



Statistik dalam Olahraga

Ali Maksum

Penulis buku *bestcitation* “Metodologi Penelitian dalam Olahraga”

@ 2018, unesa university press
ISBN: 978-602-449-123-9

Kata Pengantar

Apa yang tersaji dihadapan pembaca sejatinya merupakan ijihad akademik dalam rangka menyediakan bacaan dan sekaligus instrumen bagi para mahasiswa, termasuk peneliti muda, untuk memahami, mengolah, dan menyajikan data penelitian secara akurat dan terpercaya. Sebenarnya telah banyak buku statistik yang ditulis para ahli, hanya saja, buku-buku tersebut umumnya menggunakan perhitungan angka-angka yang terkesan “canggih” dan tidak mengaitkan langsung dengan disiplin olahraga. Akibatnya, mahasiswa tampak mengalami kesulitan ketika harus mentransformasikan ke dalam “dunia”-nya. Banyak mahasiswa yang cenderung menghindari ketika berhadapan dengan statistik. Dalam benak mereka selalu terbayang angka-angka yang “njlimet” yang memusingkan kepala. Capaian pembelajaran yang ingin diwujudkan melalui literasi buku ini bukan pada kompetensi berhitung—meski ada beberapa proses penghitungan sederhana yang masih diperlukan—melainkan statistik sebagai alat berpikir, yakni menyederhanakan masalah, mengurai masalah yang kompleks menjadi bagian-bagian yang dapat diselesaikan dengan mudah, dan menemukan pola dari suatu data atau kejadian.

Pembahasan statistik dalam buku ini lebih bersifat aplikatif dalam tataran penelitian, bukan statistik matematika yang berusaha mencari tahu asal-muasal suatu rumus atau menurunkan rumus satu ke rumus yang lain. Karena itu, penyajian materi sengaja diuraikan sesederhana mungkin untuk menghindari kesan “ruwet”. Perhitungan secara manual hanya dilakukan pada statistik deskriptif. Sementara statistik inferensial sejauh mungkin dilakukan dengan bantuan program komputer seperti SPSS (*statistical package for social sciences*).

Akhirnya, saya ingin mengutip pendapat Samuel Johnson: “*Great works are performed not by strength but by perseverance.*” Demikian pula buku ini lahir bukan dari sebuah ambisi besar melainkan dari sebuah jalinan pemikiran dan catatan kecil. Semoga dari yang “kecil” ini akan muncul sesuatu yang besar dan bermanfaat bagi banyak orang.

Surabaya, Mei 2018

Ali Maksum

Daftar Isi

Kata Pengantar

Daftar Isi

Bab 1 Pendahuluan.....	1
Bab 2 Distribusi Frekuensi.....	14
Bab 3 Tendensi Sentral	23
Bab 4 Norma Pengukuran	30
Bab 5 Variabilitas.....	36
Bab 6 Probabilitas dan Sampel.....	42
Bab 7 Skor Standar.....	48
Bab 8 Pengujian Hipotesis	54
Bab 9 Analisis Korelasi dan Regresi.....	67
Bab 10 T-Test dan Anova	86
Bab 11 Analisis Multivariat	97
Bab 12 Uji Statistik Nonparametrik	107
Daftar Pustaka	117

Bab 1 Pendahuluan

A. Pengertian Statistik dan Urgensinya dalam Olahraga

Secara etimologis, istilah statistik berasal dari kata *state* yang berarti negara atau segala sesuatu yang berhubungan dengan negara. Dalam konteks tersebut, statistik dipahami sebagai kumpulan angka-angka yang berhubungan dengan masalah pemerintahan seperti angka tentang jumlah penduduk, pendapatan masyarakat, angka kemiskinan, dan sebagainya. Pada perkembangannya, istilah statistik memiliki makna yang sangat berlainan dengan konsep awalnya. Statistik diartikan sebagai suatu metode dan prosedur yang digunakan untuk melakukan pengolahan, penafsiran, dan penarikan kesimpulan dari data hasil pengukuran.

Ada dua konteks penggunaan statistik dalam disiplin keilmuan, yang kiranya perlu diberikan pengertian lebih awal. Pertama, statistik sebagai bagian dari ilmu matematika yang membahas cara-cara mengolah data secara mendalam dan menyeluruh, termasuk darimana, mengapa, dan bagaimana suatu rumus pengolahan data diberlakukan. Pada konteks ini banyak dibahas tentang analisis matematika, aljabar linier, persamaan differensial, dan teori probabilitas. Karena itu, statistik tipe ini sering disebut teori statistik atau matematika statistik. Sementara yang kedua, penggunaan statistik yang diarahkan pada teknik dan prosedur pengolahan data, seperti rerata, simpangan baku, korelasi, regresi, dan anova. Yang terakhir ini sering disebut statistik terapan, yakni mengaplikasikan sejumlah rumus statistik ke dalam berbagai bidang keilmuan. Buku ini tidak akan membahas penggunaan statistik dalam konteks yang pertama, melainkan fokus pada penggunaan ilmu statistik dalam keolahragaan, terutama penelitian. Karena itu, tidak perlu ada kekhawatiran bagi mereka yang phobia terhadap angka-angka dan hitung-menghitung. Dalam statistik ini, tujuan utama bukan pada kompetensi berhitung—meski ada beberapa proses penghitungan sederhana yang masih diperlukan—melainkan statistik sebagai alat berpikir, yakni menyederhanakan masalah, mengurai masalah yang kompleks menjadi bagian-bagian yang dapat diselesaikan dengan mudah, dan menemukan pola dari suatu data atau kejadian.

Dalam praktik keolahragaan dewasa ini, banyak diterapkan statistik dalam pertandingan, seperti pada permainan bolabasket, sepakbola, dan bolavoli. Sebagai contoh bagaimana statistik diterapkan dalam menganalisis pertandingan sepakbola, antara Tim Nasional Indonesia dan Malaysia pada laga semifinal SEA Games 2017, 26 Agustus 2017 di Selangor Malaysia. Jika dicermati dari data statistik, Indonesia lebih dominan dibanding Malaysia. Misalnya dalam penguasaan bola pada paruh pertama, Tim Merah Putih mencapai 60%, sementara Malaysia hanya 40%, demikian juga dalam akurasi passing, Indonesia mencapai 75% sementara Malaysia 72%.

Meski demikian harus diakui, pada pertandingan tersebut Malaysia unggul dengan skor 1— 0. Sangat boleh jadi Tim Indonesia menguasai serangan, tetapi tidak efektif menghasilkan gol.

Dalam permainan bolabasket, penggunaan statistik seolah menjadi keniscayaan. Federasi bolabasket internasional (FIBA) menyediakan data statistik pertandingan di seluruh dunia yang dapat diakses melalui situs resmi mereka. Kita mengenal pemain hebat seperti Kareem Abdul Jabbar, Michael Jordan, dan Shaquille O’Neal sebagai pemain bolabasket terbaik sepanjang masa, bukan karena mereka pemain bintang NBA, tetapi karena sejumlah aksinya yang tercatat detail dalam statistik, misalnya berapa jumlah lemparan bebas, lemparan tiga angka, *rebound*, dan pelanggaran yang dilakukan pemain. Hal yang sama juga dilakukan oleh Liga Basket Nasional (NBL). Bahkan NBL memberikan penghargaan khusus kepada pemain yang telah mencapai parameter tertentu, misalnya *rebound* telah mencapai 1000 poin atau lebih, seperti Yanuar D. Priasmoro dari klub Bimasakti, Faisal J. Achmad dari Satria Muda, dan Ary Chandra dari Pelita Jaya.

Bagi pelatih, data statistik menjadi urgen guna menentukan strategi, baik dalam pelatihan maupun pertandingan. Misalnya rekor pertemuan dengan lawan, kesalahan yang seringkali dilakukan, dan poin yang banyak diciptakan. Pelatih yang jeli memanfaatkan data statistik akan berpeluang memenangkan suatu pertandingan. Kita ingat pelatih klub basket Miami Heat, Erik Spoelstra yang beberapa kali (2012—2013) membawa Miami Heat menjuarai kompetisi NBA. Spoelstra bukanlah pemain hebat dijamannya, tetapi ia merupakan ahli statistik dan videoman. Dengan data yang dimiliki, ia melakukan analisis secara mendalam untuk kemudian diterapkan dalam pertandingan. Hal yang demikian tidak banyak dilakukan oleh pelatih lain. Inilah sesungguhnya salah satu contoh sains olahraga, bagaimana melatih olahraga berbasis pada ilmu pengetahuan (*sports sciences*).

Dalam buku “*Statistical Thinking in Sport*” yang ditulis oleh Albert & Koning (2008) dijelaskan bahwa penggunaan statistik dalam olahraga begitu urgen dan sejatinya sudah berlangsung sejak tahun 1970-an, yang ditandai dengan lahirnya buku “*Management Science in Sports*” dan “*Operations Strategies in Sports*”. Sejumlah jurnal, seperti “*Chance, The American Statistician*” dan “*Quantitative Analysis in Sports*”, secara reguler juga menampilkan analisis statistik data olahraga. Bahkan dalam volume “*Anthology of Statistics in Sports*” dibahas secara khusus keberhasilan capaian rekor tertinggi dari sejumlah cabang olahraga. Memperhatikan perkembangan yang sedemikian maju—terutama kajian secara ilmiah—maka saya sampai pada kesimpulan bahwa penggunaan statistik dalam olahraga menjadi sebuah keniscayaan.

B. Fungsi Statistik dalam Penelitian

Data yang diperoleh dari pengumpulan di lapangan pada dasarnya masih berupa data mentah (*raw data*) atau skor kasar (*raw score*). Artinya, data tersebut belum diolah sehingga mudah dibaca dan memberikan gambaran yang bermakna. Dalam kaitan ini, statistik berfungsi memberikan teknik-teknik yang tepat dalam pengklasifikasian dan penyajian data penelitian. Kompas edisi 8 september 2017, sekaligus sebagai refleksi hari olahraga nasional, menerbitkan tulisan yang berjudul “Warga kurang aktivitas fisik, risiko menderita berbagai penyakit degeneratif meningkat”, mengutip hasil riset kesehatan dasar 2013, yang menyebutkan bahwa 26,1% penduduk yang berumur 10 tahun ke atas termasuk kelompok yang kurang melakukan aktivitas fisik. Sebanyak 22 dari 33 provinsi memiliki prevalensi kurang gerak di atas rerata nasional, prevalensi tertinggi dialami penduduk DKI Jakarta, yakni sebesar 44,2%. Padahal organisasi kesehatan dunia (WHO) menyatakan bahwa setiap orang harus beraktivitas fisik minimal 150 menit per minggu, dengan intensitas sedang, dan frekuensi 3-5 kali.

Sejalan dengan temuan di atas, Komnas Penjasor pada 2011 melakukan penelitian tentang aktivitas fisik terhadap 4.440 orang yang tersebar di lima kota, yakni Surabaya, Bandung, Makasar, Palembang, dan Mataram. Dari penelitian tersebut dinyatakan bahwa tingkat partisipasi olahraga sebesar 36,7% (pria 43,7% dan wanita 27,4%). Sementara tujuan berolahraga sebanyak 69,6% karena alasan kesehatan dan 16,8% untuk tujuan prestasi. Dilihat dari tingkat kebugaran jasmani, remaja putra sebesar 33,66 ml/kg/min dan putri sebesar 24,23 ml/kg/min. Adapun dilihat dari kategori, untuk putra hanya 6% yang masuk kategori baik ke atas, 23% kategori sedang, 15% rendah, dan 56% masuk kategori sangat rendah. Sementara untuk putri, hanya 2% yang masuk kategori baik ke atas, 3% kategori sedang, 30% rendah, dan 64% masuk kategori sangat rendah.

Pengumpulan data tentang aktivitas fisik yang dilakukan dalam rangka riset kesehatan dasar maupun Komnas Penjasor, sejatinya bersifat sampling. Acap kali dalam penelitian tidak mungkin melakukan analisis terhadap anggota populasi secara keseluruhan. Selain jumlahnya besar dan bahkan tidak terhitung, juga karena pertimbangan efisiensi dan efektivitas penelitian. Karena itu, analisis dilakukan melalui sebagian anggota populasi yang kemudian disebut sampel. Statistik bergerak pada tataran sampel tersebut. Dalam konteks yang demikian, statistik berfungsi memberikan parameter (ukuran) yang mencirikan populasi, menyatakan variasi, dan memberikan kecenderungan dari suatu variabel penelitian. Analisis data yang dilakukan terhadap sampel selanjutnya digunakan untuk menyimpulkan karakteristik populasi.

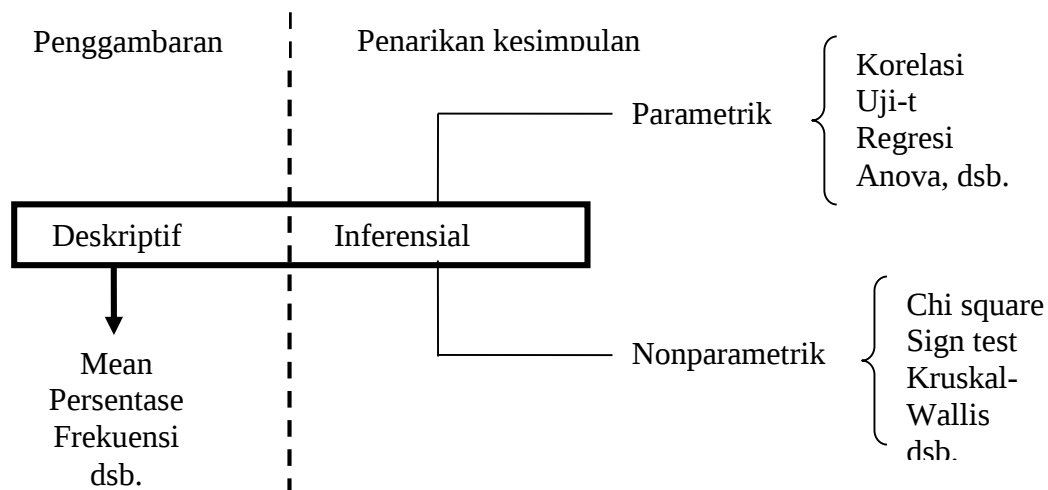
Statistik tidak hanya beroperasi pada tataran deskriptif sebagaimana diuraikan di atas, tetapi juga melakukan analisis secara mendalam. Misalnya, bagaimana kita dapat mengetahui bahwa IQ berhubungan dengan prestasi belajar? Apakah kapasitas oksigen maksimal (VO_{2max}) berhubungan dengan kemampuan berlari jauh? Kita tidak akan pernah bisa menjawab persoalan

di atas manakala kita tidak melakukan analisis terhadap data yang berhubungan dengan persoalan tersebut. Dalam kaitan ini, statistik berfungsi sebagai alat untuk menjelaskan hubungan dan tingkat hubungan antara dua variabel atau lebih. Kendati demikian perlu dicatat bahwa statistik tidak hanya bertalian dengan analisis hubungan, melainkan juga analisis perbedaan. Misalnya, perbedaan prestasi belajar antara mahasiswa reguler dan non-reguler, perbedaan motif berprestasi antara atlet yang memiliki latar belakang etnis Cina dan etnis Jawa.

C. Jenis Statistik

Berdasarkan cara yang dilakukan, statistik dapat dibagi menjadi 2 bagian, yaitu statistik deskriptif dan statistik inferensial. Statistik deskriptif adalah bagian dari statistik yang membahas mengenai penyusunan data ke dalam daftar, grafik atau bentuk lain yang dimaksudkan untuk memberikan gambaran tentang suatu variabel. Data tentang jumlah penduduk per jenis kelamin, usia, jenis pekerjaan dan data mahasiswa per program studi, angkatan, dan kelulusan adalah contoh data statistik deskriptif. Cara pengolahan data yang biasa dilakukan adalah dengan menggunakan rerata, frekuensi, persentase, bar charts, pie charts, dan sebagainya. Pengolahan data dalam jumlah besar umumnya menggunakan statistik deskriptif, dan yang paling sering adalah rerata dan persentase, seperti halnya yang dilakukan oleh Badan Pusat Statistik (BPS), Kementerian Kesehatan, dan Komnas Penjasor.

Pada tahun 2004–2007, Kementerian Pemuda dan Olahraga melakukan pengumpulan data terkait Sport Development Index (SDI), suatu instrumen untuk mengukur kemajuan pembangunan olahraga yang didasarkan pada empat aspek dasar, yakni ketersediaan ruang terbuka, partisipasi, sumberdaya manusia, dan kebugaran jasmani. Data SDI 2007 diolah berdasarkan hasil pengumpulan data di 33 propinsi di Indonesia yang mencakup 440 kabupaten/kota, dengan total subjek sebanyak 8.910 orang. Dari pengolahan data diperoleh Indeks SDI nasional sebesar 0.302. Jika dilihat dari setiap dimensi, maka diperoleh indeks ruang terbuka sebesar 0.289, indeks SDM sebesar 0.199, indeks partisipasi sebesar 0.422, dan indeks kebugaran sebesar 0.335. Indeks tersebut masuk dalam kategori rendah (norma SDI: 0,800 – 1 tinggi; 0,500 - 0,799 menengah; dan 0 - 0,499 rendah). Ibarat orang berjalan menuju suatu tujuan, angka 0,305 analog dengan 30% perjalanan. Artinya, tingkat kemajuan pembangunan olahraga nasional masih jauh dari apa yang kita harapkan.



Gambar 1.1: Pembagian statistik

Statistik inferensial – disebut juga statistik induktif - adalah bagian statistik yang membahas pengujian hipotesis dan penarikan kesimpulan berdasarkan data yang telah disusun dan diolah sebelumnya. Data tersebut biasanya berupa data sampel yang kemudian digeneralisasikan ke dalam populasi. Penyimpulan dapat berupa ada tidaknya hubungan dan/atau ada tidaknya perbedaan di antara berbagai data. Dapat juga dikatakan bahwa statistik inferensial dimaksudkan untuk menguji hipotesis, baik hipotesis nol atau pun hipotesis kerja. Teknik pengolahan data yang biasa digunakan adalah analisis korelasi, uji-t, analisis regresi, analisis varian, chi square, dan sebagainya.

Pada tahun 2005, penulis melakukan penelitian tentang dampak olahraga terhadap peningkatan kualitas hidup masyarakat ditinjau dari aspek kesehatan, psikologis, dan sosial. Dalam penelitian tersebut digunakan analisis varian (Anova), analisis statistik yang digunakan untuk menguji perbedaan antara tiga kelompok data atau lebih. Dari analisis yang dilakukan terhadap data yang diperoleh, secara keseluruhan menunjukkan bahwa ada perbedaan yang signifikan dalam hal kualitas hidup yang dicerminkan dengan tiga indikator, yaitu *subjective well-being*, *medical symptoms and vitality*, dan *pro-social behavior* antara kelompok yang secara reguler melakukan aktivitas fisik dan yang tidak dengan nilai F sebesar 25,57 pada $p < .01$. Orang yang secara teratur melakukan kegiatan olahraga, kualitas hidupnya lebih baik dibanding dengan orang yang tidak melakukan kegiatan olahraga. Lalu, bagaimana hasil analisis untuk setiap dimensi? Dari hasil analisis yang dilakukan menunjukkan bahwa ada perbedaan yang signifikan dalam hal *subjective well-being* antara kelompok reguler dan non-reguler dengan nilai F sebesar 11,77 pada $p < .01$. Artinya, aktivitas olahraga yang dilakukan secara teratur akan berdampak positif terhadap kondisi psikis seseorang yang melakukannya. Analisis yang dilakukan terhadap *medical symptoms and*

vitality antara kelompok reguler dan non-reguler menunjukkan bahwa ada perbedaan yang signifikan dengan nilai F sebesar 21,38 pada $p < .01$. Artinya, aktivitas olahraga yang dilakukan secara teratur, akan berdampak positif pada kondisi kesehatan individu yang melakukannya. Selanjutnya, terkait dengan perbedaan *pro-social behavior* antara kelompok reguler dan non-reguler, analisis statistik menunjukkan bahwa tidak ada perbedaan yang signifikan dengan nilai F sebesar 2,98 dengan $p > .05$. Artinya, seseorang yang secara teratur melakukan aktivitas olahraga, tidak dengan sendirinya perilaku sosialnya akan lebih baik dibanding orang yang tidak melakukan kegiatan olahraga.

Dalam statistik inferensial sendiri, dibagi menjadi 2 bagian, yaitu statistik parametrik dan nonparametrik. Statistik parametrik adalah suatu prosedur pengambilan kesimpulan statistik yang didasarkan pada parameter (ukuran yang mencirikan populasi). Parameter yang dimaksud berhubungan dengan asumsi-asumsi data seperti normalitas, homogenitas, dan linieritas. Statistik non-parametrik adalah suatu prosedur pengambilan kesimpulan statistik yang tidak didasarkan pada parameter. Artinya, data yang akan diolah tidak memenuhi persyaratan asumsi data sebagaimana dikemukakan di atas.

D. Dasar-Dasar Analisis Statistik

1. Data

Data dapat diartikan sebagai keterangan mengenai sesuatu. Keterangan dapat berupa angka yang lazim disebut data kuantitatif dan dapat juga berupa data bukan angka yang lazim disebut data kualitatif. Data kuantitatif diperoleh melalui proses pengukuran atau penjumlahan seperti: umur, tinggi badan, dan kecepatan lari. Sementara itu, data kualitatif biasanya diperoleh melalui wawancara, pengamatan, atau bahan tertulis seperti: jenis kelamin, jenis pekerjaan, kondisi kesejahteraan mantan atlet, gambaran interaksi siswa-guru, dan sebagainya. Statistik lebih berkenaan dengan data kuantitatif daripada data kualitatif. Data kualitatif dapat diolah dengan statistik apabila data tersebut dikuantifikasikan. Artinya, data yang semula berbentuk bukan angka diubah menjadi data berbentuk angka. Sebagai contoh, data jenis kelamin, yakni laki-laki dan perempuan. Data tersebut diubah menjadi 1 untuk laki-laki dan 2 untuk perempuan. Demikian juga data tentang kondisi kesejahteraan mantan atlet, ada yang berkecukupan (diberi kode 3), biasa-biasa saja (diberi kode 2), dan kekurangan (diberi kode 1). Data juga merupakan indikasi variabel penelitian. Sebagai contoh, variabel indeks prestasi dicerminkan oleh data berupa angka-angka seperti: 2,75; 2,50; dan 3,04. Demikian juga variabel kecepatan lari 100 meter yang dicerminkan dengan angka 11,05; 12,09, dan seterusnya.

Secara garis besar, data dapat dikelompokkan menjadi dua, yaitu data kategorik dan data kontinum. Data kategorik merujuk pada individu, objek, atau kejadian pada kategori tertentu.

Misalnya mahasiswa putra dan putri; mahasiswa yang berlatar belakang etnis madura dan jawa; jenis olahraga seperti sepakbola, bolavoli, dan bolabasket, dan sebagainya. Sementara itu data kontinum merujuk pada angka atau bilangan yang diperoleh dari pengukuran, bisa berbentuk bilangan bulat atau pecahan. Misalnya data tentang tinggi badan, berat badan, jarak lompatan, dan kebugaran.

Secara statistik, data dapat dikategorikan ke dalam empat jenis, yakni: data nominal, ordinal, interval, dan rasio. Masing-masing memiliki karakteristik yang berbeda dan karena itu juga teknik analisis statistiknya juga berbeda.

a. Data nominal

Data nominal adalah suatu data yang hanya dapat dikelompokkan secara terpisah (secara diskrit, secara kategorik) dan lebih merupakan sebuah lambang dari suatu kategori. Misalnya: data jenis kelamin dapat dikelompokkan secara terpisah menjadi laki-laki (diberi kode 1) dan perempuan (diberi kode 2); jenis pekerjaan dapat dikelompokkan menjadi pegawai negeri (kode 1), pedagang (kode 2), buruh (Kode 3), petani (kode 4), dan sebagainya. Dalam data nominal, angka 2 tidak berarti lebih tinggi dari angka satu, angka 4 tidak berarti kelipatan dari angka 2 dan begitu seterusnya. Angka hanya sekedar simbol. Beberapa teknik analisis statistik yang dapat digunakan untuk mengolah data nominal adalah: mode, korelasi kontingensi, korelasi Phi, atau Chi-square.

b. Data ordinal

Data ordinal adalah suatu data berupa angka yang menunjukkan tingkatan (jenjang) dalam sebuah urutan. Dalam data ordinal, angka tidak digunakan sebagai simbol, tetapi merupakan tingkatan. Misalnya: dalam sebuah kejuaraan Bolavoli terdapat juara 1, 2, dan 3. Angka-angka tersebut menunjukkan adanya tingkatan atau urutan, juara 1 lebih tinggi dari juara 2, dan juara 2 lebih tinggi dari juara 3. Meskipun demikian perlu dicatat bahwa jarak dari tingkatan satu ke tingkatan lain tidaklah sama. Teknik analisis statistik yang dapat digunakan untuk mengolah data ordinal antara lain: median, persentil, dan korelasi tata jenjang dari Spearman.

c. Data interval

Data interval adalah suatu data berupa angka yang batas variasi antara angka yang satu dengan yang lain sudah jelas, sehingga dapat dibandingkan. Sebagai contoh: Indeks Prestasi Mahasiswa, tingkat motivasi, dan tingkat kecerdasan. Meskipun demikian perlu dicatat bahwa nilai mutlakanya tidak dapat dibandingkan secara matematis, karena angka nol nya tidak mutlak (arbitrer). Seorang mahasiswa yang diberi nilai 0 (tidak lulus) dalam ujian bahasa Indonesia, tidak berarti mahasiswa tersebut sama sekali tidak bisa bahasa Indonesia. Demikian juga,

termometer yang menunjukkan angka 0 derajat, tidak berarti saat itu tidak ada suhu panas atau dingin. Sekalipun tidak memiliki angka nol mutlak, hampir semua teknik statistik dapat digunakan untuk menganalisis data interval, misalnya: mean, standar deviasi, varian, uji T, dan korelasi *product moment* dari Pearson.

d. Data rasio

Data rasio merupakan data pengukuran yang paling tinggi dan paling ideal. Batas intervalnya jelas, variasi nilainya tegas, dan titik nolnya mutlak. Dalam data rasio, angka nol bisa diartikan tidak ada gejala sama sekali. Misalnya: tinggi badan, berat badan, kecepatan lari, dan sebagainya. Semua teknik statistik dapat digunakan untuk menganalisis data berjenis rasio, seperti: mean, standar deviasi, varian, korelasi *product moment*, uji T, regresi, anova, dan sebagainya.

2. Variabel

Variabel adalah suatu konsep yang memiliki variabilitas atau keragaman. Sedangkan konsep sendiri adalah abstraksi atau penggambaran dari suatu fenomena atau gejala tertentu. Jika kita ingin meneliti tentang “lompat jauh”, maka lompat jauh masih merupakan sebuah konsep, karena masih berhubungan dengan pendefinisian dan penggambaran istilah lompat jauh sendiri serta tidak memiliki variasi. Akan tetapi jika diubah menjadi jarak lompatan dalam lompat jauh, maka di situ menunjukkan adanya variasi, misalnya 6,24; 7,16, dan sebagainya, maka hal tersebut dapat dikategorikan sebagai variabel.

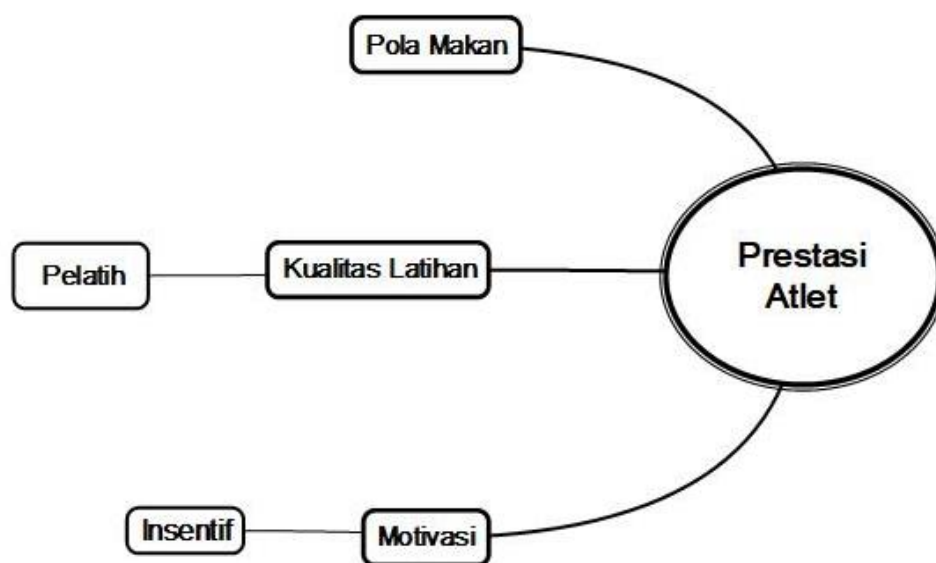
Secara garis besar, variabel ada dua macam, yaitu variabel bebas dan variabel terikat. Variabel bebas adalah variabel yang mempengaruhi, sementara variabel terikat adalah variabel yang dipengaruhi. Dalam gambar 1.2 nampak bahwa status gizi sebagai variabel bebas mempengaruhi kemampuan motorik sebagai variabel terikat.



Gambar 1.2: Variabel bebas mempengaruhi variabel terikat

Hubungan antar variabel yang demikian termasuk pola yang sederhana, padahal pola hubungan antar variabel acapkali tidak sesederhana itu, bisa lebih kompleks, langsung maupun

tidak langsung. Sebagai ilustrasi, sebagian kita beranggapan bahwa pelatih merupakan faktor penting yang mempengaruhi prestasi atlet. Hal tersebut tidak salah, namun jika hal tersebut ditempatkan dalam struktur hubungan antar variabel, maka faktor “pelatih” tidak langsung mempengaruhi prestasi atlet, melainkan melalui “kualitas latihan”. Tidak penting apakah seorang pelatih berasal dari dalam atau luar negeri, yang lebih penting apakah pelatih tersebut mampu menciptakan dan memberikan latihan yang bermutu kepada atlet. Demikian juga faktor “insentif” seperti bonus tidak langsung mempengaruhi prestasi atlet, tetapi melalui variabel “motivasi”. Dalam pola hubungan tersebut, maka variabel “kualitas pelatih” dan “motivasi” merupakan variabel moderator atau variabel perantara.



