^2_{hitung} < \chi^2_{tabel}$

H_o ditolak $\chi^2_{hitung} \geq \chi^2_{tabel}$

Dengan menggunakan angka probabilitas, dengan ketentuan:

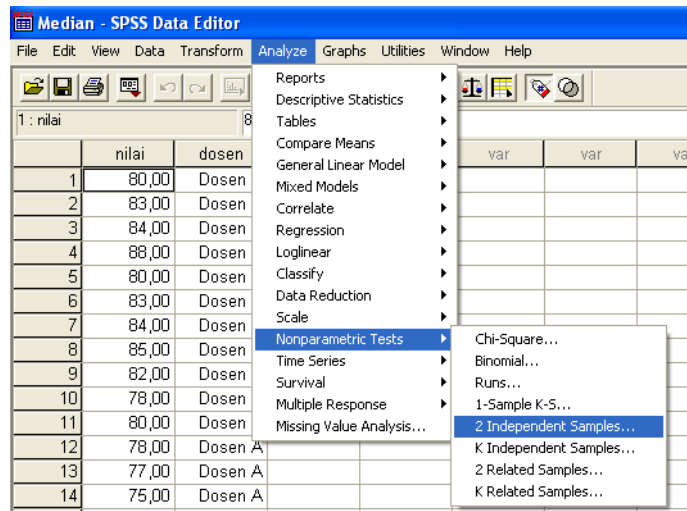
H_o diterima Probabilitas $>$ taraf nyata (α)

H_o ditolak Probabilitas $\leq$ taraf nyata (α)

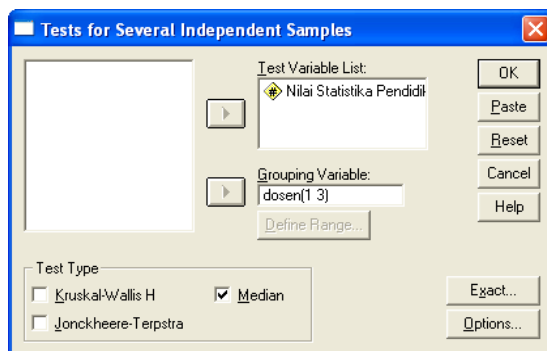
1. Aplikasi dengan SPSS

Langkah-langkah yang dapat ditempuh untuk menganalisis dengan Median test adalah sebagai berikut:

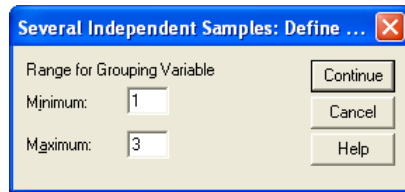
- a) Setelah data diinput, klik **Analyze ► Nonparametric Tests ► K Independent Samples**, sebagaimana gambar berikut ini.



- b) Setelah keluar gambar seperti berikut ini, destinasikan data tentang **nilai statistika pendidikan** ke **Test Variable List**, variable kategori, yaitu **dosen** ke **Grouping Variable**, Non aktifkan **Kruskal-Wallis H**, dan aktifkan **Median**, lalu klik **Define Range..**.



- c) Setelah keluar gambar seperti berikut ini, Isilah angka **1** pada kotak **Minimum**, dan angka **3** pada kotak **Maximum**, lalu klik **Continue**, lalu klik **Ok**.



d) Berikut ini adalah outputnya.

		DOSEN		
		Dosen A	Dosen B	Dosen C
Nilai Statistika	> Median	16	6	4
Pendidikan	<= Median	4	13	16

Output ini memperlihatkan bahwa terdapat 16 mahasiswa yang dibimbing oleh Dosen A, 6 dibimbing oleh Dosen B, dan 4 dibimbing oleh Dosen C yang mempunyai nilai statistika pendidikan lebih tinggi dibanding mediannya. Terdapat 4 mahasiswa dibimbing oleh Dosen A, 13 dibimbing oleh Dosen B, dan 16 dibimbing oleh Dosen C yang mempunyai nilai statistika pendidikan sama dengan mediannya atau lebih rendah.

Test Statistics ^b	
	Nilai Statistika Pendidikan
N	59
Median	77,0000
Chi-Square	16,379 ^a
df	2
Asymp. Sig.	,000

a. 0 cells (.0%) have expected frequencies less than

5. The minimum expected cell frequency is 8,4.

b. Grouping Variable: DOSEN

Output di atas memperlihatkan bahwa sampel penelitiannya ini adalah 59 mahasiswa. Median seluruh mahasiswa yang dijadikan sampel adalah 77. Skor χ^2_{hitung} adalah 16,379. Skor tersebut dibandingkan dengan tabel χ^2 dengan dk. 2:

=CHIINV(0,05;2)
5,991476

maka H_a diterima dan H_o ditolak Karena χ^2_{hitung} lebih tinggi dibanding dengan χ^2_{tabel} . Hal ini sesuai dengan skor Asymp. Signifikansi adalah 0,000 yang lebih rendah dibanding $\alpha=5\%$. Penelitian ini berkesimpulan bahwa terdapat perbedaan nilai Mata Kuliah Statistika Pendidikan mahasiswa berdasarkan perbedaan Dosen Pembimbingnya.

2. Aplikasi dengan Microsoft Excel

Untuk mengaplikasikan teknik analisis Median Extention, maka seluruh data harus diurutkan dan dicari mediannya. Skor data yang lebih tinggi dari median dipisahkan dari skor yang sama dengan median atau lebih rendah. Untuk menguji hipotesis dapat digunakan rumus χ^2 sebagai berikut:

$$\chi^2 = \sum \frac{(f_{o_{ij}} - f_{h_{ij}})^2}{f_{h_{ij}}}$$

Untuk lebih jelasnya dapat dilihat pada aplikasi Microsoft Excel berikut ini.

	A	B	C	D	E	F	G	H	I
1	Dosen A	Dosen B	Dosen C						
2	80	75	78		16	10,00	3,600		
3	83	65	76		4	10,00	3,600		
4	84	67	75						
5	88	78	76		6	9,50	1,289		
6	80	80	75		13	9,50	1,289		
7	83	72	76						
8	84	75	72		4	10,00	3,600		
9	85	72	68		16	10,00	3,600		
10	82	65	69				16,979		
11	78	60	70						
12	80	62	75						
13	78	78	78						
14	77	85	84						
15	75	76	83						
16	76	87	77						
17	76	88	69						
18	78	77	70						
19	80	72	71						
20	81	76	72						
21	80		77						
22									

formula E2	=COUNTIF(A2:A21,">77")
formula E3	=COUNTIF(A2:A21,"<=77")
formula E5	=COUNTIF(B2:B21,">77")
formula E6	=COUNTIF(B2:B21,"<=77")
formula E8	=COUNTIF(C2:C21,">77")
formula E9	=COUNTIF(C2:C21,"<=77")
formula G2	=((E2-F2)*2)/F2
formula G10	=SUM(G2:G9)

Hasil analisis Median Extention dengan Microsoft Excel ternyata menghasilkan skor yang sama dengan aplikasi SPSS.

c. Kruskal-Wallis One-way Anova

Uji Kruskal-Wallis digunakan untuk menguji hipotesis komparatif yang mempunyai 3 atau lebih sampel independen bila datanya berbentuk ordinal. Pada dasarnya, uji ini merupakan salah satu alternatif dari uji F atau Anova apabila ada asumsi yang tidak terpenuhi, misalkan data tidak bertipe interval/rasio atau distribusi datanya tidak normal. Untuk kasus terakhir ini, maka data yang semula bertipe interval atau rasio harus diubah menjadi ordinal.

Contoh:

Dilakukan penelitian perbandingan tentang prestasi kerja alumni Tarbiyah, Syari'ah, Ushuluddin, dan Dakwah. Sampel yang digunakan adalah 15 alumni Tarbiyah, 15 alumni Syari'ah, 10 alumni Ushuluddin, dan 10 alumni Dakwah.

Proses Pengambilan Keputusan.

Hipotesis:

- | | |
|----|---|
| Ho | Tidak terdapat perbedaan prestasi kerja antara alumni Tarbiyah, Syari'ah, Ushuluddin, dan Dakwah. |
| Ha | Terdapat perbedaan prestasi kerja antara alumni Tarbiyah, Syari'ah, Ushuluddin, dan Dakwah. |

Dasar Pengambilan Keputusan:

Dengan membandingkan H/χ^2_{hitung} dengan χ^2_{tabel} dengan ketentuan:

Ho diterima $H/\chi^2_{hitung} < \chi^2_{tabel}$

Ho ditolak $H/\chi^2_{hitung} \geq \chi^2_{tabel}$

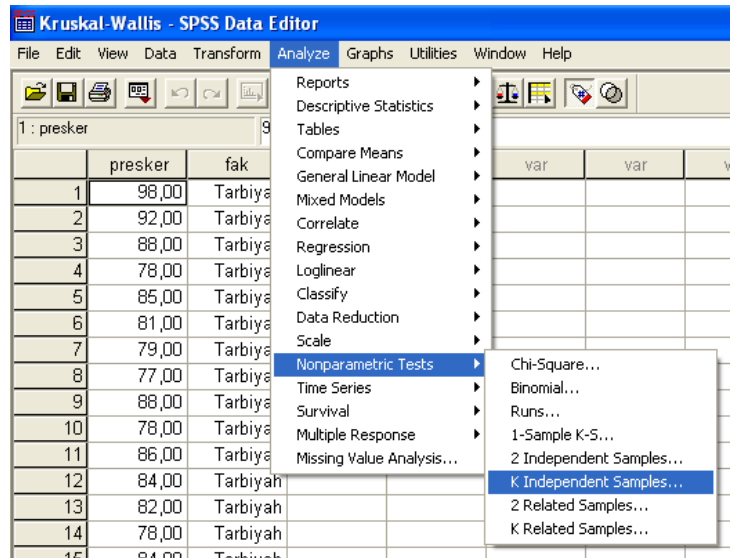
Dengan menggunakan angka probabilitas, dengan ketentuan:

