

STATISTIK PENDIDIKAN: MENGUKUR DAN MENGANALISIS DATA PEMBELAJARAN

**Hetty Elfina
Patrich Phill Edrich Papilaya
Muh. Khaedir Lutfi
Fitri Anisa Kusumastuti
Isry Laila Syathroh**



GET PRESS INDONESIA

STATISTIK PENDIDIKAN: MENGUKUR DAN MENGANALISIS DATA PEMBELAJARAN

Penulis :

Hetty Elfina
Patrich Phill Edrich Papilaya
Muh. Khaedir Lutfi
Fitri Anisa Kusumastuti
Isry Laila Syathroh

ISBN :

Editor : Ari Yanto, M.Pd.

Penyunting : Tri Putri Wahyuni., S.Pd

Desain Sampul dan Tata Letak : Atyka Trianisa, S.Pd

Penerbit : GET PRESS INDONESIA

Anggota IKAPI No. 033/SBA/2022

Redaksi :

Jln. Palarik Air Pacah No 26 Kel. Air Pacah
Kec. Koto Tangah Kota Padang Sumatera Barat
Website : www.getpress.co.id
Email : adm.getpress@gmail.com

Cetakan pertama, September 2024

Hak cipta dilindungi undang-undang
Dilarang memperbanyak karya tulis ini dalam bentuk dan
dengan cara apapun tanpa izin tertulis dari penerbit.

KATA PENGANTAR

Segala Puji dan syukur atas kehadiran Allah SWT dalam segala kesempatan. Sholawat beriring salam dan doa kita sampaikan kepada Nabi Muhammad SAW. Alhamdulillah atas Rahmat dan Karunia-Nya penulis telah menyelesaikan Buku Statistik Pendidikan: Mengukur Dan Menganalisis Data Pembelajaran ini.

Buku Ini Membahas Pengenalan Statistik Dalam Pendidikan, Analisis Deskriptif Dalam Pendidikan, Inferensi Statistik Untuk Evaluasi Pendidikan, Statistik Nonparametrik Dalam Pendidikan, Penggunaan Perangkat Lunak Statistik Dalam Analisis Pendidikan.

Proses penulisan buku ini berhasil diselesaikan atas kerjasama tim penulis. Demi kualitas yang lebih baik dan kepuasan para pembaca, saran dan masukan yang membangun dari pembaca sangat kami harapkan.

Penulis ucapkan terima kasih kepada semua pihak yang telah mendukung dalam penyelesaian buku ini. Terutama pihak yang telah membantu terbitnya buku ini dan telah mempercayakan mendorong, dan menginisiasi terbitnya buku ini. Semoga buku ini dapat bermanfaat bagi masyarakat Indonesia.

Padang, September 2024

Penulis

DAFTAR ISI

KATA PENGANTAR	i
DAFTAR ISI	ii
DAFTAR TABEL	v
DAFTAR GAMBAR	vi
BAB 1 PENGENALAN STATISTIK DALAM	
PENDIDIKAN	1
1.1 Pendahuluan	1
1.2 Pengertian Statistik dan Statistika	1
1.3 Data Statistik.....	3
1.4 Penggolongan Statistika	5
1.5 Fungsi Statistika	6
1.6 Komponen Dasar Analisis Statistika.....	7
DAFTAR PUSTAKA.....	13
BAB 2 ANALISIS DESKRIPTIF DALAM PENDIDIKAN...15	
2.1 Pendahuluan	15
2.2 Jenis-Jenis data dalam Pendidikan	18
2.2.1 Data Kuantitatif.....	18
2.2.2 Data Kualitatif.....	20
2.3 Teknik-teknik Analisis Deskriptif	22
2.3.1 Ukuran Tendensi Sentral.....	22
2.3.2 Ukuran Variabilitas.....	33
2.3.3 Distribusi Frekuensi.....	48
DAFTAR PUSTAKA.....	68
BAB 3 INFERENSI STATISTIK UNTUK EVALUASI	
PENDIDIKAN	77
3.1 Pendahuluan	77
3.2 Dasar-dasar Inferensi Statistik	78
3.2.1 Konsep Statistik Dasar	78
3.2.2 Distribusi Probabilitas	81
3.3 Pengantar Uji Hipotesis	82
3.3.1 Langkah-langkah Uji Hipotesis	83
3.3.2 Jenis Uji Statistik	83
3.3.3 Studi Kasus dalam Evaluasi Pendidikan.....	84
DAFTAR PUSTAKA.....	96

BAB 4 STATISTIK NONPARAMETRIK DALAM	
PENDIDIKAN.....	97
4.1 Pendahuluan	97
4.2 Konsep Dasar Statistik Nonparametrik.....	98
4.2.1 Pengertian Nonparametrik	98
4.2.2 Metode Nonparametrik	99
4.3 Aplikasi Statistik dalam Penelitian Pendidikan.....	100
4.3.1 Desain Penelitian	100
4.3.2 Studi Kasus	100
4.4 Keuntungan dan Keterbatasan Statistik	
Nonparametrik.....	101
4.4.1 Keuntungan Statistik Nonparametrik.....	101
4.4.2 Keterbatasan Statistik Nonparametrik.....	102
4.5 Teknik dan Metode Utama.....	103
4.5.1 Uji Chi-Square: Prinsip dasar dan penerapan.....	103
4.5.2 Uji Mann-Whitney dan Wilcoxon: Penjelasan	
dan aplikasi dalam data Pendidikan	104
4.5.3 Uji Kruskal-Wallis: Bagaimana menggunakan	
uji ini untuk lebih dari dua kelompok	104
4.6 Contoh Implementasi dalam Data Pendidikan	105
4.6.1 Studi Kasus 1: Analisis data kepuasan	
siswa menggunakan uji chi-square.....	105
4.6.2 Studi Kasus 2: Penggunaan uji Mann-Whitney	
untuk membandingkan hasil ujian antara	
dua kelas	106
4.7 Software dan Alat Statistik Nonparametrik.....	106
4.7.1 Perangkat Lunak Populer: SPSS, R, dan Python.	106
4.7.2 Tutorial Singkat: Cara melakukan analisis	
nonparametrik menggunakan alat tersebut	107
4.8 Diskusi dan Kesimpulan	109
4.8.1 Diskusi: Evaluasi hasil studi kasus dan	
metode yang digunakan.....	109
4.8.2 Kesimpulan: Implikasi untuk praktik	
pendidikan dan penelitian lebih lanjut.	109
DAFTAR PUSTAKA.....	110

BAB 5 PENGGUNAAN PERANGKAT LUNAK

STATISTIK DALAM ANALISIS PENDIDIKAN.....111

5.1 Pentingnya Penggunaan Perangkat Lunak Statistik
dalam Penelitian dan Analisis Pendidikan 111

5.2 Peran Teknologi dalam Pendidikan..... 112

5.3 Definisi dan Jenis Perangkat Lunak Statistik..... 113

5.4 Kriteria Pemilihan Perangkat Lunak 116

5.5 Tantangan dan Solusi dalam Penggunaan
Perangkat Lunak Statistik..... 118

5.6 Kesimpulan..... 119

DAFTAR PUSTAKA..... 121

BIODATA PENULIS

DAFTAR TABEL

Tabel 2.1. Tabel Frekuensinya Ujian Matematika	51
Tabel 2.2. Frekuensi Kehadiran Siswa (hari).....	52
Tabel 2.3. Jumlah Mata Pelajaran yang Ditawar Siswa ...	60
Tabel 3.1. Contoh perbedaan antara populasi dan sampel.....	78
Tabel 3.2. Contoh Tabel Skala Nominal.....	79
Tabel 3.3. Contoh Tabel Skala Ordinal.....	80
Tabel 3.4. Contoh Tabel Skala Interval.....	80
Tabel 3.5. Contoh Tabel Skala Rasio.....	81
Tabel 3.6. Karakteristik beberapa distribusi probabilitas.	82
Tabel 3.7. Beberapa jenis uji statistic.....	82
Tabel 3.8. Perbandingan antara uji z dan uji t.....	84
Tabel 3.9. Data penelitian.....	85
Tabel 3.10. Tests of Normality	86
Tabel 3.11. Test of Homogeneity of Variances.....	87
Tabel 3.12. Data yang sudah siap dimasukkan ke aplikasi SPSS	88
Tabel 3.13. Group Statistics	89
Tabel 3.14. Independent Samples Test.....	89
Tabel 3.15. Data penelitian.....	90
Tabel 3.16. Tests of Normality	92
Tabel 3.17. Data yang sudah siap dimasukkan ke aplikasi SPSS	94
Tabel 3.18. Paired Samples Statistics.....	95
Tabel 3.19. Paired Samples Test.....	95

DAFTAR GAMBAR

Gambar 1.1. Jenis Data Statistik.....	4
Gambar 1.2. Hubungan Variabel X dan Variabel Y	8
Gambar 1.3. Hubungan Antara Variabel X_1 dan X_2 secara Bersama-sama dengan Variabel Y .	8
Gambar 1.4. Hubungan Antara Variabel X_1 dengan Variabel Y dengan Variabel X_2 sebagai Moderator	9
Gambar 1.5. Hubungan Variabel X_1 dengan Variabel Y, dan Variabel X_2 sebagai Variabel Intervening	9
Gambar 1.6. Hubungan Variabel X_1 dengan Variabel Y, dan Variabel X_2 sebagai Variabel Kontrol..	10
Gambar 2.1. Histogram Nilai Matematika	57
Gambar 2.2. Diagram Batang 10 Mata Pelajaran.....	61
Gambar 2.3. Diagram Batang Jumlah Siswa 3 Jurusan...	63
Gambar 2.4. Diagram Lingkaran Jumlah Masiswa dalam empat Fakultas.....	66

BAB 1

Pengenalan Statistik dalam Pendidikan

1.1 Pendahuluan

Istilah statistik pada mahasiswa, guru maupun dosen merupakan istilah yang umum di dengar karena dalam melakukan penelitian akan dilakukan pengolahan data dengan menggunakan rumus statistik. Namun sering kali tidak menyadari bagaimana mengolah data dengan benar agar menghasilkan informasi yang bermanfaat. Dalam melakukan pengolahan data statistik yang benar bukan berarti menggunakan perhitungan yang sulit untuk dipahami (Gunawan, 2017)

Saat sekarang ini, hampir semua disiplin ilmu pengetahuan telah menggunakan pengolahan data statistik. Untuk penelitian di bidang ekonomi, akademik, dan dalam mengambil keputusan pada manajemen, penggunaan statistik dapat memberikan gambaran tentang apa yang akan diteliti dan dapat memprediksi serta memberikan rekomendasi dari keadaan yang mungkin akan muncul saat penelitian (Somantri, Muhidin, 2014)

1.2 Pengertian Statistik dan Statistika

Apakah statistik dan statistika itu sama? Saat ini masih banyak ditemukan yang tidak mampu membedakan statistik dan statistika. Statistik berasal dari bahasa latin “status” yang artinya “negara”. Awalnya statistik digunakan untuk memberikan keterangan untuk negara, seperti jumlah penduduk, keterangan pekerjaan, dan usia penduduk. Namun seiring berkembangnya jaman pengertian statistik merupakan suatu kumpulan angka-angka, seperti jumlah pengeluaran per bulan, keuntungan yang diperoleh dalam berdagang, banyaknya pengunjung yang datang dalam satu hari di satu restoran, dan lain sebagainya. Informasi

berbentuk angka-angka yang diperoleh tersebut dapat dibuat dalam bentuk tabel atau gambar (grafik). Maka **statistik** merupakan sekumpulan data, bilangan dan nonbilangan yang dibuat berbentuk tabel atau diagram untuk menjelaskan suatu masalah. Pengertian statistik dalam arti luas yaitu suatu cara keilmuan dalam mengakumulasikan, menyusun, menelaah, menampilkan keterangan berupa angka, dan menghasilkan kesimpulan.

Seorang peneliti saat melakukan penelitian kuantitatif akan bekerja dengan data berbentuk angka. Data yang diperoleh akan dikumpulkan, diolah, dan dianalisis untuk membuat kesimpulan dari penelitian yang dilakukan. Ilmu yang mempelajari bagaimana proses perhitungan akan dilakukan dinamakan dengan statistika. **Statistika** yaitu ilmu yang mengamati cara mengakumulasikan, mengolah, menyajikan, menelaah data, dan menghasilkan kesimpulan. Statistika digunakan untuk menyusun model, merumuskan jawaban sementara, mengembangkan alat pengambilan data, membuat rangka penelitian, menentukan sampel, dan menganalisis petunjuk lalu diinterpretasikan sehingga memiliki makna.

Gunawan (2013) menyatakan bahwa fungsi statistika yaitu (1) mendukung untuk pengambilan suatu ketentuan yang tepat; (2) sebagai alat dalam pengendalian kualitas; (3) mengharuskan supaya mengetahui bagaimana peluang suatu kejadian akan datang. Statistika banyak digunakan pada pendidikan, bidang keuangan, manajemen, akuntansi, dan lain sebagainya. Terdapat tiga ciri pokok statistika, yakni:

1. Beroperasi pada angka-angka. Angka-angka ini mempunyai dua makna, yakni untuk menunjukkan jumlah (frekuensi) dan menunjukkan nilai (harga). Agar dapat melakukan tugasnya maka pada statistika diperlukan bahan keterangan yang bersifat kuantitatif. Saat statistik digunakan jadi alat menelaah data kualitatif (data yang tidak menggunakan angka) maka data kualitatif yang ada wajib diganti menjadi data kuantitatif. Contoh, terdapat keterangan kualitatif mengenai kepuasan pelatihan alat peraga, yaitu “puas”, “cukup”, “kurang”. Agar data dapat dianalisis, maka data kualitatif dikonversikan ke data kuantitatif, seperti siswa

yang merasa sangat puas dengan adanya pelatihan memiliki skor 81 – 100, puas memiliki skor 61 – 80, dan kurang puas memiliki skor 30 – 60. Keterangan tersebut dapat juga diklasifikasikan sebagai siswa yang merasa puas ada 12 orang, cukup puas = 10 orang, dan tidak puas = 5 orang.

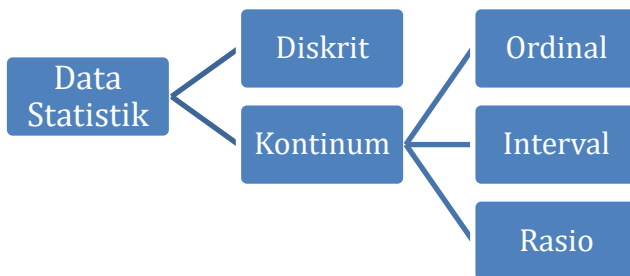
2. Bersifat objektif. Statistik sebagai penyaji fakta akan menghasilkan data sebagaimana adanya. Ini berarti statistika bekerja menurut objeknya. Kesimpulan yang dibuat dari hasil penelitian dan ramalan yang diberikan hanya berdasarkan keterangan berbentuk angka yang diperoleh dan diolah, bukan didasari pada subjektivitas atau pengaruh luar lainnya.
3. Bersifat universal (umum). Cakupan dan bidang pengolahan statistik berlaku untuk penelitian murni atau terapan. (Gunawan, 2017)

1.3 Data Statistik

Data merupakan sekumpulan angka, bukti, peristiwa, atau keadaan lain yang berdasarkan hasil observasi, penaksiran maupun pencacahan pada objek dan memiliki fungsi untuk membedakan antar objek pada variabel yang sama. Data memiliki beragam bentuk seperti jumlah siswa, nilai siswa, dan angka penjualan. Informasi berarti sekumpulan data yang lebih memiliki makna setelah melewati proses pengolahan. Pengetahuan yaitu menentukan dan mengelompokkan suatu keterangan dengan memberikan pengertian, saran, dan dasar ketetapan. Data juga diartikan sebagai petunjuk berupa kategori seperti baik, senang, puas, gagal, setuju, jarang, dan berhasil. Inilah yang disebut sebagai data statistik (Sudjana, 2002). Data yang memiliki bilangan disebut data diskrit dan data kontinu.

Data diskrit (nominal) merupakan suatu bentuk yang dapat dikategorikan berdasarkan jenisnya. Angka yang terdapat di variabel diskrit merupakan angka kuantitatif yang diperoleh berdasarkan perhitungan. Suatu angka yang menggantikan kuantitas disebut frekuensi (f) atau jumlah (N). Contoh data diskrit adalah jumlah pengangguran di usia 30 tahun ke atas; jumlah mahasiswa yang menerima bantuan berupa beasiswa di jurusan Matematika; dan jumlah wanita 20 orang.

Data kontinum merupakan data yang dikelompokkan berdasarkan tingkatan besar kecilnya yang didapat dari pengukuran. Data kontinum dikelompokkan atas data ordinal, interval, dan rasio. Angka yang digunakan dalam variabel kontinum merupakan angka kualitatif yang didapat dari hasil pengukuran. Pada statistik angka ini disebut skor, nilai, atau harga. Misalnya kecerdasan, nilai IQ 110.



Gambar 1.1. Jenis Data Statistik
(Sumber : Gunawan, 2017)

Data ordinal merupakan suatu keterangan berbentuk peringkat dan jarak antar datanya tidak harus sama. Contohnya adalah juara I dengan nilai 200, juara II dengan nilai 175, dan juara III dengan nilai 100. Data ordinal diperoleh dari data interval atau rasio.

Data interval yaitu suatu data yang memiliki jarak sama namun tidak mempunyai nilai nol absolut. Contohnya suhu 0 derajat Celcius, namun bukan berarti tidak memiliki nilai (tidak panas dan tidak dingin). Data rasio merupakan suatu data (keterangan) yang memiliki jarak sama dan mempunyai nilai nol absolut. Contohnya panjang meja 0 meter, artinya tidak ada panjang. Hasil pengukuran pada panjang dan berat merupakan contoh data rasio.

Data secara statistik terbagi atas dua, yakni data metrik dan data nonmetrik. Data metrik yaitu data yang diperoleh dari hasil menguji dan dapat mempunyai desimal. Data nonmetrik

yaitu data yang diperoleh melalui perhitungan, tidak memiliki desimal, dan memiliki kategorisasi. (Gunawan, 2017)

1.4 Penggolongan Statistika

Berdasarkan pengolahan data, dapat dibagi atas dua bagian. Pertama, aktifitas mengumpulkan data kemudian melakukan pengolahan supaya mudah dibaca oleh orang banyak. Kedua, kegiatan yang dilakukan untuk dapat ditarik kesimpulannya. Data ini diperoleh dari sampel suatu populasi lalu dianalisis dan diperoleh suatu kesimpulan yang diterapkan pada populasi. Kegiatan yang dilakukan pertama tadi dinamakan statistika deskriptif, dan kegiatan yang kedua dinamakan statistika inferensial.

Statistika deskriptif memiliki tujuan untuk mengganti sekumpulan data mentah agar dapat dipahami sebagai suatu penjelasan yang ringkas. Statistika deskriptif memiliki tujuan untuk menjelaskan objek yang diamati bagaimana adanya tanpa membuat kesimpulan. Pada statistika deskriptif diberikan penyampaian data berbentuk tabel dan diagram, menentukan rata-rata, median, modus, rentang, dan simpangan baku. Statistika deskriptif bersifat menguraikan gejala kuantitatif secara numeris dan menafsirkan informasi lebih lanjut.

Statistika inferensial memiliki tujuan dalam penyediaan dasar prakiraan yang digunakan dalam mengganti informasi menjadi pengetahuan. Statistika inferensial memiliki tujuan untuk penarikan kesimpulan pada sejumlah orang, kejadian, dan waktu secara menyeluruh. Sifat dari statistika inferensial yaitu (1) menganalisis data yang berasal dari *random sampling* (acak); (2) melakukan generalisasi serta peramalan mengenai ciri penting suatu variabel hingga kaitan antarvariabel; (3) generalisasi dan ramalan yang dibentuk berlaku untuk semua populasi didasarkan dari hasil menelaah data sampel; (4) generalisasi dan ramalan dilaksanakan melalui uji hipotesis.

Statistika inferensial terdiri atas parametrik dan nonparametrik. Statistika parametrik digunakan pada penarikan kesimpulan dari beberapa gejala dan dapat disimpulkan ke populasi. Statistika parametrik memakai dugaan tentang

populasi dan memerlukan penaksiran kuantitatif di level data interval atau rasio. Karakteristik statistika parametrik yaitu (1) menelaah berdasarkan pada dugaan bahwa data mempunyai sebaran khusus dengan kriteria yang belum ditemukan; (2) memiliki jenis data kuantitatif serta random; (3) memiliki sampel dengan jumlah terbatas sesuai kriteria keterwakilannya; (4) data berdistribusi normal; (5) data homogen, variasinya rendah. Contoh analisis statistika parametrik yaitu Uji t, Analisis Ragam (ANOVA), Uji Korelasi Pearson, dan Uji Regresi (Uji F).

Statistika nonparametrik berdasarkan parameter beberapa fenomena dan hanya terdapat kesimpulan saja di beberapa bagian dari keseluruhan. Pada statistika nonparametrik digunakan sedikit dugaan tentang populasi dan memerlukan data dengan level serendah-rendahnya ordinal. Karakteristik statistika nonparametrik adalah (1) memiliki jenis data kualitatif; (2) bukan didasari pada dugaan distribusi data; (3) analisis statistika bebas distribusi; (4) kondisi ini berlaku untuk data yang berukuran kecil dengan skala pengukuran yang bukan dari skala interval; (5) ukuran pemusatan data yang jadi fokusnya adalah median. Contoh analisis statistika nonparametrik adalah Khi Kuadrat untuk Uji Kebebasan Dua Variabel Kategori, Koefisien Korelasi Spearman, Uji Tanda Peringkat Wilcoxon, Uji Mann-Whitney, Uji Kruskal-Wallis, dan Uji Friedman.

(Somantri, Muhidin, 2014)

1.5 Fungsi Statistika

Dewasa ini hampir setiap disiplin ilmu pengetahuan melakukan perhitungan statistika dalam membuat penelitian. Penggunaan teknik analisis statistika dapat memberi dukungan yang bermakna agar mencapai tujuan kegiatan. Pada suatu penelitian, metode statistika memberikan cerminan masalah yang diteliti serta dapat memberikan prediksi dan rekomendasi pada keadaan yang dapat timbul sesuai dengan masalah yang diteliti.

Fungsi statistika adalah sebagai alat bantu. Selain itu, statistika memiliki beberapa kegunaan sebagai berikut :

1. Statistika mendukung dalam menetapkan sampel, sehingga penelitian dilakukan secara efektif dan menghasilkan data yang akurat dan valid.
2. Statistika mampu meningkatkan efisiensi, melalui pembatasan dan menentukan aturan kerja dan cara berpikir.
3. Statistika mampu menyimpulkan hasil penelitian menjadi sederhana sehingga mudah dipahami.
4. Statistika dapat melihat adatidaknya hubungan antar variabel.
5. Statistika mampu memberikan dasar untuk membuat perumusan serta mengambil kesimpulan penelitian yang tepat.
6. Statistika mampu memberikan deskripsi eksak tentang suatu dugaan yang dapat digunakan pada waktu akan datang.
7. Statistika mampu memberikan dasar pada penyusunan dugaan tentang bagaimana sesuatu akan terjadi sesuai pada keadaan yang telah diukur dan diuji.
8. Statistika mampu menelaah penyebab kausal dan perbedaan pada beberapa faktor yang bertautan dan rumit. (Mundir, 2012)

1.6 Komponen Dasar Analisis Statistika

1. Variabel penelitian

Asal kata variabel yaitu dari Bahasa Inggris “variable” yang memiliki arti ubahan, faktor tidak tetap, gejala yang dapat berubah-ubah, atau tidak satu jenis. Secara terminologi, variabel adalah suatu konsep yang memiliki keragaman atau variasi yang dapat diberi nilai. Konsep mengenai apapun yang memiliki tanda yang beragam disebut variabel. Variabel terdiri dari lima macam, yaitu variabel bebas, variabel terikat, variabel moderator, variabel intervening, dan variabel kontrol.

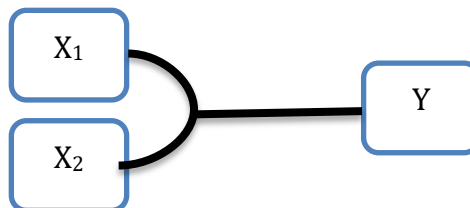
Variabel bebas (*independent*, prediktor, X) yaitu suatu variabel yang bisa memberikan dampak keragaman variabel lain yang mendampinginya. Variabel bebas merupakan variabel yang membuat alasan munculnya variabel lain.

Variabel terikat (*dependent*, Y) yaitu variabel yang terpengaruh dari adanya variabel bebas. Contoh untuk pembelajaran yaitu penggunaan video pembelajaran sebagai media ajar akan berpengaruh dengan kualitas pembelajaran, motivasi untuk berprestasi dengan hasil belajar.