Gambar 1.3: Pola hubungan variabel bebas, variabel moderator, dan variabel terikat

3. Hipotesis

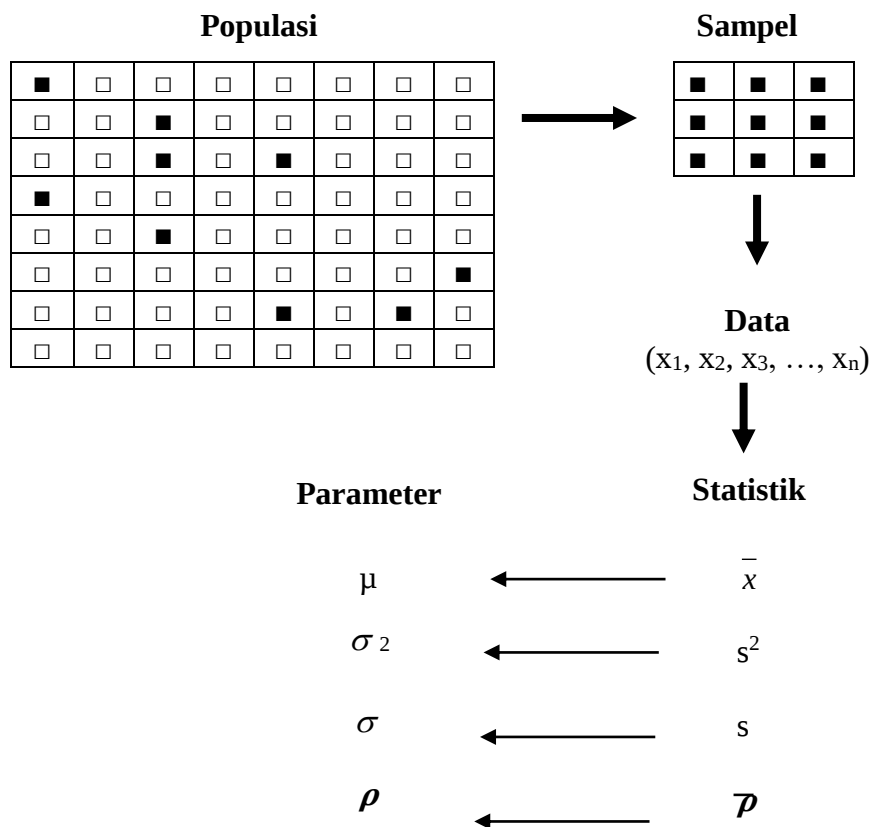
Dalam statistik, hipotesis menjadi bagian yang sangat urgen. Pengujian statistik pada dasarnya adalah pengujian terhadap hipotesis dalam suatu penelitian. Secara sederhana, hipotesis dapat diartikan sebagai dugaan sementara yang diajukan peneliti berupa pernyataan-pernyataan untuk diuji kebenarannya. Apa yang dilakukan peneliti dalam kegiatan penelitiannya adalah melakukan pembuktian hipotesis. Secara umum, ada 2 macam hipotesis, yaitu **hipotesis nihil** (disebut juga hipotesis nol, hipotesis statistik, dan disingkat H_0) dan **hipotesis kerja** (disebut juga hipotesis alternatif dan disingkat H_a). Hipotesis nihil adalah sebuah pernyataan yang menyatakan tidak adanya hubungan, tidak adanya perbedaan atau tidak adanya pengaruh antara dua variabel atau lebih. Misalnya: “tidak ada hubungan antara tingkat kecerdasan dengan prestasi belajar mahasiswa

FIK”. Hipotesis kerja adalah sebuah pernyataan yang menyatakan adanya hubungan, adanya perbedaan atau adanya pengaruh. Misalnya: “Ada perbedaan prestasi belajar antara mahasiswa PMDK dan Non-PMDK”.

Dalam penelitian, hanya ada satu hipotesis yang benar, yaitu hipotesis yang terbukti atau yang diterima. Pembuktian penerimaan hipotesis ditunjukkan oleh taraf signifikansi hasil uji statistik, lazimnya 5% atau 1%. Apabila hipotesis kerja diterima, maka hipotesis nihil ditolak. Begitu sebaliknya, jika hipotesis nihil diterima, maka hipotesis kerja ditolak. Perlu juga diberikan penjelasan di sini bahwa apabila hipotesis tidak terbukti, tidak serta merta sebuah penelitian dianggap gagal. Perlu diperiksa dulu apakah ada bagian proses penelitian yang tidak sesuai prosedur, termasuk pengumpulan datanya. Jika setelah diperiksa secara komprehensif tidak ditemukan kesalahan prosedur, maka mungkin saja penelitian tersebut justru menghasilkan kebaruan. Apabila itu terjadi, maka perlu ada diskusi yang mendalam mengapa hal tersebut bisa terjadi.

4. Populasi dan Sampel

Statistik sejatinya ingin mengkaji sejumlah parameter dalam populasi, yang merupakan sekumpulan individu, benda, atau obyek tertentu. Mengingat populasi biasanya dalam jumlah besar, seperti jumlah penduduk suatu kota, jumlah siswa sekolah, dan air satu kolam, maka akses terhadap seluruh anggota populasi menjadi tidak mudah dilakukan. Selain alasan waktu, biaya yang besar untuk mengakses seluruh anggota populasi juga menjadi pertimbangan. Atas dasar itu, pengambilan sampel menjadi penting untuk dilakukan. Tugas utama statistik inferensial adalah menarik kesimpulan terhadap variabel yang diteliti berdasarkan data yang diperoleh dari sampel untuk kemudian digeneralisasikan pada populasi. Karena itu, pemahaman mengenai konsep populasi dan sampel menjadi penting. Populasi adalah keseluruhan individu atau objek yang dimaksudkan untuk diteliti dan yang nantinya akan dikenai generalisasi. Generalisasi adalah suatu cara pengambilan kesimpulan terhadap kelompok individu atau objek yang lebih luas berdasarkan data yang diperoleh dari sekelompok individu atau objek yang lebih sedikit. Sebagian kecil individu atau objek yang dijadikan wakil dalam penelitian disebut sampel.



Gambar 1.4: Konsep populasi—sampel dalam statistik

Bagaimana suatu sampel diambil dari populasi akan berpengaruh terhadap pilihan analisis statistik yang digunakan dan keberlakuan dari kesimpulan yang diperoleh. Secara garis besar, pengambilan sampel terbagi dalam dua bagian, yakni random dan non-random. Pengambilan sampel secara random dimaksudkan bahwa subjek atau objek penelitian diambil secara acak sehingga semua anggota populasi memiliki kesempatan yang sama untuk menjadi sampel. Dalam kondisi yang demikian, hampir semua teknik statistik parametrik dapat digunakan dan generalisasi juga dapat diberlakukan. Sementara itu, pengambilan sampel non-random dimaksudkan bahwa karena adanya pertimbangan atau kriteria tertentu, maka tidak semua anggota populasi memiliki kesempatan yang sama untuk menjadi sampel. Dalam kondisi yang demikian, penggunaan teknik analisis statistik perlu lebih cermat, termasuk generalisasi yang diberlakukan. Sampel yang diambil secara non-random, keberlakuan kesimpulan terbatas pada sampel yang diteliti, kecuali jika karakteristik sampel tersebut juga dimiliki dan ada pada populasi.

Dalam logika eksperimen, penempatan subyek secara random menjadi keniscayaan. Jika dalam suatu eksperimen penempatan subyek secara random tidak dapat dilakukan, maka eksperimen tersebut bersifat semu. Cook & Campbell (1979) menyatakan bahwa “*quasi-experiments are similar to experiments except that the subjects are not randomly assigned to the*

independent variable. Quasi-experiments are used instead of experiments when random assignment is not possible...”

Masalah berikutnya terkait dengan jumlah sampel. Berapa jumlah sampel yang ideal dalam suatu penelitian. Tidak ada ukuran yang pasti, karena tergantung pada derajat homogenitas, tingkat akurasi, dan teknik analisis statistik yang digunakan. Semakin homogen suatu populasi, maka semakin sedikit jumlah sampel yang dibutuhkan. Semakin tinggi akurasi yang diinginkan, semakin besar jumlah sampel yang dibutuhkan. Secara statistik disarankan berjumlah ≥ 30 . Mengapa? Jumlah 30 dianggap sampel besar dan jika sampel diambil secara random maka data cenderung berdistribusi normal. Jika jumlah sampel kurang dari 30, maka diduga varian kesalahannya cukup besar.

Soal untuk didiskusikan dalam kelompok

1. Kelompokkan variabel-variabel berikut menurut jenis data yang dihasilkan.

Juara bulutangkis

Jarak lemparan

Usia

Kecepatan renang

agama

Kekuatan otot punggung

Motivasi berprestasi

Tinggi raihan

Vo2max

Metode mengajar

Suku bangsa

Gaya melatih

Panjang lengan

Kekuatan otot tungkai

2. Berikan 2 contoh rumusan hipotesis dan tentukan variabel-variabelnya.
3. Jika dibuat matriks, apa perbedaan mendasar dari data nominal, ordinal, interval, dan rasio.

Bab 2 Distribusi Frekuensi

Distribusi frekuensi adalah suatu cara meringkas dan menyusun sekelompok data mentah yang diperoleh dari suatu pengukuran. Acapkali kita memperoleh data dari lapangan dalam bentuk yang belum terorganisasi dengan baik sehingga tidak dapat segera digunakan. Dengan distribusi frekuensi, sekelompok data dapat dibaca dan dipahami dengan mudah. Sebagai ilustrasi, berikut ada 20 data berupa skor ujian statistik mahasiswa.

70	75	90	80	80
75	85	80	85	75
80	85	75	80	75
85	85	70	80	90

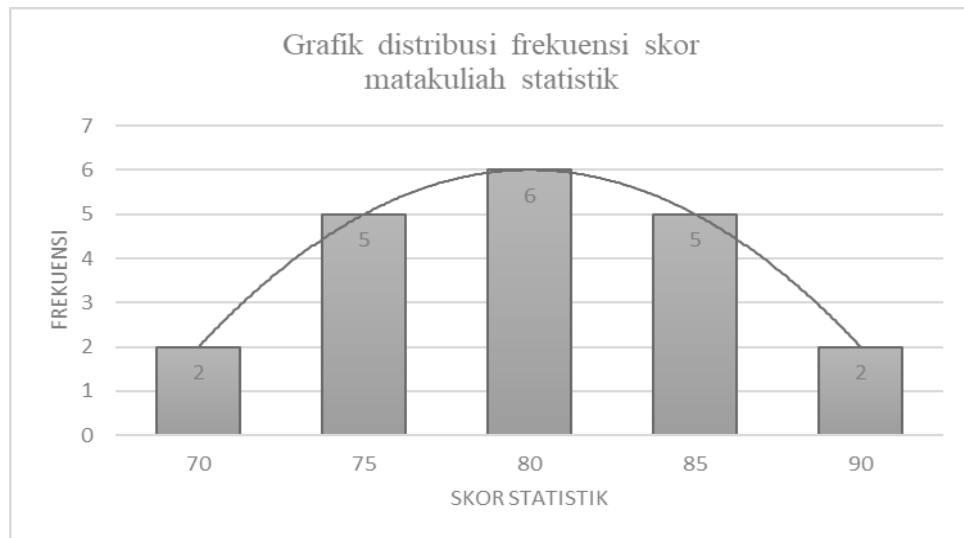
Memperhatikan data di atas, tentu kita tidak mudah untuk segera bisa membaca berapa mahasiswa yang memperoleh skor 75, berapa mahasiswa yang memperoleh skor tertinggi, dan seterusnya. Kondisi akan berbeda jika sekelompok data tersebut sudah diorganisasi dalam bentuk distribusi frekuensi seperti berikut.

Skor	Tabulasi	Frekuensi
70	II	2
75	HHH	5
80	HHH I	6
85	HHH	5
90	II	2
		N= 20

Dengan memperhatikan distribusi frekuensi di atas, kita dengan mudah mengatakan bahwa yang memperoleh skor 75 sebanyak 5 mahasiswa, yang memperoleh skor tertinggi sebanyak 2 mahasiswa, dan sebagian besar (6 mahasiswa) memperoleh skor 80. Kolom tabulasi sekadar untuk memudahkan penghitungan. Untuk selanjutnya, kolom tersebut dapat diabaikan, sehingga dalam distribusi frekuensi hanya ada 2 kolom, yakni skor dan frekuensi.

Tabel distribusi frekuensi tersebut dapat ditindaklanjuti dengan membuat grafik, dengan sumbu X berisi skor matakuliah statistik dan sumbu Y berisi frekuensi. Dengan menuangkan data tersebut dalam bentuk grafik, selain memudahkan kita membaca data tersebut, juga memberikan

informasi kepada kita mengenai normalitas distribusi data. Dari grafik tersebut tampak bahwa data terdistribusi secara normal, posisi rerata yang berada di tengah secara simetris membagi distribusi menjadi 2 bagian, kiri dan kanan. Secara sederhana, suatu data dikatakan berdistribusi normal jika sebagian besar data mengelompok ditengah, diikuti sebagian data di sebelah kiri dan kanan secara berimbang.



Gambar 2.1: Distribusi skor matakuliah statistik

Ada dua jenis distribusi frekuensi, yaitu: tunggal dan kelompok. Secara umum tujuan keduanya sama, yaitu memaparkan data sehingga mudah untuk dibaca. Distribusi frekuensi tunggal lebih sederhana karena tidak ada pengelompokan skor, sementara untuk distribusi frekuensi kelompok terdapat sejumlah interval kelas guna menggambarkan skor.

A. Distribusi Frekuensi Tunggal

Distribusi tunggal dicirikan dengan tidak adanya pengelompokan nilai-nilai dari suatu variabel. Sebagai contoh, seorang peneliti melakukan pengumpulan data tentang “tujuan masyarakat berolahraga”. Tujuan tersebut dikategorikan ke dalam 2 bagian, mereka yang melakukan olahraga untuk tujuan prestasi dan mereka yang melakukan olahraga untuk tujuan kesehatan. Dari 40 orang yang disurvei, didapat data sebagai berikut.

Tabel 2.1: Frekuensi dan Persentase Tujuan Berolahraga

Tujuan berolahraga	Frekuensi (f)	Persentase (%)
Prestasi	10	25
Kesehatan	30	75
Jumlah	40	100 %

Dalam distribusi frekuensi seperti di atas, biasanya disertakan persentase sebagai informasi tambahan. Artinya, kita tidak sekadar mengetahui frekuensi dari data tersebut, tetapi juga persentasenya.

Rumus menghitung **persentase**:

Jumlah kasus (n) dibagi dengan jumlah total (N) dikalikan 100% atau jika dirumuskan menjadi sebagai berikut:

$$\text{Persentase} = \frac{n}{N} \times 100 \%$$

Dengan mengambil contoh data pada tabel 2.1, maka kegiatan olahraga yang bertujuan prestasi, frekuensi atau jumlah kasusnya 10. Sementara jumlah frekuensi total adalah 40. Dengan menggunakan rumus di atas, maka:

$$\frac{10}{40} \times 100\% = 25\%$$

Dari data di atas, kita juga dapat menghitung lebih jauh, misalnya mengenai proporsi dan rasio. **Proporsi** adalah perbandingan jumlah kasus pada kategori tertentu dengan jumlah total, dengan menggunakan angka dasar 1.

$$\text{Proporsi} = \frac{n}{N} \times 1$$

Dari data di atas, maka dapat dihitung proporsinya sebagai berikut:

$$\frac{10}{40} \times 1 = 0,25$$

Adapun **rasio** adalah perbandingan antara jumlah kasus pada kategori tertentu dengan jumlah kasus pada kategori yang lain, tergantung pada kategori apa yang ingin kita perbandingkan, dengan menggunakan angka dasar 1.

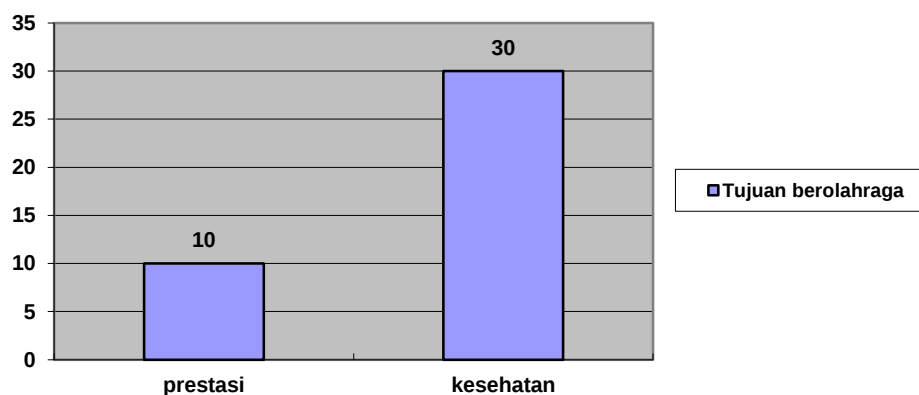
$$\text{Rasio} = \frac{n_1}{n_x} \times 1$$

Sebagai contoh, dari data pada tabel 2.1 di atas, maka rasio antara tujuan prestasi dan tujuan kesehatan adalah sebagai berikut:

$$\frac{10}{30} \times 1 = 0,333$$

Angka 0,333 mengandung arti bahwa untuk satu orang yang melakukan olahraga dengan tujuan kesehatan, terdapat 0,333 yang bertujuan prestasi. Untuk mempermudah pemahaman, angka desimal tersebut bisa dibulatkan dengan mengalikannya dengan angka 100. Dengan demikian didapatkan rasio 33. Artinya, dalam setiap 100 orang yang melakukan olahraga terdapat 33 orang berolahraga untuk tujuan prestasi. Dengan kata lain, rasio antara tujuan prestasi dan tujuan kesehatan adalah 10:30 atau 1:3.

Data sebagaimana tampak pada tabel 2.1, dapat juga disajikan dalam bentuk histogram (*bar chart*). Jika data tersebut digambarkan dalam bentuk histogram, maka dapat divisualisasikan sebagai berikut.



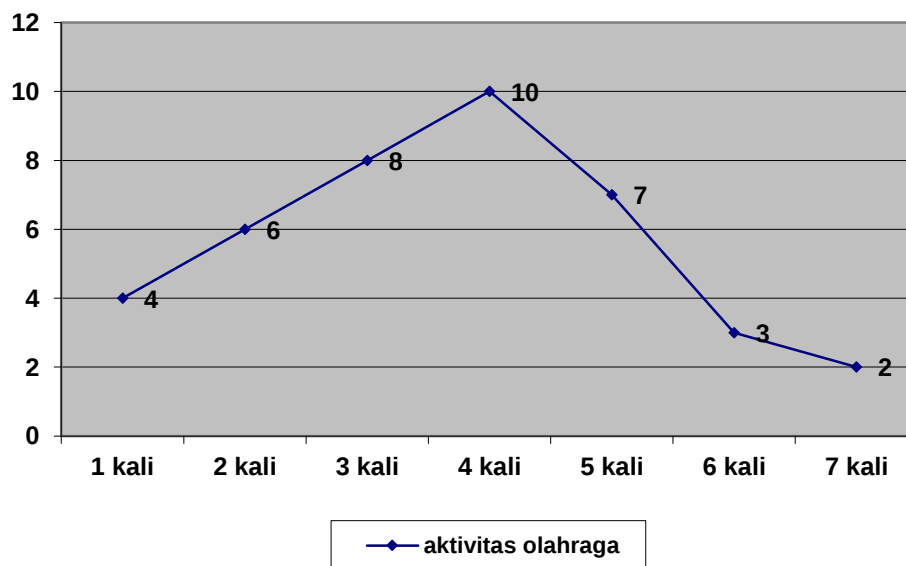
Gambar 2.2:
Histogram Frekuensi Tujuan Berolahraga

Penyajian data juga dapat dilakukan dalam bentuk poligon. Misalnya, seorang peneliti melakukan penelitian tentang “frekuensi masyarakat melakukan aktivitas olahraga dalam satu minggu”. Dari 40 orang yang disurvei, didapat data sebagai berikut.

Tabel 2.2: Frekuensi dan Persentase Aktivitas Olahraga dalam Seminggu

Aktivitas olahraga dalam satu minggu	Frekuensi (f)	Persentase (%)
7 kali	2	5
6 kali	3	7,5
5 kali	7	17,5
4 kali	10	25
3 kali	8	20
2 kali	6	15
1 kali	4	10
Jumlah	40	100 %

Jika data di atas digambarkan dalam bentuk poligon, maka dapat divisualisasikan sebagai berikut.



Gambar 2.3:

Poligon Frekuensi Olahraga Masyarakat dalam Seminggu

B. Distribusi Frekuensi Kelompok

Distribusi frekuensi kelompok dicirikan oleh penggunaan kelas interval untuk menggambarkan nilai dari variabel. Penggunaan interval dalam frekuensi kelompok dikarenakan angka-angka hasil pengukuran sangat heterogin, yang tidak dapat dikelompokkan dan disimbulkan oleh satu jenis angka tertentu sebagaimana dalam distribusi tunggal. Sebagai contoh, seorang guru pendidikan jasmani memberikan nilai kepada 40 muridnya dengan variasi skor sebagai berikut.

50	56	59	65	65	66	68	68
70	70	70	75	75	75	76	76
76	76	77	78	78	78	78	79
79	80	80	80	82	82	85	85
85	85	86	86	88	88	88	90

Bagaimana membuat distribusi frekuensi dari skor-skor tersebut? Untuk menjawab pertanyaan tersebut perlu dibuat langkah-langkah sebagai berikut.

1. Temukan skor tertinggi dan skor terendah. Dalam data di atas skor tertinggi adalah 90 dan skor terendah adalah 50.
2. Temukan *range* (jarak pengukuran), yaitu (skor tertinggi – skor terendah) + 1. Dalam data di atas *range*-nya adalah $(90 - 50) + 1 = 41$.
3. Tetapkan jumlah kelas interval, berkisar antara 5 sampai dengan 15, misalnya untuk data di atas ditetapkan 5. Usahakan jumlah kelas interval dalam bilangan ganjil agar mudah ditentukan kelompok tengahnya.
4. Hitung lebar interval, yaitu *range* dibagi jumlah kelas interval. Dalam data di atas, lebar interval adalah $41/5 = 8,2$ (dibulatkan ke atas menjadi 9). Usahakan lebar interval merupakan bilangan ganjil agar titik tengah skor merupakan bilangan bulat.
5. Susun kelompok–kelompok interval ke dalam tabel dengan memperhatikan hasil perhitungan yang telah dilakukan.

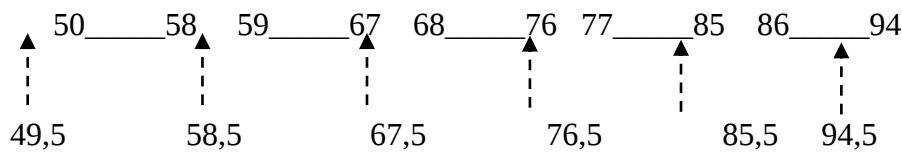
Tabel 2.3: Distribusi Frekuensi Nilai Penjas dari 40 Siswa

	Kelas Interval	Frekuensi
	86 – 94	6
	77 – 85	16
	68 – 76	12
	59 – 67	4
	50 – 58	2
	Jumlah	40

Batas bawah

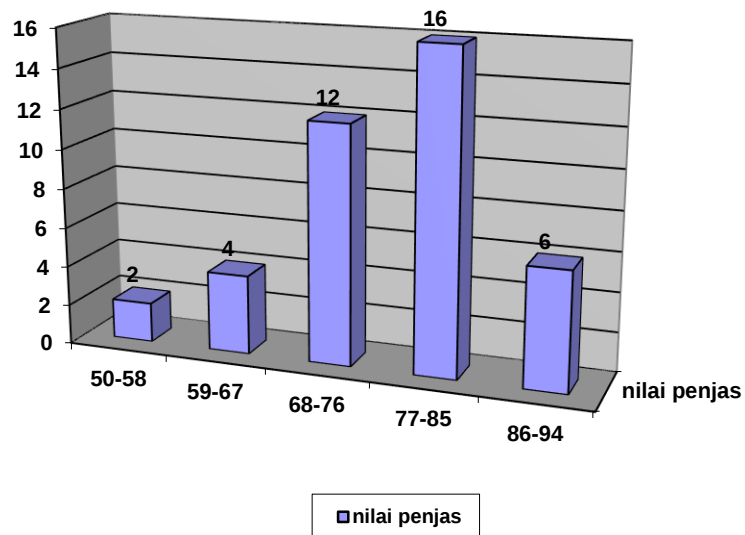
Batas atas

Skor 86, 77, 68, 59, dan 50 dalam kelas interval di atas disebut **batas bawah**; sedangkan skor 94, 85, 76, 67, dan 58 disebut **batas atas**. Kedua batas tersebut pada dasarnya merupakan **batas semu**. Mengapa? Karena kedua batas tersebut tidak dapat menghubungkan antara kelompok interval yang satu dengan lainnya. Sebagai gambaran lihatlah ilustrasi berikut:

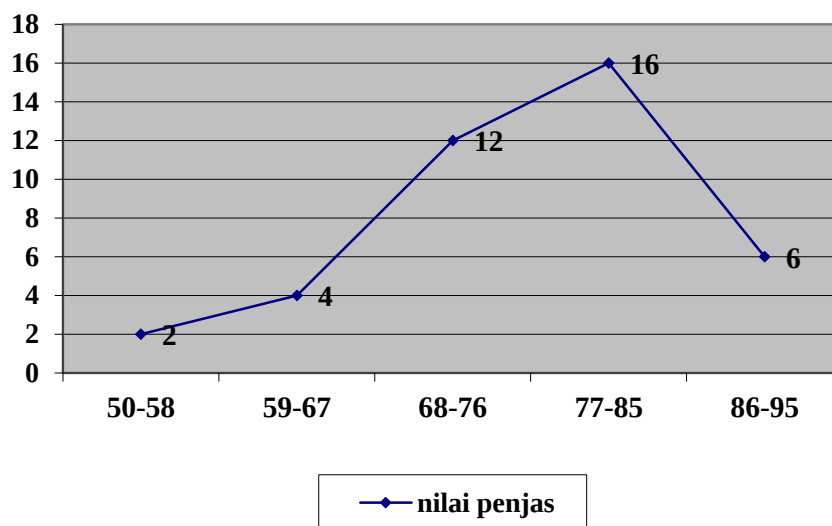


Dari ilustrasi tersebut tampak bahwa antara batas atas kelas interval pertama (58) dan batas bawah kelas interval kedua (59) tidak terhubung oleh garis. Apabila garis tersebut ingin dihubungkan, maka harus ada skor antara yang menghubungkan keduanya, yaitu 58,5. Dalam contoh tersebut skor 58,5 disebut **batas nyata**.

Dari data pada tabel 2.3, maka dapat digambarkan dalam bentuk histogram (lihat gambar 2.4) maupun dalam bentuk poligon (lihat gambar 2.5).



Gambar 2.4: Histogram frekuensi nilai Penjas 40 siswa



Gambar 2.5: Poligon frekuensi nilai Penjas 40 siswa

Soal untuk dikerjakan dan didiskusikan dalam kelompok:

1. Seorang guru memberikan skor bidang studi pendidikan jasmani kepada 20 siswa yang menjadi anak didiknya. Skor tersebut terentang sebagai berikut.

8	7	6	8	9
9	7	8	10	8
9	8	10	7	9
6	8	7	8	7

Dari data di atas, buatlah distribusi frekuensi tunggal!

2. Dalam sebuah kejuaraan bolabasket antar pelajar, terdapat 40 klub sekolah yang ikut berpartisipasi. Setelah masing-masing klub sekolah bertanding, diperoleh skor memasukkan bola sebagai berikut.

24	12	9	12	13	10
29	19	15	15	10	14
30	16	14	13	20	12
17	17	9	14	17	15
18	22	23	23	25	27

Dari data di atas, buatlah distribusi frekuensi kelompok!

Bab 3 Tendensi Sentral

Tendensi sentral adalah sebuah angka yang menunjukkan kecenderungan memusatnya sekelompok skor dalam suatu distribusi. Tendensi sentral digunakan untuk merangkum dan mendeskripsikan data dengan cara mencari angka yang dapat mewakili sekelompok data. Berapa pun jumlah suatu data, ketika kita sudah mengetahui tendensi sentralnya, maka kita dengan cepat mendapatkan gambarnya. Misalnya, dari sekitar 250 juta penduduk Indonesia, rerata lama pendidikannya 7.6 tahun, artinya sebagian besar tingkat pendidikan penduduk Indonesia masih rendah, di atas SD tetapi belum lulus SMP. Demikian juga dari hasil penelitian Komnas tahun 2011 tentang kebugaran pelajar, dari 772 pelajar yang tersebar di 6 kota (Bandung, Surabaya, Palembang, Makasar, Mataram, dan Palangkaraya), rerata Vo2max pelajar putra sebesar 30,82 dan pelajar putri sebesar 23,63. Dengan angka rerata Vo2max tersebut, kita segera mendapatkan gambaran bahwa tingkat kebugaran pelajar kita sangat rendah.

Dalam terminologi statistik, tendensi sentral terdiri dari tiga bentuk, yakni Mean, Median, dan Mode. Dari ketiga ukuran tendensi sentral tersebut, Mean paling sering digunakan untuk menggambarkan suatu data.

A. Mean

Mean atau rerata hitung adalah angka yang diperoleh dengan membagi jumlah skor keseluruhan dengan jumlah individu.

Rumus 1

$$M = \frac{\Sigma X}{N}$$

Dimana, M = mean
ΣX = jumlah total skor dalam distribusi
N = jumlah individu

Sebagai contoh, ada sejumlah angka: 9, 8, 7, 6, 5; maka jika dicari mean-nya adalah $35/5 = 7$. Rumus 1 di atas hanya cocok untuk data kasar dengan frekuensi 1. Jika skor dalam distribusi frekuensinya lebih dari satu, maka digunakan rumus sebagai berikut.