Ho diterima Probabilitas > taraf nyata (α)

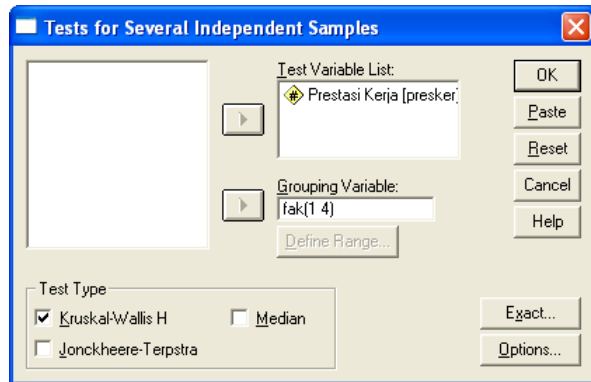
Ho ditolak Probabilitas $\leq$ taraf nyata (α)

1. Aplikasi dengan SPSS

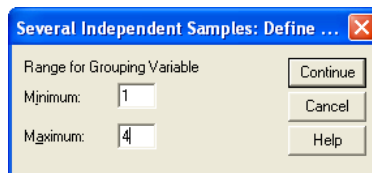
- a) Setelah data diinput, klik **Analyze ► Nonparametric Tests ► K Independent Samples**, sebagaimana gambar berikut ini.



- b) Setelah keluar gambar seperti berikut ini, destinasikan data tentang **Prestasi Kerja** ke **Test Variable List**, variable kategori, yaitu **fak** ke **Grouping Variable**, lalu klik **Define Range..**.



- c) Setelah keluar gambar seperti berikut ini, Isilah angka 1 pada kotak **Minimum**, dan angka 4 pada kotak **Maximum**, lalu klik **Continue**, lalu klik **Ok**.



d) Berikut ini adalah outputnya.

Ranks			
FAK		N	Mean Rank
Prestasi Kerja	Tarbiyah	15	36,53
	Syari'ah	15	23,07
	Ushuluddin	10	17,60
	Dakwah	10	20,50
	Total	50	

Output ini memperlihatkan bahwa sampel penelitian dari alumni Tarbiyah sebanyak 15 orang dengan rata-rata ranking 36,53, alumni Syari'ah sebanyak 15 dengan rata-rata ranking 23,07, alumni Ushuluddin sebanyak 10 dengan rata-rata ranking 17,60, alumni Dakwah sebanyak 10 dengan rata-rata ranking 20,50. Seluruh sampel berjumlah 50.

Test Statistics ^{a,b}	
Prestasi Kerja	
Chi-Square	13,197
df	3
Asymp. Sig.	,004

a. Kruskal Wallis Test

b. Grouping Variable: FAK

Output di atas memperlihatkan bahwa skor χ^2_{hitung} adalah 13,197. Skor tersebut dibandingkan dengan tabel χ^2 dengan dk. 3, yaitu: 7,814725.

=CHIINV(0,05;3)
7,814725

maka H_a diterima dan H_o ditolak karena χ^2_{hitung} lebih tinggi dibanding dengan χ^2_{tabel} . Hal ini sesuai dengan skor Asymp. Signifikansi sebesar 0,004 yang lebih rendah dibanding $\alpha=5\%$. Penelitian ini berkesimpulan bahwa terdapat perbedaan prestasi kerja antara alumni Tarbiyah, Syari'ah, Ushuluddin, dan Dakwah.

2. Aplikasi dengan Microsoft Excel

Rumus Kruskal Wallis H adalah sebagai berikut.

$$H = \frac{12}{N(N-1)} \sum_{j=1}^k \frac{R_j^2}{n_j} - 3(N+1)$$

Untuk dapat mengaplikasikan rumus tersebut, maka seluruh skor data dijadikan satu, lalu diranking. Setelah itu baru dipisahkan berdasarkan masing-masing sampel. Ranking

dari masing-masing sampel tersebut dijumlahkan. Setelah itu baru dimasukkan rumus di atas.

	A	B	C	D	E	F	G	H	
1	Tarbiyah	Ranking	Syari'ah	Ranking	Ushulud.	Ranking	Dakwah	Ranking	
2	98	50	77	17,5	79	28,5	75	10	
3	92	49	87	44,5	77	17,5	70	2	
4	88	47	79	28,5	72	4,5	73	6,5	
5	78	23,5	81	34,5	70	2	77	17,5	
6	85	42	82	38	73	6,5	79	28,5	
7	81	34,5	88	47	75	10	78	23,5	
8	79	28,5	78	23,5	77	17,5	76	13	
9	77	17,5	76	13	78	23,5	80	31,5	
10	88	47	75	10	81	34,5	81	34,5	
11	78	23,5	77	17,5	80	31,5	82	38	
12	86	43	87	44,5					
13	84	40,5	76	13					
14	82	38	72	4,5					
15	78	23,5	70	2					
16	84	40,5	74	8					
17		548		346		176		205	
18		36,53		23,07		17,60		20,50	
19									
20		0,005 =12/(50*51)							
21		35301,433 =((B17^2)/15)+((D17^2)/15)+((F17^2)/10)+((H17^2)/10)							
22		153,000 =3*51							
23		13,124 =B20*B21.B22							
24									

Hasil hitung dengan Microsoft Excel ini berbeda sedikit dengan hasil hitung SPSS. χ^2_{hitung} dengan SPSS adalah 13,197, sedangkan χ^2_{hitung} dengan Microsoft Excel adalah 13,124. Hanya saja, karena sedikitnya perbedaan tersebut tidak berakibat pada perbedaan kesimpulan yaitu menerima H_a dan menolak H_o .

DAFTAR KEPUSTAKAAN

- Ancok, Djamaludin, *Teknik Penyusunan Skala Pengukur*. (Yogyakarta: Pusat Studi Kependudukan dan Kebijakan Universitas Gadjah Mada, 2002).
- Anwar, Ali, *Cara Mudah Menulis Karya Ilmiah*. (Kediri: IAIT Press, 2009).
- Arifin, Johar, *Aplikasi Excel dalam Statistik dan Riset Terapan*. (Jakarta: Elex Media Komputindo, 2005).
- Arikunto, Suharsimi, *Manajemen Penelitian*. (Jakarta: Rineka Cipta, 2003).
- _____, *Prosedur Penelitian: Suatu Pendekatan Praktek*. (Jakarta: Rineka Cipta, 2002).
- Ary, Donald, dkk., *Pengantar Penelitian dalam Pendidikan*. Terjemahan Arief Furchan dari *Introduction to Research in Education*. (Surabaya: Usaha Nasional, 1982).
- Furqon, *Statistika Terapan untuk Penelitian*. (Bandung: Alfabeta, 2004).
- Hasan, Iqbal, *Analisis Data Penelitian dengan Statistik*. (Jakarta: Bumi Aksara, 2004).
- Kerlinger, Fred N., *Asas-asas Penelitian Behavioral*. Terjemahan Landung R. Simatupang dari *Foundation of Behavioral Research*. (Yogyakarta: UGM Press, 2006).
- Koentjaraningrat, *Metode-metode Penelitian Masyarakat*. (Jakarta: Gramedia, 1993).
- Kusnandar, Dadan, *Metode Statistik dan Aplikasinya dengan Minitab dan*

- Excel. (Yogyakarta: Madyan Press, 2004).
- Nazir, Moh., *Metode Penelitian*. (Bogor: Ghalia Indonesia, 2005).
- Riyanto, Yatim, *Metodologi Penelitian Pendidikan*. (Surabaya: SIC, 2001).
- Romus, Mahendra dkk., *Aplikasi Program SPSS dalam Analisis Data Penelitian*. (Riau: UIN Suska Press, 2007).
- Santosa, Purbayu Budi dan Ashari, *Analisis Statistik dengan Microsoft Excel dan SPSS*. (Yogyakarta: Andi, 2005).
- Santosa, R. Gunawan, *Statistik*. (Yogyakarta: Andi, 2004).
- Santoso, Singgih, *Buku Latihan SPSS Statistik Multivariat*. (Jakarta: Elex Media Komputindo, 2003).
- _____, *Buku Latihan SPSS: Statistik Non-parametrik*. (Jakarta: Elex Media Komputindo, 2003).
- _____, *Buku Latihan SPSS: Statistik Non-parametrik*. (Jakarta: Elex Media Komputindo, 2001).
- _____, *Statistik Diskriptif: Konsep dan Aplikasi dengan Microsoft Excel dan SPSS*. (Yogyakarta: Penerbit Andi, 2003).
- Saukah, Ali dkk., *Pedoman Penulisan Karya Ilmiah*. (Malang: Universitas Negeri Malang, 2000).
- Singarimbun, Masri dan Sofian Effendi (Ed.), *Metode Penelitian Survey*. (Jakarta: LP3ES, 1989).
- Somantri, Ating dan Sambas Ali Muhidin, *Aplikasi Statistik dalam Penelitian*. (Bandung: Pustaka Setia, 2006).
- Sugiyono dan Eri Wibowo, *Statistika Penelitian dan Aplikasinya dengan SPSS 10.00 for Windows*. (Bandung: CV Alfabeta, 2002).
- Sugiyono, *Metode Penelitian Bisnis*. (Bandung: CV Alfabeta, 2005).
- _____, *Metode Penelitian Pendidikan: Pendekatan Kuantitatif, Kualitatif, dan R&D*. (Bandung: CV Alfabeta, 2006).
- _____, *Statistika untuk Penelitian*. (Bandung: CV Alfabeta, 2003).
- Sulaiman, Wahid, *Analisis Regresi Menggunakan SPSS*. (Yogyakarta: Andi, 2004).
- _____, *Jalan Pintas Menguasai SPSS 10..* (Yogyakarta: Andi,

2002).