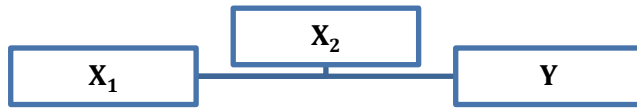


Gambar 1.2. Hubungan Variabel X dan Variabel Y
(Sumber : Mundir, 2012)

Variabel moderator yaitu variabel yang dapat memberikan pengaruh (memperkuat atau memperlemah) ikatan antara variabel bebas (X_1) dan variabel terikat (Y). Variabel moderator dinamakan variabel bebas kedua (X_2). Contohnya hubungan antara tingkat pengetahuan materi berenang dan tingkat pengalaman berenang akan semakin kuat jika didukung dengan sarana yang bagus. Saat memberikan pengaruh pada variabel terikat (Y), maka variabel bebas pertama (X_1) akan bersama dengan variabel bebas kedua (X_2), kemudian memberikan hasil dari dampak variabel bebas terhadap variabel terikat akan lebih besar. Begitupula untuk variabel bebas kedua (X_2) dapat dibuat sebagai pengontrol atas kemurnian akibat variabel bebas pertama (X_1) terhadap variabel terikat (Y).

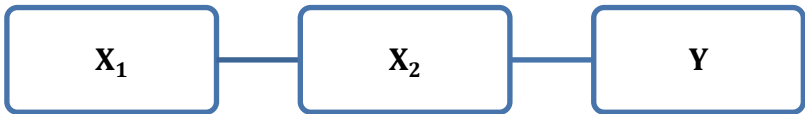


Gambar 1.3. Hubungan Antara Variabel X_1 dan X_2 secara Bersama-sama dengan Variabel Y
(Sumber : Mundir, 2012)



Gambar 1.4. Hubungan Antara Variabel X_1 dengan Variabel Y dengan Variabel X_2 sebagai Moderator
(Sumber : Mundir, 2012)

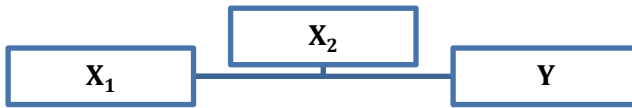
Variabel intervening yaitu variabel yang mempengaruhi (memperkuat atau memperlemah) ikatan antara variabel bebas dan variabel terikat, namun tidak dapat diukur. Contoh kepintaran seseorang akan memberikan pengaruh pada perkembangan prestasi belajar. Namun kepintaran tersebut tidak memberikan pengaruh saat seseorang itu sakit. Kondisi sakit ini disebut variabel intervening karena tingkat sakit ini sulit diukur.



Gambar 1.5. Hubungan Variabel X_1 dengan Variabel Y , dan Variabel X_2 sebagai Variabel Intervening
(Sumber : Mundir, 2012)

Variabel kontrol adalah variabel yang dikendalikan sehingga tidak memberikan pengaruh pada variabel bebas dan variabel terikat. Variabel kontrol umumnya digunakan pada penelitian eksperimen. Contohnya seorang peneliti ingin melihat perbedaan nilai dengan pembelajaran 1 dan pembelajaran 2, maka harus menetapkan variabel kontrol pada kondisi awal sebelum menjalani proses pembelajaran dan tingkat kesukaran materi. Variabel kontrol diperlukan ketika melakukan penelitian komparatif (membandingkan antar variabel) melalui penelitian eksperimen. Contohnya peneliti akan melihat pengaruh pembelajaran di suatu sekolah antara mahasiswa I dan mahasiswa II. Sebagai

variabel kontrolnya, ditetapkan kelas untuk mengajar, guru pamong, serta media yang akan digunakan.



Gambar 1.6. Hubungan Variabel X_1 dengan Variabel Y , dan Variabel X_2 sebagai Variabel Kontrol
(Sumber : Mundir, 2012)

2. Populasi dan sampel penelitian

Populasi merupakan keseluruhan objek (orang, wilayah, benda) untuk dibuat generalisasi kesimpulan dari hasil suatu penelitian. Generalisasi yaitu perlakuan dari hasil kesimpulan suatu penelitian pada semua objek menurut data yang didapat dari beberapa objek penelitian yang menjadi perwakilan. Wakil inilah yang dinamakan sebagai sampel. Sampel yang baik merupakan sampel yang mempunyai ciri-ciri, sifat-sifat, atau karakteristik yang dimilikinya kemudian dinamakan sebagai sampel yang representatif. Jika populasi terdapat di beberapa bagian, kelompok, atau wilayah, maka sampel juga wajib bersumber dari beraneka ragam populasi itu. Saat sampel tidak representatif, maka peneliti tidak dapat membuat generalisasi sebab hasilnya akan menyimpang dari kenyataan sebenarnya. Untuk kasus sampel tidak representatif ketetapan yang dihasilkan akan valid untuk sampel itu sendiri.

3. Hipotesis penelitian

Pada penelitian kuantitatif, hipotesis mempunyai kedudukan yang berarti sebab melalui hipotesis akan diberikan petunjuk yang nyata kepada peneliti agar diperoleh suatu kesimpulan. Hipotesis adalah jawaban sementara yang disusun dan dirumuskan sesuai dengan analisis teori yang terpaut, hasil ciptaan penelitian sebelumnya, atau hasil pengamatan lapangan. Sebagai jawaban sementara, maka hipotesis wajib diteliti keasliannya sesuai data yang

dikumpulkan. Usaha dalam melakukan uji kebenaran hipotesis dinamakan verifikasi.

Karena bersifat jawaban sementara, maka hipotesis tersebut didefinisikan sebanding dengan permasalahan penelitiannya. Permasalahan penelitian dapat berupa problematika deskriptif, problematika asosiatif (korelatif), atau problematika komparatif. Problematika deskriptif merupakan suatu masalah riset yang menggambarkan terdapatnya satu atau lebih variabel penelitian; problematika asosiatif (korelatif) merupakan suatu masalah riset yang menerangkan terdapatnya asosiasi (korelasi, hubungan) dari dua atau lebih variabel penelitian; Problematika komparatif merupakan suatu masalah riset yang menjelaskan terdapat perbedaan dari dua atau lebih variabel penelitian. Hipotesis kerja (hipotesis alternatif, H_a) disusun berbentuk pernyataan positif (mengiyakan), dan hipotesis nihil (hipotesis nol, H_o) disusun berbentuk pernyataan negatif (menyangkal).

a. Contoh hipotesis Deskriptif

Hipotesis Kerja: Siswa SMA memiliki waktu istirahat di rumah selama 8 jam per hari. Atau waktu istirahat siswa = 8 jam/hari.

Hipotesis Nihil: Siswa SMA memiliki waktu istirahat bukan 8 jam per hari. Atau waktu istirahat siswa $\neq$ 8 jam/hari

b. Contoh Hipotesis Asosiatif (Korelatif)

Hipotesis Kerja: Model *blanded learning* memiliki kaitan secara signifikan dengan kemampuan berhitung siswa SMA.

Hipotesis Nihil: Model *blanded learning* tidak memiliki kaitan secara signifikan dengan kemampuan berhitung siswa SMA.

c. Contoh Hipotesis Komparatif

Hipotesis Kerja: Terdapat perbedaan kemampuan penalaran matematik siswa kelas X MAN I Medan yang

berasal dari Sekolah Menengah Pertama (SMP) dan yang berasal dari Madrasah Tsanawiyah (MTs).

Hipotesis Nihil: Tidak terdapat perbedaan kemampuan penalaran matematik siswa kelas X MAN I Medan yang berasal dari Sekolah Menengah Pertama (SMP) dan yang berasal dari Madrasah Tsanawiyah (MTs) (Mundir, 2012).

DAFTAR PUSTAKA

- Gunawan, Imam. 2017. *Pengantar Statistika Inferensial*. Jakarta: PT Rajagrafindo Persada.
- Mundir. 2012. *Statistik Pendidikan: Pengantar Analisis Data untuk Penulisan Skripsi dan Tesis*. Jember : STAIN Jember Press
- Somantri, Muhidin. 2014. *Aplikasi Statistika dalam Penelitian*. Bandung: CV Pustaka Setia.
- Sudjana. 2002. *Metode Statistik*. Bandung : Tarsito

BAB 2

ANALISIS DESKRIPTIF DALAM PENDIDIKAN

2.1 Pendahuluan

Analisis deskriptif adalah metode analisis yang digunakan untuk menggambarkan dan meringkas data secara rinci. Analisis ini tidak berusaha untuk membuat inferensi atau prediksi, tetapi fokus pada penyajian karakteristik utama dari dataset yang ada. Analisis deskriptif adalah langkah awal dalam analisis data yang bertujuan untuk menggambarkan karakteristik dasar data. Ini mencakup berbagai metode statistik yang digunakan untuk meringkas, menggambarkan, dan memahami data yang dikumpulkan. Dalam konteks pendidikan, analisis deskriptif digunakan untuk memberikan gambaran umum mengenai kinerja siswa, perilaku belajar, dan hasil pendidikan lainnya.

Definisi Statistik Deskriptif Menurut Para Ahli dapat disajikan sebagai berikut :

1. William Cullen Bryant (2019)
"Statistik deskriptif adalah istilah yang digunakan untuk analisis data yang membantu menggambarkan, menunjukkan, atau meringkas data dengan cara yang bermakna sehingga pola dapat muncul dari data tersebut. Statistik deskriptif tidak memungkinkan kita untuk membuat kesimpulan di luar data yang telah kita analisis atau mencapai kesimpulan mengenai hipotesis apapun yang mungkin telah kita buat. Metode ini hanya merupakan cara untuk menggambarkan data kita." (Bryant, 2019).
2. Parampreet Kaur, J. Stoltzfus, dan Vikas Yellapu (2018)
"Statistik deskriptif digunakan untuk meringkas data secara terorganisir dengan menggambarkan hubungan antara variabel dalam sampel atau populasi. Statistik deskriptif mencakup jenis variabel (nominal, ordinal, interval, dan

rasio) serta ukuran frekuensi, tendensi sentral, dispersi/variasi, dan posisi." (Kaur et al., 2018).

3. R. Cooksey (2020)

"Statistik deskriptif adalah prosedur yang digunakan untuk memfasilitasi deskripsi dan ringkasan data. Ini mencakup penggunaan representasi grafis atau komputasi indeks atau angka yang dirancang untuk meringkas karakteristik spesifik dari variabel atau pengukuran." (Cooksey, 2020).

4. Gary Smith (2015)

"Statistik deskriptif adalah angka-angka yang digunakan untuk meringkas dan menggambarkan data. Data mengacu pada informasi yang telah dikumpulkan dari eksperimen, survei, catatan historis, dll. Misalnya, statistik deskriptif dapat berupa persentase sertifikat kelahiran yang diterbitkan di Negara Bagian New York, atau usia rata-rata ibu." (Smith, 2015).

5. J. Kadane (1978)

"Statistik deskriptif digunakan untuk menggambarkan fitur dasar data dalam sebuah studi. Hal ini memberikan ringkasan sederhana tentang sampel dan ukuran. Dengan statistik deskriptif, kita hanya menggambarkan apa yang ditunjukkan data." (Kadane, 1978).

6. M. Rendón-Macías, M. Villasís-Keever, dan M. G. Miranda-Novales (2019)

"Statistik deskriptif adalah cabang statistik yang memberikan rekomendasi tentang cara meringkas data penelitian dengan jelas dan sederhana dalam tabel, gambar, diagram, atau grafik. Sebelum melakukan analisis deskriptif, sangat penting untuk merangkum tujuannya dan mengidentifikasi skala pengukuran dari variabel yang dicatat dalam studi." (Rendón-Macías et al., 2019).

Analisis deskriptif merupakan alat penting dalam penelitian pendidikan untuk memberikan gambaran yang jelas mengenai fenomena, tren, dan masalah yang ada. Berikut ini adalah beberapa poin penting mengenai pentingnya analisis deskriptif dalam pendidikan.

- a. **Memahami Kemampuan dan Kesulitan Siswa**
Analisis deskriptif membantu dalam mengidentifikasi kemampuan dan kesulitan siswa dalam berbagai aspek pembelajaran. Misalnya, penelitian menunjukkan bahwa analisis deskriptif digunakan untuk menilai kemampuan siswa dalam menulis teks deskriptif dan mengidentifikasi kesulitan mereka dalam mengorganisasi ide dan memilih kosakata yang tepat (Yoandita, 2019).
- b. **Meningkatkan Pengembangan Kurikulum**
Analisis deskriptif juga digunakan dalam penelitian pengembangan kurikulum untuk memahami metode, alat pengumpulan data, dan tingkat pendidikan yang paling banyak diteliti. Hal ini membantu dalam merancang kurikulum yang lebih efektif dan relevan (Akdemir et al., 2015).
- c. **Evaluasi Kegiatan Pendidikan**
Dalam evaluasi pendidikan, analisis deskriptif digunakan untuk menilai perubahan sikap dan perilaku siswa, serta untuk meningkatkan minat dan motivasi mereka. Guru dapat menggunakan evaluasi deskriptif untuk memahami perilaku awal siswa dan merancang strategi pengajaran yang sesuai (Asakereh & Bahrani, 2012).
- d. **Pemetaan dan Analisis Data Pendidikan:**
Analisis statistik deskriptif memainkan peran penting dalam proses penggalan data pendidikan. Dengan menganalisis data secara deskriptif, peneliti dapat mengidentifikasi pola dan variabilitas dalam data yang dikumpulkan dari berbagai sumber pembelajaran (Dimić et al., 2019).
7. **Pengembangan dan Pemantauan Kualitas Pendidikan**
Analisis deskriptif digunakan untuk menggambarkan dan memantau atribut kualitas pendidikan, membantu dalam pengembangan produk dan proses pendidikan yang lebih baik, serta untuk jaminan dan kontrol kualitas (Gillette, 1984).

Analisis deskriptif digunakan dalam berbagai sudut pandang; antara lain :

- a. Mengidentifikasi dan menggambarkan atribut sensorik suatu produk, seperti penampilan, aroma, rasa, dan tekstur. Metode ini berguna dalam pengembangan produk, studi umur simpan, peningkatan produk, serta jaminan dan kontrol kualitas di industri makanan dan rasa (Gillette, 1984).
- b. Menganalisis produk dengan tingkat keandalan dan presisi yang tinggi, memberikan deskripsi sensoris lengkap dari produk yang diuji. Metode ini mencakup profil rasa, profil tekstur, analisis deskriptif kuantitatif (QDA), dan analisis deskriptif spektrum (Stone & Sidel, 2004).
- c. Menggambarkan data secara numerik dan grafis, seperti penyajian data dalam tabel, grafik, dan penghitungan statistik untuk tendensi sentral, variabilitas, dan distribusi. Metode grafis termasuk grafik garis, grafik batang, histogram, dan poligon frekuensi (Sonnad, 2002).

2.2 Jenis-Jenis data dalam Pendidikan

Jenis-jenis data yang digunakan dalam pendekatan statistik pendidikan sangat luas tergantung dari sudut pandang, keutamaan datanya, waktu dan sebagainya.

2.2.1 Data Kuantitatif

1. Roni, Merga, dan Morris (2019)
"Data kuantitatif mencakup data yang diukur dan dinyatakan dalam bentuk angka. Jenis data ini dapat dikategorikan menjadi rasio, interval, ordinal, dan nominal. Pemilihan jenis data yang tepat sangat penting karena menentukan uji statistik yang sesuai untuk tujuan penelitian tertentu." (Roni et al., 2019).
2. Abulela dan Harwell (2020)
"Data kuantitatif digunakan dalam penelitian pendidikan untuk melakukan analisis statistik yang memperkuat inferensi. Data kuantitatif mencakup variabel seperti jumlah

siswa, nilai ujian, dan skor keterampilan. Penting untuk memastikan kualitas data yang dikumpulkan agar hasil analisis valid." (Abulela & Harwell, 2020).

3. Gerber, Boulton-Lewis, dan Bruce (1995)
"Anak-anak semakin sering berhadapan dengan representasi grafis dari data kuantitatif dalam pendidikan formal mereka, seperti grafik dan tabel yang digunakan untuk mengukur berbagai bentuk data kuantitatif." (Gerber et al., 1995).

Beberapa contoh jenis data kuantitatif dalam konteks pendidikan yang ditemukan dalam proses pembelajaran:

1. **Nilai Ujian**

Data ini mengukur prestasi akademik siswa berdasarkan hasil ujian atau tes yang mereka ikuti. Misalnya, nilai matematika dari 0 hingga 100.

2. **Jumlah Siswa**

Data ini mencakup jumlah siswa yang terdaftar di sekolah atau kelas tertentu. Misalnya, jumlah siswa di kelas 10 adalah 30 orang.

3. **Rasio Guru-Siswa**

Data ini menunjukkan perbandingan antara jumlah guru dan jumlah siswa di suatu sekolah. Misalnya, rasio guru-siswa adalah 1:20.

4. **Tingkat Kehadiran:**

Data ini mengukur jumlah hari siswa hadir di sekolah dalam periode tertentu. Misalnya, tingkat kehadiran siswa dalam satu semester adalah 95%.

5. **Waktu Belajar:**

Data ini mengukur jumlah waktu yang dihabiskan siswa untuk belajar dalam sehari atau seminggu. Misalnya, siswa menghabiskan rata-rata 2 jam per hari untuk belajar di rumah.

6. **Skor Keterampilan:**

Data ini mengukur keterampilan tertentu yang dimiliki siswa, seperti keterampilan membaca atau kemampuan berhitung. Misalnya, skor keterampilan membaca berdasarkan tes membaca yang diberikan.

2.2.2 Data Kualitatif

Data kualitatif adalah jenis data yang bersifat deskriptif dan biasanya tidak dapat diukur secara numerik. Data ini mencakup informasi yang menggambarkan kualitas, sifat, atau karakteristik suatu fenomena. Data kualitatif sering diungkapkan dalam bentuk kata-kata, gambar, atau teks, dan digunakan untuk memahami perspektif, motivasi, dan pengalaman individu atau kelompok.

Seers, (2011), menyatakan bahwa penelitian kualitatif sering digunakan untuk memahami pengalaman dan perasaan individu terhadap suatu peristiwa atau interaksi. Selanjutnya, Libarkin & Kurdziel (2002) menyatakan bahwa data kualitatif diambil dari observasi kelas, wawancara, atau dokumen tertulis untuk memahami konteks fenomena pendidikan. Sementara, LeCompte, (2000) menyatakan bahwa analisis data kualitatif adalah proses mengubah data menjadi hasil penelitian yang dapat menjelaskan atau memprediksi sesuatu tentang apa yang dipelajari.

Data kualitatif dalam pendidikan adalah data yang bersifat deskriptif dan digunakan untuk memahami fenomena, pengalaman, dan perspektif individu atau kelompok. Berikut adalah beberapa contoh dan penjelasan tentang data kualitatif dalam konteks pendidikan:

1. Sikap Siswa

Data ini mengukur persepsi, pandangan, dan perasaan siswa terhadap berbagai aspek pendidikan, seperti metode pengajaran, materi pelajaran, atau lingkungan belajar. Misalnya, wawancara atau catatan observasi tentang bagaimana siswa merespons penggunaan teknologi di kelas. Seers, 2011 menyatakan bahwa penelitian kualitatif sering digunakan untuk memahami pengalaman dan perasaan individu terhadap suatu peristiwa atau interaksi.

2. Metode Pengajaran

- a. Data ini mengumpulkan informasi tentang berbagai metode pengajaran yang digunakan oleh guru dan bagaimana metode tersebut diterima oleh siswa. Contohnya termasuk catatan lapangan tentang interaksi di kelas, wawancara dengan guru tentang strategi

pengajaran, dan diskusi kelompok fokus dengan siswa tentang pengalaman belajar mereka.

- b. Menurut Libarkin & Kurdziel, 2002, data kualitatif diambil dari observasi kelas, wawancara, atau dokumen tertulis untuk memahami konteks fenomena pendidikan.

3. Pengalaman Belajar

Mengumpulkan data mengenai pengalaman belajar siswa, seperti tantangan yang mereka hadapi, strategi belajar yang mereka gunakan, dan dukungan yang mereka butuhkan. Hal ini dapat dilakukan melalui wawancara mendalam, jurnal reflektif siswa, atau diskusi kelompok. LeCompte, 2000 menyebutkan bahwa analisis data kualitatif adalah proses mengubah data menjadi hasil penelitian yang dapat menjelaskan atau memprediksi sesuatu tentang apa yang dipelajari.

4. Interaksi di Kelas

Data ini mencakup observasi tentang bagaimana siswa berinteraksi satu sama lain dan dengan guru selama proses pembelajaran. Ini bisa berupa catatan lapangan yang mendetail tentang dinamika kelas atau rekaman video interaksi di kelas yang dianalisis. Stojanov, 2014 menjelaskan bahwa data kualitatif sering dikumpulkan sebagai teks tidak terstruktur menggunakan berbagai teknik, yang kemudian dianalisis untuk memahami fenomena yang diteliti.

5. Pendapat Guru dan Orang Tua

Mengumpulkan data tentang pandangan dan pendapat guru serta orang tua mengenai proses pendidikan, kurikulum, dan perkembangan siswa. Ini dapat dilakukan melalui wawancara, survei terbuka, atau diskusi kelompok fokus. Brantlinger et al., 2005 menekankan pentingnya penelitian kualitatif dalam memberikan wawasan mendalam tentang berbagai aspek pendidikan yang tidak dapat diukur secara kuantitatif.

2.3 Teknik-teknik Analisis Deskriptif

2.3.1 Ukuran Tendensi Sentral

Ukuran tendensi sentral adalah nilai-nilai yang mewakili pusat atau titik tengah dari sebuah distribusi data. Dalam konteks pendidikan, ukuran ini sangat berguna untuk memahami kecenderungan umum dari prestasi siswa, efektivitas metode pengajaran, atau berbagai aspek lain dalam proses pembelajaran. Ada tiga ukuran tendensi sentral utama:

1. Nilai Mean

Nilai mean, atau rata-rata, adalah ukuran pemusatan data yang paling umum digunakan dalam statistik. Mean dihitung dengan menjumlahkan semua nilai dalam suatu kumpulan data dan kemudian membagi jumlah tersebut dengan banyaknya nilai dalam kumpulan data tersebut.

Berikut adalah beberapa definisi rata-rata menurut ahli:

- a. Rata-rata adalah nilai yang diperoleh dengan menjumlahkan semua nilai dalam dataset dan membagi jumlah tersebut dengan jumlah nilai dalam dataset tersebut (Triola, 2018).
- b. Freedman, Pisani, Purves (2007): Rata-rata adalah jumlah dari semua nilai dalam suatu set data dibagi dengan jumlah nilai dalam set tersebut (Freedman, Pisani, & Purves, 2007).
- c. Moore, McCabe, Craig (2014): Rata-rata adalah total dari semua observasi dibagi dengan jumlah observasi (Moore, McCabe, & Craig, 2014).
- d. Walpole, Myers, Myers, Ye (2012): Rata-rata adalah nilai yang merepresentasikan sekelompok data yang diperoleh dengan menjumlahkan semua data dan membagi dengan jumlah data tersebut (Walpole, Myers, Myers, & Ye, 2012).

Rumus untuk menghitung mean adalah:

$$\text{Mean}(\bar{X}) = \frac{\sum_{i=1}^n X_i}{n}$$

- $\bar{X}$ adalah mean
- $\sum$ adalah simbol penjumlahan
- X_i adalah nilai ke-i dalam kumpulan data
- n adalah jumlah nilai dalam kumpulan data

Contoh Perhitungan Mean

Misalkan kita memiliki data nilai ujian matematika dari lima siswa: 75, 80, 68, 90, dan 85.

$$\text{Mean} = (75 + 80 + 68 + 90 + 85) / 5 = 79,6$$

Jadi, nilai mean dari data nilai ujian matematika adalah 79.6.

a. Keunggulan dan Kelemahan Mean

1) Keunggulan

- Sederhana dan Mudah Dihitung:** Mean mudah dipahami dan dihitung, menjadikannya ukuran pemusatan yang populer.
- Memanfaatkan Semua Data:** Mean menggunakan semua nilai dalam kumpulan data, memberikan gambaran keseluruhan yang lengkap.

2) Kelemahan

- Sensitif terhadap Nilai Ekstrem:** Mean dapat sangat terpengaruh oleh nilai-nilai ekstrem (outlier). Misalnya, jika ada satu nilai yang sangat tinggi atau sangat rendah, itu bisa menggeser mean sehingga tidak merepresentasikan data dengan baik.
- Tidak Selalu Mewakili Data:** Dalam distribusi yang tidak simetris atau memiliki skewness, mean mungkin tidak mencerminkan pemusatan data yang sebenarnya.

b. Penerapan Mean dalam Pendidikan

Dalam konteks pendidikan, mean sering digunakan untuk mengevaluasi performa akademik siswa, misalnya:

- 1) **Rata-rata Nilai Ujian:** Menentukan rata-rata nilai ujian untuk mengevaluasi tingkat pemahaman siswa terhadap materi yang diajarkan.
- 2) **Rata-rata Kehadiran:** Menghitung rata-rata kehadiran siswa untuk menilai keterlibatan mereka dalam proses pembelajaran.

2. Median

Median adalah salah satu ukuran pemusatan dalam statistik deskriptif yang digunakan untuk menentukan nilai tengah dari suatu kumpulan data yang telah diurutkan. Median membagi data menjadi dua bagian yang sama besar, dimana separuh dari data memiliki nilai yang lebih kecil atau sama dengan median, dan separuh lainnya memiliki nilai yang lebih besar atau sama dengan median.

Median adalah ukuran pemusatan yang sangat berguna, terutama ketika data memiliki pencilan (outliers) atau distribusi yang tidak simetris, karena median tidak dipengaruhi oleh nilai-nilai ekstrim seperti mean (rata-rata).