Rumus 2

$$M = \frac{\Sigma fX}{N}$$

Dimana, M = mean
ΣfX = jumlah skor dikalikan frekuensi
N = jumlah individu

Contoh, ada sejumlah skor dengan distribusi sebagai berikut:

Tabel 3.1: Data untuk Mencari Mean pada Distribusi Tunggal

Skor (X)	Frekuensi (f)	ΣfX
9	2	18
8	3	24
7	1	7
6	2	12
5	3	15
Jumlah	11	76

Maka, jika kita ingin mencari mean-nya diperoleh:

$$\frac{76}{11} = 6,90$$

Lalu, bagaimana untuk mencari mean dalam distribusi kelompok? Hakikatnya adalah sama, hanya saja dalam distribusi kelompok, nilai X tidak mewakili nilai individu, melainkan mewakili titik tengah interval.

Contoh,

Tabel 3.2: Data untuk Mencari Mean pada Distribusi Kelompok

Kelas Interval	Titik Tengah (X)	f	ΣfX
14 - 16	15	3	45
11 - 13	12	2	24
8 - 10	9	5	45
5 - 7	6	4	24
2 - 4	3	2	6
Jumlah	-	16	144

Maka, dengan menggunakan rumus 2 didapat mean sebesar:

$$\frac{144}{16} = 9$$

Mencari Mean juga dapat dilakukan dengan cara menerka dengan rumus sebagai berikut.

Rumus 3

$$M = MT + \left(\frac{\sum fd}{N} \right) i$$

dimana;

MT = mean terkaan
d = deviasi
N = jumlah individu
i = lebar interval

Langkah-langkah untuk menghitung Mean Terkaan adalah sebagai berikut:

1. Menerka letak Mean (boleh disembarang interval) dan menjadikan titik tengahnya sebagai Mean Terkaan (MT).
2. Menghitung deviasi pada setiap interval, dengan ketentuan interval MT diberi harga 0, sedangkan di atas MT diberi harga +1 dan seterusnya, sedangkan di bawah MT diberi harga -1 dan seterusnya.
3. Mengalikan f dengan d dan mencari jumlah totalnya.
4. Menghitung N dan lebar interval.

Dengan menggunakan data pada tabel 3.2, maka MT dapat diperoleh sebagai berikut. Misalnya kita ingin menggunakan interval kedua sebagai tempat menduga Mean (lihat tabel 3.3).

Tabel 3.3: Data untuk Mencari Mean Terkaan pada Distribusi Kelompok

Kelas Interval	Titik Tengah (X)	f	d	fd
14 - 16	15	3	+1	3
11 - 13	12	2	0	0
8 - 10	9	5	-1	-5
5 - 7	6	4	-2	-8
2 - 4	3	2	-3	-6
Jumlah	-	16	-	-16

Dengan menggunakan rumus 3, maka dapat diperoleh Mean sebagai berikut:

$$\begin{aligned}
M &= 12 + \left(\frac{-16}{16} \right) 3 \\
&= 12 + (-1) 3 \\
&= 12 - 3
\end{aligned}$$

$$= 9$$

B. Median

Median – atau disebut juga rata-rata posisi – adalah angka yang terletak di tengah-tengah sederetan angka atau sebuah distribusi frekuensi. Apabila ada sekelompok data dan kemudian diurutkan mulai dari yang terkecil sampai yang terbesar, lalu dibagi menjadi dua kelompok tinggi dan separuhnya lagi kelompok rendah, maka titik tengah yang memisahkan kedua kelompok tersebut disebut Median. Sebagai contoh, ada kelompok angka yang sudah diurutkan sebagai berikut:

2 3 4 5 6 7 8 median-nya adalah 5

3 4 5 6 7 10 20 30 median-nya adalah 6,5

Cara di atas adalah cocok untuk menghitung median dari distribusi frekuensi tunggal. Untuk menghitung median dari distribusi frekuensi kelompok diperlukan rumus sebagai berikut.

$$\text{Median} = B + \left\{ \frac{\frac{1}{2}N - f_k}{f} \right\} i$$

Dimana,

B = batas bawah nyata dari interval yang mengandung median

N = Jumlah (frekuensi) individu dalam distribusi

f_k = frekuensi kumulatif di bawah interval yang mengandung median

i = lebar interval

f = Frekuensi interval yang mengandung median

Sebagai contoh,

Tabel 3.4: Data untuk mencari Median pada distribusi kelompok

Kelas Interval	f	Σf_k
14 - 16	3	16
11 - 13	2	13
<u>8 - 10</u>	<u>5</u>	<u>11</u>
5 - 7	4	6
2 - 4	2	2
Jumlah	16	-

Langkah-langkah untuk menentukan Median adalah sebagai berikut.

1. Menemukan harga $\frac{1}{2} N$, dalam data pada tabel 3.4 di atas adalah 8
2. Menentukan letak 8 pada f_k , dalam hal ini pada f_k 11 yang terletak pada interval 8–10, karena sampai dengan ini jumlah frekuensi sudah lebih dari 8
3. Menentukan batas bawah nyata interval 8-10, yaitu 7,5
4. Menentukan frekuensi pada interval 8-10, yaitu 5
5. Menentukan lebar interval, yaitu 3

Jika dimasukkan ke dalam rumus, maka:

$$\begin{aligned}\text{Median} &= 7,5 + \left\{ \frac{\frac{1}{2}16 - 6}{5} \right\} 3 \\ &= 7,5 + \left(\frac{2}{5} \right) 3 \\ &= 7,5 + 1,2 \\ &= 8,7\end{aligned}$$

Penghitungan Median dengan cara di atas lebih merupakan perkiraan, yang ada kemungkinan berbeda dengan median yang dicari dengan cara mengurutkan sebagaimana dikemukakan di awal. Median jarang digunakan untuk alat estimasi dibanding Mean, karena skor Median cenderung tidak stabil dibanding skor Mean. Median umumnya digunakan dalam data kategorik, bukan kontinum.

Mode (Modus)

Mode adalah skor yang frekuensinya paling banyak.

Contoh, sekelompok angka sebagai berikut: 2 2 3 4 5 5 5 6 6 7

Maka, mode-nya adalah 5

Demikian juga pada distribusi frekuensi kelompok seperti data di bawah ini. Mode-nya adalah 9.

Tabel 3.5: Tabel Mencari Mode dalam Distribusi Kelompok

Klas Interval	Titik Tengah (X)	f
2 – 4	3	2
5 – 7	6	4
8 – 10	<u>9</u>	<u>5</u>
11 – 13	12	2
14 – 16	15	3
Jumlah	-	16

Seperti halnya Median, Mode juga jarang dipergunakan sebagai alat estimasi. Mode biasanya digunakan pada data kategorik.

Soal untuk dikerjakan dan didiskusikan dalam kelompok.

Carilah: (1) mean, (2) median, dan (3) mode dari data berikut.

8	12	9	16
11	14	15	15
12	16	14	16
17	12	9	13
10	16	10	13

Bab 4 Norma Pengukuran

Norma pengukuran diperlukan untuk membuat suatu kategori atas sekelompok data. Berdasarkan norma tertentu, kita dapat memisahkan data menjadi bermacam-macam kategori sesuai dengan kepentingan kita. Misalnya, dengan membagi distribusi data menjadi dua bagian kita mendapatkan kategori “tinggi” dan ‘rendah’. Dengan empat kategori, kita mendapatkan kategori seperti: sangat tinggi, tinggi, sedang, dan rendah. Demikian juga seterusnya. Pembuatan norma pengukuran dengan dua kategori disebut Median; empat kategori disebut Kuartil, sepuluh kategori disebut Desil; dan seratus kategori disebut dengan Persentil. Pada bagian ini median tidak akan dibahas, mengingat telah dijelaskan pada bagian sebelumnya.

A. Kuartil

Kuartil adalah sebuah indeks yang dapat membagi distribusi data menjadi empat bagian atau kategori. Untuk membagi 4 bagian tersebut dibutuhkan 3 buah angka kuartil, yaitu: kuartil 1, kuartil 2, dan kuartil 3.

Kuartil 1 adalah suatu nilai dalam distribusi yang membagi 25% frekuensi di bagian bawah dan 75% frekuensi di bagian atas.

$$K1 = B + \left(\frac{\frac{1}{4} N - fk}{f} \right) i$$

dimana,
K1 = kuartil 1
N = jumlah individu
B = batas bawah nyata pada interval yang mengandung kuartil
fk = frekuensi kumulatif dibawah fk yang mengandung kuartil
f = frekuensi pada interval yang mengandung kuartil
i = lebar interval

Contoh,

Tabel 4.1: Tabel untuk mencari kuartil

Klas Interval	f	Σfk
28 – 32	5	23
23 - 27	2	18
18 - 22	4	16
13 – 17	3	12
<u>8 – 12</u>	<u>6</u>	9
3 - 7	3	<u>3</u>
Jumlah	23	-

Dengan data tersebut, jika dicari K1-nya adalah sebagai berikut.

$$\begin{aligned} K1 &= 7,5 + \left(\frac{\frac{1}{4} 23 - 3}{6} \right) 5 \\ &= 7,5 + \left(\frac{5,75 - 3}{6} \right) 5 \\ &= 9,79 \end{aligned}$$

Kuartil 2 adalah sebuah nilai dalam distribusi yang membagi 50% frekuensi di atas distribusi dan 50% frekuensi dibawah distribusi. Kuartil 2 hakikatnya sama dengan median.

Rumus:

$$K2 = B + \left(\frac{\frac{1}{2} N - f_k}{f} \right) i$$

Kuartil 3 adalah sebuah nilai dalam distribusi yang membagi 75% frekuensi dibagian bawah dan 25% frekuensi dibagian atas. Rumus untuk menghitung K3 adalah sebagai berikut.

$$K3 = B + \left(\frac{\frac{3}{4} N - f_k}{f} \right) i$$

B. Desil

Desil adalah sebuah indeks yang membagi distribusi data menjadi sepuluh bagian. Jika distribusi dibagi menjadi 10 bagian, maka diperlukan 9 titik batas desil, yaitu: D1, D2 sampai dengan D9.

Dasar perhitungan desil adalah persepuluhan, sebagai contoh $D1 = 1/10N$, $D2 = 2/10N$ demikian juga seterusnya. Rumus yang digunakan juga tak jauh beda dengan rumus kuartil. Hanya saja pada kuartil menggunakan per-empatan, sementara pada desil menggunakan persepuluhan. Sebagai contoh, rumus mencari D1 adalah sebagai berikut:

$$D1 = B + \left(\frac{1/10 N - f_k}{f} \right) i$$

Sebagai contoh, dengan menggunakan data pada tabel 4.1, kita mencoba mencari D3. Maka D3 dapat dihitung sebagai berikut.

$$D3 = B + \left(\frac{3/10 N - f_k}{f} \right) i$$

$$D3 = 7,5 + \left(\frac{3/10 \cdot 23 - 3}{6} \right) 5$$

$$= 7,5 + \left(\frac{6,69 - 3}{6} \right) 5$$

$$= \mathbf{10,75}$$

C. Persentil

Persentil adalah sebuah indeks yang membagi distribusi data menjadi seratus bagian. Jika distribusi dibagi menjadi 100 bagian, maka diperlukan 99 titik batas persentil, yaitu: P1, P2 sampai dengan P99.

Dasar perhitungan persentil adalah per-seratusan, sebagai contoh $P1 = 1/100N$, $P2 = 2/100N$ demikian juga seterusnya. Rumus yang digunakan juga tak jauh beda dengan rumus kuartil maupun desil. Hanya saja pada kuartil menggunakan per-empatan, pada desil menggunakan persepuluhan, sementara pada persentil menggunakan per-seratusan. Sebagai contoh, rumus mencari P1 adalah sebagai berikut:

$$P1 = B + \left(\frac{1/100 N - f_k}{f} \right) i$$

Melalui persentil, kita dapat leluasa membagi distribusi data ke dalam jumlah-jumlah yang dikehendaki. Misalnya seorang peneliti ingin membagi distribusi data tentang motivasi berprestasi

menjadi 5 kategori (misalnya: tinggi sekali, tinggi, sedang, rendah, dan rendah sekali), maka peneliti harus menemukan 4 titik persentil dengan jalan melakukan pembagian, $100/5 = 20$.

Angka 20 nantinya akan berfungsi sebagai kelipatan yang digunakan untuk menentukan dasar pembuatan kategori. Sehingga dengan demikian, 5 kategori yang diinginkan tersebut akan dibatasi oleh titik-titik P20, P40, P60, dan P80.

Persentil	Skor	Kategori
P80 -----	(-----)	Sangat Tinggi -----
P60 -----	(-----)	Tinggi -----
P40 -----	(-----)	Sedang -----
P20 -----	(-----)	Rendah -----
		Sangat Rendah

Contoh, menghitung persentil 60

Tabel 4.2: Tabel untuk Menghitung Persentil

Klas Interval	f	Σfk
28 – 32	5	23
23 - 27	2	18
<u>18 - 22</u>	<u>4</u>	16
13 – 17	3	<u>12</u>
8 – 12	6	9
3 - 7	3	3
Jumlah	23	-

Dengan data tabel 4.2 tersebut, jika dicari P60-nya adalah sebagai berikut.

$$\begin{aligned}
 P60 &= B + \left(\frac{60/100N - fk}{f} \right) i \\
 &= 17,5 + \left(\frac{60/100.23 - 12}{4} \right) 5
 \end{aligned}$$

$$\begin{aligned}
 &= 17,5 + \left(\frac{13,8 - 12}{4} \right) 5 \\
 &= \mathbf{19,75}
 \end{aligned}$$

Dari hasil tersebut dapat dikatakan bahwa $P_{60} = 19,75$. Artinya, skor yang membatasi 60% distribusi bagian bawah dengan 40% bagian atas adalah 19,75.

Soal untuk dikerjakan dalam kelompok:

Dengan menggunakan data pada tabel 4.2, carilah P20, P40, dan P80.

Bab 5 Variabilitas

Variabilitas adalah sebuah ukuran tentang derajat penyebaran nilai variabel dari suatu tendensi sentral dalam sebuah distribusi. Mungkin saja, dua kelompok data memiliki tendensi sentral yang sama, tetapi derajat penyebarannya bisa jadi berbeda. Sebagai contoh,

Kelompok data I: 4 4 5 5 5 6 6

Kelompok data II: 2 3 4 5 6 7 8

Kelompok data pertama memiliki mean $35/7 = 5$, demikian juga kelompok data kedua memiliki mean $35/7 = 5$. Nilai tendensi sentral kedua kelompok data tersebut adalah sama, yaitu 5. Akan tetapi bila dilihat dari keragaman dan penyebaran nilai dari kedua kelompok data tersebut nampak sangat berbeda. Penyebaran nilai kelompok I lebih homogen dibanding distribusi kelompok data kelompok II. Atau dengan kata lain, penyebaran nilai kelompok data II lebih beragam atau heterogen dibanding penyebaran nilai kelompok data I.

Keragaman dan penyebaran nilai dalam suatu distribusi dapat dihitung melalui: **range, mean deviasi, standar deviasi, dan varian.**

Range

Adalah jarak antara nilai tertinggi dengan nilai terendah. Rumus yang digunakan untuk menghitung range (R) adalah sebagai berikut.

$$R = (X_t - X_r) + 1 \quad \text{dimana,}$$

X_t = nilai tertinggi
 X_r = nilai terendah

Dengan menggunakan data di atas, kita dapat mencari range-nya dengan cara sebagai berikut.

Untuk kelompok data I, nilai tertinggi adalah 6, nilai terendah 4, maka range-nya adalah $(6 - 4) + 1 = 3$. Untuk kelompok data II, nilai tertinggi adalah 8, nilai terendah 2, maka range-nya adalah $(8 - 2) + 1 = 7$.

Range yang diperoleh kedua kelompok data tersebut tidak sama. Kelompok data I range-nya adalah 3, sementara kelompok data II adalah 7. Besar kecilnya range dapat digunakan sebagai

petunjuk untuk mengetahui taraf keragaman dan variabilitas suatu distribusi. Semakin tinggi range, berarti semakin beragam distribusi datanya.

Perlu dicatat di sini bahwa nilai range sangat bergantung kepada data yang ekstrim, yaitu sebuah data yang kemunculan dan ketidakmunculannya sangat berpengaruh pada tinggi rendahnya range). Karena itu, penggunaan range perlu dipahami secara hati-hati.

Mean Deviasi

Adalah skor rata-rata dari suatu penyimpangan nilai terhadap mean kelompok nilai tersebut dalam sebuah distribusi.

Rumus mencari mean deviasi (MD) adalah sebagai berikut.

$$MD = \frac{\sum d}{N} \quad \text{dimana,} \quad \begin{array}{l} \sum d = \text{jumlah deviasi} \\ N = \text{jumlah individu} \end{array}$$

Untuk mencari deviasi, masing-masing nilai dikurangi mean. Deviasi yang bertanda plus (+) menunjukkan deviasi di atas mean, sedangkan yang bertanda minus (-) berarti di bawah mean. Contoh,

Tabel 5.1: Tabel untuk menghitung mean deviasi

Nilai	Deviasi (d)
8	+3
7	+2
6	+1
5	0
4	-1
3	-2
2	-3
35	12

Untuk menghitung mean deviasi, tanda minus (-) pada deviasi diabaikan, artinya semua skor dianggap positif. Dengan demikian, pada data di atas mean deviasinya adalah $12/7 = 1,71$.

Standar Deviasi

Standar Deviasi adalah penyimpangan suatu nilai dari mean. Standar deviasi merupakan akar dari jumlah deviasi kuadrat dibagi banyaknya individu dalam distribusi. Rumus standar deviasi (SD) adalah sebagai berikut.

Rumus 1

$$SD = \sqrt{\frac{\sum d^2}{N}}$$

Dengan menggunakan data di atas, maka untuk mencari standar deviasi adalah sebagai berikut.

Nilai	Deviasi	$\sum d^2$
8	+3	9
7	+2	4
6	+1	1
5	0	0
4	-1	1
3	-2	4
2	-3	9
35	0	28

Maka,

$$\begin{aligned} SD &= \sqrt{\frac{28}{7}} \\ &= \sqrt{4} \\ &= 2 \end{aligned}$$

Rumus di atas hanya dapat digunakan untuk menghitung SD pada kelompok data dengan frekuensi tunggal. Bagaimana jika datanya frekuensinya bervariasi? Rumusnya adalah sebagai berikut.

Rumus 2

$$SD = \sqrt{\frac{\sum fd^2}{N}}$$

Contoh,

Skor	f	d	d ²	fd ²
60	2	30	900	1800
50	3	20	400	1200
40	1	10	100	100
30	2	0	0	0
20	5	-10	100	500
10	4	-20	400	1600
N=17		-	$\sum fd^2 = 5200$	

$$SD = \sqrt{\frac{\sum fd^2}{N}}$$

$$SD = \sqrt{\frac{5200}{17}}$$

$$SD = \sqrt{305,88}$$

$$= 17,49$$

Apabila datanya dalam bentuk distribusi frekuensi berkelompok, maka dapat digunakan rumus sebagai berikut.

Rumus 3

$$SD = \sqrt{\frac{\sum fX^2}{N} - \left| \frac{\sum fX}{N} \right|^2}$$

Contoh,

Interval	f	X (Titik Tengah)	ΣfX	ΣfX^2
33 – 39	2	36	72	2592
26 – 32	8	29	232	6728
19 – 25	19	22	418	9196
12 – 18	20	15	300	4500
5 - 11	11	8	88	704
Jumlah	60	-	1110	23.720

$$SD = \sqrt{\frac{\Sigma fX^2}{N} - \left| \frac{\Sigma fX}{N} \right|^2}$$

$$SD = \sqrt{\frac{23.720}{60} - \left| \frac{1110}{60} \right|^2}$$

$$SD = \sqrt{395,33 - 342,25}$$

$$SD = \sqrt{53,08}$$

$$= 7,28$$

Varian (s)

Varian adalah angka yang menunjukkan ukuran variabilitas yang dihitung dengan jalan mengkuadratkan standar deviasi. Jadi jika SD sudah diketahui, maka untuk mencari varian tinggal mengkuadratkan saja.

Akan tetapi jika standar deviasi belum diketahui, maka untuk mencari varian perlu rumus tersendiri. Rumus varian diperoleh dari dasar-dasar penghitungan standar deviasi, sebagai berikut.

$$s = \frac{\Sigma fX^2}{N} - \left[\frac{\Sigma fX}{N} \right]^2$$

Tugas untuk didiskusikan dalam kelompok:

Carilah:

- (1) range
- (2) mean deviasi
- (3) standar deviasi, dan
- (4) varian

dari data berikut.

24	12	9	12	13	10
29	19	15	15	10	14
30	16	14	13	20	12
17	17	9	14	17	15
18	22	23	23	25	27

Bab 6 Probabilitas dan Sampel

Dalam kehidupan nyata, tidak ada sesuatu yang pasti, semuanya bersifat relatif. Yang mutlak dan pasti hanyalah Tuhan Sang Maha Pencipta. Apakah seseorang yang bergelimang harta dapat dipastikan memperoleh kebahagiaan? Apakah seorang anak yang lahir dari keluarga kaya pada akhirnya juga akan menjadi kaya sebagaimana orangtua yang melahirkannya? Apakah orang sakit yang berobat ke dokter, meski sehebat apapun dokter tersebut, dapat dipastikan akan mendapatkan kesembuhan? Apakah tim yang bertabur bintang dapat dipastikan ia akan berhasil memenangkan pertandingan? Apakah atlet unggulan pertama dapat dipastikan akan menjadi pemenang dalam suatu kejuaraan? Tidak mudah untuk menjawab persoalan tersebut. Untuk mudahnya, kita bisa menjawab: belum tentu. Hal yang kurang lebih sama juga terjadi pada tataran ilmu pengetahuan. Karena itu kemudian muncul konsep probabilitas, yakni peluang terjadinya suatu peristiwa: apakah suatu peristiwa itu berpeluang terjadi atau tidak terjadi. Seseorang yang berkecukupan secara ekonomi lebih berpeluang mendapatkan kebahagiaan daripada orang yang serba kekurangan. Meski harta bukan jaminan kebahagiaan, tetapi tanpa harta orang sulit untuk bahagia. Meski sakit tanpa diobati pun ada kemungkinan sembuh, tetapi kalau berobat ke dokter, maka peluang kesembuhannya lebih tinggi. Demikian pula, atlet unggulan memiliki peluang lebih tinggi untuk menjadi juara.

Kalau kita melempar koin logam, maka kemungkinan yang muncul adalah sisi gambar atau angka. Kedua sisi tersebut bersifat eksklusif, artinya jika yang muncul gambar, maka sisi angka tidak akan muncul. Kemunculan yang satu akan meniadakan kemunculan yang lain. Ujung dari suatu permainan adalah kalah atau menang, adakalanya atlet menang dan adakalanya kalah. Peluang terjadinya suatu peristiwa yang demikian dapat diformulasikan sebagai berikut:

$$P = \frac{n}{N}$$

dimana,
P : peluang terjadinya peristiwa
n : frekuensi kejadian
N: keseluruhan peristiwa

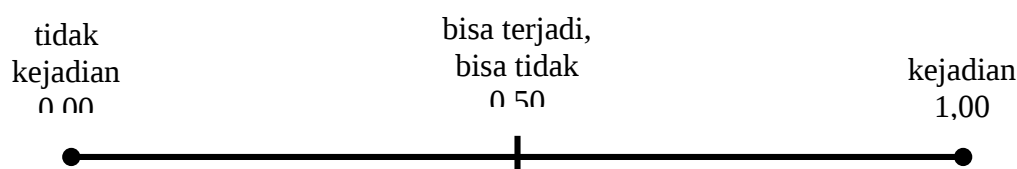
Kalau kita melempar koin logam, keseluruhan peristiwa (N) = 2, munculnya sisi gambar atau angka (n) = 1, maka peluang munculnya sisi gambar atau angka adalah $\frac{1}{2}$ atau 0,5. Demikian pula jika kita melempar dadu dengan 6 sisi (sisi 1, 2, 3, 4, 5, dan 6), keseluruhan peristiwa (N) = 6, munculnya sisi 1 atau yang lain (n) adalah 1, sehingga peluang munculnya sisi 1 atau yang lain adalah $\frac{1}{6}$ atau 0,17.

Pengertian di atas dapat diperluas dalam konteks yang lebih aplikatif. Misalnya, dari 100 atlet bulutangkis yang masuk pemusatan latihan nasional (pelatnas), terdapat 10 atlet yang menjadi juara dunia. Maka, dapat dikatakan peluang atlet pelatnas menjadi juara dunia adalah $10/100$ atau $0,10$ (10%) dan pada saat yang sama bisa dikatakan bahwa peluang kegagalan atlet pelatnas menjadi juara dunia sebesar 90%. Contoh lain, ada dua atlet bebuyutan saling mengalahkan, dari 12 kali pertemuan, atlet A memenangkan 7 kali pertandingan, sementara atlet B sebanyak 5 kali. Maka, peluang kemenangan atlet A adalah $7/12$ atau $0,58$, sementara atlet B adalah $5/12$ atau $0,42$. Fakta yang berbeda terjadi pada seorang guru pendidikan jasmani, yang mengajar 400 siswa. Pada akhir semester, dari 400 siswa tersebut, terdapat 6 siswa yang tidak lulus atau tidak mencapai standar kompetensi kelulusan. Maka, peluang ketidaklulusan siswa adalah $6/400$ atau $0,015$ (2%) atau pada saat yang sama bisa dikatakan bahwa peluang kelulusan mencapai 98%.

Dari penjelasan di atas, maka dapat disimpulkan bahwa jika angkanya mendekati 0, maka peluang keterjadiannya semakin kecil, sementara jika angkanya mendekati 1, maka peluang keterjadiannya semakin besar, sehingga dapat diformulasikan sebagai berikut:

$$0 \leq P \leq 1$$

atau bila digambarkan dalam bentuk garis kontinum tampak sebagai berikut;



Angka 0 berarti tidak kejadian (TK), sementara angka 1 berarti kejadian (K). Dengan demikian didapat;

$$P(K) = 1 - P(TK) \text{ atau berlaku hubungan;}$$

$$P(K) + P(TK) = 1$$

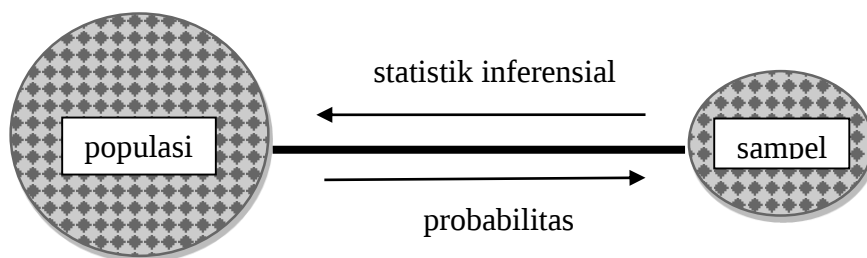
Sebagai contoh, sebuah kotak berisi 30 bola, terdiri dari 10 bola berwarna merah (M), 8 bola berwarna putih (P), dan 12 bola berwarna hijau (H). Seorang anak dengan mata tertutup diminta mengambil sebuah bola secara acak. Pertanyaannya, berapa peluang bola berwarna merah atau putih atau hijau terambil?

$$P(M) = \frac{10}{10 + 8 + 12} = 0,33$$

$$P(P) = \frac{8}{10 + 8 + 12} = 0,26$$

$$P(H) = \frac{12}{10 + 8 + 12} = 0,40$$

Dari penghitungan di atas, dapat disimpulkan bahwa dari seratus kali pengambilan bola, maka peluang terambil bola merah sebanyak 33 buah, bola putih sebanyak 26 buah, dan bola hijau sebanyak 40 buah. Lalu, bagaimana kaitan antara probabilitas dan sampel? Seperti diketahui bahwa sampel pada dasarnya diambil dari populasi, apalagi yang berjenis *probability sampling*. Setiap anggota populasi memiliki peluang yang sama untuk menjadi sampel. Sebagai contoh, di suatu sekolah terdapat 100 siswa pada kelas yang sama. Ketika seorang peneliti ingin mengambil sampel secara random dari 100 siswa tersebut, maka setiap siswa akan memiliki peluang yang sama untuk menjadi sampel. Peluang setiap siswa untuk menjadi sampel adalah 0,01 atau 1%. Dalam istilah statistik, probabilitas disimbolkan dengan p , sehingga dalam kasus di atas bisa ditulis menjadi $p = 0.01$.



Gambar 6.1
Hubungan antara probabilitas dan populasi-sampel

Relasi berikutnya terjadi pada pengambilan kesimpulan yang terjadi pada tataran sampel untuk selanjutnya digeneralisasikan pada tataran populasi. Inilah kemudian yang disebut dengan statistik inferensial. Ketika suatu kesimpulan diambil dari sampel dan selanjutnya diberlakukan pada populasi, maka dapat terjadi perbedaan akibat kesalahan sampling (*sampling error*). Kesalahan tersebut logis dan wajar, mengingat tidak pernah terjadi antara sampel dan populasi adalah identik

atau sama persis. Dalam kaitan ini, besarnya kesalahan yang dapat ditoleransi adalah di bawah 0.05 atau 5%, yang kemudian sering diformulasikan menjadi $p \leq 0.05$. Hal ini mengandung pengertian bahwa dari seratus kejadian, maka terdapat paling banyak 5 kejadian yang salah. Dari 100 orang sakit yang berobat ke dokter, maka paling banyak 5 orang yang sakitnya tidak sembuh. Standar kesalahan tidak pernah dapat diukur dengan pasti, tetapi melalui estimasi dengan rumus berikut.

$$SEM = \frac{SD}{\sqrt{n-1}}$$

Kita dapat menggunakan SEM untuk menunjukkan batas tingkat kepercayaan.