_____, *Statistik Non-parametrik: Contoh, Kasus, dan Pemecahannya dengan SPSS*. (Yogyakarta: Andi, 2003).

Supranto, J., *Teknik Sampling untuk Survey dan Eksperimen*. (Jakarta: Rineka Cipta, 2000).

Surakhmad, Winarno, *Pengantar Penelitian Ilmiah: Dasar, Metode, dan Teknik*. (Bandung: Tarsito, 1990).

Suwito dkk., *Buku Pedoman Tenaga Akademik Perguruan Tinggi Agama Islam dan PAI pada PTU*. (Jakarta: Dirjen Binbaga Islam Depag RI, 2003).

TABEL I
LUAS DI BAWAH LENGKUNGAN KURVA NORMAL
DARI 0 s.d. Z

Z	0	1	2	3	4	5	6	7	8	9
0,0	0000	0040	0080	0120	0160	0199	0239	0279	0319	0359
0,1	0398	0438	0478	0517	0557	0596	0636	0675	0714	0753
0,2	0793	0832	0871	0910	0948	0987	1026	1064	1103	1141
0,3	1179	1217	1255	1293	1331	1368	1406	1443	1480	1517
0,4	1554	1591	1628	1664	1700	1736	1772	1808	1844	1879
0,5	1915	1950	1985	2019	2054	2088	2123	2157	2190	2224
0,6	2258	2291	2324	2357	2389	2422	2454	2486	2517	2549
0,7	2580	2612	2642	2673	2703	2734	2764	2794	2823	2852
0,8	2881	2910	2939	2967	2995	3023	3051	3078	3106	3133
0,9	3159	3186	3212	3238	3264	3289	3315	3340	3365	3389
1,0	3413	3438	3461	3485	3508	3531	3554	3577	3599	3621
1,1	3643	3665	3686	3708	3729	3749	3770	3790	3810	3830
1,2	3849	3869	3888	3907	3925	3944	3962	3980	3997	4015
1,3	4032	4049	4066	4082	4099	4115	4131	4147	4162	4177
1,4	4192	4207	4222	4236	4251	4265	4279	4292	4306	4319
1,5	4332	4345	4357	4370	4382	4394	4406	4419	4429	4441
1,6	4452	4463	4474	4484	4495	4505	4515	4525	4535	4545
1,7	4554	4564	4573	4582	4591	4599	4608	4616	4625	4633
1,8	4641	4649	4656	4664	4671	4678	4686	4693	4699	4706
1,9	4713	4719	4726	4732	4738	4744	4750	4756	4761	4767
2,0	4772	4778	4783	4788	4793	4798	4808	4808	4812	4817
2,1	4821	4826	4830	4834	4838	4842	4846	4850	4854	4857
2,2	4861	4864	4868	4871	4875	4878	4881	4884	4887	4890

TABEL II
NILAI-NILAI DALAM DISTRIBUSI t

α untuk uji dua fihak (two tail test)						
	0,50	0,20	0,10	0,05	0,02	0,01
α untuk uji satu fihak (one tail test)						
dk	0,25	0,10	0,005	0,025	0,01	0,005
1	1,000	3,078	6,314	12,706	31,821	63,657
2	0,816	1,886	2,920	4,303	6,965	9,925
3	0,765	1,638	2,353	3,182	4,541	5,841
4	0,741	1,533	2,132	2,776	3,747	4,604
5	0,727	1,486	2,015	2,571	3,365	4,032
6	0,718	1,440	1,943	2,447	3,143	3,707
7	0,711	1,415	1,895	2,365	2,998	3,499
8	0,706	1,397	1,860	2,306	2,896	3,355
9	0,703	1,383	1,833	2,262	2,821	3,250
10	0,700	1,372	1,812	2,228	2,764	3,165
11	0,697	1,363	1,796	2,201	2,718	3,106
12	0,695	1,356	1,782	2,178	2,681	3,055
13	0,692	1,350	1,771	2,160	2,650	3,012
14	0,691	1,345	1,761	2,145	2,624	2,977
15	0,690	1,341	1,753	2,132	2,623	2,947
16	0,689	1,337	1,746	2,120	2,583	2,921
17	0,688	1,333	1,740	2,110	2,567	2,898
18	0,688	1,330	1,743	2,101	2,552	2,878
19	0,687	1,328	1,729	2,093	2,539	2,861
20	0,687	1,325	1,725	2,086	2,528	2,845
21	0,686	1,323	1,721	2,080	2,518	2,831
22	0,686	1,321	1,717	2,074	2,508	2,819
23	0,685	1,319	1,714	2,069	2,500	2,807
24	0,685	1,318	1,711	2,064	2,492	2,797
25	0,684	1,316	1,708	2,060	2,485	2,787
26	0,684	1,315	1,706	2,056	2,479	2,779
27	0,684	1,314	1,703	2,052	2,473	2,771

28	0,683	1,313	1,701	2,048	2,467	2,763
29	0,683	1,311	1,699	2,045	2,462	2,756
30	0,683	1,310	1,697	2,042	2,457	2,750
40	0,681	1,303	1,684	2,021	2,423	2,704
60	0,679	1,296	1,671	2,000	2,390	2,660
120	0,677	1,289	1,658	1,980	2,358	2,617
∞	0,674	1,282	1,645	1,960	2,326	2,576

Formula untuk mencari skor t_{tabel} adalah:

2,05953711	=TINV(0,05;25)
------------	----------------

TABEL III
NILAI-NILAI r PRODUCT MOMENT

N	Taraf Signif		N	Taraf Signif		N	Taraf Signif	
	5%	1%		5%	1%		5%	1%
3	0,997	0,999	27	0,381	0,487	55	0,266	0,345
4	0,950	0,990	28	0,374	0,478	60	0,254	0,330
5	0,878	0,959	29	0,367	0,470	65	0,244	0,317
6	0,811	0,917	30	0,361	0,463	70	0,235	0,306
7	0,754	0,874	31	0,355	0,456	75	0,227	0,296
8	0,707	0,834	32	0,349	0,449	80	0,220	0,286
9	0,666	0,798	33	0,344	0,442	85	0,213	0,278
10	0,632	0,765	34	0,339	0,436	90	0,207	0,270
11	0,602	0,735	35	0,334	0,430	95	0,202	0,263
12	0,576	0,708	36	0,329	0,424	100	0,195	0,256
13	0,553	0,684	37	0,325	0,418	125	0,176	0,230
14	0,532	0,661	38	0,320	0,413	150	0,159	0,210
15	0,514	0,641	39	0,316	0,408	175	0,148	0,194
16	0,497	0,623	40	0,312	0,403	200	0,138	0,181
17	0,482	0,606	41	0,308	0,398	300	0,113	0,148
18	0,468	0,590	42	0,304	0,393	400	0,098	0,128
19	0,456	0,575	43	0,301	0,389	500	0,088	0,115
20	0,444	0,561	44	0,297	0,384	600	0,080	0,105
21	0,433	0,549	45	0,294	0,380	700	0,074	0,097
22	0,423	0,537	46	0,291	0,376	800	0,070	0,091
23	0,413	0,526	47	0,288	0,372	900	0,065	0,086
24	0,404	0,515	48	0,284	0,368	1000	0,062	0,081
25	0,396	0,505	49	0,281	0,364			
26	0,388	0,496	50	0,279	0,361			

TABEL IV
HARGA-HARGA x DALAM TEST BINOMIAL
(Harga-harga dalam tabel adalah 0,...)