Median adalah nilai tengah dalam set data yang terurut dan memiliki aplikasi penting dalam teori pendidikan, termasuk dalam evaluasi hasil belajar, analisis data pendidikan, dan pemahaman konsep statistik oleh guru dan siswa.

Median merupakan konsep penting dalam statistika yang telah didefinisikan oleh berbagai ahli statistik terkemuka. Secara umum, para ahli sepakat bahwa median adalah nilai tengah yang membagi serangkaian data menjadi dua bagian yang sama besar ketika data tersebut diurutkan.

Bluman (2018) menekankan aspek pengurutan data dalam definisinya, menyatakan bahwa median adalah nilai tengah ketika data diurutkan dari terendah ke tertinggi. Sejalan dengan ini, Triola (2018) juga menyoroti pentingnya pengurutan, baik secara meningkat maupun menurun, dalam menentukan median.

Weiss (2017) memberikan perspektif yang sedikit berbeda dengan menyebut median sebagai "ukuran lokasi pusat". Definisinya menekankan fungsi median dalam membagi

distribusi data menjadi dua bagian yang sama, dengan setengah pengamatan di bawah dan setengah di atas nilai median.

Moore, McCabe, dan Craig (2017) memperkuat konsep ini dengan definisi yang serupa, menekankan posisi median sebagai nilai pengamatan yang berada tepat di tengah-tengah data yang telah diurutkan.

Walpole et al. (2016) memberikan definisi yang lebih rinci, menjelaskan bahwa jumlah data di bawah atau sama dengan median harus sama dengan jumlah data di atas atau sama dengan median. Definisi ini memperjelas fungsi median dalam membagi data secara merata.

Meskipun ada sedikit variasi dalam cara para ahli ini mendefinisikan median, inti dari konsepnya tetap konsisten. Median berfungsi sebagai titik tengah dalam serangkaian data yang telah diurutkan, membagi data menjadi dua bagian yang sama besar. Konsep ini membuatnya menjadi ukuran tendensi sentral yang berguna, terutama ketika berhadapan dengan distribusi data yang miring atau mengandung nilai-nilai ekstrem.

Pemahaman yang mendalam tentang median, sebagaimana didefinisikan oleh para ahli ini, sangat penting dalam analisis statistik. Median memberikan gambaran tentang nilai tengah suatu dataset yang sering kali lebih representatif dibandingkan mean, terutama ketika data tidak terdistribusi secara normal atau mengandung outlier. Median memiliki kelebihan dan kekurangan dalam prakteknya. Beberapa ini dapat dijelaskan sebagai berikut.

a. Kelebihan

- 1) Tidak terpengaruh oleh nilai ekstrem (*outlier*)
Median lebih tahan terhadap nilai-nilai ekstrem dibandingkan dengan rata-rata (*mean*), sehingga lebih baik dalam merepresentasikan data yang memiliki outlier (Bluman, 2018).
- 2) Cocok untuk data ordinal
Median dapat digunakan untuk data ordinal, di mana urutan penting tetapi jarak antar nilai tidak diketahui (Weiss, 2017).

- 3) Mudah diinterpretasikan
Median mudah dipahami sebagai nilai tengah yang membagi data menjadi dua bagian yang sama besar (Triola, 2018).
- 4) Berguna untuk distribusi data yang miring (*skewed*)
Pada distribusi data yang tidak simetris, median sering kali memberikan gambaran yang lebih baik tentang nilai pusat dibandingkan mean (Bluman, 2018).

b. Kekurangan

- 1) Tidak mempertimbangkan semua nilai data
Median hanya memperhatikan nilai tengah, mengabaikan informasi dari nilai-nilai lainnya dalam dataset (Weiss, 2017).
- 2) Kurang stabil untuk sampel kecil
Pada sampel yang sangat kecil, perubahan satu nilai dapat memengaruhi median secara signifikan (Triola, 2018).
- 3) Tidak cocok untuk analisis statistik lanjutan
Beberapa analisis statistik lanjutan lebih mudah dilakukan dengan menggunakan mean daripada median (Bluman, 2018).
- 4) Sulit dihitung untuk data berkelompok
Untuk data yang telah dikelompokkan, perhitungan median menjadi lebih kompleks dan kurang akurat (Weiss, 2017).

Berikut adalah langkah-langkah untuk menghitung median:

a. Mengurutkan Data: Susun data dalam urutan menaik (dari yang terkecil ke yang terbesar) atau menurun (dari yang terbesar ke yang terkecil).

b. Menentukan Posisi Median:

- 1) Jika jumlah data (n) ganjil, median adalah nilai yang berada di posisi tengah. Posisi median dapat dihitung dengan rumus:

$$\text{Posisi Median} = (n + 1)/2$$

2) Jika jumlah data (n) genap, median adalah rata-rata dari dua nilai tengah. Posisi dua nilai tengah tersebut adalah:

$$\text{Posisi Median 1} = n/2$$

$$\text{Posisi Median 2} = (n/2) + 1$$

Contoh penggunaan median dalam konteks dunia Pendidikan sebagai berikut :

Misalkan seorang guru ingin menganalisis nilai ujian matematika dari sekelompok siswa untuk mendapatkan gambaran mengenai performa kelas. Nilai-nilai ujian dari 15 siswa adalah sebagai berikut:

78, 85, 90, 67, 88, 95, 80, 85, 70, 92, 65, 77, 84, 73, 89.

Langkah-langkah untuk menghitung median:

a. **Mengurutkan Data:** Susun nilai-nilai dalam urutan menaik: 65, 67, 70, 73, 77, 78, 80, 84, 85, 85, 88, 89, 90, 92, 95.

b. **Menentukan Posisi Median:** Karena jumlah data (n) adalah 15 (ganjil), maka posisi median adalah

$$\text{Posisi Median} = \frac{n + 1}{2} = \frac{15 + 1}{2} = 8$$

c. **Menemukan Nilai Median:** Nilai di posisi ke-8 dalam data yang telah diurutkan adalah 84. Jadi, median dari nilai ujian matematika adalah 84.

Interpretasi:

Median nilai ujian sebesar 84 menunjukkan bahwa setengah dari siswa memiliki nilai di bawah 84 dan setengah lagi memiliki nilai di atas 84. Dengan kata lain, 84 adalah nilai tengah yang membagi data nilai siswa menjadi dua bagian yang sama besar.

Menggunakan median sebagai ukuran pemusatan dalam analisis nilai ujian membantu guru untuk mendapatkan gambaran yang lebih akurat mengenai performa keseluruhan kelas, terutama jika ada beberapa nilai yang sangat rendah atau sangat tinggi (pencilan) yang dapat mempengaruhi rata-rata (mean). Median lebih tahan terhadap pencilan tersebut dan memberikan representasi

yang lebih realistis tentang performa tipikal siswa dalam ujian.

3. Modus

Mode atau Modus, sebagai salah satu ukuran pemusatan dalam statistik deskriptif, memiliki filosofi informasi yang menarik dan unik.

Filosofi informasi Mode (Modus) sebagai ukuran pemusatan dalam statistik deskriptif mencakup berbagai aspek yang menarik untuk dikaji. Mode, yang merepresentasikan nilai yang paling umum atau sering muncul dalam suatu dataset, memiliki peran unik dalam pemahaman dan interpretasi data statistik.

Agresti dan Franklin (2018) dalam buku mereka "Statistics: The Art and Science of Learning from Data" menjelaskan bahwa mode mencerminkan gagasan bahwa dalam sebuah populasi atau sampel, ada kecenderungan tertentu yang dominan. Ini dapat dilihat sebagai manifestasi dari "norma" atau "standar" dalam konteks data tertentu, memberikan gambaran cepat tentang tren umum dalam dataset.

Konsep simplisitas dan intuisi mode sejalan dengan prinsip "Pisau Occam" dalam filosofi ilmu, yang dibahas oleh Baker (2016) dalam Stanford Encyclopedia of Philosophy. Prinsip ini menyatakan bahwa penjelasan paling sederhana sering kali yang terbaik. Mode, sebagai ukuran pemusatan yang relatif sederhana, menawarkan cara intuitif untuk memahami kecenderungan pusat data tanpa perhitungan kompleks.

Namun, seperti yang diungkapkan oleh Wheelan (2014) dalam bukunya "Naked Statistics", mode juga memiliki keterbatasan. Mode hanya memberikan informasi tentang nilai yang paling sering muncul, tanpa mempertimbangkan distribusi nilai-nilai lainnya. Ini mengingatkan kita bahwa setiap ukuran statistik memiliki batasan dalam menggambarkan realitas yang kompleks.

Kompleksitas dalam data sering tercermin dalam fenomena multimodalitas. DeCarlo (1997) dalam artikelnya tentang makna dan penggunaan kurtosis, menyoroti bahwa

keberadaan multiple modes (bimodal, trimodal, dll.) dalam suatu dataset menunjukkan kompleksitas dan keragaman dalam populasi. Ini mencerminkan filosofi bahwa realitas sering kali tidak dapat direduksi menjadi satu nilai tunggal. Gelman dan Hennig (2017) menekankan pentingnya konteks dalam interpretasi statistik. Dalam kasus mode, interpretasinya sangat bergantung pada konteks data. Misalnya, dalam data kategorikal, mode mungkin merupakan satu-satunya ukuran pemusatan yang bermakna.

Taleb (2007) dalam bukunya "The Black Swan" membahas konsep stabilitas dan perubahan dalam sistem kompleks. Ini relevan dengan mode, yang dapat sangat sensitif terhadap perubahan kecil dalam dataset, terutama dalam sampel kecil. Hal ini mencerminkan ketegangan antara stabilitas dan perubahan dalam sistem informasi.

Perbedaan antara data diskrit dan kontinu, yang dibahas secara mendalam oleh Jaynes (2003) dalam "Probability Theory: The Logic of Science", juga mempengaruhi interpretasi mode. Mode lebih mudah diidentifikasi dan diinterpretasikan dalam data diskrit dibandingkan data kontinu, menggambarkan perbedaan filosofis antara pandangan diskrit dan kontinu tentang realitas.

Dalam konteks teori informasi, Cover dan Thomas (2006) menjelaskan konsep entropi informasi. Mode dapat dilihat sebagai nilai dengan entropi informasi terendah - yaitu, nilai yang memberikan informasi paling sedikit ketika diketahui, karena merupakan nilai yang paling dapat diprediksi dalam dataset.

Akhirnya, O'Neil (2016) dalam "Weapons of Math Destruction" membahas implikasi etis penggunaan statistik dalam pengambilan keputusan. Penggunaan mode dalam kebijakan publik, misalnya, bisa mengabaikan kebutuhan kelompok minoritas atau outlier, menimbulkan pertanyaan tentang keadilan dan representasi dalam analisis data.

Pemahaman yang lebih mendalam tentang filosofi statistik secara umum, seperti yang dibahas oleh Bandyopadhyay dan Forster (2011) serta Mayo (2018), dapat membantu kita

menggunakan dan menginterpretasikan mode dengan lebih bijak dalam analisis statistik. Mode, sebagai ukuran pemusatan, menawarkan perspektif unik dalam memahami dan merepresentasikan data, mencerminkan keseimbangan antara simplisitas dan kompleksitas, serta pentingnya konteks dalam interpretasi data.

Kelebihan dan Kekurangan Modus

Modus, sebagai salah satu ukuran pemusatan data dalam statistik, memiliki beberapa kelebihan dan kekurangan yang perlu dipertimbangkan ketika menggunakannya dalam analisis data.

Kelebihan Modus

a. Mudah dipahami dan dihitung

- 1) Modus adalah konsep yang intuitif dan mudah dijelaskan kepada orang awam (Bluman, 2018)
- 2) Perhitungannya sederhana, terutama untuk data kategorikal atau diskrit (Wheeler, 2014)

b. Tidak terpengaruh oleh nilai ekstrem

Modus tidak sensitif terhadap outlier atau nilai ekstrem dalam dataset (Agresti, & Franklin, 2018).

c. Berguna untuk data kategorikal

Modus adalah satu-satunya ukuran pemusatan yang dapat digunakan untuk data kategorikal atau nominal (Moore et al, 2017).

d. Dapat digunakan untuk distribusi tidak simetris

Modus tetap bermakna untuk distribusi yang sangat miring atau tidak simetris (Walpole et al, 2016).

e. Membantu identifikasi cluster

Dalam distribusi multimodal, modus dapat membantu mengidentifikasi kelompok atau cluster dalam data (Everitt dan Skrondal, 2010).

f. Cepat untuk estimasi kecenderungan pusat

Dalam situasi praktis, modus dapat stimasi cepat tentang nilai yang paling umum (Bluman, 2018).

Kekurangan Modus

a. Tidak selalu unik

Dataset dapat memiliki lebih dari satu modus (bimodal, multimodal), yang dapat membuat interpretasi menjadi kompleks (Johnson dan Bhattacharyya, 2019).

b. Tidak memperhitungkan semua nilai

Modus hanya mempertimbangkan frekuensi, mengabaikan nilai aktual data lainnya (Triola, 2018).

c. Kurang stabil dalam sampel

Modus dapat berubah secara signifikan dengan perubahan kecil dalam dataset, terutama untuk sampel kecil (Howell, 2012).

d. Tidak cocok untuk perhitungan lanjutan

Modus kurang berguna untuk analisis statistik lanjutan dibandingkan mean atau median Field, 2013).

e. Bisa menyesatkan untuk data kontinu

Untuk data kontinu, modus dapat menjadi kurang bermakna dan bergantung pada bagaimana data dikelompokkan (Gravetter dan Wallnau, 2016).

f. Tidak mencerminkan keseluruhan distribusi

Modus hanya memberikan informasi tentang nilai yang paling sering muncul, bukan gambaran keseluruhan distribusi data (Urdan, 2016).

Modus adalah alat yang berguna dalam analisis statistik, terutama untuk data kategorikal dan untuk mendapatkan gambaran cepat tentang nilai yang paling umum dalam dataset (Larson dan Farber, 2014). Namun, keterbatasannya berarti bahwa modus sering digunakan bersama dengan ukuran pemusatan lain seperti mean dan median untuk memberikan gambaran yang lebih komprehensif tentang distribusi data (Weiss, 2017).

Karakteristik Mode

1. Unimodal, Bimodal, Multimodal:

- a. Data yang memiliki satu mode disebut unimodal.
- b. Data yang memiliki dua mode disebut bimodal.
- c. Data yang memiliki lebih dari dua mode disebut multimodal.

2. Tidak Unik:

- a. Dalam beberapa set data, mungkin tidak ada mode jika tidak ada nilai yang muncul lebih dari sekali.
- b. Sebaliknya, bisa ada lebih dari satu mode jika beberapa nilai muncul dengan frekuensi yang sama dan lebih sering dibandingkan nilai lainnya.

3. Aplikasi:

- a. Mode sangat berguna untuk data kategorikal, seperti warna favorit, merek produk yang paling sering dibeli, dll.
- b. Mode juga berguna untuk data numerik, terutama dalam distribusi yang tidak simetris.

Cara Menghitung Mode

Mengidentifikasi Frekuensi Terbesar:

1. Hitung frekuensi kemunculan setiap nilai dalam data.
2. Nilai yang muncul paling sering adalah mode.

Contoh Perhitungan Mode

1. **Data Kualitatif:** Misalkan sebuah survei dilakukan untuk mengetahui warna favorit dari 20 siswa dengan hasil sebagai berikut: Merah, Biru, Hijau, Biru, Kuning, Merah, Merah, Biru, Hijau, Biru, Kuning, Merah, Hijau, Biru, Merah, Biru, Kuning, Merah, Hijau, Biru.
 - a. Hitung frekuensi kemunculan setiap warna:
 - 1) Merah: 6
 - 2) Biru: 7
 - 3) Hijau: 4
 - 4) Kuning: 3
 - b. Mode adalah warna yang muncul paling sering, yaitu Biru.
2. **Data Numerik:** Misalkan nilai ujian dari 10 siswa adalah sebagai berikut: 85, 78, 92, 85, 88, 90, 85, 87, 92, 78.
 - a. Hitung frekuensi kemunculan setiap nilai:
 - 1) 85: 3
 - 2) 78: 2
 - 3) 92: 2
 - 4) 88: 1

5) 90: 1

6) 87: 1

b. Mode adalah nilai yang muncul paling sering, yaitu 85.

2.3.2 Ukuran Variabilitas

1. Range

Dalam dunia statistik deskriptif, kita mengenal sebuah konsep sederhana namun kuat yang disebut Range. Range ini bisa diibaratkan sebagai jembatan yang menghubungkan dua titik terjauh dalam sebuah lanskap data. Bayangkan Anda berdiri di tepi sungai, memandang dari satu tepi ke tepi lainnya. Jarak antara kedua tepi itulah yang kita sebut sebagai range dalam konteks data statistik.

Range memberikan kita pandangan luas tentang sebaran data dengan cara yang sangat intuitif. Ia menceritakan kisah tentang seberapa lebar rentang nilai dalam kumpulan data kita, mulai dari yang paling rendah hingga yang paling tinggi. Seperti halnya mengukur jarak antara dua kota terjauh dalam sebuah negara, range memberi kita gambaran tentang "wilayah" yang dicakup oleh data kita.

Meskipun sederhana, range memiliki kekuatan tersendiri. Ia bisa menjadi pembuka jalan dalam perjalanan kita memahami karakteristik data. Dengan cepat, range bisa memberi isyarat tentang adanya nilai-nilai yang mungkin 'nyeleneh' atau ekstrem dalam data kita. Ini seperti melihat sebuah gunung yang menjulang tinggi di kejauhan, menandakan adanya sesuatu yang berbeda dari dataran di sekitarnya.

Namun, seperti halnya melihat lanskap hanya dari dua titik terjauh, range juga memiliki keterbatasannya. Ia mungkin tidak menceritakan kisah lengkap tentang apa yang terjadi di "lembah-lembah" antara kedua ujungnya. Range tidak memberi tahu kita tentang bagaimana data tersebar di antara nilai terendah dan tertinggi. Ini seperti mengetahui jarak antara dua kota tanpa tahu apa yang ada di sepanjang perjalanan di antaranya.

Meski begitu, kesederhanaan range justru menjadi kekuatannya. Ia menjadi alat yang berguna untuk membandingkan sebaran data antar kelompok dengan cepat. Seperti membandingkan lebar dua sungai berbeda dengan sekali pandang.

Range juga sering menjadi langkah pertama dalam perjalanan kita menyelami lautan data yang lebih dalam. Ia memberi kita pijakan awal sebelum kita mulai menggunakan alat-alat yang lebih canggih untuk mengukur variabilitas data, seperti standar deviasi atau rentang antarkuartil.

Dengan demikian, range, dalam kesederhanaannya, memegang peran penting dalam statistik deskriptif. Ia menjadi penunjuk jalan yang membantu kita mulai memahami karakteristik data kita, memberi kita gambaran besar sebelum kita menyelam lebih dalam ke detail-detail yang lebih kompleks.

Definisi Range Menurut Para Pakar

Konsep range dalam statistik telah lama menjadi bagian penting dari analisis data. Berbagai pakar statistik telah memberikan definisi yang konsisten namun dengan nuansa yang sedikit berbeda sepanjang dekade.

Kita bisa memulai perjalanan ini dari tahun 1972, ketika Ronald E. Walpole dan Raymond H. Myers dalam buku mereka "Probability and Statistics for Engineers and Scientists" mendefinisikan range dengan sederhana sebagai selisih antara pengamatan terbesar dan terkecil dalam suatu sampel. Definisi ini menjadi dasar yang kuat untuk pemahaman kita tentang range.

Tiga dekade kemudian, pada tahun 2005, Robert V. Hogg, Joseph W. McKean, dan Allen T. Craig dalam "Introduction to Mathematical Statistics" memperluas pemahaman ini. Mereka tidak hanya mendefinisikan range, tetapi juga menekankan posisinya sebagai statistik yang paling sederhana dan paling tua untuk mengukur variabilitas. Ini memberikan kita perspektif historis tentang pentingnya konsep range dalam perkembangan statistik.

Lima tahun kemudian, pada 2010, Douglas C. Montgomery dan George C. Runger dalam "Applied Statistics and Probability for Engineers" memberikan definisi yang serupa namun dengan penekanan pada istilah "sampel". Mereka menyebut konsep ini sebagai "range sampel", yang mungkin untuk membedakannya dari range populasi.

Perkembangan definisi berlanjut dengan David M. Lane pada tahun 2013 dalam bukunya "Introduction to Statistics". Lane menggunakan istilah "skor" alih-alih "pengamatan" atau "nilai", yang mungkin mencerminkan fokus yang lebih besar pada aplikasi statistik dalam ilmu perilaku atau pendidikan.

Terakhir, pada tahun 2017, Sheldon M. Ross dalam "Introductory Statistics" memberikan definisi yang komprehensif, menyebutkan range sebagai perbedaan antara nilai maksimum dan minimum dari serangkaian pengamatan. Definisi Ross ini seolah-olah menjadi sintesis dari definisi-definisi sebelumnya, menggabungkan elemen-elemen kunci dari setiap definisi.

Menariknya, meskipun definisi-definisi ini muncul dalam rentang waktu 45 tahun, dari 1972 hingga 2017, inti dari konsep range tetap konsisten. Ini menunjukkan betapa fundamental dan mapannya konsep range dalam statistik. Setiap pakar, dengan caranya masing-masing, menekankan aspek yang sama: range adalah ukuran perbedaan antara nilai tertinggi dan terendah dalam suatu kumpulan data.

Perjalanan definisi range ini mencerminkan evolusi pemikiran statistik yang terus berkembang namun tetap menghormati konsep-konsep dasarnya. Ini mengingatkan kita bahwa meskipun statistik adalah bidang yang dinamis, beberapa konsep intinya tetap kokoh dan relevan sepanjang waktu.

Persamaan matematis untuk menghitung range :

Range=Nilai Maksimum–Nilai Minimum

Berikut adalah contoh soal aplikasi Range dalam dunia pendidikan:

Contoh Soal 1:

Seorang guru ingin menganalisis hasil ujian matematika dari 10 siswanya untuk memahami seberapa jauh perbedaan nilai di antara mereka. Nilai-nilai ujian yang diperoleh siswa adalah sebagai berikut:

85, 90, 78, 92, 88, 76, 95, 80, 84, 89

Hitunglah Range dari nilai ujian tersebut dan jelaskan interpretasinya dalam konteks kelas tersebut.

Penyelesaian:

1) Tentukan nilai maksimum dan minimum:

- Nilai maksimum = 95

- Nilai minimum = 76

Range=Nilai Maksimum–Nilai Minimum=95–76=19

Interpretasi:

Range sebesar 19 menunjukkan bahwa perbedaan antara nilai tertinggi dan nilai terendah di kelas tersebut adalah 19 poin. Dalam konteks kelas, ini berarti bahwa ada variasi yang cukup besar dalam kemampuan siswa dalam mengerjakan ujian matematika ini. Meskipun Range tidak memberikan informasi tentang bagaimana nilai-nilai tersebar di antara nilai maksimum dan minimum, ini memberikan petunjuk bahwa ada perbedaan yang signifikan antara siswa dengan nilai tertinggi dan terendah. Guru mungkin perlu mempertimbangkan pendekatan pengajaran yang lebih individual untuk membantu siswa yang nilainya lebih rendah mendekati nilai-nilai yang lebih tinggi.

Contoh Soal 2:

Seorang kepala sekolah ingin mengevaluasi kinerja akademik dari dua kelas yang berbeda. Berikut adalah rata-rata nilai ujian sains dari 8 siswa di Kelas A dan 8 siswa di Kelas B:

- Kelas A: 78, 85, 80, 90, 88, 92, 84, 86

- Kelas B: 65, 82, 90, 70, 75, 68, 88, 73

Hitunglah Range untuk masing-masing kelas, dan bandingkan penyebaran nilai antar kedua kelas tersebut.

Penyelesaian:

- 1) Kelas A:
 - Nilai maksimum = 92
 - Nilai minimum = 78
 - Range = $92 - 78 = 14$
- 2) Kelas B:
 - Nilai maksimum = 90
 - Nilai minimum = 65
 - Range = $90 - 65 = 25$

Interpretasi:

Range di Kelas A adalah 14, sedangkan Range di Kelas B adalah 25. Ini menunjukkan bahwa penyebaran nilai di Kelas B lebih besar dibandingkan dengan Kelas A. Dengan kata lain, di Kelas B terdapat variasi yang lebih besar antara siswa yang nilai ujiannya tertinggi dan terendah. Hal ini mungkin menunjukkan adanya perbedaan yang lebih besar dalam pemahaman materi di Kelas B dibandingkan dengan Kelas A.

Contoh Soal 3:

Seorang dosen ingin mengevaluasi hasil proyek akhir mahasiswa dalam mata kuliah Statistik. Dia mencatat nilai proyek akhir dari 6 mahasiswanya sebagai berikut:

90, 87, 95, 88, 92, 89

Dosen tersebut ingin mengetahui seberapa konsisten kinerja mahasiswa dalam proyek akhir ini. Hitunglah Range dari nilai-nilai tersebut dan berikan interpretasinya.