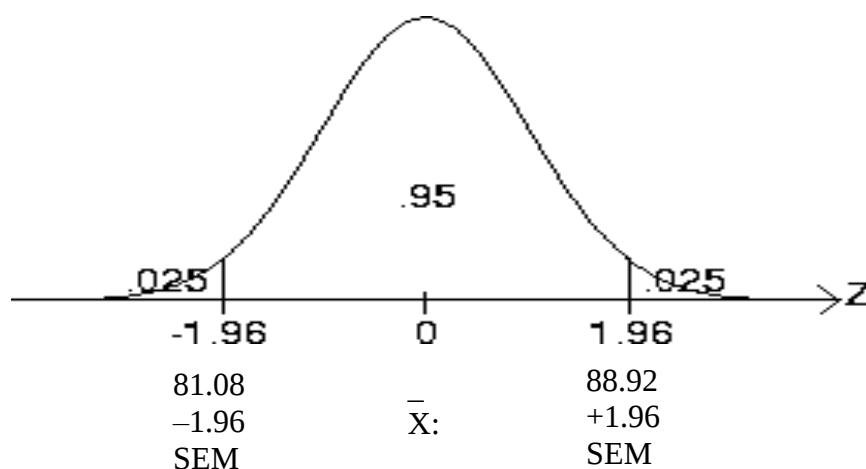
Tingkat Kepercayaan	Nilai Z
90%	1.65
95%	1.96
99%	2.58
99,9%	3.291

Sebagai ilustrasi, seorang guru memiliki 40 skor hasil ujian bidang studi pendidikan jasmani. Dari data tersebut kemudian diperoleh rerata sebesar 85 dan SEM sebesar 2, dengan tingkat kepercayaan 95%, maka rerata populasi adalah $85 \pm 1.96 (2)$

$$= 85 \pm 3.92$$

$$= 81.08 - 88.92$$

Mengapa ± 1.96 ? Karena area antara plus dan minus 1.96 nilai z sama dengan 95% total area dalam kurva normal.



Gambar 6.2
Taraf signifikansi 0,05 dalam kurva normal

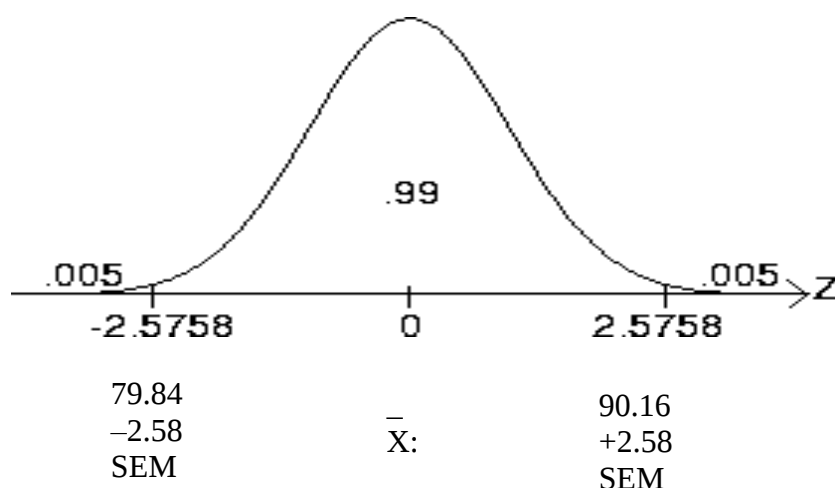
Tingkat signifikansi 5% lazim digunakan dalam ilmu-ilmu sosial, sementara ilmu-ilmu eksakta biasanya menggunakan tingkat signifikansi 1%. Meski hal tersebut bukan sebuah keniscayaan, tetapi dapat dipahami. Hal ini mengingat dalam ilmu-ilmu sosial, apalagi yang melibatkan manusia, keterjadian sesuatu melibatkan variabel yang sangat kompleks. Tidak serta merta anak-anak yang berasal dari keluarga kaya, status sosial tinggi, dan serba berkecukupan memiliki prestasi belajar yang tinggi. Sangat boleh jadi hal yang demikian justru membuat anak menjadi manja dan tidak perlu bekerja keras. Bukankah semuanya sudah tersedia? Sebaliknya, dalam ilmu eksakta, variabel-variabel lebih dapat dikendalikan. Karena obyeknya benda-benda alam dengan hukum yang tingkat kepastiannya tinggi, maka kapan pun dan di mana pun keberlakukannya akan terjadi.

Bagaimana daerah penerimaan hipotesis jika kita menggunakan tingkat kepercayaan 99%? Pada prinsipnya sama, dengan rerata sampel sebesar 85 dan SEM 2, maka diperoleh 85 ± 2.58 (SEM)

$$= 85 \pm 2.58 (2)$$

$$= 85 \pm 5.16$$

$$= 79.84 - 90.16$$



Gambar 6.3
Tarf signifikansi 0,01 dalam kurva normal

Dalam taraf signifikansi, baik 5% maupun 1%, ada daerah kritis yang digunakan untuk menerima atau menolak hipotesis nol, sebagaimana tampak pada gambar di atas. Pengujian hipotesis dikatakan signifikan apabila keterjadian tersebut hampir tidak mungkin disebabkan oleh

faktor kebetulan, sesuai dengan probabilitas yang telah ditentukan sebelumnya. Seseorang yang memberikan pupuk agar pohon mangganya berbuah harus dapat memastikan bahwa berbuahnya pohon mangga tersebut karena pupuk yang diberikan, bukan kebetulan musim, yang tanpa pupuk pun mangga tersebut akan berbuah. Seorang atlet memenangkan kejuaraan karena kesiapan fisik, mental, dan strategi yang dilakukan selama latihan, bukan karena kebetulan lawannya lemah atau sedang sakit.

Bab 7 Skor Standar

Seorang peneliti, termasuk seorang guru yang ingin mengevaluasi peserta didiknya, biasanya tertarik membandingkan antara skor individu siswa yang satu dengan individu siswa yang lain dalam kelompok. Ketika seorang siswa mendapatkan skor tertentu dalam suatu ujian, tidak serta merta kita dapat menyimpulkan bahwa skor tersebut baik atau buruk, lebih tinggi atau lebih rendah. Seorang mahasiswa yang mendapat skor 70 dalam matakuliah renang, belum tentu lebih rendah dari nilai 80 dalam matakuliah bolavoli. Apalagi jika mahasiswa diminta untuk mengerjakan soal. Menjawab soal benar 7 dari 10 soal, tentu berbeda dengan menjawab benar 7 dari 20 soal. Skor standar juga dibutuhkan jika kita memiliki data dengan satuan unit yang berbeda. Misalnya kemampuan *shooting* dalam bolabasket yang dihitung berdasarkan masuknya bola dalam keranjang, semakin banyak semakin baik. Pada saat yang sama, kita juga memiliki data tentang kecepatan lari 400 meter, yang diukur dengan waktu. Semakin cepat semakin bagus, semakin kecil waktunya semakin baik. Kedua hal tersebut, yakni banyaknya lemparan yang masuk dan kecepatan lari, memiliki satuan yang berbeda. Artinya, kedua data tersebut tidak bisa langsung dijumlahkan, apalagi dilakukan pemeringkatan. Untuk sampai pada keputusan tersebut, dibutuhkan upaya standardisasi skor. Dalam terminologi statistik, ada dua tipe skor standar, yakni Z-skor dan T-skor.

Z-Skor

Skor standar (Z) adalah suatu bilangan yang menunjukkan seberapa jauh suatu skor menyimpang dari mean (M) dalam satuan standar deviasi (SD). Rumusnya adalah sebagai berikut.

$$Z = \frac{X - M}{SD} \quad \text{dimana,} \quad \begin{array}{l} X : \text{skor individu} \\ M : \text{mean} \\ SD : \text{standar deviasi} \end{array}$$

Skor standar berfungsi untuk memberikan satuan ukuran yang baku dan sebagai ukuran untuk membandingkan dua gejala atau lebih. Sebagai contoh, seorang mahasiswa mendapat nilai 90 dalam matakuliah metodologi penelitian. Rerata dari distribusi skor metodologi penelitian adalah 70, dan nilai standar deviasi sebesar 10. Maka skor standar mahasiswa tersebut dapat dicari dengan menggunakan rumus di atas.

$$Z = \frac{X - M}{SD}$$

$$\begin{aligned}
&= \frac{90 - 70}{10} \\
&= +2
\end{aligned}$$

Ini berarti bahwa nilai matakuliah metodologi penelitian tersebut berada di atas mean sebanyak 2 SD, atau berada 2 SD di atas mean.

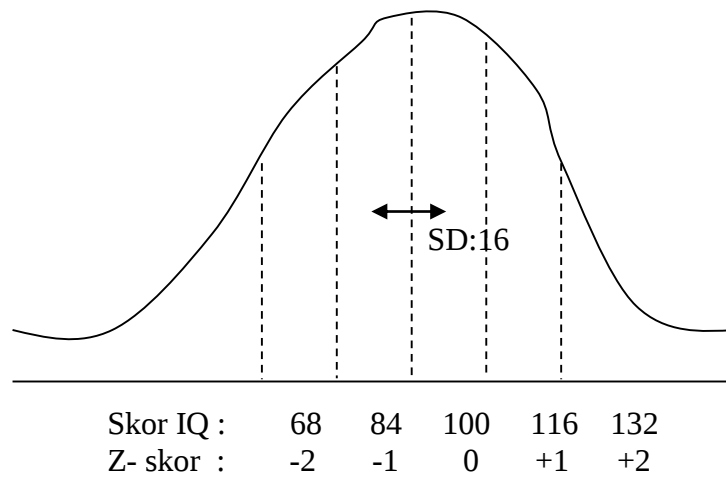
Coba kita kembangkan lagi, misalnya dalam matakuliah bolavoli seorang mahasiswa memperoleh nilai 70 dan matakuliah sepakbola memperoleh nilai 40. Dari skor tersebut, apakah dapat disimpulkan bahwa mahasiswa tersebut lebih memiliki kecakapan dibidang bolavoli daripada sepakbola? Belum tentu, marilah kita lihat lebih dalam. Pertama harus diketahui lebih dulu mean dan standar deviasi dari kedua matakuliah tersebut. Misalnya untuk matakuliah Bolavoli $M = 80$ dan $SD = 10$; sementara untuk matakuliah sepakbola $M = 30$ dan $SD = 5$, maka kita dapat membandingkannya sebagai berikut:

$$\begin{aligned}
&\text{Matakuliah bolavoli:} \\
&Z = \frac{X - M}{SD} \\
&= \frac{70 - 80}{10} \\
&= -1
\end{aligned}$$

$$\begin{aligned}
&\text{Matakuliah Sepakbola:} \\
&Z = \frac{X - M}{SD} \\
&= \frac{40 - 30}{5} \\
&= +2
\end{aligned}$$

Dari perhitungan skor standar tersebut dapat disimpulkan bahwa kemampuan mahasiswa pada matakuliah bolavoli berada 1 SD di bawah mean; sedangkan pada matakuliah sepakbola nilainya berada 2 SD di atas mean. Ini berarti bahwa mahasiswa tersebut lebih mampu pada matakuliah sepakbola daripada matakuliah bolavoli. Inilah pentingnya bagi para peneliti untuk berhati-hati dalam membandingkan suatu nilai, kelihatannya skornya lebih besar, tetapi setelah diteliti lebih dalam berlaku sebaliknya.

Contoh lain, misalnya seorang psikolog telah melakukan pengukuran IQ terhadap lima orang anak dan hasilnya adalah sebagai berikut: 68, 84, 100, 116, dan 132. Dari data tersebut ditemukan $M = 100$ dan $SD = 16$. Jika skor-skor tersebut ditempatkan dalam distribusi, maka akan tampak seperti di bawah ini.



Kita dapat mengolah Z-skor dengan bantuan SPSS. Sebagai ilustrasi, perhatikan data berikut.

Subyek	Skor MK Renang	Skor MK Bolavoli
A	50	65
B	65	78
C	55	85
D	70	80
E	67	75
F	54	90
G	58	66
H	60	78
I	64	84
J	75	82

Selanjutnya, kita lakukan standarisasi skor dengan langkah-langkah sebagai berikut.

- Entry data atau buka file data yang akan dianalisis
- Pilih menu berikut ini

Analyze

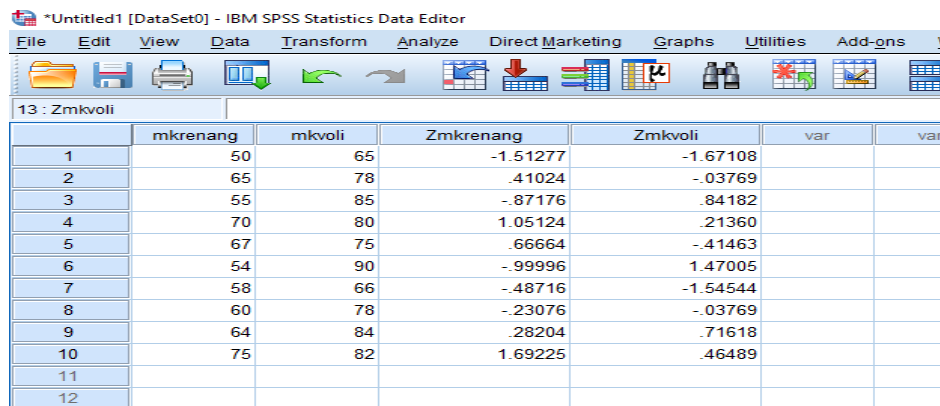
Descriptives Statistics

Descriptive

Masukkan variabel ke dalam kotak dan centang pilihan “save standardized values as variables”, maka akan diperoleh hasil sebagai berikut.

Descriptive Statistics

	N	Minimum	Maximum	Mean	Std. Deviation
Mkrenang	10	50	75	61.80	7.800
Mkvoli	10	65	90	78.30	7.959
Valid N (listwise)	10				



The screenshot shows the IBM SPSS Statistics Data Editor window. The active dataset is '13 : Zmkvoli'. It displays two columns of original data: 'mkrenang' and 'mkvoli'. Below these, there are columns for the calculated Z-scores, labeled 'Zmkrenang' and 'Zmkvoli'. The Z-scores are standardized, with a mean of 0 and a standard deviation of 1. The table also includes columns for variance ('var') for each variable.

	mkrenang	mkvoli	Zmkrenang	Zmkvoli	var	var
1	50	65	-1.51277	-1.67108		
2	65	78	.41024	-.03769		
3	55	85	-.87176	.84182		
4	70	80	1.05124	.21360		
5	67	75	.66664	-.41463		
6	54	90	-.99996	1.47005		
7	58	66	-.48716	-1.54544		
8	60	78	-.23076	-.03769		
9	64	84	.28204	.71618		
10	75	82	1.69225	.46489		
11						
12						

Dalam pengolahan tersebut muncul variabel baru, hasil proses Z skor secara otomatis. Data baru dengan Z skor tersebut memiliki rerata = 0 dan SD = 1, skor berkisar antara -3 dan +3.

Bagaimana jika datanya terkait dengan waktu seperti kecepatan lari dan kecepatan renang? Apabila kita memiliki data yang demikian, maka yang perlu dilakukan adalah *me-recode* atau mengubah tanda Z-skor dari positif ke negatif atau sebaliknya, misalnya dari 2 menjadi -2, dari -1,5 menjadi 1,5, dan seterusnya.

T-Skor

Acapkali orang ingin menghindari bilangan negatif atau desimal yang dihasilkan lewat perhitungan Z-skor. Statistik memberikan jalan keluar, yaitu dengan cara mentransformasikannya ke dalam T-skor dengan rumus sebagai berikut.

Rumus 1 **T-skor = 50 + 10z**

dimana bilangan 50 dan 10 adalah konstan. Atau bisa juga dihitung langsung berdasarkan data, yaitu:

Rumus 2
$$\text{T-skor} = 50 + \frac{(X - M)}{SD} \times 10$$

Contoh, sekelompok data diketahui memiliki $M = 80$ dan $SD = 10$. Tentukan T-skor untuk nilai 60. Maka dengan menggunakan rumus 2 dapat dihitung sebagai berikut.

$$\begin{aligned}
 T\text{-skor} &= 50 + \frac{(60 - 80)}{10} \times 10 \\
 &= 50 + \frac{-20}{10} \times 10 \\
 &= 50 + (-2) \times 10 \\
 &= 50 - 20 \\
 &= \mathbf{30}
 \end{aligned}$$

Selanjutnya, kita juga dapat melakukan standarisasi skor menggunakan SPSS dengan langkah-langkah sebagai berikut.

- Entry data atau buka file data yang akan dianalisis
- Pilih menu berikut ini

Transform

Compute Variable

Pada kolom “target variable” masukkan nama variabel baru yang akan dibuat. Pada kolom “numeric expression” masukkan rumus T skor, yaitu $50 + 10 \times \text{zmkrenang}$. Hal yang sama juga dilakukan untuk matakuliah bolavoli, maka akan diperoleh hasil sebagai berikut.

11 :

	mkrenang	mkvoli	Zmkrenang	Zmkvoli	NArenang	NAvoli	var
1	50	65	-1.51277	-1.67108	34.87	33.29	
2	65	78	.41024	-.03769	54.10	49.62	
3	55	85	-.87176	.84182	41.28	58.42	
4	70	80	1.05124	.21360	60.51	52.14	
5	67	75	.66664	-.41463	56.67	45.85	
6	54	90	-.99996	1.47005	40.00	64.70	
7	58	66	-.48716	-1.54544	45.13	34.55	
8	60	78	-.23076	-.03769	47.69	49.62	
9	64	84	.28204	.71618	52.82	57.16	
10	75	82	1.69225	.46489	66.92	54.65	
11							

Soal untuk dikerjakan dan didiskusikan dalam kelompok:

Tentukan Mean, Standar Deviasi, Z-skor, dan T-skor dari data berikut:

Subjek	Hasil Tes		Z-Skor		T-Skor	
	Lari 60m	MFT	Lari 60m	MFT	Lari 60m	MFT
1	9,66	28,3				
2	9,39	37,4				
3	10,7	23,9				
4	13,9	23,6				
5	14,9	23,6				
6	11	36,7				
7	10,3	33,6				
8	9,82	31,4				
9	12,6	28,9				
10	12,9	24,6				

Bab 8 Pengujian Hipotesis

Logika uji statistik

Pengambilan keputusan statistik selalu bicara dalam perspektif agregat data dan kecenderungan umum, bukan kasus tunggal atau spesifik. Sebagai contoh, keputusan statistik menyatakan bahwa tingkat pendidikan berkorelasi positif terhadap kesuksesan seseorang, semakin tinggi pendidikan individu semakin tinggi pula peluang keberhasilan individu tersebut. Pernyataan ini tidak bisa kemudian dibantah dengan mengatakan bahwa di kampung saya ada seseorang yang tidak lulus SD tetapi kaya raya. Bisa jadi tidak ada yang salah terhadap pernyataan tersebut, tetapi tidak berarti fakta tersebut menggugurkan pernyataan sebelumnya. Adanya fakta tersebut juga tidak bisa serta-merta dapat disimpulkan bahwa pendidikan tidak perlu, toh seseorang yang tidak lulus SD pun bisa berhasil. Demikian juga keputusan statistik yang menyatakan bahwa orang yang aktif berolahraga akan lebih sehat dan pada gilirannya memperpanjang angka harapan hidup. Sekali lagi, statistik ada pada tataran kecenderungan umum. Secara umum, individu yang berpendidikan tinggi memiliki peluang kesuksesan yang lebih tinggi dibanding mereka yang tidak berpendidikan tinggi, individu yang aktif berolahraga memiliki harapan hidup yang lebih panjang dibanding mereka yang tidak aktif.

Pengujian statistik pada dasarnya adalah pengujian atas hipotesis yang diajukan oleh peneliti, dalam bentuk hipotesis nol (H_0 atau H_{nihil}) dan hipotesis kerja (H_1 atau $H_{\text{alternatif}}$). Kalimat hipotesis nol mengandung pernyataan “tidak ada hubungan”, “tidak ada pengaruh”, atau “tidak ada perbedaan”. Misalnya intelegensi tidak berhubungan dengan kemampuan siswa dalam bermain bulutangkis, pola asuh orangtua tidak berpengaruh terhadap pola aktivitas fisik anak, dan siswa yang mengikuti ekstrakurikuler olahraga tingkat kebugarannya tidak berbeda dibanding siswa yang tidak mengikuti ekstrakurikuler olahraga. Sementara hipotesis kerja mengandung pernyataan “ada hubungan”, “ada pengaruh”, atau “ada perbedaan”. Dari sejumlah contoh rumusan hipotesis nol di atas, jika diubah ke dalam hipotesis kerja, maka berbunyi: intelegensi berhubungan dengan kemampuan siswa dalam bermain bulutangkis, pola asuh orangtua berpengaruh terhadap pola aktivitas fisik anak, dan siswa yang mengikuti ekstrakurikuler olahraga tingkat kebugarannya berbeda/lebih baik dibanding siswa yang tidak mengikuti ekstrakurikuler olahraga. Dalam proses pengujian, kedua jenis hipotesis tersebut selalu berbanding terbalik. Jika hipotesis nol benar atau terbukti, maka hipotesis alternatif dengan sendirinya salah atau tidak terbukti.

Dalam perumusan, hipotesis nol perlu ditulis pertama, baru diikuti hipotesis kerja. Mengapa hipotesis nol? Ada beberapa argumentasi yang dapat dikemukakan, pertama, sesuatu harus dinyatakan nol, tidak ada hubungan, tidak ada pengaruh, atau tidak ada perbedaan sampai bisa

dibuktikan adanya hubungan atau perbedaan tersebut. Artinya, hipotesis nol harus dianggap benar sampai ditemukan fakta yang menyatakan sebaliknya. Logika ini analog dengan pengujian kasus hukum di pengadilan. Seseorang masih dinyatakan tidak bersalah sampai ada bukti yang sah dan meyakinkan bahwa yang bersangkutan bersalah. Alasan kedua, hipotesis nol bisa dikatakan sebagai pengganggu, yang meragukan kebenaran hipotesis kerja. Seorang peneliti biasanya berangkat dari titik keraguan (skeptis) sampai kemudian ditemukan sejumlah informasi yang meyakinkan melalui data-data yang dimiliki. Lalu, bagaimana posisi hipotesis alternatif? Dalam konteks pengujian, hipotesis alternatif berfungsi sebagai petunjuk bagi peneliti dalam bekerja dan menafsirkan hasil analisis data untuk kemudian sampai pada kesimpulan.

Tabel 8.1: Pernyataan hipotesis dan simbol statistik

	Hipotesis	Pernyataan Statistik	Pernyataan Kalimat
Perbedaan	H_0	$\mu^1 = \mu^2$	Tidak ada perbedaan...
	H_a	$\mu^1 > \mu^2$	A lebih besar/tinggi/baik daripada B
		$\mu^1 < \mu^2$	A lebih kecil/rendah daripada B
		$\mu^1 \neq \mu^2$	A tidak sama dengan B
Hubungan	H_0	$\rho = 0$	Tidak ada hubungan
	H_a	$\rho \neq 0$	Ada korelasi...
		$\rho > 0$	Korelasi positif...
		$\rho < 0$	Korelasi negatif...

Kesalahan pengujian

Pengujian hipotesis akan berujung pada dua opsi, menolak hipotesis nol atau menerima hipotesis nol. Keputusan tersebut belum dapat diketahui sampai proses pengujian selesai. Karena itu, dalam pengujian hipotesis juga terbuka peluang terjadinya kesalahan. Kesalahan tipe 1 terjadi apabila peneliti menolak hipotesis nol, tetapi dalam kenyataannya setelah melalui proses pengujian, hipotesis nol tersebut benar. Sementara itu, kesalahan tipe 2 terjadi apabila peneliti menerima hipotesis nol, tetapi pada kenyataannya hipotesis nol tersebut salah.

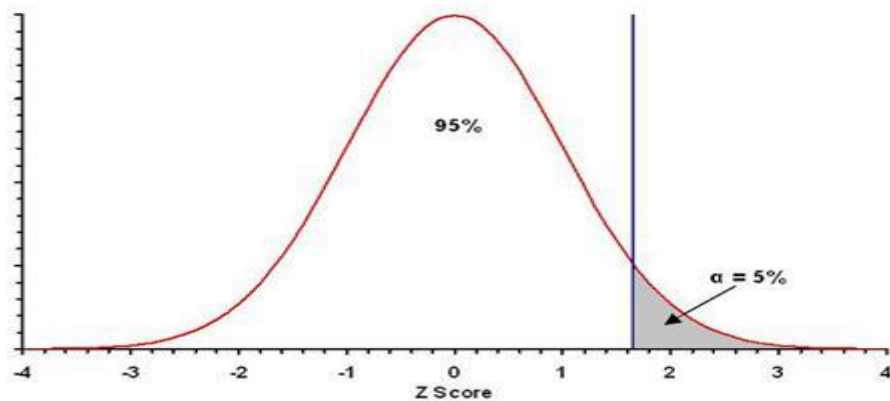
Tabel 8.2: Kesalahan Tipe 1 dan Tipe 2 dalam Pengujian Hipotesis

Klaim peneliti	Kenyataannya	
	H_0 benar	H_0 Salah
Menolak H_0	Kesalahan tipe 1	Keputusan benar
Menerima H_0	Keputusan benar	Kesalahan tipe 2

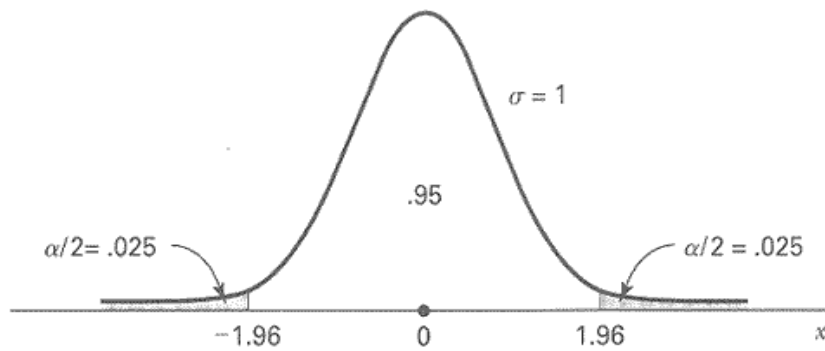
Pengujian satu arah dan dua arah

Dalam pengujian hipotesis, kita akan dihadapkan pada pilihan arah pengujian, yakni uji satu arah (*one tailed*) dan uji dua arah (*two tailed*). Uji satu arah diberlakukan apabila peneliti memiliki keyakinan yang kuat terkait dengan hipotesis yang akan diuji. Misalnya, seorang peneliti ingin mengkaji hubungan antara variabel X dan variabel Y. Apabila peneliti sudah punya keyakinan bahwa hubungan antara variabel X dan Y berarah positif atau negatif, maka peneliti dapat memutuskan menggunakan uji satu arah. Namun jika peneliti belum sampai pada keyakinan tersebut, maka dapat menggunakan uji dua arah. Contoh lain, seorang peneliti ingin menguji hipotesis: "metode A lebih efektif meningkatkan keterampilan motorik siswa dibanding metode B". Dengan pernyataan hipotesis tersebut, peneliti akan melakukan pengujian satu arah. Pertanyaannya, dari mana seorang peneliti memperoleh keyakinan tersebut? Jawabannya adalah dari kajian teori dan hasil-hasil penelitian terdahulu, yang lazimnya dipaparkan pada bab II.

Dalam pengujian satu arah, daerah penolakan hipotesis ada pada satu sisi distribusi (lihat gambar 8.1), bisa sebelah kiri atau kanan, tergantung pada rumusan hipotesis. Sebaliknya, jika pengujian bersifat dua arah, daerah penolakan hipotesis ada pada dua sisi, sebelah kiri dan kanan (lihat gambar 8.2).



Gambar 8.1: Daerah penolakan hipotesis untuk uji satu arah



Gambar 8.2: Daerah penolakan hipotesis untuk uji dua arah

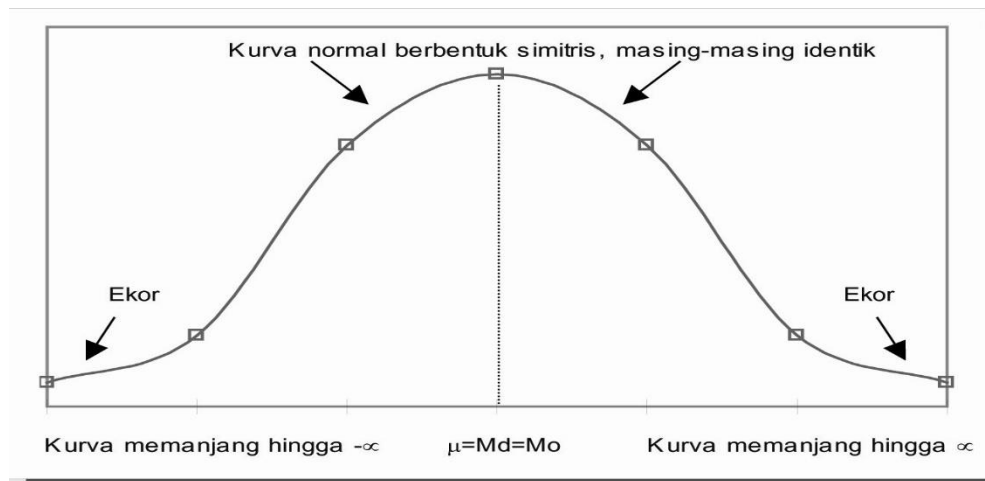
Uji persyaratan data

Sebelum kita melakukan pengujian hipotesis, terutama dalam penggunaan statistik parametrik, lazimnya dilakukan pengujian asumsi. Secara harfiah, asumsi (*assumption*) berarti *a statement accepted true without proof* atau bisa juga dimaknai sebagai *something taken for granted*. Apabila pemahaman ini juga berlaku pada pengertian asumsi dalam statistik, maka data yang akan dianalisis “dianggap” memenuhi asumsi seperti yang dipersyaratkan. Artinya, analisis dapat dilakukan tanpa harus melakukan pemeriksaan terlebih dahulu terhadap terpenuhi-tidaknya asumsi yang bersangkutan. Dalam statistik berlaku asumsi: jika sampel diambil secara random dan dalam jumlah besar (lebih dari 30 subjek), maka data akan cenderung normal dan homogen. Meski pada akhirnya data yang digunakan ternyata tidak sesuai dengan asumsi-asumsinya, tidak serta-merta hasil analisisnya dianggap salah. Hanya saja keberlakuan kesimpulan menjadi terbatas pada subjek sampel, tidak pada populasi.