Z																
N	0	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15
5	031	1	500	812	969											
6	016	188	344	656	981	984										
7	008	109	227	500	773	938	992									
8	004	062	145	363	637	855	965	996								
9	002	035	090	354	500	746	910	980	998							
10	001	020	055	172	377	623	828	945	989	999						
11		011	033	113	274	500	726	887	967	994						
12		006	019	073	194	387	613	806	927	981	997					
13		003	011	046	133	291	500	709	867	954	989	998				
14		002	006	029	090	212	395	605	788	910	971	994	999			
15		001	004	018	059	151	304	500	696	849	941	982	996			
16			002	011	038	105	227	402	598	773	895	962	989	998		
17			001	006	025	072	166	315	500	685	834	928	975	994	999	
18			001	004	015	048	119	240	407	593	760	881	952	985	996	999
19				002	010	032	084	180	324	500	676	820	916	968	990	998
20				001	00	021	058	132	252	412	588	748	868	942	879	994
21				001	004	013	039	095	192	332	500	668	808	905	961	987
22					002	008	026	067	143	262	416	584	738	857	933	974
23					001	005	017	047	105	202	339	500	661	798	895	953
24					001	003	011	032	076	154	271	419	581	729	846	924
25						002	007	022	054	115	212	345	500	655	788	885

TABEL V
HARGA FACTORIAL

N	N!
0	1
1	1
2	2
3	6
4	24
5	120
6	720
7	5040
8	40320
9	362880
10	3628800
11	39916800
12	479001600
13	6227020800
14	87178291200
15	1307674368000
16	20922789888000
17	355687428096000
18	6402373705728000
19	121645100408832000
20	2432902008176640000

Formula untuk mencari skor faktorial adalah:

39916800	=FACT(11)
----------	-----------

TABEL VI
NILAI-NILAI CHI KUADRAT

dk	Tarf signifikansi					
	50%	30%	20%	10%	5%,	1%
1	0,455	1,074	1,642	2,706	3,481	6,635
2	0,139	2,408	3,219	3,605	5,591	9,210
3	2,366	3,665	4,642	6,251	7,815	11,341
4	3,357	4,878	5,989	7,779	9,488	13,277
5	4,351	6,064	7,289	9,236	11 ,070	15,086
6	5,348	7,231	8,558	10,645	12,592	16,812
7	6,346	8,383	9,803	12,017	14,017	18,475
8	7,344	9,524	11,030	13,362	15,507	20,090
9	8,343	10,656	12,242	14,684	16,919	21,666
10	9,342	11,781	13,442	15,987	18,307	23,209
11	10,341	12,899	14,631	17,275	19,675	24,725
12	11,340	14,011	15,812	18,549	21,026	26,217
13	12,340	15,19	16,985	19,812	22,368	27,688
14	13,332	16,222	18,151	21 ,064	23,685	29,141
15	14,339	17,322	19,311	22,307	24,996	30,578
16	15,338	18,418	20,465	23,542	26,296	32,000
17	16,337	19,511	21,615	24,785	27,587	33,409
18	17,338	20,601	22,760	26,028	28,869	34,805
19	18,338	21,689	23,900	27,271	30,144	36,191
20	19,337	22,775	25,038	28,514	31,410	37,566
21	20,337	23,858	26,171	29,615	32,671	38,932
22	21,337	24,939	27,301	30,813	33,924	40,289
23	22,337	26,018	28,429	32,007	35,172	41,638
24	23,337	27,096	29,553	33,194	35,415	42,980
25	24,337	28,172	30,675	34,382	37,652	44,314
26	25,336	29,246	31,795	35,563	38,885	45,642
27	26,336	30,319	32,912	36,741	40,113	46,963
28	27,336	31,391	34,027	37,916	41,337	48,278
29	28,336	32,461	35,139	39,087	42,557	49,588
30	29,336	33,530	36,250	40,256	43,775	50,892

TABEL VII a
HARGA-HARGA KRITIS r DALAM TEST RUN
SATU SAMPEL, UNTUK $\alpha = 5 \%$

n1	n2																			
	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16	17	18	19	20	
2											2	2	2	2	2	2	2	2	2	
3				2	2	2	2	2	2	2	2	2	2	3	3	3	3	3	3	
4			2	2	2	3	3	3	3	3	3	3	3	3	4	4	4	4	4	
5			2	3	3	3	3	3	3	4	4	4	4	4	4	4	5	5	5	
6		2	2	3	3	3	3	4	4	4	4	5	5	5	5	5	5	6	6	
7		2	2	3	3	4	4	4	5	5	5	5	5	6	6	6	6	6	6	
8		2	3	3	3	4	4	5	5	5	6	6	6	6	6	7	7	7	7	
9		2	3	3	4	4	5	5	5	6	6	6	7	7	7	7	8	8	8	
10		2	3	3	4	5	5	5	6	6	7	7	7	7	8	8	8	8	9	
11		2	3	4	4	5	5	6	6	7	7	7	8	8	8	9	9	9	9	
12	2	2	3	4	4	5	6	6	7	7	7	8	8	8	9	9	9	10	10	
13	2	2	3	4	5	5	6	6	7	7	8	8	9	9	9	10	10	10	10	
14	2	2	3	4	5	5	6	7	7	8	8	9	9	9	10	10	10	11	11	
15	2	3	3	4	5	6	6	7	7	8	8	9	9	10	10	11	11	11	12	
16	2	3	4	4	5	6	6	7	8	8	9	9	10	10	11	11	11	12	12	
17	2	3	4	4	5	6	7	7	8	9	9	10	10	11	11	11	12	12	12	
18	2	3	4	5	5	6	7	8	8	9	9	10	10	11	11	12	12	13	13	
19	2	3	4	5	6	6	7	8	8	9	10	10	11	11	12	12	13	13	13	
20	2	3	4	5	6	6	7	8	9	9	10	10	11	12	12	13	13	13	14	

TABEL VII b
HARGA-HARGA KRITIS r DALAM TEST RUN
DUA SAMPEL, UNTUK $\alpha = 5 \%$

N1	N2																			
	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16	17	18	19	20	
2																				
3																				
4				9	9															
5			9	10	10	11	11													
6			9	10	11	12	12	13	13	13	13									
7				11	12	13	13	14	14	14	14	15	15	15						
8				11	12	13	14	14	15	15	16	16	16	16	17	17	17	17	17	
9					13	14	14	15	16	16	16	17	17	18	18	18	18	18	18	
10					13	14	15	16	16	17	17	18	18	18	19	19	19	20	20	
11					13	14	15	16	17	17	18	19	19	19	20	20	20	21	21	
12					13	14	16	16	17	18	19	19	20	20	21	21	21	22	22	
13						15	16	16	18	19	19	20	20	21	21	22	22	23	23	
14						15	16	18	18	19	20	20	21	22	22	23	23	23	24	
15						15	16	18	18	19	20	21	22	22	23	23	24	24	25	
16							17	18	19	20	21	21	22	23	23	2	25	25	25	
17							17	18	29	20	21	22	23	23	24	25	24	26	26	
18							17	18	19	20	21	22	23	24	25	25	26	26	26	
19							17	18	20	21	22	23	23	24	25	26	26	27	27	
20							17	18	20	21	22	23	24	25	25	26	27	27	28	

TABEL VIII
HARGA-HARGA KRITIS UNTUK TEST WILCOXON

N	Tingkat Signifikansi Untuk Test Satu Fihak (One Tail Test)		
	0,025	0,010	0,005
	Tingkat Signifikansi Untuk Test Dua Fihak (Two Tail Test)		
	0,05	0,02	0,01
6	0		
7	2	0	
8	4	2	0
9	6	3	2
10	8	5	3
11	11	7	5
12	14	10	7
13	17	13	10
14	21	16	13
15	25	20	16
16	30	24	20
17	35	28	23
18	40	33	28
19	46	38	32
20	52	43	38
21	59	49	43
22	66	56	49
23	73	62	55
24	81	69	61
25	89	77	68

TABEL IX
HARGA-HARGA KRITIS MAN-WHITNEY U TEST

	9	10	11	12	13	14	15	16	17	18	19	20
1												
2					0	0	0	0	0	0	1	1
3	1	1	1	2	2	2	3	3	4	4	4	5
4	3	3	4	5	5	6	7	7	8	9	9	10
5	5	6	7	8	9	10	11	12	13	14	15	16
6	7	8	9	11	12	13	15	16	18	19	20	22
7	9	11	12	14	16	17	19	21	23	24	26	28
8	11	13	15	17	20	22	24	26	28	30	32	34
9	14	16	18	21	23	26	28	31	33	36	38	40
10	16	19	22	24	27	30	33	36	38	41	44	47
11	18	22	25	28	31	34	37	41	44	47	50	53
12	21	24	28	31	35	38	42	46	49	53	56	60
13	23	27	31	35	39	43	47	51	55	59	63	67
14	26	30	34	38	43	47	51	56	60	65	69	73
15	28	33	37	42	47	51	56	61	66	70	75	80
16	31	36	41	46	51	56	61	66	71	76	82	87
17	33	38	44	49	55	60	66	71	77	82	88	93
18	36	41	47	53	59	65	70	76	82	88	94	100
19	38	44	50	56	63	69	75	82	88	94	101	107
20	40	47	53	60	67	73	80	87	93	100	107	114

TABEL X
TABEL HARGA-HARGA KRITIS DALAM TEST
KOLMOGOROV-SMIRNOV

N	One Tailed Test		Two Tailed Test	
	$\alpha = 0,05$	$\alpha = 0,01$	$\alpha = 0,05$	$\alpha = 0,01$
3	3			
4	4		4	
5	4	5	5	5
6	5	6	5	6
7	5	6	6	6
8	5	6	6	7
9	6	7	6	7
10	6	7	7	8
11	6	8	7	8
12	6	8	7	8
13	7	8	7	9
14	7	8	8	9
15	7	9	8	9
16	7	9	8	10
17	8	9	8	10
18	8	10	9	10
19	8	10	9	10
20	8	10	9	11
21	8	10	9	11
22	9	11	9	11
23	9	11	10	11
24	9	11	10	12
25	9	11	10	12
26	9	11	10	12
27	9	12	10	12
28	10	12	11	13
29	10	12	11	13
30	10	12	11	13
35	11	13	12	
40	11	14	13	