Penyelesaian:

- 1) Tentukan nilai maksimum dan minimum:
 - Nilai maksimum = 95
 - Nilai minimum = 87
- 2) Hitung Range:
Range = Nilai Maksimum - Nilai Minimum = $95 - 87 = 8$

Interpretasi:

Range sebesar 8 menunjukkan bahwa perbedaan antara nilai tertinggi dan nilai terendah dari proyek akhir hanya 8 poin. Ini menunjukkan bahwa kinerja mahasiswa dalam

proyek akhir tersebut cukup konsisten. Perbedaan nilai yang relatif kecil menunjukkan bahwa sebagian besar mahasiswa memiliki pemahaman yang baik dan relatif setara terhadap materi yang diajarkan.

2. Standar Defiasi

Dalam lanskap luas statistika, Standar Deviasi berdiri tegak sebagai salah satu konsep paling fundamental dan berpengaruh. Ia tidak hanya sekadar rumus matematika, tetapi juga cerminan dari cara kita memahami dan menafsirkan variabilitas dalam dunia yang penuh ketidakpastian.

Perjalanan kita dimulai dengan pemikiran Ronald A. Fisher, bapak statistika modern. Fisher, dalam karyanya "Statistical Methods for Research Workers" (1925), menekankan pentingnya variabilitas dalam pemahaman fenomena alam. Ia melihat Standar Deviasi sebagai alat untuk mengukur "scatter" atau penyebaran data, yang menurutnya sama pentingnya dengan nilai tengah dalam memahami suatu distribusi.

Konsep ini kemudian diperdalam oleh Karl Pearson, yang dalam "The Grammar of Science" (1892) menggambarkan Standar Deviasi sebagai "radius of gyration" dari sebuah distribusi. Pearson melihatnya sebagai cara untuk mengukur seberapa jauh, secara rata-rata, titik-titik data menyimpang dari pusatnya. Ini mencerminkan pemikiran bahwa untuk benar-benar memahami sebuah dataset, kita perlu melihat bukan hanya ke pusatnya, tetapi juga ke seberapa jauh data menyebar dari pusat tersebut.

George W. Snedecor dan William G. Cochran, dalam buku klasik mereka "Statistical Methods" (1989), memperluas pemahaman ini dengan menjelaskan bahwa Standar Deviasi memberikan gambaran tentang "tipikal" penyimpangan dari mean. Mereka menekankan bahwa dengan mengkuadratkan penyimpangan, kita memberikan bobot lebih pada nilai-nilai ekstrem, mencerminkan pandangan bahwa penyimpangan besar lebih signifikan dalam memahami variabilitas.

Filosofi di balik Standar Deviasi juga erat kaitannya dengan konsep distribusi normal, yang dipopulerkan oleh Carl Friedrich Gauss. Dalam "Theoria Motus Corporum Coelestium" (1809), Gauss menggambarkan kurva lonceng yang sekarang dikenal dengan namanya, di mana Standar Deviasi memainkan peran kunci dalam menentukan bentuk kurva. Ini mencerminkan pemikiran bahwa variabilitas dapat dipahami dalam kerangka probabilitas, sebuah ide yang revolusioner pada masanya dan tetap relevan hingga hari ini.

David S. Moore dan George P. McCabe, dalam "Introduction to the Practice of Statistics" (2006), menekankan peran Standar Deviasi dalam membandingkan variabilitas antar dataset. Mereka menggambarkan sebagai "penggaris" universal untuk mengukur penyebaran, memungkinkan kita membandingkan variabilitas bahkan ketika unit atau skala berbeda. Ini mencerminkan keinginan dalam statistika untuk memiliki ukuran yang tidak hanya deskriptif, tetapi juga komparatif.

Dalam konteks yang lebih luas, Nicholas Nassim Taleb, dalam bukunya "The Black Swan" (2007), membahas keterbatasan Standar Deviasi dalam menangkap risiko dan ketidakpastian di dunia nyata. Taleb mengingatkan kita bahwa meskipun Standar Deviasi adalah alat yang kuat, ia memiliki batasan, terutama dalam menghadapi peristiwa-peristiwa ekstrem atau "black swan". Ini mengingatkan kita akan pentingnya memahami tidak hanya kekuatan, tetapi juga keterbatasan dari alat-alat statistik kita.

Perjalanan filosofis Standar Deviasi mencerminkan evolusi pemikiran statistik itu sendiri - dari deskripsi sederhana ke pemahaman yang lebih dalam tentang variabilitas, ketidakpastian, dan kompleksitas dunia di sekitar kita. Ia berdiri sebagai jembatan antara kompleksitas realitas dan kebutuhan kita akan pemahaman yang dapat ditindaklanjuti, mencerminkan keinginan mendalam manusia untuk mengukur, memahami, dan pada akhirnya, mengendalikan dunia yang penuh ketidakpastian di sekitar kita.

Untuk menghitung **Standar Deviasi**, kita dapat menggunakan dua rumus yang berbeda tergantung pada apakah kita bekerja dengan populasi keseluruhan atau hanya sampel dari populasi.

a. Standar Deviasi Populasi (σ)

$$\sigma = \sqrt{\frac{1}{N} \sum_{i=1}^N (xi - \mu)^2}$$

Dimana :

σ = Standar Deviasi Populasi

N = Jumlah Total data dalam populasi

Xi = Nilai data Individu

μ = Rata-rata Populasi

Σ = Simbol penjumlahan, menjumlahkan semua nilai dari $I = 1$ hingga N

b. Standar Deviasi Sampel (s)

Jika kita menghitung Standar Deviasi dari sampel yang diambil dari populasi, rumusnya adalah:

$$s = \sqrt{\frac{1}{n-1} \sum_{i=1}^n (xi - \bar{x})^2}$$

Dimana :

S = Standar Deviasi sampel

N = Jumlah data sampel

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

Σ = Simbol Penjumlahan, menjumlahkan semua nilai dari $i=1$ hingga n

Langkah-langkah Menghitung Standar Deviasi:

- a. Hitung Rata-rata (μ atau $\bar{x}$):** Tambahkan semua nilai dalam dataset dan bagi dengan jumlah data untuk mendapatkan rata-rata.

- b. **Kurangkan Rata-rata dari Setiap Nilai ($x_i - \mu$):** Hitung selisih antara setiap nilai dalam dataset dan rata-rata.
- c. **Kuadratkan Setiap Selisih:** Kuadratkan hasil pengurangan untuk setiap nilai.
- d. **Jumlahkan Semua Kuadrat:** Tambahkan semua nilai kuadrat yang dihasilkan.
- e. **Bagi dengan Jumlah Data (Untuk Populasi) atau dengan $n-1$ (Untuk Sampel):** Jika bekerja dengan populasi, bagi jumlah kuadrat dengan N ; jika dengan sampel, bagi dengan $n-1$.
- f. **Akar Kuadrat:** Terakhir, ambil akar kuadrat dari hasil pembagian untuk mendapatkan Standar Deviasi.

Contoh Soal:

Seorang guru ingin mengetahui variasi nilai ulangan matematika dari 5 siswanya. Berikut adalah nilai yang diperoleh siswa:

80, 85, 90, 75, 95

Hitunglah Standar Deviasi dari nilai-nilai tersebut.

Langkah-langkah Perhitungan:

- a. **Hitung Rata-rata ($\bar{x}$)**

$$\bar{x} = \frac{80 + 85 + 90 + 75 + 95}{5} = \frac{425}{5} = 85$$

- b. **Kurangkan Rata-rata dari Setiap x_i atau $\bar{x}$**

$$80 - 85 = -5 ; 85 - 85 = 0 ; 90 - 85 = 5 ; 90 - 85 = 5 ; 75 - 85 = -10 ; 95 - 85 = 10$$

- c. **Kuadratkan Setiap Selisih**

$$-5^2 = 25 ; 0^2 = 0 ; 5^2 = 25 ; 10^2 = 100 ; 10^2 = 100 ;$$

- d. **Jumlahkan Semua Kuadrat**

$$25 + 0 + 25 + 100 + 100 = 250$$

$$s^2 = \frac{250}{5 - 1} = \frac{250}{4} = 62,5$$

- e. **Akar Kuadrat**

$$s = \sqrt{62,5} \approx 7,91$$

Interpretasi:

Standar Deviasi dari nilai-nilai ulangan matematika ini adalah sekitar 7.91. Ini berarti bahwa rata-rata, nilai siswa bervariasi sekitar 7.91 poin dari nilai rata-rata kelas (85). Standar Deviasi ini memberikan gambaran bahwa ada variasi moderat dalam performa siswa pada ulangan ini. Nilai ini dapat digunakan oleh guru untuk memahami seberapa konsisten siswa dalam mencapai hasil yang diharapkan.

Contoh Soal 2

Seorang dosen ingin mengevaluasi variasi nilai tugas akhir dari 6 mahasiswa di kelasnya. Berikut adalah nilai tugas akhir yang diperoleh mahasiswa: 78, 82, 88, 91, 76, 85. Hitunglah Standar Deviasi dari nilai-nilai tersebut.

Langkah-langkah Perhitungan:

a. Hitung Rata-rata ($\bar{x}$)

$$\bar{x} = \frac{78 + 82 + 88 + 91 + 76 + 85}{6} = \frac{500}{6} \approx 83,33$$

b. Kurangkan Rata-rata dari Setiap x_i atau $\bar{x}$

$$78 - 83,33 \approx -5,33 ; 82 - 83,33 \approx -1,33 ; 88 - 83,33 \approx 4,67 ; 91 - 83,33 \approx 7,67$$

$$76 - 83,33 \approx -7,33 ; 85 - 83,33 \approx 1,67$$

c. Kuadratkan Setiap Selisih

$$-5,33^2 \approx 28,41 ; -1,33^2 \approx 1,77 ; 4,67^2 \approx 21,81 ; 7,67^2 \approx 58,71 ; -7,33^2 \approx 53,71 ; 1,67^2 \approx 2,79$$

d. Jumlahkan Semua Kuadrat

$$28,41 + 1,77 + 21,81 + 58,71 + 53,71 + 2,79 \approx 167,34$$

e. Bagi dengan $n - 1$ (untuk sampel)

$$s^2 = \frac{167,34}{6 - 1} = \frac{167,34}{5} \approx 33,47$$

f. Akar Kuadrat

$$s = \sqrt{33,47} \approx 5,79$$

Interpretasi:

Standar Deviasi dari nilai tugas akhir ini adalah sekitar 5.79. Ini menunjukkan bahwa rata-rata, nilai mahasiswa bervariasi sekitar 5.79 poin dari nilai rata-rata kelas (83.33). Variasi ini bisa membantu dosen memahami perbedaan tingkat penguasaan materi di antara mahasiswa.

3. Varian

Dalam dunia statistik yang penuh dengan angka dan rumus, Varians berdiri sebagai konsep yang unik dan mendalam. Sebagai kuadrat dari standar deviasi, Varians mungkin tampak seperti langkah tambahan yang tidak perlu dalam mengukur variabilitas. Namun, di balik kesederhanaan matematisnya, tersembunyi filosofi yang kaya dan kompleks tentang bagaimana kita memahami dan mengukur ketidakpastian dalam data.

Konsep Varians pertama kali diperkenalkan oleh Ronald A. Fisher dalam karyanya "The Correlation between Relatives on the Supposition of Mendelian Inheritance" (1918). Fisher melihat Varians bukan hanya sebagai ukuran variabilitas, tetapi sebagai cara untuk memahami bagaimana sifat-sifat genetik diturunkan dari satu generasi ke generasi berikutnya (Fisher, 1981). Ini menunjukkan bahwa sejak awal, Varians dipandang sebagai alat untuk memahami proses yang mendasari variabilitas, bukan hanya deskripsi sederhana dari sebaran data.

Maurice Kendall dan Alan Stuart, dalam buku klasik mereka "The Advanced Theory of Statistics" (1969), memperdalam pemahaman ini dengan menjelaskan Varians sebagai "momen kedua" dari distribusi. Mereka menekankan bahwa dengan mengkuadratkan penyimpangan dari mean, Varians memberikan bobot yang lebih besar pada nilai-nilai ekstrem (Kendall dan Stuart, 1969). Ini mencerminkan filosofi bahwa dalam banyak situasi, penyimpangan besar dari rata-rata mungkin lebih penting atau lebih informatif daripada penyimpangan kecil.

John Tukey, bapak analisis data eksploratori, dalam karyanya "Exploratory Data Analysis" (1977), membahas

Varians dalam konteks yang lebih luas dari pemahaman data. Tukey menekankan bahwa Varians, meskipun berguna, hanyalah satu cara untuk melihat variabilitas. Ia mendorong para analis untuk "menggali" lebih dalam ke dalam data, menggunakan berbagai alat termasuk Varians, untuk mendapatkan pemahaman yang lebih kaya (Tukey, 1977). Ini mencerminkan filosofi bahwa tidak ada ukuran tunggal yang dapat menangkap sepenuhnya kompleksitas data nyata.

Dalam konteks teori probabilitas, Andrey Kolmogorov memberikan fondasi matematis yang kuat untuk konsep Varians dalam bukunya "Foundations of the Theory of Probability" (1933). Kolmogorov menunjukkan bagaimana Varians muncul secara alami dari aksioma-aksioma probabilitas, menegaskan perannya sebagai ukuran fundamental dari ketidakpastian (Kolmogorov, 1933). Ini memperkuat gagasan bahwa Varians bukan hanya alat deskriptif, tetapi juga konsep yang mendalam yang berakar pada prinsip-prinsip dasar probabilitas.

George E. P. Box, dalam artikelnya yang terkenal "Science and Statistics" (1976), membahas peran Varians dalam pemodelan statistik. Box menekankan bahwa semua model statistik pada dasarnya salah, tetapi beberapa berguna, dan Varians memainkan peran kunci dalam menilai kegunaan model (Box, 1976). Ini mencerminkan filosofi bahwa dalam dunia yang kompleks dan tidak pasti, tujuan kita bukanlah mencapai kepastian absolut, tetapi untuk mengukur dan mengelola ketidakpastian sebaik mungkin.

Nassim Nicholas Taleb, dalam bukunya "Antifragile" (2012), membawa diskusi tentang Varians ke tingkat yang lebih filosofis. Taleb berpendapat bahwa variabilitas, yang diukur oleh Varians, bukan hanya sesuatu yang harus ditoleransi atau diminimalkan, tetapi dalam banyak sistem, variabilitas sebenarnya bisa menjadi sumber kekuatan dan ketahanan (Taleb, 2012). Ini menantang pandangan konvensional tentang variabilitas sebagai sesuatu yang selalu negatif dan harus dikurangi.

Bradley Efron dan Trevor Hastie, dalam buku mereka "Computer Age Statistical Inference" (2016), membahas peran Varians dalam era big data dan machine learning. Mereka menunjukkan bagaimana konsep Varians tetap relevan bahkan dalam konteks algoritma dan metode statistik yang sangat canggih (Efron dan Hastie, 2016). Ini menegaskan daya tahan dan fleksibilitas konsep Varians dalam menghadapi perubahan dramatis dalam lanskap analisis data.

Filosofi di balik Varians mencerminkan keinginan mendalam manusia untuk memahami dan mengukur ketidakpastian. Ini bukan hanya tentang menghitung angka, tetapi tentang mengungkap struktur yang mendasari variabilitas dalam data. Varians berdiri sebagai jembatan antara realitas yang kompleks dan tidak pasti dengan kebutuhan kita akan pemahaman yang dapat diukur dan dikelola. Ia mengingatkan kita bahwa dalam dunia yang penuh dengan ketidakpastian, tugas kita bukan untuk menghilangkan variabilitas, tetapi untuk memahaminya, mengukurnya, dan terkadang bahkan memanfaatkannya.

Varians adalah ukuran yang menunjukkan seberapa jauh nilai-nilai dalam dataset tersebar dari rata-rata (mean). Varians adalah kuadrat dari Standar Deviasi, dan memberikan gambaran tentang variabilitas data dalam satuan yang sama dengan data asli, namun dalam bentuk kuadrat.

Rumus Varians:

a. Varians Populasi (σ^2)

$$\sigma^2 = \frac{1}{N} \sum_{i=1}^N (x_i - \mu)^2$$

Dimana :

σ^2 =Varians Populasi

N =Jumlah total data dalam populasi

x_i =Nilai data individu

μ = Rata – rata populasi

$\sum i = 1$ =Simbol penjumlahan, menjumlahkan semua nilai dari $i = 1$ hingga N

b. Varians Sampel (s^2)

$$s^2 = \frac{1}{n-1} \sum_{i=1}^N (x_i - \bar{x})^2$$

Dimana :

s^2 = Varians sampel

n = Jumlah data dalam sampel

x_i = Nilai data individu dalam sampel

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

$\sum i = 1$ = Simbol penjumlahan, menjumlahkan semua nilai dari $i = 1$ hingga n

Langkah-langkah Menghitung Standar Deviasi:

- a. **Hitung Rata-rata (μ atau $\bar{x}$):** Tambahkan semua nilai dalam dataset dan bagi dengan jumlah data untuk mendapatkan rata-rata.
- b. **Kurangkan Rata-rata dari Setiap Nilai ($x_i - \mu$):** Hitung selisih antara setiap nilai dalam dataset dan rata-rata.
- c. **Kuadratkan Setiap Selisih:** Kuadratkan hasil pengurangan untuk setiap nilai.
- d. **Jumlahkan Semua Kuadrat:** Tambahkan semua nilai kuadrat yang dihasilkan.
- e. **Bagi dengan Jumlah Data (Untuk Populasi) atau dengan $n-1$ (Untuk Sampel):**

1) Jika bekerja dengan populasi, bagi jumlah kuadrat dengan N ;

2) Jika dengan sampel, bagi dengan $n-1$.

Contoh 1. Perhitungan Varians

Misalkan kita memiliki data nilai ujian dari 5 siswa: 80, 85, 90, 75, 95. Kita akan menghitung Varians dari data ini (menganggapnya sebagai sampel).

Langkah-langkah Perhitungan:

- a. **Hitung Rata-rata ($\bar{x}$)**

$$\bar{x} = \frac{80 + 85 + 90 + 75 + 95}{5} = \frac{425}{5} = 85$$

b. Kurangkan Rata-rata dari Setiap x_1 atau $\bar{x}$

$$80-85=-5 ; 85-85=0 ; 90-85=5 ; 90-85=5 ; \\ 75-85=-10 ; 95-85=10$$

c. Kuadratkan Setiap Selisih

$$-5^2 = 25 ; 0^2 = 0 ; 5^2 = 25 ; 10^2 = 100 ; 10^2 = 100 ;$$

d. Jumlahkan Semua Kuadrat

$$25+0+25+100+100=250$$

$$s^2 = \frac{250}{5-1} = \frac{250}{4} = 62,5$$

Hasil:

Varians dari nilai ujian ini adalah 62.5. Ini menunjukkan bahwa nilai-nilai dalam dataset ini memiliki penyebaran rata-rata sebesar 62.5 poin kuadrat dari nilai rata-rata (85).

Contoh 2.

Seorang guru ingin menganalisis sebaran nilai ujian matematika di kelasnya. Berikut adalah nilai-nilai yang diperoleh oleh 10 siswa dalam ujian tersebut:

75, 82, 68, 90, 78, 85, 72, 88, 79, 83

Hitunglah varians dari nilai-nilai ujian tersebut.

Untuk menyelesaikan soal ini, kita akan mengikuti langkah-langkah berikut:

1. Hitung rata-rata (mean) dari nilai-nilai tersebut.
2. Hitung selisih setiap nilai dari rata-rata dan kuadratkan hasilnya.
3. Hitung rata-rata dari hasil kuadrat selisih tersebut.

a. Menghitung rata-rata:

$$\text{Jumlah semua nilai} = 75 + 82 + 68 + 90 + 78 + 85 + 72 \\ + 88 + 79 + 83 = 800 ; \text{ Rata-rata} = 800 / 10 = 80$$

b. Menghitung selisih dari rata-rata dan mengkuadratkannya:

$$(75 - 80)^2 = (-5)^2 = 25$$

$$(82 - 80)^2 = (2)^2 = 4$$

$$(68 - 80)^2 = (-12)^2 = 144$$

$$(90 - 80)^2 = (10)^2 = 100$$

$$(78 - 80)^2 = (-2)^2 = 4$$

$$(85 - 80)^2 = (5)^2 = 25$$

$$(72 - 80)^2 = (-8)^2 = 64$$

$$(88 - 80)^2 = (8)^2 = 64$$

$$(79 - 80)^2 = (-1)^2 = 1$$

$$(83 - 80)^2 = (3)^2 = 9$$

- c. Menghitung rata-rata dari hasil kuadrat selisih:

Jumlah kuadrat selisih = $25 + 4 + 144 + 100 + 4 + 25 + 64 + 64 + 1 + 9 = 440$; Varians = $440 / 10 = 44$. Jadi, varians dari nilai-nilai ujian tersebut adalah 44.

2.3.3 Distribusi Frekuensi

Distribusi frekuensi, sebuah konsep fundamental dalam statistik, mencerminkan lebih dari sekadar angka-angka dalam tabel atau grafik. Ia mewakili upaya manusia untuk memahami dan menafsirkan kompleksitas dunia di sekitar kita. Seperti yang dijelaskan oleh Stephen M. Stigler dalam bukunya "The History of Statistics" (1986), konsep ini berkembang dari kebutuhan untuk menemukan pola dan keteraturan dalam fenomena yang tampaknya acak.

Ian Hacking, dalam karyanya "The Taming of Chance" (1990), menggambarkan bagaimana distribusi frekuensi menjadi bagian integral dari "revolusi probabilistik" yang mengubah cara kita memandang dunia. Dari perspektif deterministik yang kaku, kita beralih ke pemahaman yang lebih nuansa tentang ketidakpastian dan variabilitas. Distribusi frekuensi menjadi jembatan antara data mentah dan interpretasi makna, memungkinkan kita untuk melihat hutan tanpa kehilangan pohon-pohonnya.

Gerd Gigerenzer, dalam "Reckoning with Risk" (2002), menekankan bagaimana distribusi frekuensi bukan hanya alat matematis, tetapi juga cara berpikir tentang dunia. Ia menggambarkan bagaimana kita menggunakan konsep ini untuk membuat keputusan dalam menghadapi ketidakpastian, mencerminkan keseimbangan yang rumit antara objektivitas ilmiah dan subjektivitas manusia.

Theodore M. Porter, dalam "The Rise of Statistical Thinking" (1986), menjelaskan bagaimana distribusi frekuensi menjadi bagian dari bahasa umum ilmu pengetahuan, memungkinkan disiplin ilmu yang berbeda untuk berkomunikasi dan membandingkan temuan mereka. Ini

mencerminkan gagasan bahwa di balik keragaman fenomena alam dan sosial, ada pola-pola umum yang dapat diidentifikasi dan dipelajari.

David Salsburg, dalam "The Lady Tasting Tea" (2001), menggambarkan bagaimana distribusi frekuensi dan konsep statistik lainnya mengubah cara kita melakukan penelitian ilmiah. Ia menunjukkan bagaimana konsep ini memungkinkan kita untuk mengekstrak makna dari data yang tampaknya kacau, mencerminkan kemampuan manusia untuk menemukan ketertiban dalam ketidakteraturan.

Alain Desrosières, dalam "The Politics of Large Numbers" (1998), membawa kita lebih jauh dengan mengeksplorasi bagaimana distribusi frekuensi tidak hanya alat ilmiah, tetapi juga instrumen sosial dan politik. Ia menunjukkan bagaimana cara kita mengategorikan dan mengukur fenomena melalui distribusi frekuensi dapat mempengaruhi kebijakan publik dan pemahaman sosial.

Karl Popper, meskipun tidak secara khusus membahas distribusi frekuensi dalam "The Logic of Scientific Discovery" (1959), memberikan landasan filosofis untuk memahami konsep ini dalam konteks yang lebih luas dari metode ilmiah. Gagasannya tentang falsifiabilitas dan sifat sementara pengetahuan ilmiah sangat relevan dengan cara kita menafsirkan dan menggunakan distribusi frekuensi.

Judea Pearl, dalam "Causality" (2000), membawa diskusi ini ke tingkat yang lebih dalam dengan mengeksplorasi hubungan antara distribusi frekuensi, korelasi, dan kausalitas. Karyanya mengingatkan kita akan batasan-batasan interpretasi distribusi frekuensi dan pentingnya berhati-hati dalam menarik kesimpulan kausal dari pola statistik.

R.A. Fisher, salah satu pionir statistika modern, dalam "Statistical Methods for Research Workers" (1925), meletakkan dasar-dasar teknis untuk pemahaman kita tentang distribusi frekuensi. Karyanya menunjukkan bagaimana konsep matematis yang abstrak dapat diterapkan untuk memecahkan masalah praktis dalam berbagai bidang ilmu.

Melalui lensa para pemikir ini, kita melihat bahwa distribusi frekuensi bukan hanya alat statistik, tetapi juga

cermin dari upaya manusia untuk memahami dunia. Ia mencerminkan ketegangan antara keinginan kita akan kepastian dan pengakuan akan kompleksitas realitas. Distribusi frekuensi mengingatkan kita bahwa pengetahuan kita selalu tidak sempurna dan terus berkembang, namun melalui analisis yang cermat dan interpretasi yang bijaksana, kita dapat terus memperdalam pemahaman kita tentang dunia di sekitar kita.