Dalam pengujian persyaratan data, ada 3 hal yang umumnya dilakukan, yakni: uji normalitas, uji homogenitas, dan uji linieritas. Uji normalitas dan homogenitas digunakan untuk persyaratan uji beda, sementara uji linieritas digunakan untuk persyaratan uji hubungan. Selain itu, untuk analisis regresi dibutuhkan uji homoskedastisitas.

Uji Normalitas

Uji normalitas bertujuan untuk memastikan bahwa data yang diperoleh berdistribusi simetris atau normal, yakni sebaran angka sebagian besar ada di tengah, dan semakin ke kanan atau ke kiri, sebaran angka akan semakin kecil, sehingga menyerupai bel atau kurva. Pengujian normalitas bisa dilakukan dengan Chi-Square, Kolmogorof-Smirnov, dan Shapiro-Wilks.



Gambar 8.3: Gambar kurva normal

Sebagai ilustrasi, seorang peneliti memiliki data vo2max dari 40 orang dengan sebaran data sebagai berikut.

Tabel 8.3: Data Vo2max dari 40 orang coba

Kode Subjek	Vo2max	Kode Subjek	Vo2max	Kode Subjek	Vo2max	Kode Subjek	Vo2max
w1	30,20	p8	44,90	p15	38,90	w10	34,70
p1	43,90	w4	38,20	p16	40,50	w11	31,00
p2	42,00	p9	36,40	w7	32,60	w12	30,20
p3	40,50	p10	39,20	p17	34,30	w13	34,60
p4	37,80	w5	33,20	w8	33,90	w14	34,30
p5	36,80	p11	36,80	p18	34,70	w15	32,40
w2	33,60	p12	35,40	p19	39,90	w16	27,20
p6	41,80	p13	43,90	p20	35,40	w17	33,90
p7	37,10	w6	29,80	w9	35,00	w18	30,60
w3	33,20	p14	34,70	p21	36,40	w19	25,20

Dari data di atas kemudian dianalisis dengan menggunakan SPSS (*Statistical Package for Social Science*) dengan langkah-langkah sebagai berikut.

- Entry data atau buka file data yang akan dianalisis
- Pilih menu berikut ini

Analyze

Descriptives Statistics

Explore

Plots – normality plots with tests

Maka akan diperoleh hasil sebagai berikut:

		Tests of Normality					
	gender	Kolmogorov-Smirnov ^a			Shapiro-Wilk		
		Statistic	Df	Sig.	Statistic	df	Sig.
vo2max	wanita	,144	19	,200*	,949	19	,374
	pria	,156	21	,199	,931	21	,146

a. Lilliefors Significance Correction

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

Dalam uji normalitas berlaku ketentuan: jika *p-value* lebih besar dibanding 0,05, maka data dinyatakan berdistribusi normal. Sebaliknya, jika *p-value* lebih kecil dibanding 0,05, maka data dinyatakan tidak berdistribusi normal. Dari hasil analisis di atas nampak bahwa pada kelompok wanita, uji Kolmogorov-Smirnov, *p-value* sebesar .200 (atau 0,200), sementara pada uji Shapiro-Wilk *p-value* sebesar .374. Sementara itu, pada kelompok pria, uji Kolmogorov-Smirnov, *p-value* sebesar .199, dan pada uji Shapiro-Wilk *p-value* sebesar .146. Artinya, baik pengujian melalui Kolmogorov-Smirnov maupun Shapiro-Wilk, data kedua kelompok dinyatakan berdistribusi normal karena *p-value* lebih besar dibanding 0,05.

Uji Homogenitas

Uji homogenitas bertujuan untuk memastikan bahwa varian dari setiap kelompok sama atau sejenis, sehingga perbandingan dapat dilakukan secara adil. Dari data sebagaimana tampak pada tabel 8.3, antara kelompok pria dan wanita, maka analisis homogenitas dengan menggunakan Levene test pada SPSS sebagai berikut.

- Entry data atau buka file data yang akan dianalisis
- Pilih menu berikut ini

Analyze

Compare Means

One-Way Anova

Options – homogeneity of variance test

Maka akan diperoleh hasil sebagai berikut:

Test of Homogeneity of Variances			
Vo2max			
Levene	df1	df2	Sig.
Statistic			
,731	1	38	,398

Dalam uji homogenitas berlaku ketentuan seperti pada uji normalitas, yakni: jika *p-value* lebih besar dibanding 0,05, maka data dinyatakan homogen. Sebaliknya, jika *p-value* lebih kecil dibanding 0,05, maka data dinyatakan tidak homogen. Dari hasil analisis tersebut dapat dinyatakan bahwa data kedua kelompok bersifat homogen, karena *p-value* lebih besar dibanding 0,05 atau $0,398 > 0,05$. Uji homogenitas sangat diperlukan pada analisis varian, tetapi tidak terlalu sensitif pada uji t.

Uji Linieritas

Uji linieritas dimaksudkan sebagai upaya memastikan linier tidaknya sebaran data yang ada. Uji ini dibutuhkan terutama pada analisis regresi atau korelasi yang bersifat sebab akibat. Sebagai ilustrasi, kita mengambil data pada table 8.3 dengan menambahkan variabel motivasi, sehingga tampak seperti pada tabel 8.4, di mana Vo2max sebagai variabel Y dan motivasi sebagai variabel X.

Tabel 8.4: Data Vo2max dan Motivasi dari 40 orang coba

Sby	Vo2 (Y)	Mot (X)	Sby	Vo2 (Y)	Mot (X)	Sby	Vo2 (Y)	Mot (X)	Sby	Vo2 (Y)	Mot (X)
w1	30,20	5,00	p8	44,90	10,00	p15	38,90	8,00	w10	34,70	5,00
p1	43,90	9,00	w4	38,20	7,00	p16	40,50	9,00	w11	31,00	4,00
p2	42,00	9,00	p9	36,40	7,00	w7	32,60	5,00	w12	30,20	4,00
p3	40,50	8,00	p10	39,20	8,00	p17	34,30	6,00	w13	34,60	6,00
p4	37,80	7,00	w5	33,20	6,00	w8	33,90	6,00	w14	34,30	6,00
p5	36,80	7,00	p11	36,80	6,00	p18	34,70	5,00	w15	32,40	4,00
w2	33,60	6,00	p12	35,40	5,00	p19	39,90	8,00	w16	27,20	3,00
p6	41,80	8,00	p13	43,90	9,00	p20	35,40	7,00	w17	33,90	4,00
p7	37,10	6,00	w6	29,80	3,00	w9	35,00	7,00	w18	30,60	3,00
w3	33,20	6,00	p14	34,70	5,00	p21	36,40	8,00	w19	25,20	3,00

Dari tabel 8.4 tersebut kemudian dilakukan uji linieritas dengan menggunakan Anova melalui program SPSS dengan langkah-langkah sebagai berikut.

- Entry data atau buka file data yang akan dianalisis
- Pilih menu berikut ini

Analyze

Compare Means

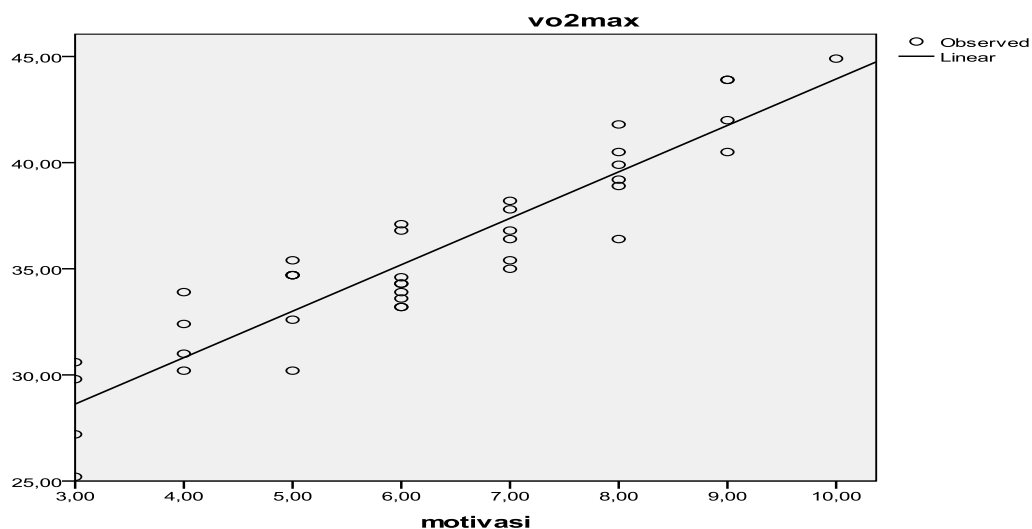
Means

Options – test for linierity

Maka akan diperoleh hasil sebagai berikut:

ANOVA Table			Sum of	df	Mean	F	Sig.
			Squares		Square		
vo2max *	Between	(Combined)	681,639	7	97,377	32,800	,000
motivasi	Groups	Linearity	662,397	1	662,397	223,122	,000
		Deviation from	19,242	6	3,207	1,080	,395
		Linearity					
	Within Groups		95,001	32	2,969		
	Total		776,640	39			

Dalam pengujian linieritas berlaku ketentuan: Jika harga F tidak signifikan atau lebih besar dari .05, maka hubungan antar prediktor dan kriterium dinyatakan linier. Sebaliknya, jika harga F signifikan atau lebih kecil dari .05, maka hubungan antara prediktor dan kriterium dinyatakan tidak linier. Dari analisis yang dilakukan diperoleh harga F (*deviation from linierity*) sebesar 1,080 pada signifikansi .395, yang berarti tidak signifikan, maka hubungan kedua variabel dinyatakan linier. Pemeriksaan linieritas juga dapat dilihat melalui diagram plot seperti tampak berikut.



Gambar 8.4: Diagram plot untuk uji linieritas

Uji Homoskedastisitas

Uji ini dimaksudkan sebagai upaya menguji kesalahan dalam model statistik, terutama pada analisis regresi, yakni apakah varian kesalahan dipengaruhi oleh faktor lain atau tidak. Kondisi

homoskedastisitas, yang mencerminkan variasi nilai residu¹ pada setiap nilai prediksi secara konstan sangat diperlukan. Sebaliknya, kita tidak menginginkan kondisi data bersifat heteroskedastisitas karena dapat mengakibatkan estimasi parameter tidak efisien sehingga variannya besar. Estimasi dianggap efisien jika memiliki varian minimum dan varian kesalahannya bersifat konstan. Dengan demikian, asumsi homoskedastisitas dapat terpenuhi.

Untuk menguji homoskedastisitas kita dapat menggunakan uji Levene atau uji korelasi Spearman. Sebagai ilustrasi, seorang peneliti ingin menguji apakah kemampuan motorik anak (Y) dapat diprediksi dengan variabel status gizi (X_1) dan frekuensi bergerak (X_2). Selanjutnya peneliti tersebut melakukan pengukuran terhadap 20 anak dan hasilnya didapatkan sebagai berikut.

Tabel 8.5: Status gizi, frekuensi bergerak, dan kemampuan motorik 20 siswa

Status Gizi (X_1)	Frekuensi Bergerak (X_2)	Kemampuan Motorik (Y)
2	4	2
2	4	1
1	4	1
1	3	1
3	6	5
4	6	4
5	3	7
5	4	6
7	3	7
6	3	8
4	5	3
3	5	3
6	9	6
6	8	6
8	6	10
9	7	9
10	5	6
9	5	6
4	7	9
4	7	10

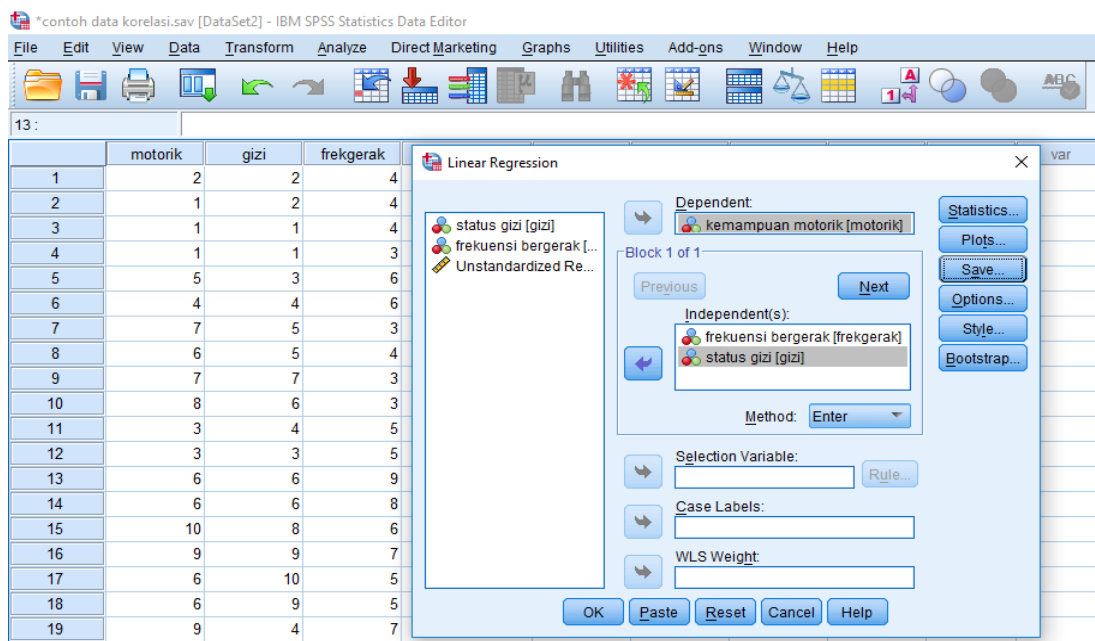
¹ Residu adalah variabel acak, variabel yang belum dijelaskan/bukan merupakan variabel yang diteliti. Apabila variabel prediksi berkorelasi dengan residu, maka sangat boleh jadi mereka adalah variabel yang sama. Jika itu terjadi, analisis regresi tidak dapat diterapkan.

Dari data tersebut selanjutnya dilakukan analisis melalui SPSS dengan langkah-langkah sebagai berikut.

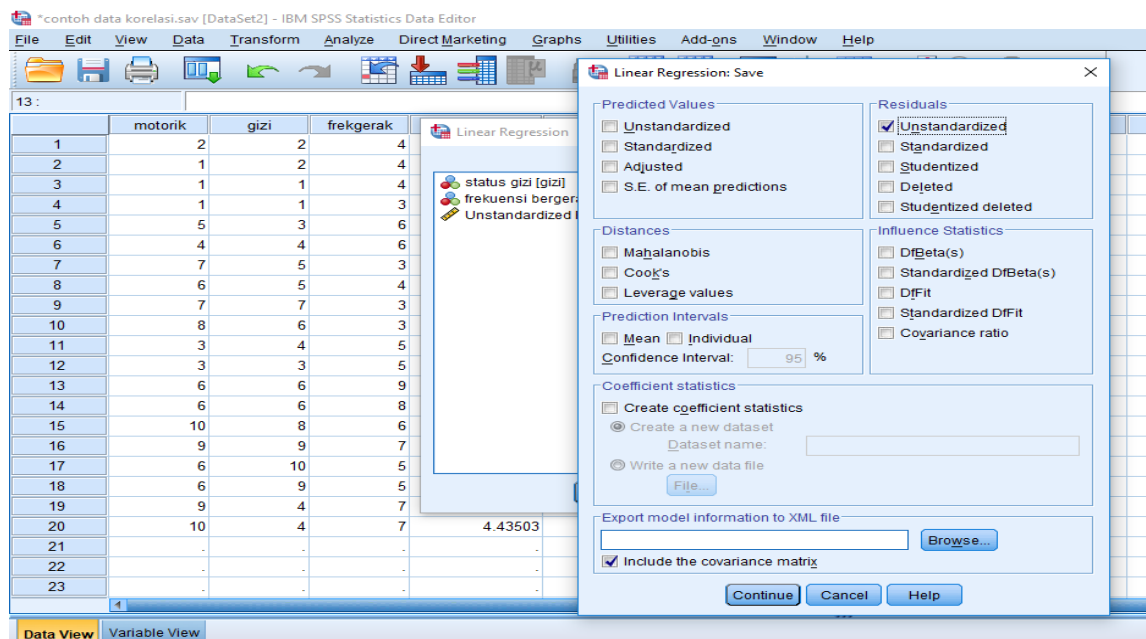
- Entry data atau buka file data yang akan dianalisis
- Pilih menu berikut ini

Analyze
Regression
Linier

Maka akan tampak laman sebagaimana gambar di bawah. Variabel kemampuan motorik dimasukkan dalam kotak “dependent”, sementara variabel status gizi dan frekuensi bergerak dimasukkan ke dalam kotak “independent”.



Selanjutnya pada menu “save” di-klik dan akan keluar laman sebagai tampak di bawah. Pada menu “residual” – unstandardized diberi tanda centang, terus di-klik “continue”.



Dari proses tersebut, maka akan diperoleh hasil sebagai berikut. Variabel residu dengan sendirinya akan muncul.

	motorik	gizi	frekgerak	RES_1	var
1	2	2	4	-1.03048	
2	1	2	4	-2.03048	
3	1	1	4	-1.35340	
4	1	1	3	-.95996	
5	5	3	6	.50556	
6	4	4	6	-1.17153	
7	7	5	3	2.33172	
8	6	5	4	.93828	
9	7	7	3	.97755	
10	8	6	3	2.65463	
11	3	4	5	-1.77808	
12	3	3	5	-1.10100	
13	6	6	9	-1.70601	
14	6	6	8	-1.31257	
15	10	8	6	2.12015	
16	9	9	7	.04963	
17	6	10	5	-2.84057	
18	6	9	5	-2.16349	
19	9	4	7	3.43503	
20	10	4	7	4.43503	
21	-	-	-	-	
22	-	-	-	-	
23	-	-	-	-	

Dengan data tersebut, selanjutnya kita akan melakukan analisis korelasi Spearman dengan langkah-langkah sebagai berikut.

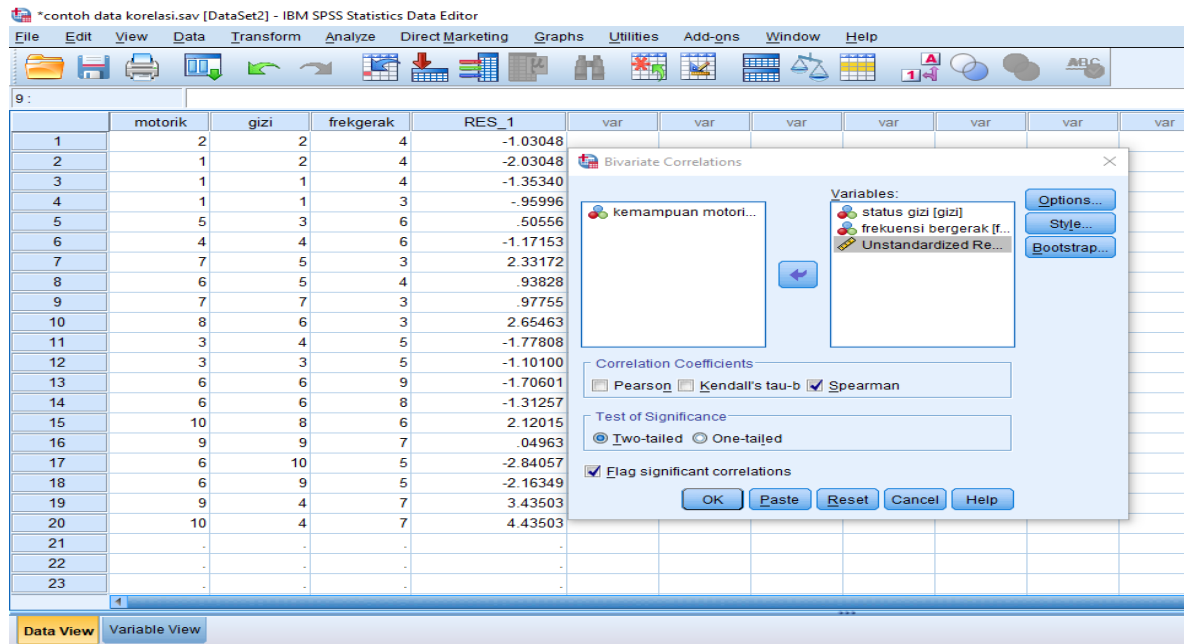
- Entry data atau buka file data yang akan dianalisis
- Pilih menu berikut ini

Analyze

Correlate

Bivariate

Perlu diingat, pada menu “correlation coefficients”, tanda centang dipindahkan dari “pearson” ke “spearman”. Selanjutnya, di-klik “ok”.



Maka akan diperoleh hasil sebagai berikut:

Correlations

			status gizi	frekuensi bergerak	Unstandardized Residual
Spearman's rho	status gizi	Correlation Coefficient	1.000	.257	-.018
		Sig. (2-tailed)	.	.274	.939
		N	20	20	20
	frekuensi bergerak	Correlation Coefficient	.257	1.000	-.079
		Sig. (2-tailed)	.274	.	.740
		N	20	20	20
	Unstandardized Residual	Correlation Coefficient	-.018	-.079	1.000
		Sig. (2-tailed)	.939	.740	.
		N	20	20	20

Dari data tersebut dapat dijelaskan bahwa taraf signifikansi variabel status gizi sebesar 0,939 dan variabel frekuensi bergerak sebesar 0,740. Dalam uji homoskedastisitas dengan menggunakan korelasi Spearman berlaku ketentuan: jika nilai signifikansi (dengan menggunakan two-tailed)

lebih besar dari 0,05, maka data bersifat homoskedastisitas. Sebaliknya, jika nilai signifikansi lebih kecil dari 0,05, maka data bersifat heteroskedastisitas. Karena nilai signifikansi kedua variabel tersebut lebih besar dari 0,05, maka dapat disimpulkan bahwa data tersebut memenuhi sifat homoskedastisitas, sehingga analisis regresi dapat dilakukan.

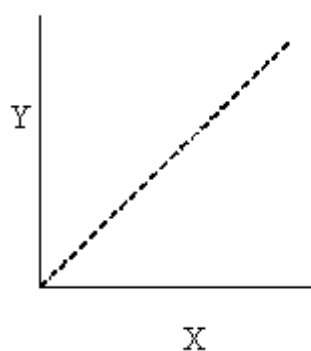
Bab 9 Analisis Korelasi dan Regresi

Secara garis besar analisis statistik dibagi dalam dua kelompok besar, yaitu analisis hubungan dan analisis perbedaan. Analisis hubungan terkait dengan sampai sejauhmana variabel yang satu berhubungan dengan variabel yang lain. Sementara analisis perbedaan terkait dengan sampai sejauhmana suatu variabel berbeda dari kondisi yang satu ke kondisi yang lain.

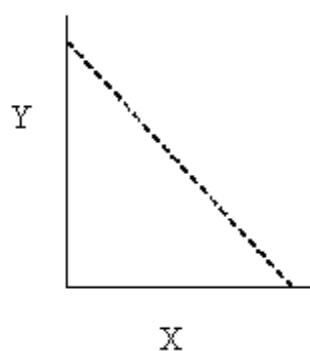
Desain	Jenis Data	Analisis Statistik
Hubungan	Nominal	Korelasi Phi
	Ordinal	Korelasi Tata-Jenjang
	Interval	Korelasi Produk Momen, Regresi
	Rasio	Korelasi Produk Momen, Regresi
Perbedaan/Perbandingan	Frekuensi	Chi-Square
	Interval	Uji Beda Mean (Uji-t, Anova)

Pengertian Korelasi

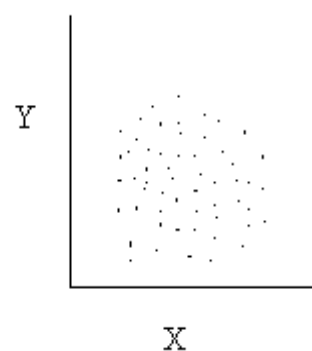
Korelasi adalah sebuah teknik analisis statistik yang digunakan untuk mencari hubungan (korelasi) antara dua variabel atau lebih. Dua variabel yang akan dicari hubungannya tersebut masing-masing disebut variabel bebas (variabel X) dan variabel terikat (variabel Y). Sebagai contoh, jika kita akan meneliti tentang “status gizi dengan kemampuan gerak anak”, maka “status gizi” merupakan variabel X dan “kemampuan gerak anak” merupakan variabel Y.



Korelasi Positif



Korelasi Negatif



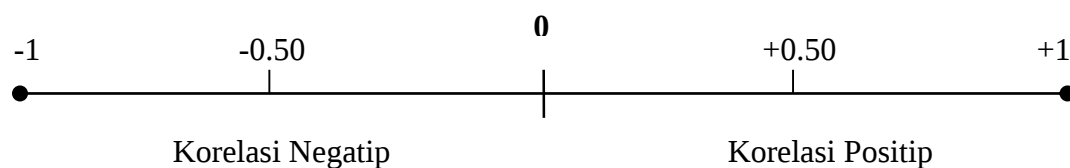
Korelasi Nihil

Dilihat dari arahnya, korelasi dapat dikategorisasikan ke dalam tiga macam, yaitu: korelasi positif, korelasi negatif, dan korelasi nihil. **Korelasi positif** terjadi jika kenaikan nilai pada

variabel X diikuti oleh kenaikan nilai variabel Y, atau bisa juga terjadi sebaliknya, penurunan nilai pada variabel X diikuti dengan penurunan nilai pada variabel Y. Sebagai contoh, hubungan antara motivasi dan prestasi belajar. Semakin tinggi motivasi semakin tinggi pula prestasi belajar, dan sebaliknya, semakin rendah motivasi semakin rendah pula prestasi belajar.

Korelasi negatip terjadi apabila kenaikan nilai pada variabel X diikuti dengan penurunan nilai pada variabel Y, atau sebaliknya, penurunan nilai pada variabel X diikuti dengan kenaikan nilai pada variabel Y. Sebagai contoh, hubungan antara berat badan dengan kelincahan. Semakin tinggi berat badan maka semakin rendah kelincahan. Sedangkan **korelasi nihil** terjadi apabila variabel X dan Y tidak memiliki hubungan yang sistematis.

Arah korelasi tersebut ditunjukkan oleh suatu nilai yang disebut koefisien korelasi. Koefisien korelasi berkisar antara -1 hingga 1, nilai negatip menunjukkan korelasi negatip dan nilai positip menunjukkan korelasi positip. Besarnya koefisien yang bergerak dari titik nol mencerminkan kuatnya hubungan yang ditunjukkan di antara variabel. Semakin dekat dengan angka -1 atau +1, maka semakin kuat korelasi diantara variabel tersebut dan jika semakin mendekati angka 0, maka korelasi diantara variabel tersebut semakin lemah.



Jenis-Jenis Korelasi

1. Korelasi Produk Momen

Teknik korelasi ini ditemukan oleh Karl Pearson, seorang matematikawan yang lahir di Inggris tahun 1857. Korelasi ini digunakan untuk menganalisis hubungan antara dua variabel yang datanya berjenis interval atau rasio.

2. Korelasi Tata Jenjang (*rank order correlation*)

Teknik korelasi ini dikembangkan oleh Charles Spearman, seorang ahli psikologi sekaligus ahli statistik kelahiran Inggris. Korelasi ini digunakan untuk menganalisis hubungan antara dua variabel yang datanya ordinal.

3. Korelasi Phi

Teknik korelasi ini digunakan untuk menganalisis hubungan antara dua variabel yang datanya berjenis nominal.

4. Korelasi Point Serial

Teknik korelasi ini digunakan untuk menganalisis hubungan antara dua variabel, di mana variabel X datanya berjenis ordinal dan variabel Y datanya berjenis interval atau rasio.

Dalam penelitian olahraga, jenis analisis korelasi produk momen yang paling sering digunakan. Ini mengingat dalam olahraga, data yang muncul lebih banyak dalam bentuk rasio seperti kecepatan lari, jarak lemparan, berat badan, tinggi badan, dan sebagainya. Terkait dengan ini, ada dua model korelasi, yaitu korelasi sederhana dan korelasi ganda.

Korelasi Sederhana

Korelasi sederhana adalah suatu analisis korelasi yang menghubungkan antara satu variabel bebas (X) dan satu variabel terikat (Y).

Rumus

$$r = \frac{N \cdot \sum XY - \sum X \cdot \sum Y}{\sqrt{\{N \cdot \sum X^2 - (\sum X)^2\} \{N \cdot \sum Y^2 - (\sum Y)^2\}}}$$

Korelasi *product moment* termasuk dalam statistik inferensial yang umumnya digunakan untuk menguji hipotesis. Sebagai contoh, hipotesis yang dirumuskan adalah: ada hubungan yang signifikan antara berat badan (variabel X) dan kemampuan pull-ups (variabel Y). Untuk kepentingan pengujian dilakukanlah pengukuran terhadap sepuluh orang yang hasilnya seperti tampak pada tabel 9.1. Pertanyaannya, berapa koefisien korelasi antara variabel X dan variabel Y?

Tabel 9.1: Skor Pengukuran Berat Badan dan Pull-Ups dari 10 Subjek

No	Berat Badan (X)	Pull-Ups (Y)	X ²	Y ²	XY
1	104	4	10.816	16	416
2	86	2	7.396	4	172
3	92	6	8.464	36	552
4	112	1	12.544	1	112
5	96	4	9.216	16	384
6	98	13	9.604	169	1.274
7	110	0	12.100	0	0
8	86	9	7.396	81	774
9	105	1	11.025	1	105
10	91	10	8.281	100	910
$\Sigma=980$		$\Sigma=50$	$\Sigma=96.842$	$\Sigma=424$	$\Sigma=4.699$

Dengan menggunakan rumus di atas, maka diperoleh koefisien korelasi sebagai berikut.

$$r = \frac{10.4699 - 980.50}{\sqrt{\{(10.96842) - (980)^2\} \{(10.424 - (50)^2)\}}}$$

$$r = \frac{46990 - 49000}{\sqrt{\{(968420 - 960400) (4240 - 2500)\}}}$$

$$r = \frac{-2010}{\sqrt{\{8020.1740\}}}$$

$$r = \frac{-2010}{896.417}$$

$$r = \frac{-2010}{37363}$$

$$r = -.538$$

$$r = -.54 \text{ (dibulatkan)}$$

Untuk menguji koefisien korelasi (r) yang diperoleh tersebut, ia harus dikonsultasikan dengan tabel nilai r produk momen (lihat lampiran 1). Sebelumnya harus ditentukan derajat kebebasan (*degree of freedom*), yaitu $df = N-2$. Dalam tabel r produk momen, df 8 adalah 0.632 untuk taraf signifikansi 5% dan 0.765 untuk taraf signifikansi 1%. Jadi, koefisien korelasi yang diperoleh tidak signifikan, baik pada taraf 5% maupun 1%, mengingat r hitung lebih kecil dibanding r tabel. Dengan demikian hipotesis yang berbunyi: “Ada hubungan yang signifikan antara berat badan dan kemampuan pull-ups” ditolak. Artinya, berdasarkan data empirik sebagai hasil pengujian lapangan, setidaknya untuk data di atas, menunjukkan bahwa tidak ada hubungan antara berat badan dengan kemampuan pull-ups.