TABEL XI
HARGA-HARGA z UNTUK TEST RUN WALD-WOLFOWITZ

z	0,00	0,01	0,02	0,03	0,04	0,05	0,06	0,07	0,08	0,09
0,0	0,5000	0,4960	0,4920	0,4880	0,4840	0,4801	0,4761	0,4721	0,4681	0,4641
0,1	0,4602	0,4562	0,4522	0,4483	0,4443	0,4404	0,4364	0,4325	0,4286	0,4247
0,2	0,4207	0,4168	0,4129	0,4090	0,4052	0,4013	0,3974	0,3936	0,3897	0,3859
0,3	0,3821	0,3783	0,3745	0,3707	0,3669	0,3632	0,3594	0,3557	0,3520	0,3483
0,4	0,3446	0,3409	0,3372	0,3336	0,3300	0,3264	0,3228	0,3192	0,3156	0,3121
0,5	0,3086	0,3050	0,3015	0,2981	0,2946	0,2912	0,2877	0,2843	0,2810	0,2776
0,6	0,2743	0,2709	0,2676	0,2643	0,2611	0,2578	0,2546	0,2514	0,2483	0,2451
0,7	0,2420	0,2389	0,2358	0,2327	0,2297	0,2266	0,2236	0,2206	0,2177	0,2148
0,8	0,2119	0,2090	0,2061	0,2033	0,2005	0,1977	0,1949	0,1922	0,1894	0,1867
0,9	0,1841	0,1814	0,1788	0,1762	0,1736	0,1711	0,1685	0,1660	0,1635	0,1611
1,0	0,1587	0,1562	0,1535	0,1515	0,1492	0,1469	0,1446	0,1423	0,1401	0,1379
1,1	0,1357	0,1335	0,1314	0,1292	0,1271	0,1251	0,1230	0,1210	0,1190	0,1170
1,2	0,1151	0,1131	0,1112	0,1093	0,1075	0,1056	0,1038	0,1020	0,1003	0,0985
1,3	0,0968	0,0951	0,0934	0,0918	0,0901	0,0885	0,0869	0,0853	0,0838	0,0823
1,4	0,0808	0,0793	0,0778	0,0764	0,0749	0,0735	0,0721	0,0708	0,0694	0,0681
1,5	0,0668	0,0655	0,0643	0,0630	0,0618	0,0606	0,0594	0,0581	0,0571	0,0559
1,6	0,0548	0,0537	0,0526	0,0516	0,0505	0,0495	0,0485	0,0475	0,0465	0,0455
1,7	0,0445	0,0436	0,0427	0,0418	0,0409	0,0401	0,0392	0,0384	0,0375	0,0367
1,8	0,0359	0,0351	0,0344	0,0336	0,0329	0,0322	0,0314	0,0307	0,0301	0,0294
1,9	0,0287	0,0281	0,0274	0,0268	0,0262	0,0256	0,0250	0,0244	0,0239	0,0233
2,0	0,0228	0,0222	0,0217	0,0212	0,0207	0,0202	0,0197	0,0192	0,0188	0,0183
2,1	0,0179	0,0174	0,0170	0,0166	0,0162	0,0158	0,0154	0,0150	0,0146	0,0143
2,2	0,0139	0,0136	0,0132	0,0129	0,0125	0,0122	0,0119	0,0116	0,0113	0,0110
2,3	0,0107	0,0104	0,0102	0,0099	0,0096	0,0094	0,0091	0,0089	0,0087	0,0084
2,4	0,0082	0,0080	0,0078	0,0075	0,0073	0,0071	0,0069	0,0068	0,0066	0,0064
2,5	0,0062	0,0060	0,0059	0,0057	0,0055	0,0054	0,0052	0,0051	0,0049	0,0048
2,5	0,0047	0,0045	0,0044	0,0043	0,0041	0,0040	0,0039	0,0038	0,0037	0,0036
2,7	0,0035	0,0034	0,0033	0,0032	0,0031	0,0030	0,0029	0,0028	0,0027	0,0026
2,8	0,0026	0,0025	0,0024	0,0023	0,0023	0,0022	0,0021	0,0021	0,0020	0,0019
2,9	0,0019	0,0018	0,0013	0,0017	0,0016	0,0016	0,0015	0,0015	0,0014	0,0014