Demikianlah, distribusi frekuensi berdiri sebagai monumen intelektual yang menggabungkan matematika, filosofi, dan pemahaman manusia tentang alam semesta - sebuah jembatan antara yang diketahui dan yang tidak diketahui, antara kepastian dan ketidakpastian, yang terus membentuk cara kita memandang dan berinteraksi dengan dunia.

1. Tabel Frekuensi

Misalkan seorang guru ingin menganalisis distribusi nilai ujian matematika dari 30 siswa di kelasnya. Berikut adalah nilai-nilai yang diperoleh (dalam skala 0-100):

75, 82, 68, 90, 78, 85, 72, 88, 79, 83,
76, 81, 69, 92, 77, 86, 71, 87, 80, 84,
73, 89, 70, 91, 74, 88, 67, 93, 82, 85

Langkah-langkah untuk membuat tabel frekuensi:

- 1) Tentukan rentang nilai (nilai tertinggi - nilai terendah)
- 2) Tentukan jumlah kelas (biasanya menggunakan aturan Sturges: $1 + 3.3 \log n$, di mana n adalah jumlah data)
- 3) Hitung lebar interval kelas (rentang nilai / jumlah kelas)
- 4) Buat tabel dengan kolom untuk interval kelas, frekuensi, dan frekuensi relatif

Penyelesaian :

- 1) Rentang nilai: $93 - 67 = 26$
- 2) Jumlah kelas: $1 + 3.3 \log 30 \approx 5.87$ (dibulatkan menjadi 6)
- 3) Lebar interval kelas: $26 / 6 \approx 4.33$ (dibulatkan menjadi 5 untuk kemudahan)

Tabel 2.1. Tabel Frekuensinya Ujian Matematika

Interval Nilai	Frekuensi	Frekuensi Relatif
65-69	2	6,67 %
70-74	5	16,67 %
75-79	6	20,00 %
80-84	7	23,33 %
85-90	6	20.00 %
91-94	4	13,33 %
Total	30	100 %

Penjelasan tabel:

- 1) Interval Nilai: Rentang nilai untuk setiap kelas.
- 2) Frekuensi: Jumlah siswa yang mendapat nilai dalam interval tersebut.
- 3) Frekuensi Relatif: Persentase siswa dalam setiap interval (frekuensi dibagi total siswa, dikali 100%).

Interpretasi tabel:

- 1) Mayoritas siswa (23.33%) mendapat nilai antara 80-84.
- 2) Distribusi nilai cenderung normal, dengan lebih sedikit siswa di ujung-ujung distribusi.
- 3) Hanya sedikit siswa (6.67%) yang mendapat nilai di bawah 70.
- 4) 13.33% siswa mendapat nilai di atas 90, menunjukkan performa yang sangat baik.

Tabel frekuensi ini memberikan gambaran yang jelas tentang distribusi nilai ujian di kelas tersebut. Guru dapat menggunakan informasi ini untuk:

- 1) Menilai efektivitas pengajaran
- 2) Mengidentifikasi area yang mungkin memerlukan perhatian khusus
- 3) Merencanakan intervensi untuk siswa yang berada di ujung bawah distribusi

- 4) Memberikan penghargaan atau tantangan tambahan untuk siswa di ujung atas distribusi

Tabel frekuensi adalah langkah awal dalam analisis statistik deskriptif. Dari sini, guru bisa melanjutkan dengan membuat visualisasi seperti histogram atau diagram batang untuk representasi grafis dari distribusi nilai.

Contoh 2.

Hasil pencatatan persentase kehadiran 200 murid di sekolah Inti mandiri data dipaparkan sebagai berikut :

Data Kehadiran Siswa SMA Inti Mandiri (200 siswa)

95, 98, 92, 100, 88, 97, 93, 99, 91, 96, 94, 97, 90, 100, 89, 98, 92, 99, 93, 95, 96, 94, 97, 91, 99, 93, 95, 98, 92, 96, 94, 97, 90, 98, 92, 95, 93, 99, 91, 96, 97, 93, 98, 92, 100, 89, 96, 94, 99, 91, 95, 97, 93, 98, 90, 96, 94, 97, 92, 95, 98, 92, 97, 93, 99, 91, 95, 98, 94, 96, 90, 97, 93, 98, 92, 95, 94, 99, 91, 96, 95, 98, 92, 97, 94, 99, 90, 96, 93, 98, 91, 95, 97, 94, 99, 92, 96, 93, 98, 91, 97, 93, 98, 92, 95, 94, 99, 90, 96, 93, 97, 91, 98, 94, 95, 92, 99, 93, 96, 91, 98, 94, 97, 92, 95, 93, 99, 91, 96, 94, 97, 90, 98, 93, 95, 92, 99, 91, 96, 94, 95, 97, 93, 98, 92, 96, 94, 99, 90, 97, 93, 95, 91, 98, 94, 96, 92, 97, 93, 99, 96, 98, 92, 97, 94, 95, 93, 99, 91, 96, 98, 94, 97, 92, 95, 93, 99, 90, 96, 94, 97, 93, 98, 91, 95, 94, 99, 92, 96, 93, 97, 90, 98, 94, 95, 92, 99, 91, 96, 93

Perhitungan untuk Tabel Frekuensi

- 1) Rentang data: $100 - 88 = 12$ hari
- 2) Jumlah kelas: Menggunakan aturan Sturges: $1 + 3.3 \log 200 \approx 8.59$ (dibulatkan menjadi 8)
- 3) Lebar interval kelas: $12 / 8 = 1.5$ (dibulatkan menjadi 2 untuk kemudahan)

Tabel 2.2. Frekuensi Kehadiran Siswa (hari)

Jumlah Kehadiran	Frekuensi	Frekuensi Relatif
87 - 88	1	0,5 %
89 - 90	12	6,0 %
91 - 92	32	16,00 %
93 - 94	46	23,0 %
95 - 96	49	24,5 %
97 - 98	41	20,5 %

Jumlah Kehadiran	Frekuensi	Frekuensi Relatif
99 – 100	19	9,5 %
Total	200	100 %

Penjelasan :

- 1) Distribusi Keseluruhan:
 - a. Mayoritas siswa (24.5%) hadir antara 95-96 hari dari 100 hari sekolah.
 - b. Kehadiran secara umum sangat baik, dengan 77.5% siswa hadir 93 hari atau lebih.
- 2) Kehadiran Tinggi:
 - a) 9.5% siswa memiliki kehadiran sangat tinggi (99-100 hari).
 - b) 54.5% siswa hadir 95 hari atau lebih, menunjukkan komitmen yang kuat terhadap sekolah.
- 3) Kehadiran Rendah:
 - a) Hanya 0.5% siswa (1 siswa) yang hadir kurang dari 89 hari.
 - b) 6.5% siswa hadir kurang dari 91 hari, yang mungkin memerlukan perhatian khusus.
- 4) Implikasi untuk Kebijakan Sekolah:
 - a) Sekolah dapat menyelidiki faktor-faktor yang berkontribusi pada kehadiran tinggi untuk 30% siswa yang hadir 97 hari atau lebih.
 - b) Perlu ada strategi intervensi untuk 6.5% siswa dengan kehadiran terendah (kurang dari 91 hari).
- 5) Tren Kehadiran:
 - a) Distribusi cenderung sedikit miring ke kiri (*slight negative skew*), menunjukkan bahwa lebih banyak siswa memiliki kehadiran tinggi.
 - b) Ini menunjukkan budaya kehadiran yang sangat positif di sekolah.

Penggunaan Analisis dalam Konteks Pendidikan:

- 1) Identifikasi Pola: Sekolah dapat mengidentifikasi pola kehadiran dan menyelidiki faktor-faktor yang mempengaruhi kehadiran tinggi dan rendah.

- 2) Intervensi Terarah: Fokus pada 6.5% siswa dengan kehadiran terendah untuk meningkatkan kehadiran mereka.
- 3) Penghargaan dan Motivasi: Memberi penghargaan kepada 9.5% siswa dengan kehadiran sangat tinggi sebagai contoh positif.
- 4) Perencanaan Program: Merancang program untuk mempertahankan tingkat kehadiran yang tinggi dan mendorong peningkatan bagi mereka yang di bawah rata-rata.
- 5) Benchmarking: Menggunakan data ini sebagai benchmark untuk semester-semester mendatang atau perbandingan dengan sekolah lain.
- 6) Analisis Korelasi: Menyelidiki korelasi antara tingkat kehadiran dan prestasi akademik untuk memperkuat pentingnya kehadiran reguler.
- 7) Komunikasi dengan Pemangku Kepentingan: Menggunakan data ini untuk menginformasikan orang tua, dewan sekolah, dan pembuat kebijakan tentang keberhasilan kebijakan kehadiran saat ini dan area yang memerlukan perbaikan.

Data yang lebih lengkap ini memberikan gambaran yang lebih akurat tentang kehadiran siswa di sekolah, memungkinkan untuk analisis yang lebih mendalam dan pengambilan keputusan yang lebih tepat sasaran.

2. Histogram

Histogram adalah representasi grafis dari distribusi data yang diatur ke dalam interval-interval atau "bins." Histogram terdiri dari batang (bars) di mana tinggi setiap batang menggambarkan jumlah data yang berada dalam suatu interval tertentu.

Histogram didasarkan pada filosofi visualisasi data, yang bertujuan untuk menyederhanakan kompleksitas data numerik sehingga pola dan tren dapat dilihat dengan lebih jelas. Dengan mengelompokkan data ke dalam interval, histogram memberikan gambaran cepat tentang bagaimana

data tersebut terdistribusi, apakah normal, miring, atau memiliki outlier.

Kegunaan

- 1) Memvisualisasikan Distribusi Data:** Histogram digunakan untuk menunjukkan bagaimana data terdistribusi, membantu dalam mengidentifikasi pola distribusi seperti normalitas atau kecondongan (skewness) (Doane, D. P., & Seward, L. E., 2011).
- 2) Identifikasi Outlier:** Histogram memudahkan dalam mengidentifikasi nilai-nilai ekstrem atau outlier yang mungkin mempengaruhi analisis lebih lanjut (Chambers, J. M., Cleveland, W. S., Kleiner, B., & Tukey, P. A., 1983).
- 3) Pengambilan Keputusan:** Histogram dapat digunakan dalam berbagai bidang, seperti kualitas manufaktur atau penelitian ilmiah, untuk membantu pengambilan keputusan berdasarkan distribusi data (Montgomery, D. C., 2017).
- 4) Analisis Kualitas:** Dalam manajemen kualitas, histogram sering digunakan untuk menganalisis proses dan mengidentifikasi area yang memerlukan perbaikan (Ishikawa, K., 1985).

Kelebihan

- 1) Mudah Dipahami:** Histogram merupakan alat visual yang sederhana dan mudah dipahami oleh khalayak luas, termasuk mereka yang tidak memiliki latar belakang statistik (Everitt, B. S., & Skrondal, A., 2010).
- 2) Memberikan Gambaran Keseluruhan:** Histogram memberikan gambaran komprehensif tentang distribusi data, termasuk karakteristik pusat dan penyebaran (Scott, D. W., 2015).
- 3) Mengidentifikasi Pola:** Histogram efektif dalam mengidentifikasi pola seperti simetri, kecondongan, dan adanya modus dalam data (Cleveland, W. S., 1993).
- 4) Fleksibilitas:** Histogram dapat diterapkan pada berbagai jenis data numerik, baik kontinu maupun diskrit (Freedman, D., Pisani, R., & Purves, R., 2007).

Kekurangan

- 1) **Kehilangan Detail Data:** Ketika data dikelompokkan ke dalam interval, informasi detail dari data asli mungkin hilang, mengurangi resolusi analisis (Silverman, B. W., 1986).
- 2) **Sensitif terhadap Pemilihan Bin:** Hasil histogram sangat dipengaruhi oleh lebar dan jumlah bin yang dipilih, yang dapat menyebabkan distorsi informasi jika tidak dipilih dengan tepat (Wand, M. P., & Jones, M. C., 1995).
- 3) **Tidak Efektif untuk Data Kecil:** Pada dataset yang kecil, histogram mungkin tidak memberikan gambaran yang akurat, karena kurangnya jumlah data yang cukup untuk mengisi setiap bin (Tukey, J. W., 1977).
- 4) **Tidak Efektif untuk Data Kategorikal:** Histogram lebih cocok untuk data numerik; data kategorikal sebaiknya divisualisasikan menggunakan bar chart (Agresti, A., 2018).

Contoh: Distribusi Nilai Ujian

Misalkan sebuah kelas terdiri dari 50 siswa yang telah mengikuti ujian matematika, dan guru ingin melihat bagaimana distribusi nilai siswa tersebut. Nilai ujian berkisar dari 0 hingga 100, dan guru memutuskan untuk mengelompokkan nilai-nilai tersebut ke dalam interval atau bin dengan rentang 10 poin.

Langkah-Langkah:

- 1) **Kumpulkan Data:** Misalnya, berikut adalah nilai-nilai ujian dari 50 siswa:
 - 95, 82, 67, 88, 75, 62, 78, 90, 84, 70, 92, 61, 85, 73, 76, 81, 55, 69, 80, 77, 74, 93, 86, 58, 79, 91, 64, 89, 83, 87, 66, 60, 82, 72, 95, 68, 71, 59, 90, 88, 85, 63, 87, 92, 60, 65, 71, 89, 56, 81.
- 2) **Tentukan Bin:** Tentukan bin untuk mengelompokkan nilai. Misalnya:
 - Bin 1: 50-59
 - Bin 2: 60-69
 - Bin 3: 70-79

- Bin 4: 80-89
- Bin 5: 90-100

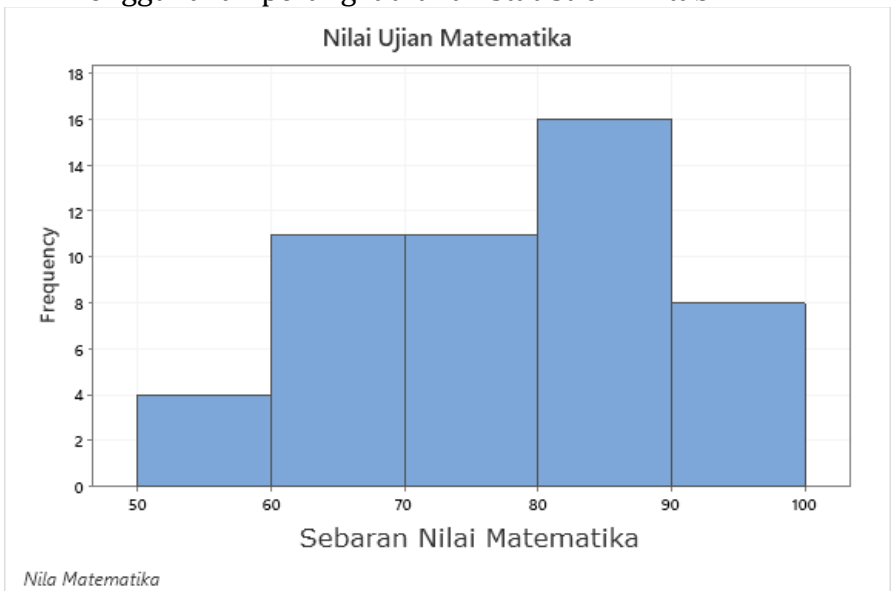
3) Hitung Frekuensi: Hitung berapa banyak nilai yang masuk ke dalam setiap bin:

- Bin 1 (50-59): 5 siswa
- Bin 2 (60-69): 11 siswa
- Bin 3 (70-79): 12 siswa
- Bin 4 (80-89): 14 siswa
- Bin 5 (90-100): 8 siswa

4) Buat Histogram:

- Pada sumbu horizontal (x-axis), letakkan interval bin (50-59, 60-69, dst.).
- Pada sumbu vertikal (y-axis), letakkan frekuensi (jumlah siswa).
- Gambarlah batang (bars) untuk setiap bin dengan tinggi yang sesuai dengan frekuensi yang dihitung.

Hasilnya seperti ditampilkan pada gambar 2.1, dengan menggunakan perangkat lunak statistic Minitab.



Gambar 2.1. Histogram Nilai Matematika

Interpretasi Histogram:

- 1) **Simetri:** Jika histogram simetris, ini berarti nilai ujian tersebar merata di sekitar rata-rata, menunjukkan distribusi yang seimbang.
- 2) **Kecondongan (*Skewness*):** Jika histogram miring ke kanan (*positif skewness*), lebih banyak siswa yang mendapat nilai rendah. Sebaliknya, jika miring ke kiri (*negatif skewness*), lebih banyak siswa yang mendapat nilai tinggi.
- 3) **Modus:** Bin dengan frekuensi tertinggi menunjukkan modus, atau nilai yang paling sering muncul. Dalam contoh ini, bin 4 (80-89) adalah yang paling banyak, menunjukkan bahwa sebagian besar siswa mendapat nilai di antara 80 dan 89.

Penggunaan Lain Histogram dalam Pendidikan:

- 1) **Kehadiran Siswa:** Visualisasi data kehadiran untuk melihat pola absensi dalam kelas tertentu.
- 2) **Evaluasi Pengajaran:** Menggunakan histogram untuk memvisualisasikan hasil survei kepuasan siswa terhadap pengajaran, dengan berbagai interval nilai seperti sangat tidak puas hingga sangat puas.
- 3) **Distribusi Waktu Belajar:** Menggunakan histogram untuk melihat berapa banyak waktu yang dihabiskan siswa untuk belajar, dengan interval waktu seperti 0-1 jam, 2-3 jam, dan seterusnya.

Histogram seperti ini sangat berguna bagi guru dan administrator sekolah untuk mendapatkan wawasan tentang kinerja siswa, menilai efektivitas pengajaran, dan membuat keputusan yang berbasis data untuk meningkatkan hasil pendidikan.

Diagram Batang (*Bar Chart*)

Diagram Batang adalah representasi grafis yang menggunakan batang untuk menunjukkan nilai kuantitatif. Setiap batang mewakili kategori data, dan tinggi atau panjang batang sesuai dengan nilai yang diwakili. Diagram

batang sering digunakan untuk membandingkan jumlah atau frekuensi data di antara beberapa kategori.

Filosofi di balik diagram batang adalah mempermudah perbandingan visual antara kategori yang berbeda. Dengan mewakili data secara grafis, diagram batang memungkinkan pengguna melihat dengan cepat kategori mana yang memiliki nilai tertinggi, terendah, atau bagaimana distribusi data di antara berbagai kategori (Cleveland, W. S., 1994).

Kegunaan

- 1) Perbandingan Kategori:** Diagram batang sangat berguna untuk membandingkan nilai-nilai di antara beberapa kategori, seperti jumlah penjualan per produk, kehadiran siswa per kelas, atau frekuensi tanggapan dalam survei (Kelleher, C., & Wagener, T., 2011).
- 2) Visualisasi Data Kategorikal:** Diagram batang efektif untuk data yang bersifat kategorikal (nominal atau ordinal), di mana kategori tidak memiliki hubungan numerik yang intrinsik (Everitt, B. S., 2002).
- 3) Pengambilan Keputusan:** Diagram batang membantu dalam analisis data yang diperlukan untuk pengambilan keputusan di berbagai bidang seperti bisnis, pendidikan, dan penelitian (Few, S., 2006).
- 4) Monitoring Perubahan:** Digunakan untuk memantau perubahan data dari waktu ke waktu, terutama jika data dikumpulkan secara periodik (Tuft, E. R., 1983).

Kelebihan

- 1) Mudah Dipahami:** Diagram batang mudah dibaca dan dipahami oleh semua orang, termasuk yang tidak memiliki latar belakang teknis (Chambers, J. M., Cleveland, W. S., Kleiner, B., & Tukey, P. A., 1983).
- 2) Fleksibilitas:** Dapat digunakan untuk data diskrit, baik kategorikal maupun numerik, dan dapat menampilkan data dalam bentuk horizontal atau vertikal (Wainer, H., 1997).

- 3) **Perbandingan Jelas:** Membuat perbandingan antar kategori menjadi sangat jelas dan mudah dikenali (Cleveland, W. S., 1994).
- 4) **Tampilan Beragam:** Diagram batang dapat diadaptasi dalam berbagai format seperti diagram batang berkelompok (grouped bar chart) atau diagram batang bertumpuk (stacked bar chart) (Kelleher, C., & Wagener, T., 2011).

Kekurangan

- 1) **Keterbatasan pada Data Kontinu:** Diagram batang tidak ideal untuk data kontinu, karena lebih cocok untuk data yang dikelompokkan ke dalam kategori yang berbeda (Tuft, E. R., 1983).
- 2) **Kelemahan dalam Visualisasi Detail:** Ketika jumlah kategori atau batang terlalu banyak, diagram batang bisa menjadi sulit dibaca dan membingungkan (Few, S., 2006).
- 3) **Sensitif terhadap Skala:** Skala pada sumbu y (atau x) sangat penting. Jika skala tidak dipilih dengan hati-hati, perbandingan antar kategori bisa menjadi menyesatkan (Wainer, H., 1997).
- 4) **Ruang yang Dibutuhkan:** Diagram batang dengan banyak kategori memerlukan ruang yang lebih luas, terutama jika menggunakan batang horizontal (Kelleher, C., & Wagener, T., 2011).

Contoh :

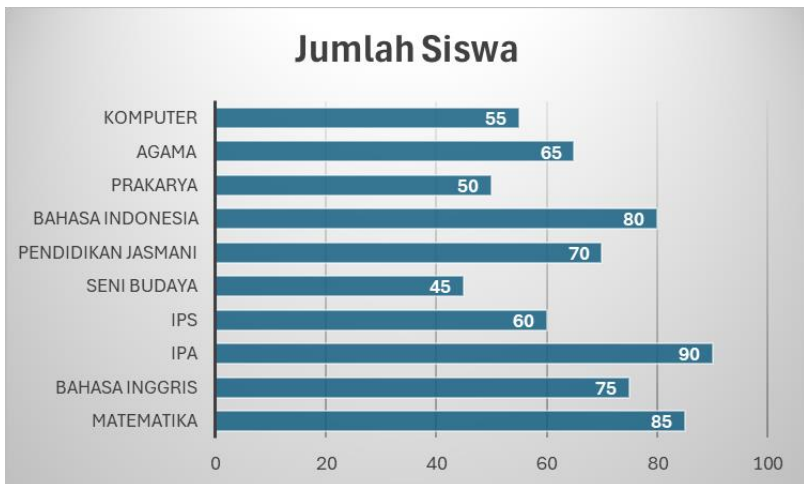
Jumlah 10 Mata Pelajaran yang diikuti dari Siswa Sekolah Mengengah Atas Negeri 1 tahun 2023. ditampilkan pada tabel 2.3.

Tabel 2.3. Jumlah Mata Pelajaran yang Ditawar Siswa

Mata Pelajaran	Jumlah Siswa
Matematika	85
Bahasa Inggris	75
IPA	90
IPS	60

Mata Pelajaran	Jumlah Siswa
Seni Budaya	45
Pendidikan Jasmani	70
Bahasa Indonesia	80
Prakarya	50
Agama	65
Komputer	55

Hasil dari 10 mata pelajaran ditampilkan dalam diagram batang pada gambar 2.2.



Gambar 2.2. Diagram Batang 10 Mata Pelajaran

3. Diagram Lingkaran dalam Statistik Deskriptif

Diagram lingkaran, atau yang sering disebut sebagai diagram pie, adalah representasi grafis dari data dalam bentuk lingkaran yang dibagi menjadi beberapa bagian sesuai dengan proporsi masing-masing kategori atau nilai. Setiap bagian mewakili satu kategori data, dan luasnya sebanding dengan persentase dari total keseluruhan.

Kelebihan Diagram Lingkaran

1) Visualisasi yang Jelas:

Diagram lingkaran memberikan representasi visual yang langsung terlihat mengenai proporsi masing-masing kategori data. Hal ini memudahkan pembaca untuk memahami perbandingan antar kategori secara sekilas (Cleveland, dan McGill, 1984).

2) Mudah Dipahami:

Diagram lingkaran sangat populer dan mudah dipahami oleh banyak orang, bahkan mereka yang tidak memiliki latar belakang statistik. (Everitt, 1998)."

3) Penggunaan yang Familiar:

Diagram ini sering digunakan dalam laporan, presentasi, dan media massa, membuatnya mudah dikenali dan dipahami oleh audiens yang luas (Kosslyn, 2006).

Kekurangan Diagram Lingkaran

1) Kesulitan dalam Membandingkan Bagian:

Diagram lingkaran kurang efektif untuk membandingkan bagian yang memiliki ukuran yang sangat mirip atau ketika jumlah bagian banyak, karena perbedaan ukuran bisa menjadi sulit dilihat (Tufte, 1983).

2) Informasi yang Terbatas:

Diagram lingkaran hanya menunjukkan distribusi data untuk satu variabel kategoris pada satu waktu. Jika ada lebih dari satu variabel atau ingin menampilkan perubahan data dari waktu ke waktu, diagram ini tidak cukup informatif (Few, 2007).

3) Tidak Cocok untuk Data dengan Banyak Kategori:

Jika diagram lingkaran memiliki terlalu banyak bagian, diagram ini menjadi berantakan dan sulit dibaca, membuat informasi yang disampaikan kurang efektif (Ware, 2012).