Tidak adanya korelasi yang signifikan dalam data di atas dapat disebabkan karena kecilnya jumlah sampel, yakni hanya 10. Sangat boleh jadi jika sampelnya ditambah, misalnya menjadi 30, maka besar kemungkinan hasilnya akan berbeda.

Dalam contoh berikut kita akan mencoba melakukan analisis korelasi dengan jumlah data yang lebih besar. Kita tidak akan menghitung dengan cara manual, melainkan melalui software statistik program SPSS (*statistical package for social science*).

Tabel 9.2: Data vo2max dan motivasi dari 40 subyek

Sby	Vo2 (Y)	Mot (X)	Sby	Vo2 (Y)	Mot (X)	Sby	Vo2 (Y)	Mot (X)	Sby	Vo2 (Y)	Mot (X)
1	30,20	5,00	11	44,90	10,00	21	38,90	8,00	31	34,70	5,00
2	43,90	9,00	12	38,20	7,00	22	40,50	9,00	32	31,00	4,00
3	42,00	9,00	13	36,40	7,00	23	32,60	5,00	33	30,20	4,00
4	40,50	8,00	14	39,20	8,00	24	34,30	6,00	34	34,60	6,00
5	37,80	7,00	15	33,20	6,00	25	33,90	6,00	35	34,30	6,00
6	36,80	7,00	16	36,80	6,00	26	34,70	5,00	36	32,40	4,00
7	33,60	6,00	17	35,40	5,00	27	39,90	8,00	37	27,20	3,00
8	41,80	8,00	18	43,90	9,00	28	35,40	7,00	38	33,90	4,00
9	37,10	6,00	19	29,80	3,00	29	35,00	7,00	39	30,60	3,00
10	33,20	6,00	20	34,70	5,00	30	36,40	8,00	40	25,20	3,00

Dengan menggunakan program SPSS, maka kita dapat mengolah data di atas dengan langkah-langkah sebagai berikut.

- Entry data atau buka file data yang akan dianalisis
- Pilih menu berikut ini

Analyze
Correlate
Bivariate

Maka akan diperoleh hasil sebagai berikut:

Descriptive Statistics

	Mean	Std. Deviation	N
vo2max	35.6275	4.46249	40
motivasi	5.9750	1.95445	40

Correlations

		vo2max	motivasi
vo2max	Pearson Correlation	1	.645**
	Sig. (2-tailed)		.000
	N	40	40
motivasi	Pearson Correlation	.645**	1
	Sig. (2-tailed)	.000	
	N	40	40

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

Dari hasil pengolahan tersebut nampak bahwa rerata vo2max sebesar 35,63 dengan standar deviasi 4,46. Sementara itu, rerata motivasi sebesar 5,97 dengan standar deviasi 1,95. Hasil analisis korelasi dari kedua variabel tersebut menunjukkan bahwa ada korelasi yang signifikan antara variabel motivasi dan vo2max dengan koefisien korelasi sebesar 0,64 dengan *p-value* 0.01. Jika kita ingin mengetahui lebih jauh terkait kontribusi variabel motivasi terhadap vo2max, maka kita dapat menggunakan koefisien determinasi dengan rumus sebagai berikut.

$$\text{Koefisien determinasi} = r^2 \times 100\%$$

Dengan menggunakan data di atas, maka koefisien determinasi dari motivasi terhadap vo2max adalah $0,64^2 \times 100\% = 40,96\%$. Dalam konteks relasi antara motivasi dan vo2max, maka dapat dikatakan bahwa 40,96% dari variabel vo2max dapat dijelaskan dari variabel motivasi. Selebihnya, atau 59,04% dari variabel lain yang belum bisa dijelaskan. Perlu diingat bahwa ketika relasi antar variabel bersifat bivariat, hanya dua variabel, maka koefisien korelasi dan koefisien

determinasinya begitu tinggi. Namun jika dimasukkan variabel yang lain, sangat mungkin koefisien tersebut akan mengecil. Karena itu, dalam tingkat tertentu, sangat disarankan analisis bersifat multivariat.

Korelasi Ganda

Korelasi ganda (*multiple correlation*) adalah suatu analisis korelasi yang menghubungkan antara sekelompok variabel bebas (bisa dua atau lebih, misalnya X_1 dan X_2) dengan satu variabel terikat (Y).

Rumus

$$R_{y.12} = \sqrt{\frac{(r_{1.y})^2 + (r_{2.y})^2 - 2(r_{1.y})(r_{2.y})(r_{1.2})}{1 - (r_{1.2})^2}}$$

Keterangan:

$R_{y.12}$ = korelasi antara X_1 dan X_2 dengan Y

$r_{1.y}$ = korelasi antara X_1 dan Y

$r_{2.y}$ = korelasi antara X_2 dan Y

$r_{1.2}$ = korelasi antara X_1 dan X_2

Dari rumus tersebut tampak bahwa untuk menghitung korelasi ganda, harus ditemukan dulu korelasi tunggalnya. Anggap saja telah dilakukan perhitungan, $r_{1.y} = 0,840$; $r_{2.y} = 0,605$; dan $r_{1.2} = 0,539$. Berdasarkan harga-harga korelasi tunggal yang sudah diperoleh, maka koefisien korelasi ganda dapat dihitung sebagai berikut.

$$\begin{aligned} R_{y.12} &= \sqrt{\frac{(r_{1.y})^2 + (r_{2.y})^2 - 2(r_{1.y})(r_{2.y})(r_{1.2})}{1 - (r_{1.2})^2}} \\ &= \sqrt{\frac{(0,840)^2 + (0,605)^2 - 2(0,840)(0,605)(0,539)}{1 - (0,539)^2}} \end{aligned}$$

$$= \sqrt{\frac{(0,7056 + 0,366025) - 2(0,2739198)}{0,709479}}$$

$$= 0,859$$

Sebelum digunakan untuk mengambil kesimpulan, harga korelasi ganda sebesar 0,859 tersebut harus diuji signifikansinya terlebih dahulu, yaitu dengan menggunakan rumus F berikut.

$$F = \frac{R^2/m}{(1-R^2)/(N-m-1)}$$

Keterangan:

R^2 = korelasi kuadrat atau koefisien determinasi

m = jumlah variabel bebas

N = jumlah sampel

Korelasi Tata Jenjang (*rank order correlation* - Spearman)

Analisis korelasi ini digunakan untuk menganalisis hubungan dua variabel yang datanya ordinal.

Adapun rumus korelasi tata jenjang adalah sebagai berikut.

$$\rho = 1 - \frac{6\sum D^2}{N(N^2 - 1)}$$

di mana,

ρ = koefisien korelasi tata jenjang

D = perbedaan antar urutan

N = jumlah sampel

1&6 = bilangan konstan

Sebagai contoh, ada kejuaran senam artistik di sekolah yang diikuti oleh 11 siswa. Selain dinilai oleh guru olahraganya sendiri, juga dinilai oleh guru olahraga dari sekolah yang lain. Dari penilaian yang dilakukan, maka diperoleh skor sebagai berikut.

Tabel 9.3: Hasil penilaian kejuaraan senam dari dua juri

Siswa	Urutan penilaian Juri 1	Urutan penilaian Juri 2	Perbedaan (D)	D ²
A	1	4	-3	9
B	2	3	-1	1
C	3	1	+2	4
D	4	2	+2	4
E	5	5	0	0
F	6	6	0	0
G	7	8	-1	1
H	8	9	-1	1
I	9	7	+2	4
J	10	11	-1	1
K	11	10	+1	1
			0	26

Dengan menggunakan rumus di atas, maka

$$\begin{aligned}
 \rho &= 1 - \frac{6\sum D^2}{N(N^2 - 1)} \\
 &= 1 - \frac{(6)(26)}{11(121 - 1)} \\
 &= \mathbf{0,88}
 \end{aligned}$$

Jika kita menghitung secara manual, maka untuk mengetahui tingkat signifikansi perlu dikonfirmasi dengan tabel korelasi tata jenjang (*rho correlation*).

Dengan menggunakan program SPSS, maka kita dapat mengolah data di atas dengan langkah-langkah sebagai berikut.

- Entry data atau buka file data yang akan dianalisis
- Pilih menu berikut ini

Analyze
 Correlate
 Bivariate

Maka akan diperoleh hasil sebagai berikut:

Correlations				
			juri1	juri2
Spearman's rho	juri1	Correlation Coefficient	1.000	.882**
		Sig. (2-tailed)	.	.000
		N	11	11
	juri2	Correlation Coefficient	.882**	1.000
		Sig. (2-tailed)	.000	.
		N	11	11

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

Ada kalanya kita memiliki dua data, yang satu bersifat ordinal dan yang lainnya bersifat interval dan kita ingin mengkorelasikan keduanya. Bagaimana caranya? Pertama, kita perlu menurunkan skala data, dari interval ke ordinal, mengingat kita tidak mungkin menaikkan skala data, dari ordinal ke interval. Proses selanjutnya sama dengan pengolahan sebelumnya. Hanya saja perlu diingat, jika ada data yang sama, maka konversi urutannya diambil dari rerata dari data tersebut.

Tabel 9.4: Hasil penilaian bermain bolabasket

Siswa	Urutan skor bermain	Skor tes lemparan basket	Urutan skor tes	Perbedaan (D)	D ²
A	1	19	1	0	0
B	2	17	3,5	-1,5	2,25
C	3	18	2	+1	1
D	4	17	3,5	+0,5	0,25
E	5	15	6	-1	1
F	6	14	8	-2	4
G	7	15	6	+1	1
H	8	15	6	+2	4
I	9	12	10	-1	1
J	10	13	9	+1	1
K	11	8	11	0	0
L	12	5	12	0	0
				0	15,5

Dengan menggunakan rumus di atas, maka

$$\begin{aligned}\rho &= 1 - \frac{6\sum D^2}{N(N^2 - 1)} \\ &= 1 - \frac{(6)(15.5)}{12(144 - 1)} \\ &= \mathbf{0,945}\end{aligned}$$

Jika data tersebut diolah dengan menggunakan SPSS, maka diperoleh hasil sebagai berikut.

Correlations			rank bermain	skor tes basket
Spearman's rho	rank bermain	Correlation	1.000	-.945**
		Coefficient		
		Sig. (2-tailed)	.	.000
		N	12	12
	skor tes basket	Correlation	-.945**	1.000
		Coefficient		
		Sig. (2-tailed)	.000	.
		N	12	12

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

Korelasi Phi

Korelasi ini dikembangkan oleh Karl Pearson, digunakan untuk menganalisis hubungan antara dua variabel yang memiliki skala nominal atau kategorik. Korelasi ini disimbulkan dengan ϕ (phi) dengan rumus sebagai berikut.

$$\phi = \frac{ad - bc}{\sqrt{(a + b)(c + d)(a + c)(b + d)}}$$

Untuk dapat menggunakan rumus di atas, kita memerlukan tabel kontingensi 2x2 sebagai berikut.

X	Y		Jumlah
	1	2	
1	a	b	(a+b)
2	c	d	(c+d)
Jumlah	(a+c)	(b+d)	N

Sebagai ilustrasi, kita ingin mengetahui apakah ada hubungan antara jenis kelamin dan pilihan cabang olahraga. Jenis kelamin dikategorikan ke dalam laki-laki (L) dan perempuan (P), sementara jenis olahraga dikategorikan ke dalam olahraga terukur (OT) dan olahraga tidak terukur (OTT). Kita memiliki data 100 mahasiswa dengan pilihan sebagai berikut.

Jenis Kelamin	Pilihan Cabor		Jumlah
	OT	OTT	
L	30	20	(50)
P	15	35	(50)
Jumlah	(45)	(55)	100

Dengan menggunakan rumus di atas, maka dapat diperoleh hasil sebagai berikut.

$$\begin{aligned}
 \phi &= \frac{ad - bc}{\sqrt{(a+b)(c+d)(a+c)(b+d)}} \\
 &= \frac{30.35 - 20.15}{\sqrt{(50)(50)(45)(55)}} \\
 &= \frac{750}{2487} \\
 &= \mathbf{0,30}
 \end{aligned}$$

Untuk memastikan apakah koefisien tersebut signifikan atau tidak, kita perlu mengonversi ke dalam nilai chi-square, dengan rumus:

$$\chi^2 = (\phi)^2 \times N$$

Maka dengan demikian dapat diperoleh nilai chi-square sebagai berikut;

$$(0,30)^2 \times 100 = 9$$

Selanjutnya, angka tersebut dikonfirmasi ke dalam tabel nilai chi-square, dengan $df = 1$ dan taraf signifikansi 5%, maka menunjukkan angka 3,841. Artinya, angka 9 melampaui nilai chi-square sebesar 3,841. Dengan demikian dapat disimpulkan bahwa ada hubungan yang signifikan antara jenis kelamin dan pilihan cabang olahraga. Mahasiswa cenderung memilih cabang olahraga terukur, sementara mahasiswi cenderung memilih cabang olahraga tidak terukur.

Dengan menggunakan program SPSS, maka kita dapat mengolah data di atas dengan langkah-langkah sebagai berikut.

- Entry data atau buka file data yang akan dianalisis
- Pilih menu berikut ini

Analyze

Descriptive Statistics

Crosstabs

Maka akan diperoleh hasil sebagai berikut:

jenis kelamin * pilihan cabang Crosstabulation

Count

		pilihan cabang		Total
		olahraga terukur	olahraga tidak terukur	
jenis kelamin	laki-laki	30	20	50
	perempuan	15	35	50
Total		45	55	100

Symmetric Measures

		Value	Approx. Sig.
Nominal by Nominal	Phi	.302	.003
	Cramer's V	.302	.003
N of Valid Cases		100	

Korelasi Point Serial

Korelasi ini juga dikembangkan oleh Karl Pearson, digunakan untuk menghubungkan dua variabel yang bersifat kategorik-kontinum (*artificial dichotomy*) dan interval/rasio. Istilah serial merujuk pada jumlah variabel X, jika variabel X dibagi menjadi dua disebut korelasi biserial, jika dibagi tiga disebut korelasi triserial, dan jika dibagi empat disebut quartoserial. Sebagai ilustrasi, seorang peneliti ingin menguji apakah ada hubungan antara aktivitas latihan dan tingkat kebugaran. Variabel aktivitas latihan dikategorikan ke dalam aktif latihan (1) dan tidak aktif latihan (0), sementara variabel kebugaran berjenis data interval. Adapun data yang dimiliki adalah sebagai berikut.

Tabel 9.5: Data aktivitas latihan dan kebugaran

Subyek yang aktif latihan	Kebugaran	Subyek yang tidak aktif latihan	Kebugaran
A1	48	B16	36
A2	50	B17	25
A3	46	B18	34
A4	48	B19	31
A5	40	B20	30
A6	38	B21	26
A7	42	B22	36
A8	40	B23	33
A9	43	B24	32
A10	44	B25	26
A11	38	B26	30
A12	38	B27	28
A13	42	B28	28
A14	37	B29	34
A15	35	B30	31

Dengan menggunakan program SPSS, maka kita dapat mengolah data di atas dengan langkah-langkah sebagai berikut.

- Entry data atau buka file data yang akan dianalisis
- Pilih menu berikut ini

Analyze

Correlate

Bivariate

Maka akan diperoleh hasil sebagai berikut:

Correlations

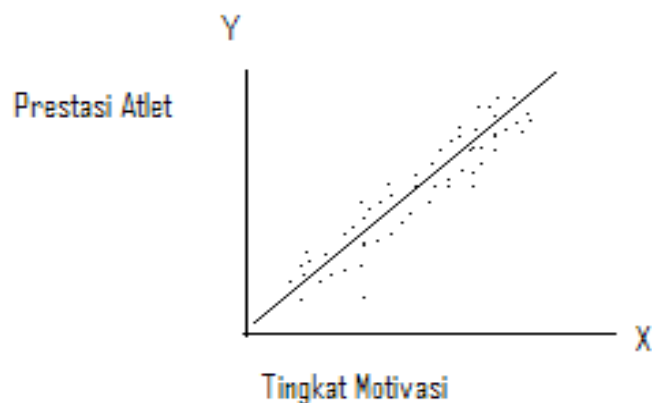
		latihan	kebugaran
latihan	Pearson Correlation	1	.820**
	Sig. (2-tailed)		.000
	N	30	30
kebugaran	Pearson Correlation	.820**	1
	Sig. (2-tailed)	.000	
	N	30	30

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

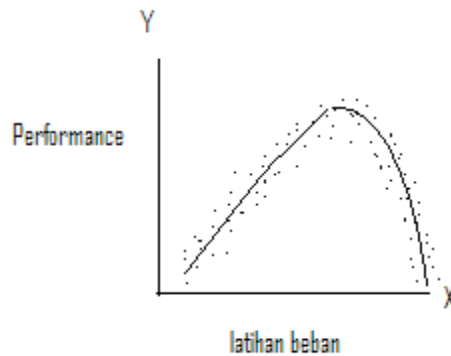
Analisis Regresi

Analisis regresi adalah suatu teknik analisis statistik yang digunakan untuk mengadakan peramalan (prediksi) besarnya variasi yang terjadi pada variabel Y (kriterium) berdasarkan variabel X (prediktor). Analisis regresi juga dapat digunakan untuk menentukan bentuk hubungan, arah, dan besarnya koefisien korelasi antara variabel X dan Y. Untuk bisa dilakukan analisis regresi, jenis datanya harus interval atau rasio.

Regresi bisa berbentuk linier dan non-linier. Regresi linier terjadi apabila variasi peningkatan nilai pada variabel kriterium diikuti oleh peningkatan nilai pada variabel prediktor, demikian juga sebaliknya. Misalnya. seorang peneliti ingin mengkaji prestasi atlet berdasarkan motivasinya. Semakin tinggi motivasi, semakin tinggi prestasi yang dicapai.



Sedangkan bentuk hubungan yang non-linier terjadi apabila peningkatan variasi nilai dari salah satu variabel tidak serta merta diikuti oleh peningkatan nilai pada variabel yang lain. Pada titik tertentu justru mengalami penurunan. Sebagai contoh pengaruh latihan beban terhadap peningkatan *performance*.



Terkait dengan bentuk regresi di atas, peneliti harus dapat memastikan terlebih dahulu mengenai distribusi data yang akan dianalisis, apakah linier atau justru non-linier. Jika distribusi data berbentuk linier, maka harus diolah dengan menggunakan analisis regresi linier. Sebaliknya, jika distribusi data berbentuk non-linier, maka harus diolah dengan analisis regresi non-linier. Prosedur yang dilakukan untuk mengetahui apakah suatu data linier atau tidak adalah dengan melakukan uji linieritas. Pada uji linieritas, harga F empirik diharapkan lebih kecil daripada F teoretik; dan jika F empirik lebih besar daripada F teoretik, berarti distribusi data tersebut tidak linier.

Regresi Sederhana

Analisis regresi sederhana digunakan untuk menentukan dasar ramalan dari suatu distribusi data yang terdiri dari variabel kriterium (Y) dan satu variabel prediktor (X) yang memiliki bentuk hubungan linier. Variabel kriterium disebut juga variabel terikat dan variabel prediktor disebut juga sebagai variabel bebas. Ada perbedaan mendasar antara korelasi dengan regresi, meski keduanya merupakan uji hubungan antar variabel. Analisis korelasi sekadar menguji hubungan antar variabel tanpa harus mengetahui variabel mana yang menjadi sebab dan variabel mana yang menjadi akibat. Sementara itu, pada analisis regresi sampai pada keputusan relasi sebab-akibat di antara variabel tersebut.

Adapun rumus untuk regresi sederhana adalah sebagai berikut.

$$Y = a + bX + e$$

Dimana.

Y = kriterium

a = intersep (konstanta regresi) atau harga yang memotong sumbu Y

b = koefisien regresi atau sering disebut slope, gradien, atau kemiringan garis

e = *error* atau residual

Untuk menemukan harga a dan b digunakan rumus sebagai berikut.

$$a = \frac{\Sigma Y \cdot \Sigma X^2 - \Sigma X \cdot \Sigma XY}{N \cdot \Sigma X^2 - (\Sigma X)^2}$$

$$b = \frac{N \cdot \Sigma XY - \Sigma X \cdot \Sigma Y}{N \cdot \Sigma X^2 - (\Sigma X)^2}$$

Sebagai ilustrasi, seorang peneliti ingin memprediksi tingkat kebugaran (Y) berdasarkan frekuensi latihan (X). Secara teoretik, tingkat kebugaran seseorang dipengaruhi oleh frekuensi latihan dari yang bersangkutan. Peneliti kemudian melakukan pengumpulan data terhadap 30 subjek seperti tampak pada tabel berikut.

Tabel 9.6: Data frekuensi latihan dan kebugaran dari 30 orang subjek

<i>Subjek</i>	<i>Frekuensi Latihan (X)</i>	<i>Kebugaran (Y)</i>	<i>Subjek</i>	<i>Frekuensi Latihan (X)</i>	<i>Kebugaran (Y)</i>
1	6	65	16	2	25
2	4	40	17	4	40
3	3	30	18	2	30
4	3	35	19	3	30
5	6	60	20	3	35
6	2	30	21	6	60
7	2	25	22	5	50
8	5	50	23	5	55
9	4	45	24	6	65
10	5	55	25	4	45
11	6	60	26	2	25
12	5	50	27	4	40
13	5	55	28	2	30
14	4	45	29	3	30
15	6	65	30	3	35

Dari data tersebut selanjutnya dilakukan analisis melalui SPSS dengan langkah-langkah sebagai berikut.

- Entry data atau buka file data yang akan dianalisis
- Pilih menu berikut ini

Analyze
Regression
Linier

Maka akan diperoleh hasil sebagai berikut:

Model Summary^b					
Model		R	R Square	Adjusted R Square	Std. Error of the Estimate
1	dimension0	,975 ^a	,952	,950	2,97309
a. Predictors: (Constant), Frek.lat					
b. Dependent Variable: kebugaran					

ANOVA^b						
Model		Sum of Squares	df	Mean Square	F	Sig.
1	Regression	4860,000	1	4860,000	549,818	,000 ^a
	Residual	247,500	28	8,839		
	Total	5107,500	29			
a. Predictors: (Constant), Frek.lat						
b. Dependent Variable: kebugaran						

Coefficients^a					
Model		Unstandardized Coefficients		Standardized Coefficients	T
		B	Std. Error	Beta	Sig.
1	(Constant)	7,500	1,628		4,606
	Frek.lat	9,000	,384	,975	23,448
a. Dependent Variable: kebugaran					

Dari hasil analisis di atas, dapat diinterpretasi bahwa korelasi antara variabel prediktor (frekuensi latihan) dan variabel kriterium (kebugaran), yang disimbulkan dengan R (simbol korelasi pada regresi²) sebesar .975 dengan standar kesalahan (*standard error of estimate*) sebesar

² Ada perbedaan rentang koefisien antara R dan r, pada r korelasi koefisien bergerak antara -1 sampai dengan +1, sementara pada R regresi koefisien bergerak antara 0 sampai dengan 1

2,97 . Artinya, hubungan kedua variabel sangat tinggi bahkan mendekati sempurna. Dari analisis regresi yang dilakukan nampak bahwa harga F sebesar 549,818 pada signifikansi .000. Besarnya sumbangan relatif variabel prediktor terhadap variabel kriterium sebesar 95%, diperoleh dari koefisien beta dikuadratkan kemudian dikalikan 100%. Dari analisis tersebut juga diperoleh konstanta sebesar 7,500 dan koefisien regresi prediktor sebesar 9,000. Jika harga-harga tersebut dibuat dalam persamaan garis regresi, maka diperoleh persamaan sebagai berikut:

$$\begin{aligned} Y &= a + bX + e \\ &= 7,500 + 9,000X + e \end{aligned}$$

Bab 10 T-Test dan Anova

Jika pada bagian sebelumnya kita telah mempelajari analisis korelasi dan regresi, yang merupakan bentuk uji hubungan antar variabel, maka pada bagian ini akan dibahas tentang T-test dan Anova, yang merupakan bentuk uji beda antar kelompok. T-test digunakan untuk menganalisis perbedaan antara dua buah distribusi. Sementara itu, Anova digunakan untuk menguji perbedaan data yang lebih dari dua kelompok.

T-Test

T-test dikembangkan oleh William Gossett, seorang ahli statistik berkebangsaan Inggris yang meninggal tahun 1937. T-test merupakan teknik analisis statistik yang dipergunakan untuk menguji signifikansi perbedaan dua buah mean yang berasal dari dua buah distribusi. Ada dua jenis T-Test, yaitu untuk sampel yang berbeda (*independent sample*) dan sampel yang sejenis (*dependent sample*).

T-Test untuk Sampel Berbeda

Sampel berbeda dimaksudkan bahwa distribusi data yang dibandingkan berasal dari dua kelompok yang berbeda. Sebagai contoh, seorang peneliti ingin mengkaji pengaruh latihan beban terhadap peningkatan kekuatan otot tungkai. Maka, peneliti akan mencari dua kelompok subjek. Kelompok pertama diberi latihan beban (sebagai kelompok eksperimen) dan kelompok kedua tidak diberi latihan beban (sebagai kelompok kontrol). Setelah perlakuan berupa pemberian latihan beban diberikan dalam jangka waktu tertentu, peneliti kemudian mengukur kekuatan otot tungkai kedua kelompok. Data kedua kelompok tersebut kemudian dibandingkan (dianalisis) dengan menggunakan t-test untuk sampel berbeda.

Adapun rumus yang digunakan adalah sebagai berikut.

$$t = \frac{M_1 - M_2}{\sqrt{\frac{s^2}{N_1} + \frac{s^2}{N_2}}}$$

di mana.

M_1 = Mean pada distribusi sampel 1

M_2 = Mean pada distribusi sampel 2

s_1^2 = Nilai varian pada distribusi sampel 1

s_2^2 = Nilai varian pada distribusi sampel 2

N_1 = Jumlah individu pada sampel 1

N_2 = Jumlah individu pada sampel 2

Untuk dapat mengerjakan rumus tersebut, perlu diketahui lebih dulu rumus menentukan varian populasi (s^2), yaitu:

$$s^2 = \frac{\left\{ \frac{\sum X_1^2 - (\sum X_1)^2}{N_1} \right\} + \left\{ \frac{\sum X_2^2 - (\sum X_2)^2}{N_2} \right\}}{(N_1 + N_2) - 2}$$

Untuk lebih jelasnya, marilah kita lihat contoh berikut. Seorang peneliti ingin membandingkan nilai matakuliah fisiologi dari dua kelompok mahasiswa yang berbeda, yakni UMPTN dan PMDK. Peneliti tersebut mengajukan hipotesis: “Terdapat perbedaan prestasi belajar fisiologi yang signifikan antara mahasiswa dari jalur UMPTN dan jalur PMDK”. Peneliti tersebut selanjutnya melakukan pengumpulan data dan divisualisasikan sebagaimana tampak pada tabel 10.1.

Tabel 10.1: Nilai Matakuliah Fisiologi pada Dua Kelompok Mahasiswa

Nilai Fisiologi Mahasiswa UMPTN				Nilai Fisiologi Mahasiswa PMDK			
Nilai (X_1)	f	fX_1	fX_1^2	Nilai (X_2)	f	fX_2	fX_2^2
80	1	80	6400	75	1	75	5625
75	4	300	22500	73	1	73	5329
72	5	360	25920	70	3	210	14700
70	10	700	49000	66	5	330	21780
68	15	1020	69360	65	9	585	38025
65	13	845	54925	62	15	930	57660
63	6	378	23814	60	9	540	32400
60	4	240	14400	58	4	232	13456
55	2	110	6050	55	3	165	9075
				50	2	100	5000
	N = 60	$\sum fX_1 = 4033$	$\sum fX_1^2 = 272369$		N=52	$\sum fX_2 = 3240$	$\sum fX_2^2 = 203050$
$M_1 = 4033/60 = 67.217$				$M_2 = 3240/52 = 62.308$			

Sebelum kita memasukkan sejumlah besaran tersebut ke dalam rumus t-tes, perlu dicari dulu varian populasinya.

$$s^2 = \frac{\{ 272.369 - \frac{(4.033)^2}{60} \} + \{ 203.050 - \frac{(3.240)^2}{52} \}}{(60 + 52) - 2}$$

$$= \frac{(272.369 - 271.084.82) + (203.050 - 201.876.92)}{110}$$

$$= 22.3387279$$

Setelah besarnya varian populasi diketahui, maka langkah berikutnya adalah memasukkan angka-angka di atas ke dalam rumus t-tes sebagai berikut.

$$t = \frac{M_1 - M_2}{\sqrt{s^2 \left[\frac{1}{N_1} + \frac{1}{N_2} \right]}}$$

$$t = \frac{67.217 - 62.308}{\sqrt{\frac{22.3387279}{60} + \frac{22.3387279}{52}}}$$

$$t = \frac{4.909}{\sqrt{0.3723121 + 0.4295903}}$$

$$t = 5.4819131$$

$$t = \mathbf{5.48} \text{ (dibulatkan)}$$

Apakah koefisien t hasil perhitungan tersebut signifikan atau dengan kata lain hipotesis yang diajukan terbukti? Untuk menjawabnya, t hitung perlu dikonsultasikan dengan t tabel (lihat lampiran 2). Dengan $df (N_1 + N_2) - 2$, maka didapat df sebesar 110. Mengingat df 110 tidak ada dalam tabel, maka kita perlu melakukan interpolasi. Caranya, pada analisis dua ekor, kita peroleh

df 60 pada taraf signifikansi 5% adalah 2.000, sedangkan df 120 adalah 1.980. Untuk mudahnya, df 60 dan df 120 dijumlah kemudian dibagi dua, yakni $(2.000 + 1.980) : 2 = 1.990$.