TABEL XII
NILAI-NILAI UNTUK DISTRIBUSI F

Baris atas untuk 5%
Baris bawah untuk 1%

$v_1 = dk$ penyebut	$v_2 = dk$ pembilang																			
	1	2	3	4	5	6	7	8	9	10	11	12	14	16	20	24	30	40	50	75
1	181	200	216	225	230	234	237	239	241	242	243	244	245	246	248	249	250	251	252	253
2	4.052	4.899	5.403	5.625	5.764	5.859	5.928	5.981	6.022	6.056	6.082	6.108	6.142	6.169	6.208	6.234	6.258	6.286	6.302	6.323
3	18.51	19.00	19.16	19.25	19.30	19.33	19.36	19.37	19.38	19.39	19.40	19.41	19.42	19.43	19.44	19.45	19.46	19.47	19.47	19.48
4	88.49	90.01	90.17	90.25	90.30	90.33	90.34	90.36	90.38	90.40	90.41	90.42	90.43	90.44	90.45	90.46	90.47	90.48	90.48	90.50
5	10.13	9.55	9.28	9.12	9.01	8.94	8.88	8.84	8.81	8.78	8.76	8.74	8.71	8.69	8.66	8.64	8.62	8.60	8.58	8.57
6	34.12	30.81	29.48	28.71	28.24	27.91	27.67	27.49	27.34	27.23	27.13	27.05	26.92	26.83	26.69	26.60	26.50	26.41	26.30	26.27
7	7.71	6.94	6.59	6.39	6.28	6.18	6.09	6.04	6.00	5.96	5.93	5.91	5.87	5.84	5.80	5.77	5.74	5.71	5.70	5.68
8	21.20	18.00	16.69	15.88	15.52	15.21	14.98	14.80	14.68	14.54	14.45	14.37	14.24	14.15	14.02	13.93	13.83	13.74	13.69	13.61
9	6.61	5.79	5.41	5.19	5.06	4.96	4.88	4.82	4.78	4.74	4.70	4.68	4.64	4.60	4.56	4.53	4.50	4.48	4.44	4.42
10	18.28	13.27	12.06	11.39	10.97	10.67	10.45	10.27	10.15	10.05	9.96	9.89	9.77	9.68	9.55	9.47	9.38	9.29	9.24	9.17
11	5.90	5.14	4.78	4.53	4.39	4.28	4.21	4.15	4.10	4.06	4.03	4.00	3.98	3.92	3.87	3.84	3.81	3.77	3.76	3.72
12	13.74	10.92	9.78	9.15	8.76	8.47	8.26	8.10	7.98	7.87	7.78	7.72	7.60	7.52	7.39	7.31	7.23	7.14	7.08	7.02
13	12.26	8.55	8.45	7.95	7.46	7.19	7.00	6.84	6.71	6.62	6.54	6.47	6.35	6.27	6.15	6.07	5.98	5.90	5.85	5.78
14	5.32	4.46	4.07	3.84	3.60	3.50	3.44	3.39	3.34	3.31	3.28	3.23	3.23	3.20	3.15	3.12	3.08	3.05	3.03	3.00
15	11.26	8.65	7.69	7.01	6.63	6.37	6.19	6.03	5.91	5.82	5.74	5.67	5.58	5.48	5.38	5.28	5.20	5.11	5.06	5.00
16	5.12	4.26	3.88	3.63	3.48	3.37	3.29	3.23	3.18	3.13	3.10	3.07	3.02	2.98	2.93	2.90	2.86	2.82	2.80	2.77
17	10.58	8.02	6.90	6.42	6.06	5.80	5.62	5.47	5.36	5.28	5.18	5.11	5.00	4.92	4.80	4.73	4.64	4.56	4.51	4.45
18	4.94	3.98	3.59	3.36	3.20	3.09	3.01	2.96	2.90	2.86	2.82	2.79	2.74	2.70	2.66	2.61	2.57	2.53	2.50	2.47
19	8.65	7.20	6.22	5.67	5.32	5.07	4.88	4.74	4.63	4.54	4.46	4.40	4.29	4.21	4.10	4.02	3.94	3.88	3.80	3.74
20	4.84	3.98	3.59	3.36	3.20	3.09	3.01	2.96	2.90	2.86	2.82	2.79	2.74	2.70	2.66	2.61	2.57	2.53	2.50	2.47
21	4.75	3.88	3.49	3.26	3.11	3.00	2.92	2.85	2.80	2.76	2.72	2.69	2.64	2.60	2.54	2.50	2.46	2.42	2.40	2.38
22	4.67	3.80	3.41	3.18	3.02	2.92	2.84	2.77	2.72	2.68	2.63	2.60	2.55	2.51	2.46	2.42	2.38	2.34	2.32	2.28
23	4.60	3.74	3.34	3.11	2.96	2.85	2.77	2.70	2.65	2.60	2.56	2.53	2.48	2.44	2.39	2.36	2.31	2.27	2.24	2.21
24	4.48	3.61	3.21	2.98	2.82	2.71	2.63	2.56	2.51	2.47	2.43	2.40	2.36	2.32	2.28	2.24	2.20	2.16	2.13	2.10
25	4.40	3.53	3.13	2.90	2.74	2.63	2.55	2.48	2.43	2.39	2.35	2.32	2.28	2.24	2.20	2.16	2.12	2.08	2.04	2.00
26	4.32	3.45	3.05	2.82	2.66	2.55	2.47	2.40	2.35	2.31	2.27	2.24	2.20	2.16	2.12	2.08	2.04	2.00	1.96	1.92
27	4.24	3.37	2.97	2.74	2.58	2.47	2.39	2.32	2.27	2.23	2.19	2.16	2.12	2.08	2.04	2.00	1.96	1.92	1.88	1.84
28	4.16	3.29	2.89	2.66	2.50	2.39	2.31	2.24	2.19	2.15	2.11	2.08	2.04	2.00	1.96	1.92	1.88	1.84	1.80	1.76
29	4.08	3.21	2.81	2.58	2.42	2.31	2.23	2.16	2.11	2.07	2.03	2.00	1.96	1.92	1.88	1.84	1.80	1.76	1.72	1.68
30	4.00	3.13	2.73	2.50	2.34	2.23	2.15	2.08	2.03	1.99	1.95	1.92	1.88	1.84	1.80	1.76	1.72	1.68	1.64	1.60
31	3.92	3.05	2.65	2.42	2.26	2.15	2.07	2.00	1.95	1.91	1.87	1.84	1.80	1.76	1.72	1.68	1.64	1.60	1.56	1.52
32	3.84	2.97	2.57	2.34	2.18	2.07	1.99	1.92	1.87	1.83	1.79	1.76	1.72	1.68	1.64	1.60	1.56	1.52	1.48	1.44
33	3.76	2.89	2.49	2.26	2.10	1.99	1.91	1.84	1.79	1.75	1.71	1.68	1.64	1.60	1.56	1.52	1.48	1.44	1.40	1.36
34	3.68	2.81	2.41	2.18	2.02	1.91	1.83	1.76	1.71	1.67	1.63	1.60	1.56	1.52	1.48	1.44	1.40	1.36	1.32	1.28
35	3.60	2.73	2.33	2.10	1.94	1.83	1.75	1.68	1.63	1.59	1.55	1.52	1.48	1.44	1.40	1.36	1.32	1.28	1.24	1.20
36	3.52	2.65	2.25	2.02	1.86	1.75	1.67	1.60	1.55	1.51	1.47	1.44	1.40	1.36	1.32	1.28	1.24	1.20	1.16	1.12
37	3.44	2.57	2.17	1.94	1.78	1.67	1.59	1.52	1.47	1.43	1.39	1.36	1.32	1.28	1.24	1.20	1.16	1.12	1.08	1.04
38	3.36	2.49	2.09	1.86	1.70	1.59	1.51	1.44	1.39	1.35	1.31	1.28	1.24	1.20	1.16	1.12	1.08	1.04	1.00	0.96
39	3.28	2.41	2.01	1.78	1.62	1.51	1.43	1.36	1.31	1.27	1.23	1.20	1.16	1.12	1.08	1.04	1.00	0.96	0.92	0.88
40	3.20	2.33	1.93	1.70	1.54	1.43	1.35	1.28	1.23	1.19	1.15	1.12	1.08	1.04	1.00	0.96	0.92	0.88	0.84	0.80
41	3.12	2.25	1.85	1.62	1.46	1.35	1.27	1.20	1.15	1.11	1.07	1.04	1.00	0.96	0.92	0.88	0.84	0.80	0.76	0.72
42	3.04	2.17	1.77	1.54	1.38	1.27	1.19	1.12	1.07	1.03	0.99	0.96	0.92	0.88	0.84	0.80	0.76	0.72	0.68	0.64
43	2.96	2.09	1.69	1.46	1.30	1.19	1.11	1.04	0.99	0.95	0.91	0.88	0.84	0.80	0.76	0.72	0.68	0.64	0.60	0.56
44	2.88	2.01	1.61	1.38	1.22	1.11	1.03	0.96	0.91	0.87	0.83	0.80	0.76	0.72	0.68	0.64	0.60	0.56	0.52	0.48
45	2.80	1.93	1.53	1.30	1.14	1.03	0.95	0.88	0.83	0.79	0.75	0.72	0.68	0.64	0.60	0.56	0.52	0.48	0.44	0.40
46	2.72	1.85	1.45	1.22	1.06	0.95	0.87	0.80	0.75	0.71	0.67	0.64	0.60	0.56	0.52	0.48	0.44	0.40	0.36	0.32
47	2.64	1.77	1.37	1.14	0.98	0.87	0.79	0.72	0.67	0.63	0.59	0.56	0.52	0.48	0.44	0.40	0.36	0.32	0.28	0.24
48	2.56	1.69	1.29	1.06	0.90	0.79	0.71	0.64	0.59	0.55	0.51	0.48	0.44	0.40	0.36	0.32	0.28	0.24	0.20	0.16
49	2.48	1.61	1.21	0.98	0.82	0.71	0.63	0.56	0.51	0.47	0.43	0.40	0.36	0.32	0.28	0.24	0.20	0.16	0.12	0.08
50	2.40	1.53	1.13	0.90	0.74	0.63	0.55	0.48	0.43	0.39	0.35	0.32	0.28	0.24	0.20	0.16	0.12	0.08	0.04	0.00