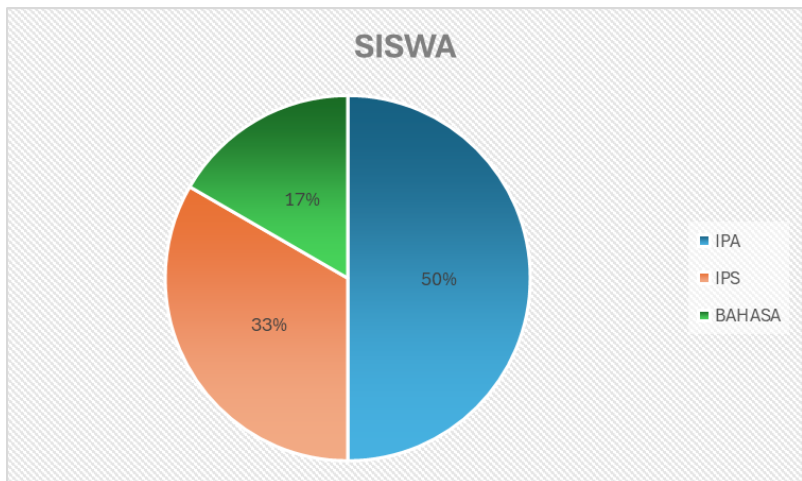
Diagram lingkaran memang memiliki tempat dalam statistik deskriptif, terutama ketika kita ingin menunjukkan perbandingan proporsi dari beberapa kategori. Namun,

perlu berhati-hati dalam penggunaannya, terutama dalam situasi di mana perbandingan yang lebih rinci diperlukan.

Contoh Kasus

Misalnya, sebuah SMA memiliki tiga jurusan: IPA (Ilmu Pengetahuan Alam), IPS (Ilmu Pengetahuan Sosial), dan Bahasa. Jumlah total siswa di sekolah tersebut adalah 600 siswa. Data mengenai jumlah siswa di masing-masing jurusan adalah sebagai berikut:

- 1) **IPA:** 300 siswa (50% dari total siswa)
- 2) **IPS:** 200 siswa (33.33% dari total siswa)
- 3) **Bahasa:** 100 siswa (16.67% dari total siswa)



Gambar 2.3. Diagram Batang Jumlah Siswa 3 Jurusan

Penerapan Diagram Lingkaran:

Diagram lingkaran dapat digunakan untuk menggambarkan distribusi persentase siswa di setiap jurusan. Dengan menggunakan diagram lingkaran, bagian-bagian lingkaran akan dibagi sesuai dengan proporsi jumlah siswa di masing-masing jurusan.

- 1) **Bagian pertama:** Mewakili jurusan IPA dengan 50% dari lingkaran.
- 2) **Bagian kedua:** Mewakili jurusan IPS dengan 33.33% dari lingkaran.

- 3) **Bagian ketiga:** Mewakili jurusan Bahasa dengan 16.67% dari lingkaran.

Manfaat Penggunaan Diagram Lingkaran:

- 1) **Visualisasi Proporsi:** Diagram ini membantu pihak sekolah, orang tua, dan siswa untuk dengan mudah memahami sebaran siswa di masing-masing jurusan tanpa harus melihat angka-angka secara langsung.
- 2) **Pengambilan Keputusan:** Dengan melihat diagram lingkaran, pihak manajemen sekolah bisa mengambil keputusan yang lebih tepat, misalnya dalam hal alokasi sumber daya, penentuan jumlah guru per jurusan, dan perencanaan kegiatan belajar mengajar.

Contoh 2

Sebuah universitas memiliki beberapa fakultas dengan berbagai program studi. Berikut adalah data jumlah mahasiswa berdasarkan fakultas dan program studi untuk semester ini:

1) Fakultas Ilmu Komputer

- a) Teknik Informatika: 300 mahasiswa (15% dari total mahasiswa)
- b) Sistem Informasi: 200 mahasiswa (10% dari total mahasiswa)

2) Fakultas Ekonomi:

- a) Manajemen: 250 mahasiswa (12.5% dari total mahasiswa)
- b) Akuntansi: 200 mahasiswa (10% dari total mahasiswa)

3) Fakultas Teknik:

- a) Teknik Sipil: 150 mahasiswa (7.5% dari total mahasiswa)
- b) Teknik Elektro: 100 mahasiswa (5% dari total mahasiswa)

4) Fakultas Sastra:

- a) Sastra Inggris: 200 mahasiswa (10% dari total mahasiswa)

- b) Sastra Jepang: 150 mahasiswa (7.5% dari total mahasiswa)

5) Fakultas Ilmu Kesehatan:

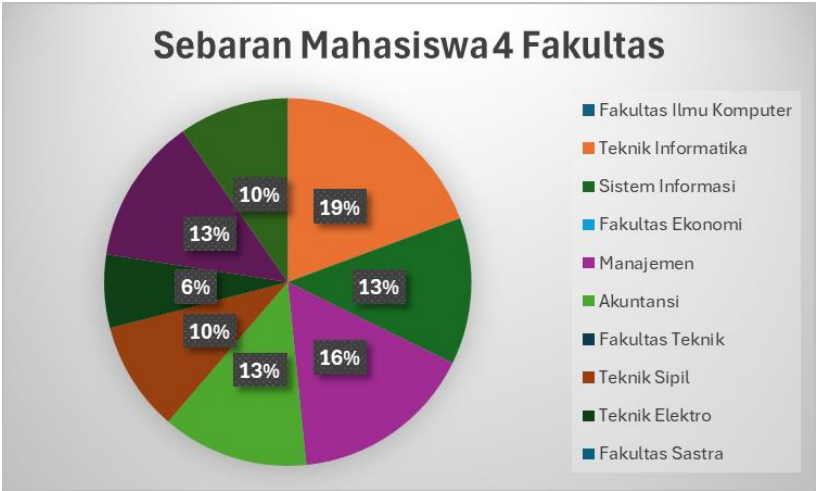
- a) Keperawatan: 250 mahasiswa (12.5% dari total mahasiswa)
- b) Farmasi: 200 mahasiswa (10% dari total mahasiswa)

Penerapan Diagram Lingkaran

Dalam hal ini, kita dapat menggunakan beberapa diagram lingkaran untuk menggambarkan distribusi data:

- 1. Diagram Lingkaran 1:** Menampilkan persentase mahasiswa di setiap fakultas.
 - a. **Fakultas Ilmu Komputer:** 25%
 - b. **Fakultas Ekonomi:** 22.5%
 - c. **Fakultas Teknik:** 12.5%
 - d. **Fakultas Sastra:** 17.5%
 - e. **Fakultas Ilmu Kesehatan:** 22.5%
- 2. Diagram Lingkaran 2:** Menampilkan persentase mahasiswa berdasarkan program studi di Fakultas Ilmu Komputer.
 - a. **Teknik Informatika:** 60%
 - b. **Sistem Informasi:** 40%
- 3. Diagram Lingkaran 3:** Menampilkan persentase mahasiswa berdasarkan program studi di Fakultas Ekonomi.
 - a. **Manajemen:** 55.6%
 - b. **Akuntansi:** 44.4%
- 4. Diagram Lingkaran 4:** Menampilkan persentase mahasiswa berdasarkan program studi di Fakultas Teknik.
 - a. **Teknik Sipil:** 60%
 - b. **Teknik Elektro:** 40%
- 5. Diagram Lingkaran 5:** Menampilkan persentase mahasiswa berdasarkan program studi di Fakultas Sastra.
 - a. **Sastra Inggris:** 57.1%
 - b. **Sastra Jepang:** 42.9%
- 6. Diagram Lingkaran 6:** Menampilkan persentase mahasiswa berdasarkan program studi di Fakultas Ilmu Kesehatan.

- a. Keperawatan: 55.6%
- b. Farmasi: 44.4%



Gambar 2.4. Diagram Lingkaran Jumlah Masiswa dalam empat Fakultas

Manfaat Penggunaan Diagram Lingkaran

- Analisis Perbandingan:** Dengan beberapa diagram lingkaran, pihak universitas bisa membandingkan proporsi mahasiswa antar fakultas dan program studi secara visual. Misalnya, diagram 1 memberikan gambaran umum mengenai sebaran mahasiswa di seluruh fakultas, sementara diagram 2 hingga 6 memberikan detail mengenai distribusi mahasiswa di setiap program studi.
- Pengambilan Keputusan:** Informasi ini bisa digunakan oleh pihak universitas untuk menentukan alokasi anggaran, merencanakan jumlah dosen, atau melakukan strategi pemasaran untuk program studi tertentu yang kurang diminati.

Penggunaan dalam Presentasi

Dalam sebuah rapat senat universitas, rektor bisa menggunakan rangkaian diagram lingkaran ini untuk mempresentasikan distribusi mahasiswa kepada dekan-dekan

fakultas. Dengan visualisasi yang jelas, setiap dekan bisa memahami posisi fakultas mereka dan program studi yang memerlukan perhatian lebih lanjut.

Diagram lingkaran dalam konteks ini memberikan berbagai sudut pandang mengenai data mahasiswa, mulai dari distribusi antar fakultas hingga ke detail per program studi, memudahkan pengambilan keputusan yang lebih berbasis data.

DAFTAR PUSTAKA

- Abulela, M. A. A., & Harwell, M. (2020). Data Analysis: Strengthening Inferences in Quantitative Education Studies Conducted by Novice Researchers. *Educational Sciences: Theory & Practice*, 20(1), 59–78.
<https://doi.org/10.12738/jestp.2020.1.005>
- Agresti, A. (2018). *Statistical methods for the social sciences*. Upper Saddle River Pearson.
- Agresti, A., Franklin, C. A., & Bernhard Klingenberg. (2018). *Statistics : the art and science of learning from data*. Pearson Education Limited.
- Agresti, A., Franklin, C., & Klingenberg, B. (2018). *Statistics : The Art and Science of Learning from Data, 4/E*. Open Library; Pearson.
<https://openlibrary.telkomuniversity.ac.id/pustaka/165852/statistics-the-art-and-science-of-learning-from-data-4-e.html>
- Akdemir, E., Karameşe, E. N., & Arslan, A. (2015a). Descriptive Analysis of Researches on Curriculum Development in Education. *Procedia - Social and Behavioral Sciences*, 174, 3199–3203.
<https://doi.org/10.1016/j.sbspro.2015.01.1062>
- Akdemir, E., Karameşe, E. N., & Arslan, A. (2015b). Descriptive Analysis of Researches on Curriculum Development in Education. *Procedia - Social and Behavioral Sciences*, 174, 3199–3203.
<https://doi.org/10.1016/j.sbspro.2015.01.1062>
- Alain Desrosières. (1998). *The politics of large numbers : a history of statistical reasoning*. Harvard University Press.
- Aylmer, R. (1925). *Statistical Methods for Research Workers*.
- Aylmer, R., Alfred, P., & Austen, C. (1966). *Commentary on R.A. Fisher's Paper on The Correlation Between Relatives on the Supposition of Mendelian Inheritance. Transactions of the Royal Society of Edinburgh*, 52, 1918 ... By P.A.P. Moran ... and C.A.B. Smith. [With the Text.]. Cambridge University Press.

- Baker, A. (2016). *Simplicity* (*Stanford Encyclopedia of Philosophy*). Stanford.edu.
<https://plato.stanford.edu/entries/simplicity/>
- Bandyopadhyay, P. S. (2011). *Philosophy of Statistics (Handbook of Philosophy of Science)*. North Holland.
- Bluman, A. G. (2008). *Bluman, Elementary Statistics: A Step by Step Approach*, © 2009, 7e, Student Edition (Reinforced Binding) with Formula Card. McGraw-Hill Education.
- Box, G. E. P. (1976). Science and Statistics. *Journal of the American Statistical Association*, 71(356), 791–799.
<https://doi.org/10.2307/2286841>
- Brantlinger, E., Jimenez, R., Klingner, J., Pugach, M., & Richardson, V. (2005). Qualitative Studies in Special Education. *Exceptional Children*, 71(2), 195–207.
<https://doi.org/10.1177/001440290507100205>
- Carl Friedrich Gauss. (1809). *Theoria motus corporum coelestium in sectionibus conicis solem ambientium*.
- Chambers, J. M. (1983). *Graphical Methods for Data Analysis*. Duxbury Resource Center.
- Cleveland, W. S. (1993). *Visualizing Data*. Hobart Press.
- Cleveland, W. S. (1994). *The Elements of Graphing Data*. Hobart Press.
- Cleveland, W. S., & McGill, R. (1984). Graphical Perception: Theory, Experimentation, and Application to the Development of Graphical Methods. *Journal of the American Statistical Association*, 79(387), 531–554.
<https://doi.org/10.2307/2288400>
- Cooksey, R. W. (2020). Descriptive Statistics for Summarising Data. *Illustrating Statistical Procedures: Finding Meaning in Quantitative Data*, 61–139. https://doi.org/10.1007/978-981-15-2537-7_5
- Cover, T. M., & Thomas, J. A. (2006). *Elements of information theory*. Wiley-Interscience.
- Dabić, T., & Stojanov, Ž. (2014). Techniques for Collecting Qualitative Field Data in Education Research: Example of Two Studies in Information Technology Field. *Proceedings of the 1st International Scientific Conference - Sinteza 2014*. <https://doi.org/10.15308/sinteza-2014-362-367>

- David , D. (2013). *David M. Lane*. Apple Books.
<https://books.apple.com/us/author/david-m-lane/id684002458>
- DeCarlo, L. T. (1997). On the meaning and use of kurtosis. *Psychological Methods*, 2(3), 292–307.
<https://doi.org/10.1037/1082-989x.2.3.292>
- Doane, D. P., & Seward, L. E. (2011). Measuring Skewness: A Forgotten Statistic? *Journal of Statistics Education*, 19(2).
<https://doi.org/10.1080/10691898.2011.11889611>
- Edgar, T. F., & Manz, D. O. (2017). Descriptive Study. *Elsevier EBooks*, 131–151. <https://doi.org/10.1016/b978-0-12-805349-2.00005-4>
- Efron, B., & Hastie, T. (2016). *Computer age statistical inference*. Cambridge University Press.
- Everitt, B. (1998). *The Cambridge Dictionary of Statistics*. Cambridge University Press.
- Everitt, B. (1999). *Chance Rules*. Springer Science & Business Media.
- Everitt, B. S. (2006). *A Handbook of Statistical Analyses Using R*. CRC Press.
- Everitt, B., & Skrondal, A. (2010). *The Cambridge dictionary of statistics*. Cambridge University Press.
- Few, S. (2006). *Information Dashboard Design*. Oreilly & Associates Incorporated.
- Field, A. P. (2013). *Discovering statistics using IBM SPSS statistics* (4th ed.). Sage Publications Ltd.
- Fisher, R. A. (1925). *Fisher, R.A. (1925) Statistical Methods for Research Workers. Oliver and Boyd, Edinburgh, Scotland. - References - Scientific Research Publishing. Scirp.org.*
<https://www.scirp.org/reference/referencespapers?referenceid=2056938>
- Freedman, D., Pisani, R., & Purves, R. (2007). *Statistics*. W.W. Norton & Company.
- Freedman, D., Pisani, R., & Purves, R. (2009). *Statistics*. Viva Books.
- Gabrijela Dimić, Dejan Rančić, O. Pronic Rancic, & Petar Spalević. (2019). *Descriptive Statistical Analysis in the Process of Educational Data Mining*.

- <https://doi.org/10.1109/telsiks46999.2019.9002177>
- Gelman, A., & Hennig, C. (2017). Beyond subjective and objective in statistics. *Journal of the Royal Statistical Society: Series a (Statistics in Society)*, 180(4), 967–1033. <https://doi.org/10.1111/rssa.12276>
- Gerber, R., Boulton-Lewis, G., & Bruce, C. (1995). Children's understanding of graphic representations of quantitative data. *Learning and Instruction*, 5(1), 77–100. [https://doi.org/10.1016/0959-4752\(95\)00001-j](https://doi.org/10.1016/0959-4752(95)00001-j)
- GILLETTE, M. (1984). Applications of Descriptive Analysis. *Journal of Food Protection*, 47(5), 403–409. <https://doi.org/10.4315/0362-028x-47.5.403>
- Grami, A. (2019). Descriptive Statistics. *Probability, Random Variables, Statistics, and Random Processes: Fundamentals & Applications*, 239–258. <https://doi.org/10.1002/9781119300847.ch8>
- Gravetter, F. J., & Wallnau, L. B. (2017). *Statistics for the behavioral sciences*. Cengage.
- Gunter, B., & Wainer, H. (1998). Visual Revelations: Graphical Tales of Fate and Deception from Napoleon Bonaparte to Ross Perot. *The American Statistician*, 52(1), 83. <https://doi.org/10.2307/2685574>
- Hacking, I. (1990). *The taming of chance*. Cambridge University Press.
- Hogg, R. V., Mckean, J. W., & Craig, A. T. (2020). *Introduction to mathematical statistics*. Pearson.
- Howell, D. C. (2012). *Statistical Methods for Psychology*. Cengage Learning.
- Jabr, F. (2021a). John A. Long - Publications List. *Publicationslist.org*, 14(6). <http://publicationslist.org/jlong>
- Jabr, F. (2021b). John A. Long - Publications List. *Publicationslist.org*, 14(6). <http://publicationslist.org/jlong>
- Jaynes, E. T. (2003). *Probability Theory*. Cambridge University Press.
- John Wilder Tukey. (1977). *Exploratory Data Analysis*. Addison-Wesley Publishing Company.

- Johnson, R. A., & Bhattacharyya, G. K. (2019). *Statistics : principles and methods*. Wiley.
- Karl Raimund Popper. (1963). *Conjectures and Refutations*.
- Kaur, P., Stoltzfus, J., & Yellapu, V. (2018). Descriptive Statistics. *International Journal of Academic Medicine*, 4(1), 60–63. ResearchGate. https://doi.org/10.4103/ijam.ijam_7_18
- Kelleher, C., & Wagener, T. (2011). Ten guidelines for effective data visualization in scientific publications. *Environmental Modelling & Software*, 26(6), 822–827. <https://doi.org/10.1016/j.envsoft.2010.12.006>
- Kendall, M. G., & Stuart, A. (1969). *Kendall, M.G. and Stuart, A. (1969) The Advanced Theory of Statistics, Vol. 1, Chapter 3. Hafner Publishing Company, New York. - References - Scientific Research Publishing. Scirp.org. https://www.scirp.org/reference/referencespapers?referenceid=3253483*
- Kolmogorov, A. N., Bharucha-Reid, A. T., & Morrison, N. (2018). *Foundations of the theory of probability*. Dover Publications, Inc.
- Kosslyn, S. M. (2006). *Graph Design for the Eye and Mind*. Oxford University Press, USA.
- Larson, R., & Farber, B. (2014). *e Book Instant Access for Elementary Statistics: Picturing the World, Global Edition*. Pearson Higher Ed.
- LeCompte, M. D. (2000). Analyzing Qualitative Data. *Theory into Practice*, 39(3), 146–154. https://doi.org/10.1207/s15430421tip3903_5
- Libarkin, J. C., & Kurdziel, J. P. (2018). Research Methodologies in Science Education: The Qualitative-Quantitative Debate. *Journal of Geoscience Education*, 50(1), 78–86. <https://doi.org/10.1080/10899995.2002.12028053>
- Loyer, M., & Triola, M. F. (2006). *Student's Solutions Manual for Essentials of Statistics*. Addison Wesley Longman.
- Mackenzie, D. A. (1981). *Statistics in Britain, 1865-1930 : the social construction of scientific knowledge*. Edinburgh University Press.
- Mat Roni, S., Merga, M. K., & Morris, J. E. (2019). Data Types and Samples. *Conducting Quantitative Research in Education*,

- 39–45. https://doi.org/10.1007/978-981-13-9132-3_4
- Mayo, D. G. (2018). *Statistical inference as severe testing : how to get beyond the statistics wars*. Cambridge University Press.
- Montgomery, D. C. (2020). *Introduction To Statistical Quality Control*. Wiley.
- Montgomery, D. C., & Runger, G. C. (2011). *Applied statistics and probability for engineers*. Wiley.
- Moore, D. S., & McCabe, G. P. (2005). *Moore, D. S., & McCabe, G. P. (2005). Introduction to the Practice of Statistics (5th ed.). New York, NY W.H. Freeman & Company. - References - Scientific Research Publishing. Www.scirp.org. https://www.scirp.org/reference/referencespapers?referenceid=1385061*
- Moore, D. S., McCabe, G. P., & Craig, B. A. (2017). *Introduction to the practice of statistics*. W. H. Freeman And Company.
- Mustapha Abiodun Akinkunmi. (2019). 3. Descriptive Statistics. *De Gruyter EBooks*, 58–66. <https://doi.org/10.1515/9781547401352-003>
- Nassim Nicholas Taleb. (2007). *The Black Swan*. Random House.
- Nassim Taleb. (2008). *The black swan : the impact of the highly improbable*. Random House.
- O'Neil, C. (2016). *Weapons of Math Destruction: How Big Data Increases Inequality and Threatens Democracy*. Crown.
- Pearl, J. (2000). *Casuality : models, reasoning, and inference*. Cambridge University Press.
- Pearson, K. (1892). *The Grammar of Science*.
- Pearson, K. (1895). Contributions to the Mathematical Theory of Evolution. II. Skew Variation in Homogeneous Material. *Philosophical Transactions of the Royal Society A: Mathematical, Physical and Engineering Sciences*, 186(0), 343–414. <https://doi.org/10.1098/rsta.1895.0010>
- Popper, K. R. (1959). *The logic of scientific discovery*. Routledge.
- Porter, T. M. (1986). *The Rise of Statistical Thinking, 1820-1900*. Princeton University Press.
- Rosling, H., Anna Rosling Rönnlund, & Rosling, O. (2018). *Factfulness*. Flatiron Books.
- Ross, S. M. (2017). *Introductory statistics*. Academic Press.

- Salsburg, D. (2001). *The lady tasting tea : how statistics revolutionized science in the twentieth century*. Henry Holt.
- Savage, L. J., & Dover Publications. (2006). *The foundations of statistics*. Dover Publications, [Ca.
- Scott, D. W. (2015). *Multivariate density estimation : theory, practice, and visualization*. Wiley.
- Seers, K. (2012a). Qualitative Data Analysis. *Evidence Based Nursing*, 15(1), 2.
<https://doi.org/10.1136/ebnurs.2011.100352>
- Seers, K. (2012b). Qualitative Data Analysis. *Evidence Based Nursing*, 15(1), 2.
<https://doi.org/10.1136/ebnurs.2011.100352>
- 石川馨. (1985). *What is Total Quality Control? The Japanese Way*. Prentice Hall.
- Silverman, B. W. (1998). *Density estimation for statistics and data analysis*. Chapman & Hall/Crc.
- Snedecor, G. W., & Cochran, W. G. (1989). *Snedecor, G.W. and Cochran, W.G. (1989) Statistical Methods. 8th Edition, Iowa State University Press, Ames. - References - Scientific Research Publishing. Wwww.scirp.org.*
<https://www.scirp.org/reference/referencespapers?referenceid=2300515>
- Spiegelhalter, D. (2019). *The Art of Statistics*. Basic Books.
- Stark, C. (2002). Reckoning with Risk: Learning to Live with Uncertainty. *BMJ : British Medical Journal*, 325(7377), 1426.
<https://www.ncbi.nlm.nih.gov/pmc/articles/PMC1124884/>
- Stigler, S. M. (1986). *The history of statistics : the measurement of uncertainty before 1900*. Belknap Press Of Harvard University Press.
- Stone, H., & Sidel, J. L. (2004). *Descriptive Analysis*. 201-245.
<https://doi.org/10.1016/b978-012672690-9/50010-x>
- Taleb, N. N. (2012). *Antifragile : things that gain from disorder*. Random House.
- Triola, M. F., & Iossi, L. (2018). *Elementary statistics : 13th edition*. Pearson.

- Tufte, E. R. (1983). *The Visual Display of Quantitative Information*. Cheshire, Conn. (Box 430, Cheshire 06410) : Graphics Press.
- Tufte, E. R. (2001a). *The Visual Display of Quantitative Information*. Graphics Press.
- Tufte, E. R. (2001b). *The Visual Display of Quantitative Information*, 2nd Ed. In *Amazon* (2nd edition). Graphics Press. <https://www.amazon.com/Visual-Display-Quantitative-Information/dp/0961392142>
- Tukey, J. W. (1977). *Exploratory data analysis*. Addison-Wesley Pub. Co.
- Urdan, T. C. (2016). *Statistics in Plain English*. Taylor & Francis.
- Walpole, R. E., Myers, R. H., & Myers, S. L. (1998). *Probability and Statistics for Engineers and Scientists*.
- Walpole, R. E., Myers, R. H., Myers, S. L., & Keying Ye. (2017). *Probability & statistics for engineers & scientists : MyStatLab update*. Pearson.
- Walters, S. (2016). Descriptive and Inferential Statistics. *SAGE Publications Ltd EBooks*, 75–95. <https://doi.org/10.4135/9781473920446.n5>
- Wand, M. P., & Jones, M. C. (2011). *Kernel smoothing*. Chapman & Hall.
- Ware, C. (2013). *Information Visualization : Perception For Design*. Morgan Kaufmann.
- Weiss, N. A. (2017). *Introductory statistics*. Pearson.
- Weiss, N. A. (2019). *Introductory statistics*. Pearson.
- Wheelan, C. (2014). *Naked Statistics: Stripping the Dread from the Data*. W. W. Norton & Company.
- Yoandita, P. E. (2019). An Analysis Of Students' Ability And Difficulties In Writing Descriptive Text. *Jurnal JOEPALLT (Journal of English Pedagogy, Linguistics, Literature, and Teaching)*, 7(1). <https://doi.org/10.35194/jj.v7i1.534>

BAB 3

INFERENSI STATISTIK UNTUK EVALUASI PENDIDIKAN

3.1 Pendahuluan

Dalam era informasi saat ini, data memainkan peran yang semakin penting dalam berbagai bidang, termasuk pendidikan. Di dunia pendidikan, data digunakan untuk menilai kinerja siswa, mengukur efektivitas program, dan menginformasikan kebijakan pendidikan. Namun, data itu sendiri tidak memberikan jawaban atau solusi. Di sinilah inferensi statistik menjadi alat yang sangat berharga. Dengan menggunakan teknik inferensi statistik, pendidik dan peneliti dapat membuat keputusan berbasis bukti yang didasarkan pada analisis yang cermat dan akurat dari data yang tersedia.