Ternyata t hitung (5.48) lebih besar daripada t tabel (1.990). Dengan demikian, hipotesis yang menyatakan: “Terdapat perbedaan prestasi belajar fisiologi yang signifikan antara mahasiswa dari jalur UMPTN dan jalur PMDK” dapat diterima. Artinya terdapat perbedaan yang signifikan antara mahasiswa UMPTN dan PMDK dalam hal prestasi belajar fisiologi.

Data tersebut juga dapat dianalisis melalui SPSS dengan langkah-langkah sebagai berikut.

- Entry data atau buka file data yang akan dianalisis
- Pilih menu berikut ini

Analyze

Compare Means

Independent samples T-test

Maka akan diperoleh hasil sebagai berikut:

Group Statistics					
	jalur masuk	N	Mean	Std. Deviation	Std. Error Mean
fisiologi	jalur umptn	60	67.2167	4.66539	.60230
	jalur pmdk	52	62.3077	4.79599	.66508

Independent Samples Test									
		Levene's Test for Equality of Variances		t-test for Equality of Means					
		F	Sig.	t	df	Sig. (2-tailed)	Mean Difference	Std. Error Difference	95% Confidence Interval of the Difference
									Lower Upper
fisiologi	Equal variances assumed	.065	.799	5.482	110	.000	4.90897	.89549	3.13432 6.68363
	Equal variances not assumed			5.471	106.839	.000	4.90897	.89727	3.13020 6.68775

Dari analisis data menggunakan SPSS, tampak tidak ada perbedaan hasil dengan pengolahan data secara manual. Ini menunjukkan bahwa analisis data menggunakan cara manual sama akuratnya dengan menggunakan SPSS. Tetapi tentu, dari segi kecepatan, mengolah data menggunakan *software* jauh lebih efisien dibanding menggunakan cara manual.

T-Test untuk Sampel Sejenis

Sampel sejenis dimaksudkan bahwa distribusi data yang dibandingkan berasal dari kelompok subjek yang sama. Misalnya, bila kita ingin menganalisis perbedaan antara hasil *pretest* dan *posttest* pada kelompok tertentu, maka dapat kita gunakan T-Test sampel sejenis. Adapun rumus yang digunakan adalah sebagai berikut.

$$t = \frac{\sum D}{\sqrt{\frac{(N\sum D^2 - (\sum D)^2)}{N - 1}}}$$

di mana.

D = Perbedaan setiap pasangan skor (*pretest-posttest*)

N = Jumlah sampel

Sebagai contoh, seorang peneliti ingin mengkaji pengaruh latihan beban terhadap kemampuan daya ledak otot tungkai. Peneliti tersebut kemudian mengumpulkan data kepada 10 orang coba. Pertama dilakukan *pretest* berupa “jump and reach” kepada subjek. Selanjutnya subjek diberi latihan beban untuk beberapa waktu, kemudian diukur lagi (*posttest*). Hasilnya sebagai berikut.

Tabel 10.2: Skor Pretest dan Posttest “Jump and Reach” pada 10 Subjek

Subjek	Jump and Reach Test (cm)		D	D ²
	Pretest	Posttest		
1	12	16	4	16
2	15	21	6	36
3	13	15	2	4
4	20	22	2	4
5	21	21	0	0
6	19	23	4	16
7	14	16	2	4
8	17	18	1	1
9	16	22	6	36
10	18	23	5	25
$\sum$	165	197	32	142
M	16.5	19.7	3.2	

Dari data sebagaimana tampak pada tabel 10.2 tersebut kemudian dianalisis dengan menggunakan rumus berikut.

$$t = \frac{\sum D}{\sqrt{\frac{(N\sum D^2 - (\sum D)^2)}{N - 1}}}$$

$$t = \frac{32}{\sqrt{\frac{1.420 - 1.024}{9}}}$$

$$t = \frac{32}{\sqrt{44}}$$

$$t = \frac{32}{6.63}$$

$$t = 4.826$$

$$t = \mathbf{4.83} \text{ (dibulatkan)}$$

Hasil perhitungan tersebut selanjutnya dikonsultasikan dengan tabel (lihat lampiran 2). Dengan $df = N - 1 = 9$, maka didapat $t(9) = 1.83$. $p < .05$. Artinya t hitung lebih besar dibanding t tabel. Dengan demikian, terdapat perbedaan yang signifikan antara hasil *pretest* dan *posttest*. Untuk mengetahui peningkatannya dapat diketahui dengan cara berikut.

$$\text{Peningkatannya} = \frac{M_D}{M_{\text{pre}}} \times 100$$

$$\begin{aligned} &= \frac{3.2}{16.5} \times 100 \\ &= \mathbf{19.4\%} \end{aligned}$$

Data tersebut juga dapat dianalisis melalui SPSS dengan langkah-langkah sebagai berikut.

- Entry data atau buka file data yang akan dianalisis
- Pilih menu berikut ini

Analyze

Compare Means Paired samples T-test

Maka akan diperoleh hasil sebagai berikut:

Paired Samples Statistics					
		Mean	N	Std. Deviation	Std. Error Mean
Pair 1	posttest	19.7000	10	3.12872	.98939
	pretest	16.5000	10	3.02765	.95743

Paired Samples Test								
	Paired Differences					t	df	Sig. (2-tailed)
	Mean	Std. Deviation	Std. Error Mean	95% Confidence Interval of the Difference				
				Lower	Upper			
Pair 1 posttest - pretest	3.20000	2.09762	.66332	1.69945	4.70055	4.824	9	.001

Anova Satu Jalur (One-Way Anova)

Jika uji-t digunakan untuk menguji perbedaan antara dua kelompok data, maka Anova (*analysis of variance*) digunakan untuk menguji perbedaan antara tiga kelompok data atau lebih. Dengan demikian dapat dikatakan bahwa Anova pada dasarnya adalah ekstensi dari Uji T. Sesuai dengan namanya, Anova selalu berkaitan dengan variasi angka yang disebut varian. Dalam statistik deskriptif, varian dikenal sebagai bentuk kuadrat standar deviasi.

Anova tidak saja mampu menguji perbedaan antara tiga kelompok data atau lebih dari satu variabel bebas, tetapi juga dapat untuk menyelesaikan kelompok data yang berasal dari dua variabel bebas atau lebih. Anova dengan satu variabel bebas disebut Anova satu jalur, sedangkan Anova dengan dua variabel bebas atau lebih disebut Anova ganda atau Anova faktorial.

Dasar pemikiran Anova adalah bahwa harga varian total pada populasi dalam suatu pengamatan dapat dianalisis menjadi dua sumber, yaitu varian antar kelompok (*between*) dan varian dalam kelompok (*within*). Skor varian antar kelompok akan dijadikan pembilang, sedangkan skor varian dalam kelompok dijadikan penyebut. Sebagai salah satu bentuk statistik parametrik, Anova memiliki asumsi-asumsi antara lain; (1) sampel harus berasal dari populasi yang terdistribusi

secara normal, (2) nilai varian dalam kelompok harus homogen, (3) data berjenis interval atau rasio, dan (4) sampel diambil secara random.

Untuk melakukan pengolahan data dengan Anova, kita membutuhkan Mean Square antar kelompok (MS_{ak}) dan Mean Square dalam kelompok (MS_{dk}) sebagai berikut.

$$F = \frac{MS_{ak}}{MS_{dk}} = \frac{V_{ak} / df_{ak}}{V_{dk} / df_{dk}}$$

Sebelum sampai pada rumus di atas, kita memerlukan sejumlah skor yang diperlukan, yakni: varian total, varian antar kelompok, dan varian dalam kelompok sebagai berikut.

$$\text{Varian total} = V_t = \frac{\sum X_t^2 - (\sum X_t)^2 / N}{(N - 1)}$$

$$\text{Varian antar kelompok} = V_{ak} = \frac{(\sum X_1)^2}{(n_1)} + \frac{(\sum X_2)^2}{(n_2)} + \dots - \frac{(\sum X_t)^2}{N}$$

$$\text{Varian dalam kelompok} = V_{dk} = V_t - V_{ak}$$

Sebagai ilustrasi, seorang peneliti ingin menguji apakah kondisi lingkungan yang stress dapat berpengaruh terhadap kemampuan pemain bolabasket memasukan bola dalam keranjang. Selanjutnya peneliti mengambil 30 subyek secara random dan kemudian dibagi secara random pula ke dalam tiga kelompok: 10 subyek masuk dalam kelompok dengan stress tinggi, 10 subyek masuk dalam kelompok dengan stress sedang, dan 10 subyek masuk dalam kelompok dengan stress rendah, dengan hasil sebagai berikut.

Tabel 10.3: Skor memasukkan bolabasket dalam tiga kondisi stress yang berbeda

Kelompok 1 Stress tinggi	Kelompok 2 Stress sedang	Kelompok 3 Stress rendah
11	12	10
12	14	11
13	15	11
14	16	12
15	16	12
15	17	13
16	18	13
17	19	14
18	20	14
19	22	15

Mengingat Anova pengolahannya melibatkan banyak rumus sehingga lebih kompleks dibanding Uji T, maka pada kesempatan ini proses analisisnya dilakukan dengan SPSS, meski tetap terbuka peluang untuk melakukannya secara manual.

Data tersebut dianalisis melalui SPSS dengan langkah-langkah sebagai berikut.

- Entry data atau buka file data yang akan dianalisis
- Pilih menu berikut ini

Analyze

Compare Means

One-Way Anova

Maka akan diperoleh hasil sebagai berikut:

Descriptives

stress

	N	Mean	Std. Deviation	Std. Error	95% Confidence Interval for Mean		Minimum	Maximum
					Mean			
					Lower Bound	Upper Bound		
stress tinggi	10	15.0000	2.58199	.81650	13.1530	16.8470	11.00	19.00
stress sedang	10	16.9000	2.96086	.93630	14.7819	19.0181	12.00	22.00
stress rendah	10	12.5000	1.58114	.50000	11.3689	13.6311	10.00	15.00
Total	30	14.8000	2.98733	.54541	13.6845	15.9155	10.00	22.00

ANOVA

Stress

	Sum of Squares	df	Mean Square	F	Sig.
Between Groups	97.400	2	48.700	8.147	.002
Within Groups	161.400	27	5.978		
Total	258.800	29			

Dari hasil analisis di atas, maka dapat ditafsirkan bahwa terdapat perbedaan yang signifikan dalam kemampuan subyek memasukkan bolabasket pada tiga kondisi stress yang berbeda. Hal tersebut ditunjukkan dengan harga F sebesar 8,147 dengan signifikansi .002.

Contoh lain, seorang peneliti ingin menguji efektivitas tiga metode latihan terhadap peningkatan kapasitas aerobik. Dari hasil pengukuran diperoleh data sebagai berikut.

Tabel 10.4: Data Pengukuran Kapasitas Aerobik dari Penerapan Tiga Metode

<i>Metode A</i>	<i>Metode B</i>	<i>Metode C</i>
60,00	43,00	32,00
58,00	47,00	35,00
56,00	49,00	32,00
55,00	46,00	36,00
56,00	45,00	35,00
54,00	44,00	36,00
60,00	46,00	34,00
56,00	47,00	36,00
54,00	45,00	32,00
57,00	42,00	34,00

Dari data tersebut selanjutnya dilakukan analisis melalui SPSS dengan langkah-langkah sebagai berikut.

- Entry data atau buka file data yang akan dianalisis
- Pilih menu berikut ini

Analyze

Compare Means

Oneway Anova

Maka akan diperoleh hasil sebagai berikut:

Oneway Anova					
Kapasitas Aerobik					
	Sum of Squares	df	Mean Square	F	Sig.
Between Groups	2508,800	2	1254,400	318,316	,000
Within Groups	106,400	27	3,941		
Total	2615,200	29			

Multiple Comparisons							
Kapasitas Aerobik							
Scheffe							
(I) metode latihan	(J) metode latihan	Mean Difference (I-J)	Std. Error	Sig.	95% Confidence Interval		
					Lower Bound	Upper Bound	
Metode A	Metode B	11,20000*	,88778	,000	8,9006	13,4994	
	Metode C	22,40000*	,88778	,000	20,1006	24,6994	
Metode B	Metode A	-11,20000*	,88778	,000	-13,4994	-8,9006	
	Metode C	11,20000*	,88778	,000	8,9006	13,4994	
Metode C	Metode A	-22,40000*	,88778	,000	-24,6994	-20,1006	
	Metode B	-11,20000*	,88778	,000	-13,4994	-8,9006	

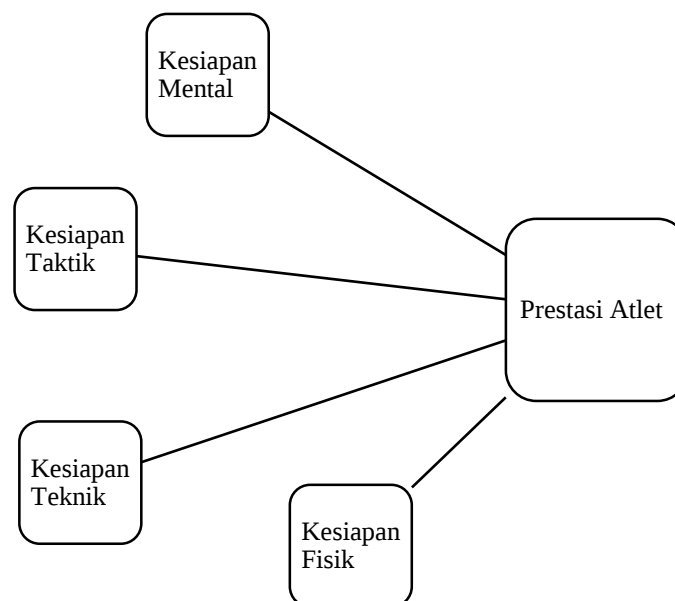
*. The mean difference is significant at the 0.05 level.

Dari hasil analisis di atas dapat diinterpretasi bahwa terdapat perbedaan yang signifikan di antara ketiga metode dalam mempengaruhi kapasitas aerobik. Hal ini ditunjukkan dengan harga F sebesar 318,316 pada signifikansi .000. Setelah dipastikan terdapat perbedaan yang signifikan, analisis dilanjutkan dengan *Post Hoc Comparisons*³ dengan menggunakan uji Scheffe. Hasil analisis menunjukkan bahwa kombinasi perbedaan antara metode A, metode B, dan metode C, kesemuanya signifikan pada .000.

³ Post hoc comparisons hanya dapat dilakukan ketika perbedaan antar kelompok yang ditunjukkan oleh harga F signifikan. Analisis ini tidak dimaksudkan untuk menguji hipotesis, tetapi lebih merupakan upaya melihat kombinasi perbedaan mean

Bab 11 Analisis Multivariat

Dalam pengembangan keilmuan, analisis multivariat menjadi sangat dibutuhkan. Ini mengingat, hampir semua fenomena ilmu pengetahuan tidak pernah beralur tunggal. Sebagai ilustrasi, ketika seseorang mengalami sakit, apalagi yang bersifat komplikasi, maka ia membutuhkan penanganan dokter dengan sejumlah keahlian, seperti dokter bedah, dokter penyakit dalam, dokter anastasi, dan sebagainya. Demikian juga ketika kita mengkaji tentang prestasi atlet, maka banyak melibatkan variabel, mencakup aspek fisik, teknik, taktik, dan psikologis (Bompa & Haff, 2009). Seorang atlet yang memiliki kualitas fisik yang memadai—misalnya power yang bagus dan vo2max yang tinggi—tidak serta merta akan memenangkan pertandingan, tetapi masih tergantung dari variabel yang lain. Sebuah tim yang bertabur bintang—yang secara teknik memiliki keunggulan—tidak otomatis akan memenangkan pertandingan, ada faktor taktik dan psikologis yang juga berpengaruh. Karena itu, penggunaan analisis multivariat dalam pengembangan keilmuan menjadi sebuah keniscayaan.



Gambar 11.1: Logika multivariat yang mempengaruhi prestasi atlet

Ada sejumlah teknik statistik yang dapat dikategorikan ke dalam analisis multivariat, seperti Regresi Ganda, Anova faktorial, Manova, Analisis Faktor, dan Analisis Diskriminan. Selain itu, ada juga teknik analisis yang lebih *sophisticated* seperti *Structural Equation Modelling* (SEM) dengan Lisrel, Prelis, dan Simplis. Dalam logika SEM, relasi antar variabel tidak saja bersifat langsung, tetapi juga bersifat tidak langsung. Pada bagian ini tidak akan dijelaskan keseluruhan analisis multivariat, melainkan fokus pada Regresi Ganda, Anova Faktorial, dan Manova.

Regresi Ganda (*Multiple Regression*)

Dalam analisis multivariat, regresi menjadi bagian analisis statistik yang utama. Karena itulah, Pedhazur, seorang ahli statistik dari New York University dan Kerlinger seorang ahli psikologi menulis buku khusus setebal 822 halaman yang berjudul “*Multiple Regression in Behavioral Research*”. Buku tersebut telah menjadi rujukan utama pembelajaran statistik di perguruan tinggi ternama di seluruh dunia. Mengapa regresi ganda menjadi urgen dalam analisis statistik? Ini karena regresi ganda berusaha menggabungkan kebutuhan inferensial dalam logika statistik dan penjelasan rasional-teoretik dalam memahami relasi suatu variabel. Seperti diketahui bahwa tugas utama ilmu pengetahuan adalah menjelaskan suatu fenomena, misalnya fenomena prestasi belajar, prestasi atlet, dan kinerja guru, yang dalam konstruksi metode penelitian disebut variabel terikat. Tanpa ada penjelasan yang rasional atas relasi suatu variabel, maka tidak ada makna yang berarti dalam suatu penelitian.

Ketika seorang peneliti ingin mengkaji “kualitas latihan”, maka diyakini banyak faktor yang mempengaruhinya, mulai dari faktor pelatih, sarana prasarana, motivasi atlet, kompetisi, sampai pada dukungan ilmiah (Bompa & Haff, 2009). Sebuah proses penelitian yang baik harus mampu mengidentifikasi sumber-sumber variabel yang diduga atau dominan mempengaruhi kualitas latihan. Pendek kata, variabel yang mempengaruhi kualitas latihan sangat kompleks, lalu bagaimana mengidentifikasi variabel-variabel tersebut? Dalam konteks yang demikian, analisis regresi ganda menjadi penting untuk menguji relasi antara variabel tersebut.

Tabel 11.1: Data Frekuensi, Intensitas dan Kebugaran dari 30 Subjek

Subjek	Frek. Latihan (X ₁)	Int. Latihan (X ₂)	Kebugaran (Y)	Subjek	Frek. Latihan (X ₁)	Int. Latihan (X ₂)	Kebugaran (Y)
1	6	5	65	16	2	1	25
2	4	3	40	17	4	3	40
3	3	2	30	18	2	1	30
4	3	3	35	19	3	2	30
5	6	5	60	20	3	3	35
6	2	1	30	21	6	5	60
7	2	1	25	22	5	4	50
8	5	4	50	23	5	4	55
9	4	4	45	24	6	5	65
10	5	4	55	25	4	4	45
11	6	5	60	26	2	1	25
12	5	4	50	27	4	3	40
13	5	4	55	28	2	1	30
14	4	4	45	29	3	2	30
15	6	5	65	30	3	3	35

Jika regresi sederhana hanya memiliki satu variabel prediktor, maka pada regresi ganda, variabel prediktornya lebih dari satu, bisa dua, bisa tiga, dan seterusnya. Bentuk pertanyaan mendasar dari setiap analisis regresi adalah: Perubahan apa yang diharapkan dapat terjadi dalam variabel terikat sebagai hasil dari perubahan dalam variabel bebas. Sebagai ilustrasi, seorang peneliti ingin mengkaji fenomena kebugaran sebagai variabel terikat. Secara teoretik, variabel kebugaran antara lain dipengaruhi oleh frekuensi latihan dan intensitas latihan. Selanjutnya, proses pengukuran dilakukan dan memperoleh data sebagai berikut (*dummy data*).

Ide dasar regresi pada hakikatnya sama, yakni:

$$Y = a + bX + e$$

di mana,

Y = skor variabel terikat

a = intercept

b = koefisien regresi

X = skor variabel bebas

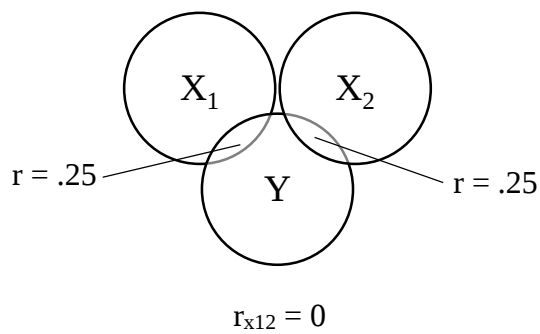
e = kesalahan atau residu

Rumus tersebut dapat diperluas dengan menambah variabel bebas menjadi sebagai berikut.

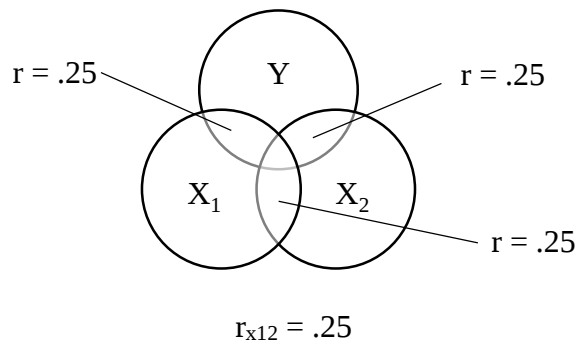
$$Y = a + b_1X_1 + b_2X_2 + \dots + b_kX_k + e$$

Jika menggunakan dua variabel bebas, maka ada X_1 dan X_2 , jika tiga variabel bebas berarti ada X_1 , X_2 , dan X_3 , dan begitu seterusnya. Dalam logika regresi, relasi variabel X terhadap Y yang disimbulkan dengan R mengandung pengertian proporsi varian pada Y yang dapat dijelaskan oleh variabel X , selebihnya adalah varian kesalahan atau residual. Pendek kata, $1 - r^2$ merupakan varian yang belum dapat dijelaskan. Penambahan variabel X akan mengurangi varian kesalahan yang akan terjadi.

Selain itu, relasi antar variabel bebas juga perlu mendapat perhatian. Adakalanya antara variabel bebas tidak ada korelasi yang signifikan sebagaimana tampak pada gambar a, tetapi banyak terjadi ada korelasi yang signifikan di antara variabel bebas seperti tampak pada gambar b. Apabila di antara variabel bebas terdapat korelasi, maka dapat dipastikan ada varian sama yang dijelaskan oleh kedua variabel tersebut. Inilah yang kemudian disebut dengan *multicollinearity*. Apabila antar variabel bebas terjadi *multicollinearity*, maka estimasi persamaan regresi menjadi sulit untuk ditentukan. Meski para ahli statistik sendiri belum sepakat mengenai seberapa besar koefisien korelasi yang dianggap sebagai *multicollinearity*. Dalam penelitian eksperimen, sangat mungkin tidak ada korelasi antar variabel bebas. Namun dalam penelitian noneksperimen, relasi antar variabel bebas menjadi sulit dihindari.



Gambar (a)



Gambar (b)

Dengan memperhatikan data pada tabel 11.1, maka dapat dilakukan analisis regresi ganda melalui SPSS dengan langkah-langkah sebagai berikut.

- Entry data atau buka file data yang akan dianalisis
- Pilih menu berikut ini

Analyze
 Regression
 Linier

Maka akan diperoleh hasil sebagai berikut:

Model Summary^b					
Model		R	R Square	Adjusted R Square	Std. Error of the Estimate
dimension0	1	,976 ^a	,953	,950	2,97992
a. Predictors: (Constant), Int.lat, Frek.lat					
b. Dependent Variable: kebugaran					

ANOVA^b						
Model		Sum of Squares	df	Mean Square	F	Sig.
1	Regression	4867,742	2	2433,871	274,087	,000 ^a
	Residual	239,758	27	8,880		
	Total	5107,500	29			
a. Predictors: (Constant), Int.lat, Frek.lat						
b. Dependent Variable: kebugaran						

		Coefficients ^a				
Model		Unstandardized Coefficients		Standardized Coefficients	t	Sig.
		B	Std. Error	Beta		
1	(Constant)	8,274	1,831		4,520	,000
	Frek.lat	7,774	1,368	,843	5,683	,000
	Int.lat	1,290	1,382	,138	,934	,359
a. Dependent Variable: kebugaran						

Dari hasil analisis di atas, dapat diinterpretasi bahwa korelasi antara variabel prediktor 1 (frekuensi latihan) dan variabel prediktor 2 (intensitas latihan) terhadap kriterium (kebugaran) sebesar .976. Dari analisis regresi yang dilakukan nampak bahwa harga F sebesar 274,087 pada signifikansi .000. Korelasi prediktor 1 terhadap kriterium sebesar .843 dan korelasi prediktor 2 terhadap kriterium sebesar .138. Data ini sekaligus juga menguatkan pendapat bahwa kuatnya hubungan yang bersifat bivariat, satu variabel bebas dan satu variabel terikat, akan melemah ketika disertakan variabel bebas berikutnya. Dalam kasus di atas, hubungan antara frekuensi latihan dan kebugaran sebesar .975. Namun setelah dimasukkan variabel bebas lain berupa intensitas latihan, maka korelasi keduanya melemah menjadi .843.

Dari analisis tersebut juga diperoleh konstanta sebesar 8,274 dan koefisien regresi prediktor 1 sebesar 7,774 dan prediktor 2 sebesar 1,290. Jika harga-harga tersebut dibuat dalam persamaan garis regresi, maka diperoleh persamaan sebagai berikut:

$$Y = a + b_1X_1 + b_2X_2 + e$$

$$= 8,274 + 7,774X_1 + 1,290X_2 + e$$

Anova Faktorial

Dari sisi istilah, Anova faktorial sering juga disebut Anova dua jalur, yaitu suatu teknik statistik yang digunakan untuk menguji perbedaan di antara kelompok data yang berasal dari dua variabel bebas atau lebih. Istilah “jalur” sebenarnya lebih mengacu pada jumlah variabel bebas yang dianalisis, sehingga kalau variabel bebasnya ada tiga, maka logikanya dapat disebut Anova tiga jalur. Dari analisis ini, peneliti dapat menguji pengaruh secara bersama-sama sejumlah variabel bebas terhadap variabel terikat, baik secara terpisah maupun gabungan (interaksi).

Dalam konteks Anova, variabel bebas disebut sebagai faktor. Sebagai sebuah variabel, faktor bisa berupa metode latihan, metode mengajar, gaya kepemimpinan, tingkat motivasi, jenis kelamin, dan sebagainya. Apabila faktor tersebut terbagi dalam dua atau tiga bagian, maka bagian-bagian tersebut disebut sel. Perlu diingat bahwa dalam pembagian faktorial tersebut, setiap sel

harus terpisah secara jelas dan semua data atau kasus dapat terbagi habis dalam sel yang ada, tidak boleh ada yang tersisa. Dengan pola $p \times q$ faktorial, ada banyak kemungkinan yang dapat dilakukan, misalnya 2×2 , 2×3 , 3×3 , dan seterusnya. Berikut adalah contoh desain faktorial dengan dua variabel bebas, A dan B. Variabel A terdiri dari dua set: A_1 dan A_2 . Sementara untuk variabel B terdiri dari tiga set: B_1 , B_2 , dan B_3 . Sel yang dihasilkan dari pembagian silang tersebut adalah A_1B_1 , A_1B_2 , dan seterusnya.

	B_1	B_2	B_3
A_1	A_1B_1	A_1B_2	A_1B_3
A_2	A_2B_1	A_2B_2	A_2B_3

Anova faktorial lazimnya digunakan untuk menganalisis data pada penelitian eksperimen, yakni menguji pengaruh beberapa variabel bebas terhadap sebuah variabel terikat. Ada beberapa keuntungan ketika kita menggunakan analisis Anova faktorial, yaitu (1) ada peluang untuk menguji interaksi antar variabel bebas dalam mempengaruhi variabel terikat, (2) melakukan kontrol terhadap varian kesalahan mengingat variabelnya tidak tunggal, (3) lebih efisien karena dapat menguji pengaruh variabel, baik secara terpisah maupun kombinasi, secara simultan, dan (4) pengaruh perlakuan dapat dikaji dalam kondisi yang berbeda sehingga generalisasinya menjadi lebih luas.

Sebagai contoh, seorang peneliti ingin menguji pengaruh metode dan interval latihan terhadap kapasitas aerobik. Ada 30 subyek yang ditempatkan secara acak dalam sel yang tersedia dan setelah dilakukan perlakuan, maka diperoleh data sebagaimana tabel berikut.

Tabel 11.2: Data Kapasitas Aerobik Menurut Jenis Interval dan Metode

<i>Jenis Interval</i>	<i>Jenis Metode</i>		
	Metode A	Metode B	Metode C
Interval A	60, 58, 56, 55, 56	43, 47, 49, 46, 45	32, 35, 32, 36, 35
Interval B	54, 60, 56, 54, 57	44, 40, 47, 45, 42	36, 34, 36, 32, 34

Dari data tersebut selanjutnya dilakukan analisis melalui SPSS dengan langkah-langkah sebagai berikut.