$v_1 = dk$ penyebut	$v_1 = dk$ pembilang																			
	1	2	3	4	5	6	7	8	9	10	11	12	14	16	20	24	30	40	60	75
36	4.11	3.26	2.80	2.53	2.48	2.36	2.28	2.21	2.15	2.10	2.06	2.03	1.89	1.93	1.87	1.82	1.78	1.72	1.69	1.65
38	7.39	5.25	4.38	3.68	3.58	3.35	3.18	3.04	2.94	2.86	2.78	2.72	2.62	2.54	2.43	2.35	2.26	2.17	2.12	2.04
40	4.10	3.25	2.85	2.62	2.46	2.35	2.26	2.18	2.14	2.08	2.05	2.02	1.96	1.92	1.85	1.80	1.76	1.71	1.67	1.63
42	7.35	5.21	4.34	3.65	3.54	3.32	3.15	3.02	2.91	2.82	2.75	2.68	2.58	2.51	2.40	2.32	2.22	2.14	2.08	2.00
44	4.08	3.23	2.84	2.61	2.45	2.34	2.25	2.18	2.12	2.07	2.04	2.00	1.95	1.90	1.84	1.79	1.74	1.69	1.65	1.61
46	7.31	5.18	4.31	3.63	3.51	3.29	3.12	2.99	2.88	2.80	2.73	2.66	2.56	2.49	2.37	2.29	2.20	2.11	2.05	1.97
48	4.07	3.22	2.83	2.59	2.44	2.32	2.24	2.17	2.11	2.06	2.03	1.99	1.94	1.89	1.82	1.78	1.73	1.68	1.64	1.60
50	7.24	5.12	4.25	3.57	3.46	3.24	3.07	2.94	2.84	2.75	2.68	2.62	2.52	2.44	2.32	2.24	2.15	2.06	2.00	1.92
55	4.06	3.21	2.82	2.58	2.43	2.31	2.23	2.16	2.10	2.05	2.01	1.98	1.92	1.87	1.81	1.76	1.72	1.66	1.63	1.58
60	7.21	5.10	4.24	3.56	3.44	3.22	3.05	2.92	2.82	2.73	2.66	2.60	2.50	2.42	2.30	2.22	2.13	2.04	1.96	1.88
65	4.04	3.19	2.80	2.56	2.41	2.30	2.21	2.14	2.08	2.03	1.99	1.95	1.90	1.85	1.78	1.74	1.70	1.64	1.61	1.56
70	7.19	5.08	4.22	3.74	3.42	3.20	3.04	2.90	2.80	2.71	2.64	2.56	2.46	2.38	2.28	2.20	2.11	2.02	1.96	1.88
75	4.03	3.18	2.79	2.55	2.40	2.29	2.20	2.13	2.07	2.02	1.98	1.95	1.90	1.85	1.78	1.71	1.69	1.63	1.60	1.55
80	7.17	5.06	4.20	3.72	3.41	3.18	3.02	2.88	2.78	2.70	2.62	2.56	2.46	2.38	2.28	2.18	2.10	2.00	1.91	1.86
85	4.02	3.17	2.78	2.54	2.39	2.27	2.18	2.11	2.05	2.00	1.97	1.93	1.88	1.83	1.76	1.72	1.67	1.61	1.58	1.52
90	7.12	5.01	4.15	3.68	3.37	3.15	2.98	2.83	2.75	2.66	2.59	2.53	2.43	2.35	2.23	2.15	2.00	1.96	1.90	1.82
95	4.00	3.15	2.76	2.52	2.37	2.23	2.17	2.10	2.01	1.99	1.95	1.92	1.86	1.81	1.75	1.70	1.63	1.59	1.56	1.50
100	7.08	4.98	4.13	3.65	3.31	3.12	2.95	2.82	2.72	2.63	2.56	2.50	2.40	2.32	2.20	2.12	2.03	1.93	1.87	1.79
105	4.00	3.14	2.75	2.51	2.36	2.24	2.15	2.08	2.02	1.98	1.94	1.90	1.85	1.80	1.74	1.68	1.63	1.57	1.54	1.49
110	7.04	4.95	4.10	3.62	3.34	3.09	2.93	2.78	2.70	2.61	2.54	2.47	2.37	2.30	2.18	2.09	2.00	1.90	1.84	1.76
115	4.00	3.15	2.76	2.52	2.37	2.23	2.17	2.10	2.01	1.99	1.95	1.92	1.86	1.81	1.75	1.70	1.63	1.59	1.56	1.50
120	7.01	4.92	4.08	3.60	3.28	3.07	2.91	2.77	2.67	2.59	2.51	2.45	2.35	2.28	2.15	2.07	1.98	1.88	1.82	1.74
125	4.00	3.14	2.75	2.51	2.36	2.24	2.15	2.08	2.02	1.98	1.94	1.90	1.84	1.79	1.72	1.67	1.62	1.56	1.54	1.47
130	7.00	4.93	4.10	3.62	3.34	3.09	2.93	2.78	2.70	2.61	2.54	2.47	2.37	2.30	2.18	2.09	2.00	1.90	1.84	1.76
135	4.00	3.15	2.76	2.52	2.37	2.23	2.17	2.10	2.01	1.99	1.95	1.92	1.86	1.81	1.75	1.70	1.63	1.59	1.56	1.50
140	7.00	4.93	4.10	3.62	3.34	3.09	2.93	2.78	2.70	2.61	2.54	2.47	2.37	2.30	2.18	2.09	2.00	1.90	1.84	1.76
145	4.00	3.15	2.76	2.52	2.37	2.23	2.17	2.10	2.01	1.99	1.95	1.92	1.86	1.81	1.75	1.70	1.63	1.59	1.56	1.50
150	7.00	4.93	4.10	3.62	3.34	3.09	2.93	2.78	2.70	2.61	2.54	2.47	2.37	2.30	2.18	2.09	2.00	1.90	1.84	1.76
155	4.00	3.15	2.76	2.52	2.37	2.23	2.17	2.10	2.01	1.99	1.95	1.92	1.86	1.81	1.75	1.70	1.63	1.59	1.56	1.50
160	7.00	4.93	4.10	3.62	3.34	3.09	2.93	2.78	2.70	2.61	2.54	2.47	2.37	2.30	2.18	2.09	2.00	1.90	1.84	1.76
165	4.00	3.15	2.76	2.52	2.37	2.23	2.17	2.10	2.01	1.99	1.95	1.92	1.86	1.81	1.75	1.70	1.63	1.59	1.56	1.50
170	7.00	4.93	4.10	3.62	3.34	3.09	2.93	2.78	2.70	2.61	2.54	2.47	2.37	2.30	2.18	2.09	2.00	1.90	1.84	1.76
175	4.00	3.15	2.76	2.52	2.37	2.23	2.17	2.10	2.01	1.99	1.95	1.92	1.86	1.81	1.75	1.70	1.63	1.59	1.56	1.50
180	7.00	4.93	4.10	3.62	3.34	3.09	2.93	2.78	2.70	2.61	2.54	2.47	2.37	2.30	2.18	2.09	2.00	1.90	1.84	1.76
185	4.00	3.15	2.76	2.52	2.37	2.23	2.17	2.10	2.01	1.99	1.95	1.92	1.86	1.81	1.75	1.70	1.63	1.59	1.56	1.50
190	7.00	4.93	4.10	3.62	3.34	3.09	2.93	2.78	2.70	2.61	2.54	2.47	2.37	2.30	2.18	2.09	2.00	1.90	1.84	1.76
195	4.00	3.15	2.76	2.52	2.37	2.23	2.17	2.10	2.01	1.99	1.95	1.92	1.86	1.81	1.75	1.70	1.63	1.59	1.56	1.50
200	7.00	4.93	4.10	3.62	3.34	3.09	2.93	2.78	2.70	2.61	2.54	2.47	2.37	2.30	2.18	2.09	2.00	1.90	1.84	1.76
205	4.00	3.15	2.76	2.52	2.37	2.23	2.17	2.10	2.01	1.99	1.95	1.92	1.86	1.81	1.75	1.70	1.63	1.59	1.56	1.50
210	7.00	4.93	4.10	3.62	3.34	3.09	2.93	2.78	2.70	2.61	2.54	2.47	2.37	2.30	2.18	2.09	2.00	1.90	1.84	1.76
215	4.00	3.15	2.76	2.52	2.37	2.23	2.17	2.10	2.01	1.99	1.95	1.92	1.86	1.81	1.75	1.70	1.63	1.59	1.56	1.50
220	7.00	4.93	4.10	3.62	3.34	3.09	2.93	2.78	2.70	2.61	2.54	2.47	2.37	2.30	2.18	2.09	2.00	1.90	1.84	1.76
225	4.00	3.15	2.76	2.52	2.37	2.23	2.17	2.10	2.01	1.99	1.95	1.92	1.86	1.81	1.75	1.70	1.63	1.59	1.56	1.50
230	7.00	4.93	4.10	3.62	3.34	3.09	2.93	2.78	2.70	2.61	2.54	2.47	2.37	2.30	2.18	2.09	2.00	1.90	1.84	1.76
235	4.00	3.15	2.76	2.52	2.37	2.23	2.17	2.10	2.01	1.99	1.95	1.92	1.86	1.81	1.75	1.70	1.63	1.59	1.56	1.50
240	7.00	4.93	4.10	3.62	3.34	3.09	2.93	2.78	2.70	2.61	2.54	2.47	2.37	2.30	2.18	2.09	2.00	1.90	1.84	1.76
245	4.00	3.15	2.76	2.52	2.37	2.23	2.17	2.10	2.01	1.99	1.95	1.92	1.86	1.81	1.75	1.70	1.63	1.59	1.56	1.50
250	7.00	4.93	4.10	3.62	3.34	3.09	2.93	2.78	2.70	2.61	2.54	2.47	2.37	2.30	2.18	2.09	2.00	1.90	1.84	1.76
255	4.00	3.15	2.76	2.52	2.37	2.23	2.17	2.10	2.01	1.99	1.95	1.92	1.86	1.81	1.75	1.70	1.63	1.59	1.56	1.50
260	7.00	4.93	4.10	3.62	3.34	3.09	2.93	2.78	2.70	2.61	2.54	2.47	2.37	2.30	2.18	2.09	2.00	1.90	1.84	1.76
265	4.00	3.15	2.76	2.52	2.37	2.23	2.17	2.10	2.01	1.99	1.95	1.92	1.86	1.81	1.75	1.70	1.63	1.59	1.56	1.50
270	7.00	4.93	4.10	3.62	3.34	3.09	2.93	2.78	2.70	2.61	2.54	2.47	2.37	2.30	2.18	2.09	2.00	1.90	1.84	1.76
275	4.00	3.15	2.76	2.52	2.37	2.23	2.17	2.10	2.01	1.99	1.95	1.92	1.86	1.81	1.75	1.70	1.63	1.59	1.56	1.50
280	7.00	4.93	4.10	3.62	3.34	3.09	2.93	2.78	2.70	2.61	2.54	2.47	2.37	2.30	2.18	2.09	2.00	1.90	1.84	1.76
285	4.00	3.15	2.76	2.52	2.37	2.23	2.17	2.10	2.01	1.99	1.95	1.92	1.86	1.81	1.75	1.70	1.63	1.59	1.56	1.50
290	7.00	4.93	4.10	3.62	3.34	3.09	2.93	2.78	2.70	2.61	2.54	2.47	2.37	2.30	2.18	2.09	2.00	1.90	1.84	1.76
295	4.00	3.15	2.76	2.52	2.37	2.23	2.17	2.10	2.01	1.99	1.95	1.92	1.86	1.81	1.75	1.70	1.63	1.59	1.56	1.50
300	7.00	4.93	4.10	3.62	3.34	3.09	2.93	2.78	2.70	2.61	2.54	2.47	2.37	2.30	2.18	2.09	2.00	1.90	1.84	1.76
305	4.00	3.15	2.76	2.52	2.37	2.23	2.17	2.10	2.01	1.99	1.95	1.92	1.86	1.81	1.75	1.70	1.63	1.59	1.56	1.50
310	7.00	4.93	4.10	3.62	3.34	3.09	2.93	2.78	2.70	2.61	2.54	2.47	2.37	2.30	2.18	2.09	2.00	1.90	1.84	1.76
315	4.00	3.15	2.76	2.52	2.37	2.23	2.17	2.10	2.01	1.99	1.95	1.92	1.86	1.81	1.75	1.70	1.63	1.59	1.56	1.50
320	7.00	4.93	4.10	3.62	3.34	3.09	2.93	2.78	2.70	2.61	2.54									

TABEL XIII
TABEL NILAI-NILAI RHO

N	Taraf	Signif	N	Taraf	Signif
	5%	1%		5%	1%
5	1,000		16	0,506	0,665
6	0,886	1,000	18	0,475	0,626
7	0,786	0,929	20	0,450	0,591
8	0,738	0,881	22	0,428	0,562
9	0,683	0,833	24	0,409	0,537
10	0,648	0,794	26	0,392	0,515
12	0,591	0,777	28	0,377	0,496
14	0,544	0,715	30	0,364	0,478