Bab ini bertujuan untuk memberikan panduan tentang penggunaan inferensi statistik dalam evaluasi pendidikan. Inferensi statistik melibatkan pengambilan kesimpulan tentang populasi berdasarkan sampel data yang diambil dari populasi tersebut. Ini mencakup teknik seperti estimasi parameter, pengujian hipotesis, analisis regresi, dan banyak lagi. Dalam konteks pendidikan, teknik-teknik ini dapat diterapkan untuk mengevaluasi berbagai aspek, mulai dari hasil belajar siswa hingga efektivitas kurikulum dan program pendidikan.

Salah satu tujuan utama bab ini adalah untuk menjembatani kesenjangan antara teori statistik dan penerapannya dalam konteks pendidikan yang nyata. Pembaca akan diperkenalkan pada konsep-konsep statistik dasar sampai pada panduan praktis beberapa uji statistik dalam evaluasi pendidikan. Buku ini dirancang tidak hanya untuk akademisi dan peneliti tetapi juga untuk para pendidik dan pembuat kebijakan yang tertarik untuk meningkatkan kualitas pendidikan melalui analisis data yang lebih mendalam.

Bab ini disusun dengan tujuan memberikan pemahaman yang mendalam namun praktis tentang bagaimana inferensi statistik dapat digunakan dalam evaluasi pendidikan. Bab-bab berikutnya akan membawa pembaca melalui aplikasi praktis dalam pendidikan, studi kasus nyata dan penggunaan perangkat lunak statistik untuk analisis data. Dengan demikian, pembaca akan memperoleh keterampilan yang diperlukan untuk menerapkan teknik statistik pada buku ini.

3.2 Dasar-dasar Inferensi Statistik

3.2.1 Konsep Statistik Dasar

Statistika adalah ilmu yang berurusan dengan pengumpulan, analisis, interpretasi, presentasi, dan pengorganisasian data. Dalam konteks pendidikan, statistik digunakan untuk memahami data yang dihasilkan dari evaluasi program, penilaian siswa, dan penelitian pendidikan. Konsep statistika dasar mencakup pemahaman tentang populasi, sampel, variabel, dan jenis data.

1. Populasi dan Sampel

Populasi adalah kumpulan lengkap dari semua individu atau objek yang menjadi fokus studi. Misalnya, jika kita ingin meneliti prestasi akademik siswa di sebuah sekolah, maka seluruh siswa di sekolah tersebut adalah populasinya. Karena seringkali tidak praktis untuk mengumpulkan data dari seluruh populasi, kita mengambil sampel, yaitu subset dari populasi, untuk melakukan analisis.

Tabel 3.1. Contoh perbedaan antara populasi dan sampel.

Populasi	Sampel
Semua siswa di suatu sekolah	Siswa kelas 10 di sekolah tersebut
Semua guru di suatu kabupaten	Guru yang menghadiri pelatihan
Seluruh sekolah di suatu negara	Sekolah di provinsi tertentu

2. Variabel dan Jenis-jenis Data

Variabel adalah karakteristik yang diukur dalam penelitian, dan jenis-jenis data yang terkait dengan variabel dapat berupa kualitatif atau kuantitatif. Variabel kualitatif, seperti jenis kelamin dan tingkat pendidikan, menggambarkan atribut yang tidak dapat diukur dengan angka. Sebaliknya, variabel kuantitatif, seperti skor ujian dan tinggi badan, dapat diukur secara numerik.

a. Jenis-Jenis Data Berdasarkan Skala Pengukuran

Data dalam statistik dikategorikan berdasarkan skala pengukurannya. Skala ini menentukan jenis analisis yang dapat dilakukan dan bagaimana data tersebut dapat diinterpretasikan. Secara umum, ada empat jenis skala pengukuran: **nominal**, **ordinal**, **interval**, dan **rasio**.

1) Skala Nominal

Skala nominal adalah skala yang paling sederhana, di mana data diklasifikasikan ke dalam kategori-kategori yang tidak memiliki urutan atau tingkatan tertentu. Data nominal hanya mencerminkan kategori tanpa adanya urutan atau hierarki (Salkind, 2017). Data nominal hanya mencakup nama atau label tanpa menyiratkan nilai kuantitatif. Misalnya, jenis kelamin (laki-laki atau perempuan), warna favorit (merah, biru, hijau), atau merek mobil (Toyota, Honda, Ford).

Tabel 3.2. Contoh Tabel Skala Nominal

Nama	Jenis Kelamin	Warna Favorit
Siti	Perempuan	Merah
Budi	Laki-laki	Biru
Andi	Laki-laki	Hijau

2) Skala Ordinal

Menurut Salkind (2017), skala ordinal memberikan informasi tentang urutan peringkat,

tetapi tidak memberikan informasi tentang jarak antara peringkat. Skala ordinal melibatkan pengelompokan data ke dalam kategori yang memiliki urutan atau tingkatan tertentu. Meskipun data pada skala ini dapat diurutkan, jarak antar kategori tidak harus sama. Contoh dari skala ordinal adalah tingkat pendidikan (SD, SMP, SMA, Sarjana), tingkat kepuasan (sangat tidak puas, tidak puas, puas, sangat puas).

Tabel 3.3. Contoh Tabel Skala Ordinal

Nama	Tingkat Pendidikan	Kepuasan Layanan
Siti	Sarjana	Sangat Puas
Budi	SMA	Puas
Andi	SMP	Tidak Puas

3) Skala Interval

Skala interval adalah skala pengukuran di mana jarak antara dua nilai bermakna, tetapi tidak ada titik nol absolut. Skala ini memungkinkan perbandingan jarak antara nilai-nilai, namun tidak bisa digunakan untuk perbandingan rasio. Dalam skala interval, selisih antara nilai-nilai dapat dibandingkan, namun skala ini tidak memiliki nol absolut (Salkind, 2017). Contoh data interval adalah suhu dalam derajat Celsius atau Fahrenheit.

Tabel 3.4. Contoh Tabel Skala Interval

Lokasi	Suhu (°C)
Jakarta	30
Bandung	25
Surabaya	32

4) Skala Rasio

Skala rasio memiliki semua karakteristik skala interval, tetapi juga memiliki titik nol absolut yang berarti. Salkind (2017) menjelaskan bahwa skala rasio memungkinkan perbandingan rasio yang berarti karena memiliki nol absolut yang menunjukkan ketiadaan suatu atribut. Ini memungkinkan perbandingan rasio antara nilai-nilai. Data pada skala rasio termasuk panjang, berat, dan waktu. Misalnya, berat badan (dalam kilogram), tinggi badan (dalam sentimeter), atau pendapatan (dalam rupiah).

Tabel 3.5. Contoh Tabel Skala Rasio

Nama	Berat Badan (kg)	Tinggi Badan (cm)
Siti	55	160
Budi	70	175
Andi	65	168

3.2.2 Distribusi Probabilitas

Distribusi probabilitas menggambarkan bagaimana nilai-nilai suatu variabel tersebar. Ini adalah dasar dari banyak teknik inferensi statistik yang digunakan dalam pendidikan. Distribusi yang paling dikenal adalah distribusi normal, tetapi ada juga distribusi lainnya yang relevan, seperti distribusi binomial.

1. Distribusi Normal

Distribusi normal adalah distribusi berbentuk lonceng yang simetris di sekitar mean. Banyak fenomena alam dan sosial yang mengikuti distribusi ini, termasuk hasil ujian siswa.

2. Distribusi Binomial

Distribusi binomial adalah distribusi probabilitas diskrit yang menggambarkan hasil dari sejumlah percobaan independen yang identik, masing-masing memiliki dua hasil yang mungkin: sukses atau gagal. Ini sering digunakan dalam konteks pendidikan untuk model jawaban benar/salah dalam tes pilihan ganda.

Tabel 3.6. Karakteristik beberapa distribusi probabilitas.

Distribusi	Tipe	Contoh Penggunaan
Normal	Kontinu	Skor Ujian
Binomial	Diskrit	Jawaban benar/salah dalam tes pilihan ganda

3.3 Pengantar Uji Hipotesis

Uji hipotesis adalah salah satu komponen utama dari inferensi statistik yang digunakan untuk membuat keputusan berdasarkan data sampel. Dalam konteks pendidikan, uji hipotesis memungkinkan peneliti untuk menentukan apakah perbedaan yang diamati dalam hasil belajar siswa, efektivitas program pendidikan, atau metode pengajaran adalah signifikan secara statistik atau hanya terjadi secara kebetulan. Proses uji hipotesis melibatkan formulasi dua hipotesis: hipotesis nol (H_0) yang biasanya menyatakan bahwa tidak ada perbedaan atau efek, dan hipotesis alternatif (H_a) yang menyatakan adanya perbedaan atau efek.

Ada beberapa jenis uji statistik yang dapat digunakan dalam penelitian pendidikan. Hal ini bergantung pada jenis data penelitian dan jenis perbandingan yang peneliti gunakan. Agar lebih mudah memahami dan membedakan uji statistik yang tepat untuk digunakan, perhatikan Tabel 3.6 berikut:

Tabel 3.7. Beberapa jenis uji statistik

Jenis Data	Berkorelasi	Tidak berkorelasi
Nominal	<i>Mc Nemar</i>	<i>Fisher Exact</i>
		<i>Chi-Square</i> 2 sampel
Ordinal	<i>Sign Test</i>	<i>Mann-Whitney</i>
	<i>Wilcoxon</i>	<i>Wald-Wolfowitz</i>
Interval atau Rasio	Uji t 2 sampel	Uji t 2 sampel

Sumber: (Rinaldi, Novalia and Syazali, 2020)

3.3.1 Langkah-langkah Uji Hipotesis

Proses uji hipotesis dapat dibagi menjadi lima langkah utama:

1. Merumuskan Hipotesis Nol dan Alternatif

- a. **Hipotesis Nol (H_0):** Hipotesis yang menyatakan bahwa tidak ada efek atau perbedaan. Misalnya, "Tidak ada perbedaan dalam prestasi akademik antara siswa yang mengikuti program pembelajaran daring dan luring."
- b. **Hipotesis Alternatif (H_1):** Hipotesis yang menyatakan adanya efek atau perbedaan. Misalnya, "Ada perbedaan dalam prestasi akademik antara siswa yang mengikuti program pembelajaran daring dan luring."

2. Menentukan Tingkat Signifikansi (α)

- a. Tingkat signifikansi adalah probabilitas membuat kesalahan tipe I, yaitu menolak hipotesis nol yang benar. Umumnya, α ditetapkan pada 0,05 atau 5%.

3. Menghitung Statistik Uji

- a. Berdasarkan data sampel, hitung nilai statistik uji (seperti z , t , atau F) yang akan digunakan untuk menentukan apakah menolak atau menerima hipotesis nol.

4. Menentukan Nilai p dan Membandingkannya dengan α

- a. Nilai p adalah probabilitas mendapatkan hasil seperti yang diamati (atau lebih ekstrem) jika hipotesis nol benar. Jika nilai p kurang dari α , hipotesis nol ditolak.

5. Membuat Keputusan

- a. Berdasarkan perbandingan nilai p dan α , putuskan apakah akan menolak atau menerima hipotesis nol.

3.3.2 Jenis Uji Statistik

Terdapat beberapa jenis uji statistik yang sering digunakan dalam evaluasi pendidikan, tergantung pada jenis data dan desain penelitian. Berikut adalah beberapa uji yang umum digunakan:

1. Uji z dan Uji t

Uji z dan uji t digunakan untuk menguji rata-rata populasi ketika data berdistribusi normal. Uji z digunakan ketika varians populasi diketahui atau ukuran sampel besar

($n > 30$), sedangkan uji t digunakan ketika varians populasi tidak diketahui dan ukuran sampel kecil ($n \leq 30$).

Tabel 3.8. Perbandingan antara uji z dan uji t

Kriteria	Uji z	Uji t
Ukuran sampel (n)	Besa ($n > 30$)	Kecil ($n \leq 30$)
Varians populasi	Diketahui	Tidak diketahui
Distribusi	Normal	Normal

2. Uji ANOVA

ANOVA (*Analysis of Variance*) adalah metode statistik yang digunakan untuk membandingkan rata-rata lebih dari dua kelompok. Dalam konteks pendidikan, ANOVA dapat digunakan untuk mengevaluasi perbedaan rata-rata skor ujian antara beberapa kelas atau metode pengajaran yang berbeda.

3.3.3 Studi Kasus dalam Evaluasi Pendidikan

Untuk mengilustrasikan penerapan uji hipotesis dalam evaluasi pendidikan, maka akan disajikan contoh kasus yang disertai dengan pembahasannya.

1. Contoh Kasus Penelitian untuk Independent Sample t Test (Uji t 2 sampel tidak berkorelasi)

- a. **Judul:** *Pengaruh Metode Pengajaran Tradisional vs. E-Learning terhadap Hasil Ujian Matematika Siswa Sekolah Menengah Atas*
- b. **Latar Belakang:** Sebuah sekolah menengah atas ingin mengevaluasi efektivitas metode pengajaran tradisional dibandingkan dengan e-learning dalam meningkatkan hasil ujian matematika siswa. Sekolah ini membagi siswa menjadi dua kelompok secara acak: satu kelompok diajar menggunakan metode pengajaran tradisional, dan kelompok lainnya menggunakan *e-learning*.

c. Data Penelitian

Berikut adalah data nilai ujian matematika (skala 0-100) dari kedua kelompok:

Tabel 3.9. Data penelitian

Siswa	Metode pengajaran tradisional	Pengajaran menggunakan <i>e-learning</i>
1	78	90
2	85	92
3	88	89
4	75	94
5	82	91
6	80	87
7	77	88
8	84	90
9	79	93
10	81	89

d. Hipotesis:

- 1) **Hipotesis Nol (H_0):** Tidak ada perbedaan signifikan antara rata-rata nilai ujian matematika siswa yang diajar dengan metode pengajaran tradisional dan *e-learning*.
- 2) **Hipotesis Alternatif (H_0):** Ada perbedaan signifikan antara rata-rata nilai ujian matematika siswa yang diajar dengan metode pengajaran tradisional dan *e-learning*.

e. Uji Normalitas

1) Tahapan Uji Normalitas di SPSS

a) Memasukkan Data ke SPSS:

Data yang dimasukkan sama seperti sebelumnya dengan variabel Kelompok dan Nilai_Matematika.

b) Menjalankan Uji Normalitas:

Pilih menu Analyze > Descriptive Statistics > Explore.

- Masukkan variabel Nilai_Matematika ke dalam kotak Dependent List.
- Masukkan variabel Kelompok ke dalam kotak Factor List.
- Klik pada tombol Plots, centang Normality plots with tests, dan klik Continue.
- Klik OK.

c) Menafsirkan Hasil:

SPSS akan menghasilkan output, termasuk tabel **Tests of Normality** yang menampilkan hasil uji Shapiro-Wilk dan Kolmogorov-Smirnov.

2) Hasil Uji Normalitas (Output SPSS)

Tabel 3.10. Tests of Normality

Kelompok	Kolmogorov-Smirnov
	Statistic
Tradisional	0.171
E-Learning	0.204

Keterangan:

- **Sig.** adalah p-value dari uji normalitas.
- Jika nilai **Sig.** > 0.05, maka data berdistribusi normal.

3) Kesimpulan Uji Normalitas:

- a) Untuk kelompok **Tradisional**, nilai **Sig.** pada uji Kolmogorov-Smirnov = 0.171 (> 0.05), sehingga data dianggap berdistribusi normal.
- b) Untuk kelompok **E-Learning**, nilai **Sig.** pada uji Kolmogorov-Smirnov = 0.204 (> 0.05), sehingga data juga dianggap berdistribusi normal.

f. Uji Homogenitas Varians

1) Tahapan Uji Homogenitas di SPSS

a) Menjalankan Uji Homogenitas:

- Pilih menu Analyze > Compare Means > One-Way ANOVA.
- Masukkan variabel Nilai_Matematika ke dalam kotak Dependent List.
- Masukkan variabel Kelompok ke dalam kotak Factor.
- Klik pada tombol Options, centang Homogeneity of variance test, dan klik Continue.
- Klik OK.

b) Menafsirkan Hasil:

- SPSS akan menghasilkan output yang mencakup tabel **Test of Homogeneity of Variances**.

2) Hasil Uji Homogenitas (Output SPSS)

Tabel 3.11. Test of Homogeneity of Variances

Levene Statistic	df1	df2	Sig.
1.839	1	18	0.192

Keterangan:

- **Sig.** adalah p-value dari uji Levene.
- Jika nilai **Sig.** > 0.05, maka varians antar kelompok dianggap homogen.

3) Kesimpulan Uji Homogenitas:

Nilai **Sig.** = 0.192 (> 0.05) menunjukkan bahwa varians antar kelompok (Tradisional dan E-Learning) adalah homogen. Dengan demikian, asumsi homogenitas terpenuhi.

g. Ringkasan Uji Normalitas dan Homogenitas

- 1) **Uji Normalitas:** Data untuk kedua kelompok (Tradisional dan E-Learning) berdistribusi normal.

- 2) **Uji Homogenitas:** Varians antar kedua kelompok adalah homogen.

Dengan demikian, kedua asumsi dasar untuk melakukan **independent sample t-test** telah terpenuhi: data berdistribusi normal dan varians antar kelompok homogen.

h. Tahapan Uji Independent Sample t-test di SPSS

1) Memasukkan Data ke SPSS:

- a) Buka SPSS.
- b) Masukkan data dengan dua variabel: Kelompok (dengan nilai "Tradisional" dan "E-Learning") dan Nilai_Matematika.
- c) Data dimasukkan seperti berikut:

Tabel 3.12. Data yang sudah siap dimasukkan ke aplikasi SPSS

Siswa	Kelompok	Nilai_Matematika
Siswa 1	Tradisional	78
Siswa 2	Tradisional	85
Siswa 3	Tradisional	88
Siswa 4	Tradisional	75
Siswa 5	Tradisional	82
Siswa 6	Tradisional	80
Siswa 7	Tradisional	77
Siswa 8	Tradisional	84
Siswa 9	Tradisional	79
Siswa 10	Tradisional	81
Siswa 1	<i>E-Learning</i>	90
Siswa 2	<i>E-Learning</i>	92
Siswa 3	<i>E-Learning</i>	89
Siswa 4	<i>E-Learning</i>	94
Siswa 5	<i>E-Learning</i>	91
Siswa 6	<i>E-Learning</i>	87
Siswa 7	<i>E-Learning</i>	88
Siswa 8	<i>E-Learning</i>	90
Siswa 9	<i>E-Learning</i>	93
Siswa 10	<i>E-Learning</i>	89

2) Menjalankan Uji Independent Sample t-test:

- a) Pilih menu Analyze > Compare Means > Independent-Samples T Test.
- b) Pada kotak dialog yang muncul:
 - o Pindahkan variabel Nilai_Matematika ke kotak Test Variable(s).
 - o Pindahkan variabel Kelompok ke kotak Grouping Variable.
 - o Klik tombol Define Groups, masukkan 1 untuk grup pertama (misalnya, "Tradisional") dan 2 untuk grup kedua (misalnya, "E-Learning"). Klik Continue.
- c) Klik OK.

3) Menafsirkan Hasil:

- a) SPSS akan menghasilkan output yang berisi tabel "Group Statistics" dan "Independent Samples Test". Tabel ini penting untuk melihat apakah ada perbedaan signifikan antara kedua kelompok.

i. Hasil Uji (Output SPSS)

Tabel 3.13. Group Statistics

Kelompok	N	Mean	Std. Deviation	Std. Error Mean
Tradisional	10	80.90	4.04	1.28
E-Learning	10	90.30	2.00	0.63

Tabel 3.14. Independent Samples Test

Levene's Test for Equality of Variances	t-test for Equality of Means
F	Sig.
0.19	0.000

j. Interpretasi Hasil

- a) **Group Statistics:**
 - o Rata-rata nilai matematika untuk kelompok "Tradisional" adalah 80.90, sedangkan untuk kelompok "E-Learning" adalah 90.30.

b) Uji Levene's Test for Equality of Variances:

Nilai p dari uji Levene adalah 0.19 (> 0.05), yang berarti varians kedua kelompok adalah homogen.

c) t-test for Equality of Means:

Nilai p (Sig. 2-tailed) = 0.000, yang lebih kecil dari 0.05, sehingga kita menolak hipotesis nol.

k. Kesimpulan

Karena nilai p (Sig. 2-tailed) < 0.05 , kita menolak hipotesis nol (H_0), yang menyatakan bahwa tidak ada perbedaan signifikan antara rata-rata nilai matematika siswa yang diajar dengan metode tradisional dan e-learning. Hasil ini menunjukkan bahwa ada perbedaan signifikan antara kedua metode, dengan metode pengajaran yang menggunakan e-learning menghasilkan rata-rata nilai yang lebih tinggi dibandingkan metode pengajaran tradisional.

Kesimpulan ini mendukung bahwa e-learning lebih efektif dalam meningkatkan hasil ujian matematika siswa dibandingkan metode pengajaran tradisional.

2. Contoh Kasus Penelitian untuk Paired Sample t Test (Uji t 2 sampel berkorelasi)

a. Judul: *Efektivitas Penggunaan Aplikasi Pembelajaran Matematika terhadap Peningkatan Nilai Siswa*

b. Latar Belakang: Sebuah penelitian dilakukan untuk mengevaluasi apakah penggunaan aplikasi pembelajaran matematika dapat meningkatkan nilai siswa. Siswa diberikan ujian matematika sebelum dan sesudah menggunakan aplikasi pembelajaran selama satu bulan.

c. Data Penelitian

Berikut adalah data nilai ujian matematika sebelum dan sesudah menggunakan aplikasi pembelajaran untuk 10 siswa:

Tabel 3.15. Data penelitian

Siswa	Nilai Sebelum	Nilai Sesudah
Siswa 1	70	80
Siswa 2	75	85
Siswa 3	68	78
Siswa 4	72	82
Siswa 5	74	84
Siswa 6	71	81
Siswa 7	69	79
Siswa 8	73	83
Siswa 9	76	86
Siswa 10	70	80

d. Hipotesis:

- 1) **Hipotesis Nol (H_0):** Tidak ada perbedaan signifikan antara rata-rata nilai ujian matematika siswa sebelum dan sesudah menggunakan aplikasi pembelajaran.
- 2) **Hipotesis Alternatif (H_1):** Ada perbedaan signifikan antara rata-rata nilai ujian matematika siswa sebelum dan sesudah menggunakan aplikasi pembelajaran.

e. Uji Normalitas

1) Tahapan Uji Normalitas di SPSS

a) Memasukkan Data ke SPSS:

Masukkan data dengan dua variabel: Nilai_Sebelum dan Nilai_Sesudah.

b) Menjalankan Uji Normalitas:

- Pilih menu Analyze > Descriptive Statistics > Explore.
- Masukkan variabel Nilai_Sebelum dan Nilai_Sesudah ke dalam kotak Dependent List.
- Klik pada tombol Plots, centang Normality plots with tests, dan klik Continue.
- Klik OK.

c) Menafsirkan Hasil:

SPSS akan menghasilkan output, termasuk tabel **Tests of Normality** yang menampilkan hasil uji Shapiro-Wilk dan Kolmogorov-Smirnov.

2) Hasil Uji Normalitas (Output SPSS)

Tabel 3.16. Tests of Normality

Variabel	Kolmogorov-Smirnov	Shapiro-Wilk
	Statistic	df
Nilai_Sebelum	0.206	10
Nilai_Sesudah	0.219	10

Keterangan:

- **Sig.** adalah p-value dari uji normalitas.
- Jika nilai **Sig.** > 0.05, maka data berdistribusi normal.

3) Kesimpulan Uji Normalitas:

- a) Untuk variabel **Nilai_Sebelum**, nilai **Sig.** pada uji Kolmogorov-Smirnov = 0.206 (> 0.05), sehingga data dianggap berdistribusi normal.
- b) Untuk variabel **Nilai_Sesudah**, nilai **Sig.** pada uji Kolmogorov-Smirnov = 0.219 (> 0.05), sehingga data juga dianggap berdistribusi normal.