- Entry data atau buka file data yang akan dianalisis
- Pilih menu berikut ini

Analyze

General Linier Model

Univariate

Maka akan diperoleh hasil sebagai berikut:

Univariate Analysis of Variance

Tests of Between-Subjects Effects					
Dependent Variable: Kapasitas Aerobik					
Source	Type III Sum of Squares	Df	Mean Square	F	Sig.
Corrected Model	2514,400 ^a	5	502,880	119,733	,000
Intercept	61834,800	1	61834,800	14722,571	,000
metode	2508,800	2	1254,400	298,667	,000
interval	2,133	1	2,133	,508	,483
metode * interval	3,467	2	1,733	,413	,666
Error	100,800	24	4,200		
Total	64450,000	30			
Corrected Total	2615,200	29			
a. R Squared = ,961 (Adjusted R Squared = ,953)					

Dari hasil analisis di atas dapat diinterpretasi bahwa metode latihan berpengaruh signifikan terhadap peningkatan kapasitas aerobik. Hal ini ditunjukkan dengan harga F sebesar 296,667 pada signifikansi .000. Sementara itu, variabel interval tidak mempengaruhi peningkatan kapasitas aerobik. Hal ini ditunjukkan dengan harga F sebesar .508 pada signifikansi .483. Selain itu, hasil analisis juga menunjukkan bahwa tidak ada interaksi antara variabel metode dan variabel interval dalam mempengaruhi peningkatan kapasitas aerobik. Hal ini ditunjukkan oleh harga F sebesar .413 dengan signifikansi .666.

Manova (Multivariate Analysis of Variance)

Pada dasarnya Manova merupakan perluasan dari Anova. Jika Anova menguji perbedaan satu variabel terikat terhadap beberapa variabel bebas, maka Manova menguji perbedaan beberapa variabel terikat terhadap beberapa variabel bebas. Pedhazur (1982) menyatakan bahwa “*Manova is extension of univariate analysis of variance designed to test simultaneously differences among groups on multiple dependent variables (p.710).*” Terkait dengan pengujian perbedaan antar kelompok yang melibatkan beberapa variabel terikat, Manova sering dipertukarkan dengan

analisis Diskriminan. Meski demikian, disarankan Manova digunakan lebih dulu untuk menguji signifikansi perbedaan keseluruhan kelompok. Jika hipotesis nol ditolak, maka analisis diskriminan dapat digunakan untuk mengidentifikasi variabel yang dominan membedakan antar kelompok.

Tabel 11.3: Nilai Rerata dan Standar Deviasi Untuk Setiap Dimensi dari Tiga Kelompok Responden

Dimensi Traits	Kelompok Responden		
	ABT	ABR	KBA
Ambisi Prestatif	21.42 (2.49)	19.93 (3.02)	17.90 (3.73)
Kerja Keras	34.10 (3.30)	33.20 (4.56)	31.60 (5.56)
Gigih	28.62 (2.66)	27.68 (4.44)	25.85 (3.66)
Mandiri	32.88 (3.37)	30.78 (4.29)	31.03 (4.98)
Komitmen	26.03 (2.45)	23.35 (4.49)	20.35 (3.30)
Cerdas	23.75 (3.07)	21.95 (3.23)	22.45 (4.06)
Swakendali	30.48 (2.81)	29.43 (4.43)	29.52 (6.00)
Total	197.28 (15.58)	186.30 (22.25)	178.70 (24.91)

Keterangan: angka dalam kurung adalah standar deviasi

Untuk tujuan tersebut, bentuk uji hipotesis yang banyak digunakan adalah Wilks' λ (lamda). Jika hasil analisis Wilks' λ menunjukkan perbedaan yang signifikan, maka analisis dilanjutkan dengan melihat kombinasi perbedaan di antara kelompok. Dalam pengujian berlaku ketentuan: jika signifikansi $> 0,05$, maka hipotesis nol diterima dan jika signifikansi $< 0,05$, maka hipotesis nol ditolak.

Sebagai ilustrasi, akan dipaparkan hasil penelitian terhadap tiga kelompok subyek yang diukur berdasarkan intensitas *traits* yang dimiliki. Dari data seperti tampak pada tabel 11.3 ditunjukkan bahwa terdapat perbedaan nilai rata-rata pada setiap dimensi di tiga kelompok responden. Misalnya pada dimensi ambisi prestatif, nilai rata-rata pada kelompok atlet berprestasi tinggi (ABT) = 21.42; kelompok atlet berprestasi rendah (ABR) = 19.93; dan kelompok bukan atlet (KBA) = 17.90. Pada dimensi kerja keras, nilai rata-rata ABT= 34.10; ABR= 33.20; dan KBA= 31.60. Apakah perbedaan-perbedaan tersebut signifikan secara statistik? Inilah yang masih akan diuji dengan menggunakan Analisis Multivariat.

Tabel 11.4: Hasil Analisis Multivariat

Hipotesis	Arah hipotesis	Nilai F	Signifikansi	Uji Wilks' Λ (Lambda)
Mayor (Seluruh dimensi)	abt>abr>kba	27.696	.000	*
Minor				
Ambisi prestatif	abt>abr>kba	8.981	.000	
Kerja keras	abt>abr>kba	6.174	.003	
Gigih	abt>abr>kba	6.986	.001	
Mandiri	abt>abr>kba	5.843	.004	
Komitmen	abt>abr>kba	20.657	.000	
Cerdas	abt>abr>kba	4.508	.013	
Swakendali	abt>abr>kba	3.074	.050	

Tabel 11.5: Perbandingan Kelompok ABT, ABR, dan KBA Berdasarkan Perbedaan Mean

Dimensi	Kelompok Responden		Perbedaan Mean	Signifikansi
	Utama	Pembanding		
Ambisi Prestatif	ABT	ABR	1.418*	.048
		KBA	3.964*	.000
Kerja Keras	ABT	ABR	0.468	.646
		KBA	4.819*	.001
Gigih	ABT	ABR	0.689	.404
		KBA	4.172*	.000
Mandiri	ABT	ABR	1.665	.081
		KBA	4.184*	.001
Komitmen	ABT	ABR	2.486*	.002
		KBA	6.689*	.000
Cerdas	ABT	ABR	1.513	.054
		KBA	2.838*	.008
Swakendali	ABT	ABR	0.588	.565
		KBA	3.430*	.015
Total	ABT	ABR	8.827	.064
		KBA	30.096*	.000

Secara umum, berdasarkan analisis statistik multivariat diperoleh informasi bahwa terdapat perbedaan yang signifikan pada dimensi total dengan nilai $F = 27.696$ pada signifikansi 0.000. Artinya, kelompok ABT memiliki *trait* kepribadian dengan intensitas yang lebih kuat dibanding dengan kelompok ABR dan KBA. Selain itu, dari analisis multivariat tersebut juga dapat diketahui bahwa dimensi ambisi prestatif dan komitmen merupakan variabel dominan yang membedakan kualitas kepribadian antara kelompok atlet berprestasi tinggi dengan kelompok atlet berprestasi rendah maupun kelompok bukan atlet. Setelah dipastikan terdapat perbedaan yang signifikan, maka analisis dilanjutkan dengan *Pairwise Comparisons* yang bertujuan untuk menemukan *traits* mana yang dominan ada pada kelompok ABT. Dari hasil analisis tersebut tampak bahwa terdapat

perbedaan signifikan antara kelompok ABT dan ABR terjadi pada dimensi ambisi prestatif dengan perbedaan rata-rata sebesar 1.418 pada $p .048$ dan pada dimensi komitmen dengan perbedaan rata-rata sebesar 2.486 pada $p .002$. Sementara itu, jika dibandingkan dengan kelompok KBA, seluruh dimensi berbeda secara signifikan.

Bab 12 Uji Statistik Nonparametrik

Analisis statistik nonparametrik diberlakukan apabila persyaratan data, terutama dari aspek normalitas tidak dapat terpenuhi. Dalam terminologi statistik, normalitas data bertalian dengan hasil pengujian asumsi data atau dapat juga berkaitan dengan jumlah anggota sampel. Sampel dalam jumlah besar (≥ 30) diasumsikan data akan berdistribusi normal. Karena itu, uji nonparametrik biasanya juga digunakan untuk menganalisis data dengan jumlah sampel kecil. Selain itu, uji nonparametrik biasanya juga dipakai untuk menganalisis data dengan skala ordinal dan nominal. Sebagai peneliti tentu sangat menginginkan suatu data berdistribusi normal, namun tidak selamanya hal tersebut dapat terpenuhi. Itu sebabnya, meski uji nonparametrik merupakan solusi, tetapi mutu analisisnya tidak sekuat statistik parametrik. Dari sisi pengolahan data, teknik analisis statistik nonparametrik lebih sederhana, tidak membutuhkan rumus dan angka-angka yang rumit. Berikut dipaparkan sejumlah teknik analisis statistik nonparametrik.

Tabel 12.1: Pilihan uji nonparametrik ketika uji parametrik tidak dapat dilakukan

Kondisi sampel	Uji parametrik	Uji nonparametrik
Satu sampel	Uji T satu sampel	Wilcoxon signed rank test, Chi-square
Sampel berpasangan	Uji T sampel sejenis	Sign test, Wilcoxon signed-rank test
Dua sampel berbeda	Uji T sampel berbeda	Mann-Whitney test, Wilcoxon Rank Sum Test, Kolmogorov-Smirnov
Tiga atau lebih sampel berbeda	One-way Anova	Kruskal-Wallis test

Chi-square

Chi-square (χ^2) dikembangkan oleh Karl Pearson, ahli statistik matematika dari Inggris. Chi-square merupakan teknik analisis statistik yang digunakan untuk menguji perbedaan frekuensi. Sebagai contoh, seorang mahasiswa ingin meneliti tentang sikap mahasiswa terhadap kewajiban menyusun skripsi sebagai prasyarat kelulusan. Secara random diambil 100 mahasiswa sebagai sampel penelitian. Sikap yang dimaksud dinyatakan dalam dua pernyataan, yaitu: setuju dan tidak setuju. Dari data yang dikumpulkan didapat bahwa 60 orang menyatakan setuju dan 40 orang menyatakan tidak setuju. Data ini kemudian diolah dengan menggunakan Chi-square. Rumus untuk mencari nilai Chi-square adalah sebagai berikut.

$$\chi^2 = \sum \left[\frac{(fo - fe)^2}{fe} \right]$$

dimana.

χ^2 = nilai chi-square

fo = frekuensi yang diperoleh

fe = frekuensi yang diharapkan

Koefisien fo dalam kasus di atas adalah 60 dan 40, sedangkan koefisien fe didapat dari jumlah frekuensi kedua kelompok, dibagi jumlah kelompok. Jadi $(60 + 40) : 2 = 50$. Untuk memudahkan, penghitungan biasanya dilakukan dalam bentuk tabel seperti berikut.

Tugas Skripsi	Fo	fe	fo-fe	$(fo-fe)^2$	$(fo-fe)^2 : fe$
Setuju	60	50	10	100	2
Tidak setuju	40	50	-10	100	2
					$\chi^2 = 4$

Tabel di atas menunjukkan hasil perhitungan nilai Chi-square terhadap perbedaan frekuensi antara mahasiswa yang setuju dengan tugas skripsi dan yang tidak setuju, yaitu sebesar 4. Untuk memastikan apakah perbedaan tersebut signifikan, angka tersebut perlu dikonsultasikan dengan nilai chi-square dalam tabel (lampiran 3), dengan memperhatikan derajat kebebasan, yaitu jumlah kelompok dikurangi satu. Dalam chi-square berlaku ketentuan, jika nilai empirik hasil perhitungan lebih besar dibandingkan dengan nilai teoretiknya dalam tabel, maka dapat disimpulkan terdapat perbedaan yang signifikan diantara kelompok yang ada. Dalam contoh di atas, dengan db = 1, pada taraf signifikansi 5% nilai teoretiknya sebesar 3,84 dan pada taraf signifikansi 1% sebesar 6,63.

Memperhatikan hasil penghitungan di atas, maka nilai 4 lebih besar dibandingkan 3,84 pada signifikansi 5%, tetapi lebih kecil dibandingkan 6,63 pada signifikansi 1%. Dengan demikian dapat disimpulkan bahwa pada taraf signifikansi 5%, terdapat perbedaan yang bermakna di kalangan mahasiswa dalam menyikapi skripsi. Sementara itu, pada taraf signifikansi 1%, tidak terdapat perbedaan yang bermakna.

Data tersebut juga dapat dianalisis melalui SPSS dengan langkah-langkah sebagai berikut.

- Entry data atau buka file data yang akan dianalisis
- Pilih menu berikut ini

Analyze
Nonparametric tests
Chi square

Maka akan diperoleh hasil sebagai berikut:

Sikap terhadap skripsi

	Observed N	Expected N	Residual
Setuju	60	50.0	10.0
tidak setuju	40	50.0	-10.0
Total	100		

Test Statistics

	sikap thp skripsi
Chi-Square	4.000 ^a
Df	1
Asymp. Sig.	.046

a. 0 cells (0.0%) have expected frequencies less than 5. The minimum expected cell frequency is 50.0.

Dari hasil penghitungan melalui SPSS tampak sejalan dengan penghitungan manual. Chi-square sebesar 4.000 dengan df (*degree of freedom*) 1 signifikan pada .046.

Sign Test (Uji Tanda)

Sign test digunakan untuk menguji perbedaan dua rerata dari kelompok sampel berpasangan ketika normalitas data tidak dapat terpenuhi. Seperti yang sudah kita bahas sebelumnya, apabila data berdistribusi normal, maka kita tentu menggunakan T Test. Sesuai dengan namanya, pengujian Sign test menggunakan tanda positif dan negatif di antara pasangan data pengamatan. Jika ada data yang sama, artinya tidak ada perubahan, disimbolkan dengan nol, maka data tersebut diabaikan. Sebagai ilustrasi, seorang peneliti ingin menguji apakah “pendekatan modifikasi olahraga” yang menekankan pada kesenangan dan keterlibatan dalam aktivitas fisik, dapat meningkatkan kemampuan motorik siswa. Untuk tujuan tersebut, dilakukan pembelajaran terhadap 10 siswa (kelas kecil) dengan data sebelum dan sesudah perlakuan.

Kemampuan motorik siswa		Perbedaan + atau –
sebelum	sesudah	(sesudah – sebelum)
100	120	+
125	117	–
101	96	–
114	111	–
100	100	0
98	98	0
95	106	+
128	140	+
120	131	+
104	110	+
		5+, 3–

Dalam uji tanda berlaku ketentuan: Hipotesis nol diterima apabila probabilitas hasil sampel ≥ 0.05 dan hipotesis nol ditolak apabila probabilitas hasil sampel < 0.05 . Sehubungan dengan data penelitian di atas, maka uji statistiknya dapat dinyatakan sebagai berikut.

H0: $p = 0.5$ (tidak ada peningkatan kemampuan motorik)

H1: $p \neq 0.5$ (ada peningkatan kemampuan motorik), dengan tingkat signifikansi 0.05 atau 5%

Data tersebut selanjutnya dianalisis melalui SPSS dengan langkah-langkah sebagai berikut.

- Entry data atau buka file data yang akan dianalisis
- Pilih menu berikut ini

Analyze

Nonparametric tests

Related samples – sign test

Maka akan diperoleh hasil sebagai berikut:

Hypothesis Test Summary

	Null Hypothesis	Test	Sig.	Decision
1	The median of differences between sebelum and sesudah equals 0.	Related-Samples Sign Test	.727 ¹	Retain the null hypothesis.

Asymptotic significances are displayed. The significance level is .05.

¹Exact significance is displayed for this test.

Berdasarkan hasil analisis tersebut nampak bahwa nilai signifikansi sebesar 0,727, lebih besar dari 0,05. Artinya, hipotesis nol diterima dan hipotesis kerja ditolak. Dengan demikian dapat disimpulkan bahwa tidak ada peningkatan kemampuan motorik anak setelah mereka diberikan pembelajaran modifikasi.

Wilcoxon Test

Secara fungsional, uji Wilcoxon relatif sama dengan uji tanda. Dari sisi mutu analisis, uji Wilcoxon lebih baik karena tidak saja memperhatikan tanda dalam penghitungan, melainkan juga selisih di antara skor. Sebagai contoh, kita dapat menggunakan data pada uji tanda di atas. Selain kita perlu menghitung selisih skor, kita juga perlu mengurutkan atau merangking selisih skor tersebut tanpa memperhatikan tandanya. Skor terkecil diberi nomor urut atau rangking 1, skor terkecil kedua diberi rangking 2, dan begitu seterusnya. Jika terdapat selisih skor yang sama, maka digunakan rerata dari rangking tersebut. Pengujian dilakukan terhadap jumlah skor yang paling sedikit di antara tanda positif atau negatif. Jika skor yang bertanda positif jumlahnya lebih kecil, maka kita menggunakan skor positif, begitu sebaliknya.

Kemampuan motorik siswa			Rangking	
sebelum	sesudah	selisih	Selisih +	Selisih -
100	120	+20	8	
125	117	-8		4
101	96	-5		2
114	111	-3		1
100	100	0		
98	98	0		
95	106	+11	5.5	
128	140	+12	7	
120	131	+11	5.5	
104	110	+6	3	
			29	7

Catatan: jika ada data yang selisihnya nol, maka data tersebut diabaikan, tidak disertakan dalam analisis

Dalam uji Wilcoxon berlaku ketentuan menerima hipotesis nol jika probabilitas hasil sampel ≥ 0.05 dan menolak hipotesis nol apabila probabilitas hasil sampel < 0.05 . Secara statistik dapat dirumuskan sebagai berikut.

$H_0: \mu_1 = \mu_2$ (tidak ada perbedaan antar kelompok)

$H_1: \mu_1 \neq \mu_2$ (ada perbedaan antar kelompok), dengan tingkat signifikansi 0.05 atau 5%

Data tersebut selanjutnya dianalisis melalui SPSS dengan langkah-langkah sebagai berikut.

- Entry data atau buka file data yang akan dianalisis
- Pilih menu berikut ini

Analyze

Nonparametric tests

Related samples – Wilcoxon signed rank test

Maka akan diperoleh hasil sebagai berikut:

Hypothesis Test Summary

	Null Hypothesis	Test	Sig.	Decision
1	The median of differences between sebelum and sesudah equals 0.	Related- Samples Wilcoxon Signed Rank Test	.123	Retain the null hypothesis.

Asymptotic significances are displayed. The significance level is .05.

Berdasarkan hasil analisis tersebut nampak bahwa nilai signifikansi sebesar 0,123, lebih besar dari 0,05. Artinya, hipotesis nol diterima dan hipotesis kerja ditolak. Dengan demikian dapat disimpulkan bahwa tidak ada peningkatan kemampuan motorik anak setelah mereka diberikan pembelajaran modifikasi. Ternyata pengujian yang dilakukan dengan uji tanda maupun uji Wilcoxon mendapatkan hasil yang sama.

Mann-Whitney U Test

Uji ini disebut juga dengan Wilcoxon Rank Sum Test, digunakan untuk menguji perbedaan dua median dari kelompok yang berbeda. Uji ini diberlakukan apabila variabel terikat berskala ordinal atau interval/rasio namun tidak berdistribusi normal. Sebagai ilustrasi, seorang peneliti ingin menguji dampak latihan olahraga terhadap berat badan seseorang. Untuk kepentingan tersebut, peneliti melakukan pengukuran berat badan kepada 10 orang yang rutin berolahraga dan 10 orang yang tidak pernah berolahraga. Hasil pengukuran tersebut selanjutnya dipaparkan sebagai berikut.

Berat badan (kg)		Rangking	
Tidak berolahraga	Rutin berolahraga	Tidak berolahraga	Rutin berolahraga
60	59	2.5	1
65	60	4.5	2.5
67	68	6	8
75	76	16	17
74	68	14.5	8
80	72	18	13
89	86	20	19
74	70	14.5	11
68	70	8	11
70	65	11	4.5

Dalam uji Mann-Whitney berlaku ketentuan menerima hipotesis nol jika probabilitas hasil sampel ≥ 0.05 dan menolak hipotesis nol apabila probabilitas hasil sampel < 0.05 . Secara statistik dapat dirumuskan sebagai berikut.

H0: $\mu_1 = \mu_2$ (tidak ada perbedaan antar kelompok)

H1: $\mu_1 \neq \mu_2$ (ada perbedaan antar kelompok), dengan tingkat signifikansi 0.05 atau 5%

Data tersebut selanjutnya dianalisis melalui SPSS dengan langkah-langkah sebagai berikut.

- Entry data atau buka file data yang akan dianalisis
- Pilih menu berikut ini

Analyze

Nonparametric tests

Independent samples – Mann-whitney U test

Maka akan diperoleh hasil sebagai berikut:

Ranks				
	kondisi	N	Mean Rank	Sum of Ranks
beratbadan	tidak berolahraga	10	11.50	115.00
	rutin berolahraga	10	9.50	95.00
	Total	20		

Test Statistics ^a	
	beratbadan
Mann-Whitney U	40.000
Wilcoxon W	95.000
Z	-.759
Asymp. Sig. (2-tailed)	.448
Exact Sig. [2*(1-tailed Sig.)]	.481 ^b

a. Grouping Variable: kondisi

b. Not corrected for ties.

Berdasarkan hasil analisis tersebut nampak bahwa nilai signifikansi sebesar 0,448, lebih besar dari 0,05. Artinya, hipotesis nol diterima dan hipotesis kerja ditolak. Dengan demikian dapat disimpulkan bahwa latihan olahraga, setidaknya untuk data di atas, tidak berpengaruh terhadap berat badan seseorang.

Kruskal-Wallis Test

Teknik statistik ini digunakan untuk menguji perbedaan antar kelompok, dua atau lebih, di mana variabel terikatnya berskala ordinal, atau interval/rasio tetapi tidak berdistribusi normal. Kruskal-Wallis Test pada dasarnya merupakan perluasan dari Mann-Whitney U Test, yang hanya dapat menguji perbedaan dua buah kelompok.

Sebagai ilustrasi, seorang guru ingin mengajarkan lempar lembing kepada siswa dengan tiga metode yang berbeda, sebutlah metode A, B, dan C. Sebanyak 12 siswa dibagi ke dalam tiga kelompok, masing-masing kelompok empat siswa. Setelah dilakukan pembelajaran, kemampuan lempar lembing siswa diukur dengan hasil sebagai berikut.

Kemampuan siswa melempar lembing (meter)					
Metode A		Metode B		Metode C	
X	Rank	X	Rank	X	Rank
28	12	19	8	16	6
26	11	17	7	14	5
22	10	13	4	10	2
20	9	11	3	8	1
$\Sigma = 42$		$\Sigma = 22$		$\Sigma = 14$	

Dalam Kruskal-Wallis Test berlaku ketentuan menerima hipotesis nol jika probabilitas hasil sampel ≥ 0.05 dan menolak hipotesis nol apabila probabilitas hasil sampel < 0.05 . Secara statistik dapat dirumuskan sebagai berikut.

$H_0: \mu_1 = \mu_2 = \mu_3$ (tidak ada perbedaan antar kelompok)

$H_1: \mu_1 \neq \mu_2 \neq \mu_3$ (ada perbedaan antar kelompok), dengan tingkat signifikansi 0.05 atau 5%

Data tersebut selanjutnya dianalisis melalui SPSS dengan langkah-langkah sebagai berikut.

- Entry data atau buka file data yang akan dianalisis
- Pilih menu berikut ini

Analyze

Nonparametric tests

K Independent samples – Kruskal-Wallis Test

Maka akan diperoleh hasil sebagai berikut:

Ranks

	metode	N	Mean Rank
lempar	metode A	4	10.50
lembing	metode B	4	5.50
	metode C	4	3.50
	Total	12	

Test Statistics^{a,b}

	lempar lembing
Chi-Square	8.000
df	2
Asymp. Sig.	.018

a. Kruskal Wallis Test

b. Grouping Variable:
metode

Berdasarkan hasil analisis tersebut dapat dinyatakan bahwa angka signifikansi sebesar 0,018, lebih kecil dibanding 0,05. Artinya hipotesis nol ditolak dan menerima hipotesis kerja. Dengan demikian dapat disimpulkan bahwa terdapat perbedaan yang signifikan dalam hal hasil lempar lembing, antara pembelajaran dengan menggunakan metode A, B, dan C.

Daftar Pustaka

- Albert, J.H. & Koning, R.H. (2008). *Statistical thinking in sports*. New York: Taylor & Francis Group
- Ary, D., Jacobs, L.C. & Razavieh, A. (1990). *Introduction to research in education* (fourth edition). New York: Harcourt Brace College Publishers.
- Bompa, T.O & Haff, G.G. (2009). *Periodization: Theory and methodology of training* (fifth edition). Champaign, IL: Human Kinetics
- Cohen, J. (1992). *Statistical power analysis*. New York: Sage Publication Inc.
- Cook & Campbell (1979). *Quasi-experimentation: Design & analysis issues for field settings*. Chicago: Rand McNally College Publishing Company.
- Dixon, W. At al.(1985). *Introduction to statistical analysis*. McGraw-Hill Book Company.
- Hair, J.,F., Anderson, R.E., Tatham, R.L. & Black, W.C. (1998). *Multivariate data analysis* (fifth edition). New Jersey: Prentice Hall
- Larson, R. & Farber, B. (2012). *Elementary statistics* (5th edition). Boston: Prentice Hall
- Maksum, A. (2012). *Metodologi penelitian dalam olahraga*. Surabaya: Unesa University Press
- Newell, J., Aitchison, T. & Grant, S. (2010). *Statistics for sports and exercise science: A practical approach*. London: Routledge
- Pedhazur, E. J., (1982). *Multiple regression in behavioral research* (2nd ed.). New York: CBS College Publishing.
- Pedhazur, E.J. & Schmelkin, L.P. (1991). *Measurement, design, and analysis: An integrated approach*. New Jersey: Lawrence Erlbaum Associates Publisher
- Syayib, M.A. (2013). *Applied statistics* (2nd edition). London: bookboon.com
- Thomas & Nelson (1990). *Research method in physical activity*. 2nd ed.. Champaign. IL: Human Kinetics.
- Tilley. A. (1990). *An introduction to psychological research and statistics*. Brisbane: Pineapple Press.
- Winarsunu, T. (1996). *Statistik: Teori dan aplikasinya dalam penelitian*. Malang: UMM Press

Lampiran 1

Tabel Nilai Kritik r Product Moment *

df (N-2)	Taraf Signifikansi		df (N-2)	Taraf Signifikansi		df (N-2)	Taraf Signifikansi	
	5%	1%		5%	1%		5%	1%
1	0.997	0.999	26	0.374	0.478	60	0.250	0.325
2	0.950	0.990	27	0.367	0.470	70	0.232	0.302
3	0.878	0.959	28	0.361	0.463	80	0.217	0.283
4	0.811	0.917	29	0.355	0.456	90	0.205	0.267
5	0.754	0.874	30	0.349	0.449	100	0.195	0.254
6	0.707	0.834	31	0.344	0.442			
7	0.666	0.798	32	0.339	0.436			
8	0.632	0.765	33	0.334	0.430			
9	0.602	0.735	34	0.329	0.424			
10	0.576	0.708	35	0.325	0.418			
11	0.553	0.684	36	0.320	0.413			
12	0.532	0.661	37	0.316	0.408			
13	0.514	0.641	38	0.312	0.403			
14	0.497	0.623	39	0.308	0.398			
15	0.482	0.606	40	0.304	0.393			
16	0.468	0.590	41	0.401	0.389			
17	0.456	0.575	42	0.297	0.384			
18	0.444	0.561	43	0.294	0.380			
19	0.433	0.549	44	0.291	0.376			
20	0.423	0.537	45	0.288	0.372			
21	0.413	0.526	46	0.284	0.368			
22	0.404	0.515	47	0.281	0.364			
23	0.396	0.505	48	0.279	0.361			
24	0.388	0.496	49	0.275	0.357			
25	0.381	0.487	50	0.273	0.354			

* *two-tailed test*

Lampiran 2

Tabel Nilai Kritik t

<i>df</i>	Tes Satu Ekor (<i>one-tailed test</i>)		Tes Dua Ekor (<i>two-tailed test</i>)	
	Signifikansi 5%	Signifikansi 1%	Signifikansi 5%	Signifikansi 1%
1	6.314	31.821	12.706	63.657
2	2.920	6.965	4.303	9.925
3	2.353	4.541	3.182	5.841
4	2.132	3.747	2.776	4.604
5	2.015	3.365	2.571	4.032
6	1.943	3.143	2.447	3.707
7	1.895	2.998	2.365	3.499
8	1.860	2.896	2.306	3.355
9	1.833	2.821	2.262	3.250
10	1.812	2.764	2.228	3.169
11	1.796	2.718	2.201	3.106
12	1.782	2.681	2.179	3.053
13	1.771	2.650	2.160	3.012
14	1.761	2.624	2.145	2.977
15	1.753	2.602	2.131	2.947
16	1.746	2.583	2.120	2.921
17	1.740	2.567	2.110	2.898
18	1.734	2.552	2.101	2.878
19	1.729	2.539	2.093	2.861
20	1.725	2.528	2.086	2.845
21	1.721	2.518	2.080	2.831
22	1.717	2.508	2.074	2.819
23	1.714	2.500	2.069	2.807
24	1.711	2.492	2.064	2.797
25	1.708	2.485	2.060	2.787
26	1.706	2.479	2.056	2.779
27	1.703	2.473	2.052	2.771
28	1.701	2.467	2.048	2.763
29	1.699	2.462	2.045	2.756
30	1.697	2.457	2.042	2.750
40	1.684	2.423	2.021	2.704
60	1.671	2.390	2.000	2.660
120	1.658	2.358	1.980	2.617
∞	1.282	2.326	1.960	2.576

Tabel Distribusi Chi-Square (χ^2)

d.b	Tingkat Signifikansi	
	5%	1%
1	3,84	6,63
2	5,99	9,21
3	7,81	11,3
4	9,49	13,3
5	11,1	15,1
6	12,6	16,8
7	14,1	18,5
8	15,5	20,1
9	16,9	21,7
10	18,3	23,2
11	19,7	24,7
12	21,0	26,2
13	22,4	27,7
14	23,7	29,1
15	25,0	30,6
16	26,3	32,0
17	27,6	33,4
18	28,9	34,8
19	30,1	36,2
20	31,4	37,6
21	32,7	38,9
22	33,9	40,3
23	35,2	41,6
24	35,4	43,0
25	37,7	44,3
26	38,9	45,6
27	40,1	47,0
28	41,3	48,3
29	42,6	49,6
30	43,8	50,9
40	55,8	53,7

50	67,5	76,2
60	79,1	88,4
70	90,5	100,4
80	101,9	112,3
90	113,1	124,1
100	124,3	135,8