TABEL XIV
TABEL HARGA-HARGA KRITIS Z DALAM OBSERVASI
DISTRIBUSI NORMAL

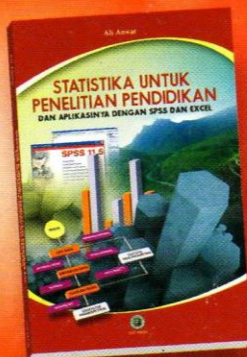
Z	0,00	0,01	0,02	0,03	0,04	0,05	0,06	0,07	0,08	0,09
0,0	,5000	,4960	,4920	,4880	,4840	,4801	,4761	,4721	,4681	,4641
0,1	,4602	,4562	,4522	,4483	,4443	,4404	,4364	,4325	,4286	,4247
0,2	,4207	,4168	,4129	,4090	,4052	,4013	,3974	,3936	,3897	,3859
0,3	,3821	,3783	,3745	,3707	,3669	,3632	,3594	,3557	,3520	,3483
0,4	,3446	,3409	,3372	,3336	,3300	,3264	,3228	,3192	,3156	,3121
0,5	,3085	,3050	,3015	,2981	,2946	,2912	,2877	,2843	,2810	,2776
0,6	,2743	,2709	,2676	,2643	,2611	,2578	,2546	,2514	,2483	,2451
0,7	,2420	,2389	,2358	,2327	,2296	,2266	,2236	,2206	,2177	,2148
0,8	,2119	,2090	,2061	,2033	,2005	,1977	,1949	,1922	,1894	,1867
0,9	,1841	,1814	,1788	,1762	,1736	,1711	,1685	,1660	,1635	,1611
1,0	,1587	,1562	,1539	,1515	,1492	,1469	,1446	,1423	,1401	,1379
1,1	,1357	,1335	,1314	,1292	,1271	,1251	,1230	,1210	,1190	,1170
1,2	,1151	,1131	,1112	,1093	,1075	,1056	,1038	,1020	,1003	,0985
1,3	,0968	,0951	,0934	,0918	,0901	,0885	,0869	,0853	,0838	,0823
1,4	,0808	,0793	,0778	,0764	,0749	,1735	,0721	,0708	,0694	,0681
1,5	,0668	,0655	,0643	,0630	,0618	,0606	,0594	,0582	,0571	,0559
1,6	,0548	,0537	,0526	,0516	,0505	,0495	,0485	,0475	,0465	,0455
1,7	,0446	,0436	,0427	,0418	,0409	,0410	,0392	,0384	,0375	,0367
1,8	,0359	,0351	,0344	,0336	,0329	,0322	,0314	,0307	,0301	,0294
1,9	,0287	,0281	,0274	,0268	,0262	,0256	,0250	,0244	,0239	,0233
2,0	,0228	,0222	,0217	,0212	,0207	,0202	,0197	,0192	,0188	,0183
2,1	,0179	,0174	,0170	,0166	,0162	,0158	,0154	,0150	,0146	,0143
2,2	,0139	,0136	,0132	,0129	,0125	,0122	,0119	,0116	,0113	,0110
2,3	,0107	,0104	,0102	,0099	,0096	,0094	,0091	,0089	,0087	,0084
2,4	,0082	,0080	,0078	,0075	,0073	,0071	,0069	,0068	,0066	,0064
2,5	,0062	,0060	,0059	,0057	,0055	,0054	,0052	,0051	,0049	,0048
2,6	,0047	,0045	,0044	,0043	,0041	,0040	,0039	,0038	,0037	,0036
2,7	,0035	,0034	,0033	,0032	,0031	,0030	,0029	,0028	,0027	,0026
2,8	,0026	,0025	,0024	,0023	,0023	,0022	,0021	,0021	,0020	,0019
2,9	,0019	,0013	,0018	,0017	,0016	,0016	,0015	,0015	,0014	,0014

**CARA MENGHITUNG NILAI TABEL: t , F , χ^2 , DAN z
DENGAN APLIKASI MICROSOFT EXCEL**

- 01 Untuk mencari nilai tabel t adalah: `TINV(Probability,Deg_freedom)`
Contoh:
Diketahui pada sebuah penelitian bahwa degree of freedomnya =276, pada $\alpha = 5\%$, maka Tabel t -nya adalah: 1,968596735
=`TINV(0,05;276)`
- 02 Untuk mencari nilai tabel F adalah: `FINV(Probability,Deg_freedom1,Deg_freedom2)`
Contoh:
Diketahui pada sebuah penelitian bahwa degree of freedom1 =103, degree of freedom2 = 74 pada $\alpha = 5\%$, maka Tabel F -nya adalah: 1,328514188
=`FINV(0,05;103;174)`
- 03 Untuk mencari nilai tabel χ^2 adalah: `CHIINV(Probability,Deg_freedom)`
Contoh:
Diketahui pada sebuah penelitian bahwa degree of freedomnya =3, pada $\alpha = 5\%$, maka Tabel χ^2 nya adalah: 7,814724703
=`CHIINV(0,05;3)`
- 04 Untuk mencari nilai tabel z adalah: `NORMADIST(z)`
Contoh:
Diketahui pada sebuah penelitian bahwa skor $z = 0,47$, maka Tabel z -nya adalah: 0,680822481
=`NORMSDIST(0,47)`

CATATAN

[illegible]



STATISTIKA UNTUK PENELITIAN PENDIDIKAN

DAN APLIKASINYA DENGAN SPSS DAN EXCEL

Buku ini merupakan buku statistik terapan yang menyajikan aplikasi dalam menguji instrumen penelitian, menentukan besarnya sampel, mendeskripsikan data dan menganalisisnya, serta menguji hipotesis dengan menggunakan SPSS dan Excel. Sebagai buku statistik terapan buku ini tidak menjelaskan bagaimana rumus-rumus statistik itu disusun dan dihasilkan.

Selama ini, jika orang mendengar kata statistik, maka asosiasi mereka adalah tentang sesuatu yang ruwet, memusingkan, penuh dengan rumus-rumus yang rumit, membosankan, dan sebagainya. Mahasiswa yang menulis skripsi dengan pendekatan kuantitatif ternyata sering menemukan kesulitan untuk menentukan analisis statistik yang tepat dan mengaplikasikan analisis yang telah dipilih.

Dengan semakin meluasnya ketersediaan perangkat keras dan lunak komputer, ia telah memberikan banyak kemudahan kepada peneliti dalam menganalisis hasil penelitian dengan statistik. Bila dahulu ada anggapan bahwa penelitian yang meneliti tiga variabel atau lebih hanya dapat dilakukan oleh peneliti setingkat mahasiswa S2 atau bahkan S3, dikarenakan sulitnya analisis multivariat, sekarang, hal tersebut dapat dilakukan oleh mahasiswa S1 atau oleh peneliti pemula. Buku ini berusaha menunjukkan kemudahan statistik dengan memberikan contoh dan aplikasinya dengan SPSS dan Microsoft Excel.



Penulis buku ini, Ali Anwar, lahir di Demak, Jawa Tengah, pada tanggal 3 Mei 1964, adalah Direktur dan Dosen Program Pascasarjana IAIT Kediri, Lektor Kepala dalam Sejarah Pendidikan Islam di Jurusan Tarbiyah STAIN Kediri, Wakil Ketua Asosiasi Peneliti Sosial Keagamaan Indonesia (APSKI), dan Ketua Asosiasi Analisis Data Kuantitatif STAIN Kediri. Penulis sekarang bertempat tinggal di Jl. Sunan Ampel RT 01 RW 02 No. 18 Rejomulyo Kota Kediri, Kode Pos: 64129. E-mail: ali_anwar03@yahoo.co.id

Sejak 29 Agustus 2008, mantan Pembantu Rektor I IAIT Kediri periode 2000-2004, mantan Asisten Direktur I Program Pascasarjana IAIT Kediri periode 2004-2008, dan alumnus program doktor UIN Jakarta yang mendapatkan predikat Cumlaude/Terpuji ini berjuang untuk dapat lulus dari ujian dan cobaan yang diberikan oleh sahabatnya yang sedang berkuasa berupa berbagai hambatan untuk dapat mengajukan kenaikan jabatan fungsionalnya menjadi Guru Besar.

Buku ini adalah buku keempat penulis yang ber-ISBN. Yang pertama, "Eksistensi Pendidikan Tradisional di Tengah Arus Modernisasi Pendidikan: Studi terhadap Kelangsungan Madrasah Hidayatul Mubtadi'in di Pondok Pesantren Lirboyo Kediri", dalam Irwan Abdullah dkk. (Ed.), *Agama, Pendidikan Islam, dan Tanggung Jawab Sosial Pesantren*, (Yogyakarta: Sekolah Pascasarjana UGM dan Pustaka Pelajar, 2008). Kedua, *Pembaruan Pendidikan di Pesantren Lirboyo Kediri*, (Kediri: IAIT Press, 2008). Dan ketiga, *Cara Mudah Menulis Karya Ilmiah*, (Kediri: IAIT Press, 2009). Sejak 2007, penulis buku ini telah menerbitkan 4 (empat) hasil penelitiannya pada berbagai jurnal terakreditasi nasional dan 3 (tiga) hasil penelitian lainnya pada Jurnal at-Tarbawi STAIN Surakarta, Realita, Empirisma STAIN Kediri.



Diterbitkan oleh: **IAIT Press**
Program Pascasarjana IAIT Kediri
Jl. KH. Wahid Hasyim 62 Kediri 64114
Telp/Faks.: (0354) 777239, 772879
E-mail: iait.press@yahoo.co.id

ISBN 978-979-18633-2-2