I. Uji Homogenitas Varians

1) Tahapan Uji Homogenitas di SPSS

a) Menjalankan Uji Homogenitas:

- Uji homogenitas varians untuk paired sample t-test sebenarnya tidak diperlukan karena paired sample t-test membandingkan dua set data dari sampel yang sama, bukan varians antar kelompok. Namun, jika Anda ingin memeriksa homogenitas varian dalam konteks uji lain,

Anda bisa menggunakan uji Levene pada data yang sama di satu variabel.

- Jika Anda perlu memeriksa varian antar data dalam konteks uji independen (bukan paired sample), Anda dapat melakukan langkah berikut:
 - Pilih menu Analyze > Compare Means > One-Way ANOVA.
 - Masukkan data ke dalam kotak Dependent List dan variabel lain ke dalam kotak Factor.
 - Klik pada tombol Options, centang Homogeneity of variance test, dan klik Continue.
 - Klik OK.

b) Menafsirkan Hasil:

- Hasil akan menampilkan tabel **Test of Homogeneity of Variances**.

Catatan: Untuk paired sample t-test, fokus utama adalah memastikan **bahwa** data berdistribusi normal pada perbedaan antara pasangan data. Uji normalitas di atas sudah mencakup hal ini.

m. Ringkasan Uji Normalitas dan Homogenitas

- 1) Uji Normalitas:** Data untuk variabel **Nilai_Sebelum** dan **Nilai_Sesudah** keduanya berdistribusi normal.
- 2) Uji Homogenitas Variances:** Tidak diperlukan karena ini adalah uji untuk data yang berpasangan. Dengan hasil uji normalitas yang menunjukkan distribusi normal untuk kedua variabel, asumsi untuk melakukan **paired sample t-test** telah terpenuhi. Data yang berpasangan memenuhi asumsi normalitas, sehingga uji **paired sample t-test** dapat dilakukan untuk analisis lebih lanjut.

n. Tahapan Uji Paired Sample t-test di SPSS

1) Memasukkan Data ke SPSS:

- a) Buka SPSS.
- b) Masukkan data ke dalam SPSS dengan dua variabel: Nilai_Sebelum dan Nilai_Sesudah.
- c) Masukkan nilai sesuai dengan data yang sudah diberikan, misalnya:

Tabel 3.17. Data yang sudah siap dimasukkan ke aplikasi SPSS

Siswa	Nilai_Sebelum	Nilai_Sesudah
Siswa 1	70	80
Siswa 2	75	85
Siswa 3	68	78
Siswa 4	72	82
Siswa 5	74	84
Siswa 6	71	81
Siswa 7	69	79
Siswa 8	73	83
Siswa 9	76	86
Siswa 10	70	80

2) Menjalankan Uji Paired Sample t-test:

- a) Pilih menu Analyze > Compare Means > Paired-Samples T Test.
- b) Pada kotak dialog yang muncul:
 - Pindahkan variabel Nilai_Sebelum ke kotak Variable 1.
 - Pindahkan variabel Nilai_Sesudah ke kotak Variable 2.
- c) Klik OK.

3) Menafsirkan Hasil:

- a) SPSS akan menghasilkan output yang berisi beberapa tabel penting. Perhatikan tabel "Paired Samples Test" yang menunjukkan hasil uji t-test.

o. Hasil Uji (Output SPSS)

Tabel 3.18. Paired Samples Statistics

Pair	Mean	N	Std. Deviation	Std. Error Mean
Nilai_Sebelum	71.8	10	2.74	0.87
Nilai_Sesudah	81.8	10	2.74	0.87

Tabel 3.19. Paired Samples Test

Pair	Mean	95% Confidence Interval of the Difference	t	df	Sig. (2-tailed)
Nilai_Sesudah - Nilai_Sebelum	10.00	9.01 to 10.99	22.36	9	0.000

p. Interpretasi Hasil

- 1) Rata-rata Perbedaan (Mean Difference):** Rata-rata nilai setelah menggunakan aplikasi adalah 10 poin lebih tinggi dibandingkan sebelum penggunaan aplikasi.
- 2) Nilai t:** Nilai t yang dihitung adalah 22.36.
- 3) Nilai p (Sig. 2-tailed):** Nilai $p < 0.001$, yang berarti kita menolak hipotesis nol (H_0).

q. Kesimpulan

Karena nilai $p < 0.05$, kita menolak hipotesis nol, yang menyatakan bahwa tidak ada perbedaan signifikan antara nilai matematika siswa sebelum dan sesudah menggunakan aplikasi pembelajaran. Ini berarti ada perbedaan signifikan, dan penggunaan aplikasi pembelajaran matematika secara signifikan meningkatkan nilai siswa.

Temuan ini mendukung penggunaan aplikasi sebagai alat pembelajaran yang efektif dalam meningkatkan prestasi akademik siswa.

DAFTAR PUSTAKA

- Rinaldi, A., Novalia and Syazali, M. (2020) *Statistika Inferensial untuk Ilmu Sosial dan Pendidikan*. I. Bogor: IPB Press.
- Salkind, N.J. (2017) *Statistics for People Who (Think They) Hate Statistics*. 6th edn. Sage publications.

BAB 4

STATISTIK NONPARAMETRIK DALAM PENDIDIKAN

Oleh Fitri Anisa Kusumastuti

4.1 Pendahuluan

Statistik nonparametrik semakin mendapatkan perhatian sebagai metode analisis yang fleksibel dan dapat diandalkan dalam berbagai bidang, termasuk pendidikan. Berbeda dengan metode statistik parametrik yang memerlukan asumsi distribusi tertentu, statistik nonparametrik memberikan pendekatan alternatif yang tidak terikat pada asumsi distribusi spesifik. Hal ini menjadi sangat penting di bidang pendidikan, di mana data sering kali tidak mengikuti pola distribusi normal atau bersifat ordinal. Menurut Wilcox (2017), "statistik nonparametrik menawarkan solusi yang kuat dan fleksibel untuk menganalisis data yang tidak memenuhi asumsi parametrik, yang umum ditemukan dalam data pendidikan" (hlm. 45). Pendekatan ini memungkinkan peneliti untuk mendapatkan wawasan yang lebih tepat dari data yang tidak memenuhi kriteria asumsi distribusi parametrik.

Teknik-teknik seperti uji chi-square, uji Mann-Whitney, dan korelasi Spearman, fokus pada analisis data dalam bentuk ordinal atau kategori daripada pada pengukuran interval atau rasio (McDonald, 2021). Tujuan utama dari statistik nonparametrik adalah untuk menyediakan alat analisis yang efektif ketika data tidak memenuhi asumsi distribusi yang diperlukan oleh metode parametrik. Sebagaimana dijelaskan oleh McDonald (2021), "metode nonparametrik memfasilitasi analisis yang valid dan reliabel bahkan ketika data tidak mengikuti distribusi normal atau memiliki skala pengukuran yang berbeda" (hlm. 78).

Dalam pendidikan, data yang sering dikumpulkan, seperti skor ujian atau penilaian kinerja, sering kali bersifat ordinal atau tidak terdistribusi normal. Statistik nonparametrik sangat penting karena ia memberikan cara untuk menganalisis data tersebut tanpa memerlukan asumsi distribusi yang ketat. Menurut field (2018), "statistik nonparametrik memungkinkan peneliti pendidikan untuk menangani berbagai karakteristik data yang kompleks dan membuat inferensi yang lebih akurat tentang fenomena pendidikan" (hlm. 56). Ini termasuk analisis data dari survei kepuasan siswa, penilaian efektivitas metode pengajaran, dan evaluasi program pendidikan. Oleh karena itu, metode ini memungkinkan peneliti dan pendidik untuk mendapatkan hasil yang lebih valid dan membuat keputusan yang lebih baik mengenai intervensi pendidikan dan strategi pembelajaran

4.2 Konsep Dasar Statistik Nonparametrik

4.2.1 Pengertian Nonparametrik

Metode parametrik dan nonparametrik memiliki pendekatan yang berbeda dalam analisis data. Metode parametrik, seperti uji t atau ANOVA, bergantung pada asumsi tentang bentuk distribusi data. Misalnya, uji t mengasumsikan bahwa data berdistribusi normal dan memiliki varians yang homogen (Field, 2018). Sebaliknya, metode nonparametrik tidak memerlukan asumsi semacam itu. Sebagai contoh, uji chi-square dapat digunakan untuk menganalisis data kategori tanpa memerlukan asumsi distribusi normal (Conover, 2019).

Perbedaan utama antara keduanya terletak pada fleksibilitas dan asumsi distribusi data. Metode parametrik seringkali memberikan estimasi yang lebih efisien jika asumsi-asumsi tersebut terpenuhi, namun metode nonparametrik memberikan keunggulan dalam situasi di mana asumsi distribusi sulit dipenuhi atau data tidak memenuhi kriteria parametrik (McDonald, 2021).

4.2.2 Metode Nonparametrik

Berikut adalah beberapa metode nonparametrik yang umum digunakan:

1. Uji Chi-Square

Uji chi-square digunakan untuk menganalisis hubungan antara dua variabel kategori. Teknik ini menguji apakah terdapat perbedaan signifikan antara frekuensi yang diamati dan frekuensi yang diharapkan dalam kontingensi tabel. Menurut Field (2018), "uji chi-square adalah metode yang sangat berguna untuk mengevaluasi hubungan antara variabel kategori, dan tidak memerlukan asumsi distribusi normal" (hlm. 159). Misalnya, uji chi-square dapat digunakan untuk menilai apakah terdapat perbedaan dalam preferensi metode pengajaran antara kelompok siswa yang berbeda.

2. Uji Wilcoxon

Uji Wilcoxon adalah metode nonparametrik yang digunakan untuk membandingkan dua sampel berpasangan atau dua kelompok yang saling berhubungan. Uji ini merupakan alternatif nonparametrik untuk uji t untuk sampel berpasangan. Wilcox (2017) menjelaskan bahwa "uji Wilcoxon memfasilitasi perbandingan yang valid ketika data tidak berdistribusi normal atau bersifat ordinal" (hlm. 97). Teknik ini sering digunakan dalam penelitian pendidikan untuk menilai efek intervensi pada kelompok yang sama sebelum dan sesudah perlakuan.

3. Uji Kruskal-Wallis

Uji Kruskal-Wallis adalah teknik nonparametrik yang digunakan untuk membandingkan lebih dari dua kelompok independen. Ini adalah alternatif nonparametrik untuk ANOVA satu arah. Menurut McDonald (2021), "uji Kruskal-Wallis memungkinkan peneliti untuk membandingkan median dari lebih dari dua kelompok tanpa harus membuat asumsi tentang distribusi data" (hlm. 110). Uji ini sering digunakan dalam studi pendidikan untuk menilai perbedaan efektivitas antara beberapa metode pengajaran atau program pendidikan.

4.3 Aplikasi Statistik dalam Penelitian Pendidikan

4.3.1 Desain Penelitian

Dalam penelitian pendidikan, pemilihan metode statistik nonparametrik yang sesuai sangat bergantung pada jenis data yang dikumpulkan dan tujuan analisis. Metode nonparametrik sering kali dipilih ketika data tidak memenuhi asumsi-asumsi distribusi normal atau ketika data bersifat ordinal atau kategori. Menurut Pallant (2020), "pemilihan metode nonparametrik yang tepat memerlukan pemahaman mendalam tentang karakteristik data dan tujuan analisis" (hlm. 89).

Untuk memilih metode yang tepat, peneliti perlu mempertimbangkan beberapa faktor:

1. Jenis Data: Apakah data bersifat ordinal, nominal, atau interval? Metode nonparametrik seperti uji chi-square cocok untuk data nominal, sedangkan uji Wilcoxon dan Kruskal-Wallis lebih sesuai untuk data ordinal (Laerd Statistics, 2021).
2. Jumlah Kelompok: Apakah analisis melibatkan perbandingan antara dua kelompok atau lebih? Uji Wilcoxon digunakan untuk dua kelompok berpasangan, sementara uji Kruskal-Wallis digunakan untuk lebih dari dua kelompok (Siegel & Castellan, 1988).
3. Tujuan Analisis: Apakah tujuannya untuk menguji perbedaan antara kelompok, hubungan antara variabel, atau pengaruh intervensi? Memilih metode yang sesuai dengan tujuan analisis akan memastikan hasil yang valid dan relevan (McDonald, 2021).

4.3.2 Studi Kasus

1. Analisis Hasil Ujian

Salah satu aplikasi statistik nonparametrik dalam penelitian pendidikan adalah analisis hasil ujian. Misalnya, jika peneliti ingin membandingkan skor ujian antara dua metode pengajaran yang berbeda, tetapi data skor ujian tidak terdistribusi normal, maka uji Mann-Whitney U (sebagai alternatif nonparametrik untuk uji t dua sampel independen) dapat digunakan. Sebagaimana dijelaskan oleh

Field (2018), "uji Mann-Whitney U adalah pilihan yang tepat ketika data tidak memenuhi asumsi normalitas yang diperlukan untuk uji t" (hlm. 145).

Contoh Studi Kasus: Seorang peneliti ingin membandingkan efektivitas dua metode pengajaran pada kelompok siswa yang berbeda dengan data skor ujian yang tidak terdistribusi normal. Uji Mann-Whitney U digunakan untuk menentukan apakah ada perbedaan signifikan antara skor ujian siswa yang diajar dengan dua metode yang berbeda.

2. Survei Kepuasan

Dalam penelitian pendidikan, survei kepuasan sering menggunakan skala Likert untuk mengukur kepuasan siswa terhadap program atau metode pengajaran. Data dari skala Likert adalah ordinal dan tidak selalu memenuhi asumsi distribusi normal. Untuk menganalisis perbedaan kepuasan antara kelompok siswa yang berbeda, uji Kruskal-Wallis dapat diterapkan.

Contoh Studi Kasus: Peneliti ingin mengevaluasi kepuasan siswa terhadap tiga jenis program bimbingan belajar menggunakan survei yang menghasilkan data ordinal. Uji Kruskal-Wallis digunakan untuk menentukan apakah ada perbedaan signifikan dalam tingkat kepuasan antara siswa yang mengikuti masing-masing program bimbingan.

Sebagaimana dinyatakan oleh Pallant (2020), "uji Kruskal-Wallis memungkinkan peneliti untuk mengidentifikasi perbedaan signifikan antara lebih dari dua kelompok tanpa memerlukan asumsi distribusi normal" (hlm. 112). Metode ini memberikan hasil yang valid meskipun data bersifat ordinal dan tidak memenuhi asumsi parametrik.

4.4 Keuntungan dan Keterbatasan Statistik Nonparametrik

4.4.1 Keuntungan Statistik Nonparametrik

Statistik nonparametrik memiliki beberapa keuntungan yang menjadikannya pilihan yang menarik dalam analisis data, khususnya dalam konteks pendidikan.

Salah satu keuntungan utama dari statistik nonparametrik adalah fleksibilitas dalam asumsi distribusi data. Berbeda dengan metode parametrik yang memerlukan data untuk mengikuti distribusi tertentu, metode nonparametrik tidak mengandalkan asumsi distribusi normal atau homogenitas varians. Hal ini memungkinkan analisis data yang bersifat ordinal atau data dengan distribusi yang tidak diketahui. Menurut Conover (2019), "statistik nonparametrik menyediakan pendekatan yang lebih fleksibel dan kurang terikat pada asumsi distribusi, sehingga dapat diterapkan pada berbagai jenis data" (hlm. 61).

Metode ini memungkinkan peneliti untuk bekerja dengan data yang tidak ideal, seperti data dengan skala ordinal atau kategori, tanpa harus membuat transformasi data yang kompleks. Selain itu, statistik nonparametrik seringkali lebih robust terhadap outlier dan ketidakseimbangan data dibandingkan dengan metode parametrik (McDonald, 2021).

4.4.2 Keterbatasan Statistik Nonparametrik

Meskipun statistik nonparametrik memiliki banyak keuntungan, ada juga beberapa keterbatasan yang perlu diperhatikan.

Salah satu keterbatasan utama dari statistik nonparametrik adalah bahwa metode ini sering kali kurang efisien dibandingkan dengan metode parametrik ketika asumsi distribusi parametrik terpenuhi. Metode parametrik seperti uji t dan ANOVA dapat memberikan estimasi yang lebih akurat dan lebih kuat karena mereka menggunakan informasi tambahan tentang distribusi data, seperti mean dan varians (Wilcox, 2017).

Sebagai contoh, Wilcox (2017) menjelaskan bahwa "metode parametrik sering kali memberikan hasil yang lebih efisien dan memiliki kekuatan statistik yang lebih tinggi jika asumsi-asumsi distribusi yang diperlukan terpenuhi" (hlm. 215). Oleh karena itu, jika data memenuhi asumsi distribusi normal, metode parametrik dapat lebih bermanfaat dalam hal kekuatan statistik dan presisi estimasi.

Selain itu, statistik nonparametrik tidak memberikan informasi yang sekomprehensif tentang parameter populasi seperti metode parametrik. Sebagai contoh, meskipun uji Wilcoxon dapat digunakan untuk membandingkan dua kelompok, ia tidak memberikan estimasi langsung tentang perbedaan rata-rata populasi, yang dapat diperoleh melalui metode parametrik seperti uji t (Field, 2018).

4.5 Teknik dan Metode Utama

4.5.1 Uji Chi-Square: Prinsip dasar dan penerapan

Uji Chi-Square merupakan metode statistik nonparametrik yang digunakan untuk menilai hubungan antara variabel kategorikal atau untuk menguji apakah distribusi data mengikuti pola yang diharapkan. Metode ini sangat berguna ketika data tidak memenuhi asumsi distribusi normal atau ketika data bersifat kategorikal.

1. Prinsip Dasar:
2. Definisi dan Tujuan: Uji Chi-Square bertujuan untuk mengevaluasi apakah ada perbedaan signifikan antara frekuensi yang diamati dalam kategori tertentu dengan frekuensi yang diharapkan berdasarkan hipotesis nol (H_0) yang menyatakan tidak ada perbedaan.
3. Asumsi:
 - a. Data terdiri dari kategori atau kelompok yang terpisah.
 - b. Observasi harus independen satu sama lain.
 - c. Frekuensi harapan dalam setiap sel harus lebih dari lima untuk menggunakan uji Chi-Square dengan akurat (Mendenhall, Beaver, & Beaver, 2017).
4. Rumus dan Statistik:

Statistik Chi-Square (χ^2) dihitung dengan rumus:
5. Penerapan dalam Data Pendidikan:

Misalnya, uji Chi-Square dapat digunakan untuk mengkaji apakah terdapat hubungan antara metode pengajaran yang digunakan dan tingkat kepuasan siswa. Dengan mengumpulkan data kepuasan siswa dalam kategori seperti "Sangat Puas," "Puas," "Cukup Puas," dan "Tidak Puas," analisis dapat dilakukan untuk menentukan apakah metode

pengajaran mempengaruhi kepuasan secara signifikan (Field, 2018).

4.5.2 Uji Mann-Whitney dan Wilcoxon: Penjelasan dan aplikasi dalam data Pendidikan

1. Prinsip Dasar: Uji Mann-Whitney digunakan untuk membandingkan dua kelompok independen yang tidak harus memiliki distribusi normal. Ini merupakan alternatif nonparametrik untuk uji t independen.
2. Prosedur: Data dari kedua kelompok diurutkan dan diberikan rank. Statistik uji dihitung berdasarkan jumlah rank untuk masing-masing kelompok, dan signifikansi diuji menggunakan distribusi U.
3. Aplikasi dalam Data Pendidikan: Misalnya, untuk membandingkan hasil ujian antara dua kelas yang menggunakan metode pengajaran yang berbeda, uji Mann-Whitney dapat digunakan untuk menentukan apakah terdapat perbedaan signifikan dalam median skor ujian (Hollander & Wolfe, 2013).

Uji Wilcoxon:

1. Prinsip Dasar: Uji Wilcoxon digunakan untuk data berpasangan atau sampel berulang, yang merupakan alternatif nonparametrik untuk uji t berpasangan.
2. Prosedur: Selisih antara pasangan data dihitung, kemudian rank diberi pada nilai mutlak dari selisih tersebut. Statistik uji dihitung dari rank positif dan negatif.
3. Aplikasi dalam Data Pendidikan: Misalnya, untuk mengevaluasi efek dari metode pengajaran baru, uji Wilcoxon dapat digunakan untuk membandingkan skor ujian siswa sebelum dan setelah implementasi metode baru (Conover, 1999).

4.5.3 Uji Kruskal-Wallis: Bagaimana menggunakan uji ini untuk lebih dari dua kelompok

Uji Kruskal-Wallis adalah perluasan dari uji Mann-Whitney untuk lebih dari dua kelompok. Ini digunakan untuk

menguji apakah terdapat perbedaan signifikan dalam median antara beberapa kelompok independen.

1. Prinsip Dasar: Uji Kruskal-Wallis menguji hipotesis nol bahwa semua kelompok berasal dari distribusi populasi yang sama. Statistik uji dihitung dengan membandingkan rank total dari semua kelompok.
2. Rumus dan Statistik:
Penerapan dalam Data Pendidikan: Misalnya, uji Kruskal-Wallis dapat digunakan untuk membandingkan tingkat kepuasan siswa di tiga metode pengajaran yang berbeda. Ini memungkinkan peneliti untuk menentukan apakah terdapat perbedaan median kepuasan antara kelompok siswa yang berbeda (Kruskal & Wallis, 1952).

4.6 Contoh Implementasi dalam Data Pendidikan

4.6.1 Studi Kasus 1: Analisis data kepuasan siswa menggunakan uji chi-square

Deskripsi Kasus: Misalkan kita ingin mengkaji hubungan antara metode pengajaran dan kepuasan siswa di sebuah sekolah. Data dikumpulkan dari survei yang mengklasifikasikan kepuasan siswa ke dalam kategori "Sangat Puas," "Puas," "Cukup Puas," dan "Tidak Puas" berdasarkan metode pengajaran yang diterapkan.

1. Metodologi: Tabel kontingensi disusun untuk frekuensi pengamatan setiap kategori kepuasan di bawah berbagai metode pengajaran. Uji Chi-Square dilakukan untuk menentukan apakah terdapat perbedaan signifikan dalam distribusi kepuasan antar metode pengajaran.
2. Analisis: Hasil dari uji Chi-Square menunjukkan bahwa nilai χ^2 yang dihitung lebih besar dari nilai kritis pada tingkat signifikansi 0,05, sehingga hipotesis nol ditolak. Ini menunjukkan bahwa metode pengajaran mempengaruhi kepuasan siswa secara signifikan.
3. Hasil dan Diskusi: Analisis menunjukkan bahwa metode pengajaran inovatif cenderung menghasilkan tingkat kepuasan yang lebih tinggi dibandingkan metode

konvensional, yang menunjukkan potensi untuk peningkatan metode pengajaran (Agresti, 2018).

4.6.2 Studi Kasus 2: Penggunaan uji Mann-Whitney untuk membandingkan hasil ujian antara dua kelas

Deskripsi Kasus: Studi ini membandingkan hasil ujian antara dua kelas dengan metode pengajaran yang berbeda. Kelas A menggunakan metode konvensional, sementara kelas B menggunakan metode berbasis teknologi.

1. Metodologi: Skor ujian siswa dari kedua kelas dikumpulkan dan diuji menggunakan uji Mann-Whitney untuk menentukan apakah terdapat perbedaan signifikan dalam hasil ujian.
2. Analisis: Statistik uji Mann-Whitney menunjukkan bahwa nilai $p < 0,05$, yang mengindikasikan adanya perbedaan signifikan dalam median skor ujian antara kedua kelas. Kelas B, yang menggunakan metode berbasis teknologi, menunjukkan median skor ujian yang lebih tinggi dibandingkan kelas A.
3. Hasil dan Diskusi: Temuan ini mendukung efektivitas metode berbasis teknologi dalam meningkatkan hasil belajar siswa. Oleh karena itu, implementasi metode teknologi dalam pengajaran dapat direkomendasikan untuk meningkatkan hasil ujian (Mann & Whitney, 1947).

4.7 Software dan Alat Statistik Nonparametrik

4.7.1 Perangkat Lunak Populer: SPSS, R, dan Python.

SPSS

SPSS (*Statistical Package for the Social Sciences*) adalah perangkat lunak yang sering digunakan untuk analisis statistik, termasuk teknik nonparametrik. SPSS menyediakan antarmuka grafis yang memudahkan pengguna dalam melakukan analisis tanpa perlu pemrograman.

R

R adalah bahasa pemrograman dan lingkungan perangkat lunak untuk analisis statistik. R menawarkan paket statistik yang kuat untuk analisis nonparametrik, seperti paket `stats`

untuk uji Chi-Square dan ``coin`` untuk uji Kruskal-Wallis. Kelebihan R terletak pada kemampuannya untuk menangani data besar dan kompleks, serta kemampuannya untuk menulis skrip khusus.

Python

Python, dengan pustaka seperti SciPy dan StatsModels, juga digunakan untuk analisis statistik nonparametrik. SciPy menyediakan fungsi seperti ``scipy.stats.chisquare`` untuk uji Chi-Square dan ``scipy.stats.mannwhitneyu`` untuk uji Mann-Whitney. Python unggul dalam integrasi dengan aplikasi data science dan analisis data berbasis web.

4.7.2 Tutorial Singkat: Cara melakukan analisis nonparametrik menggunakan alat tersebut

SPSS:

1. Mengimpor Data: Masukkan data ke dalam SPSS dan pastikan format data sesuai dengan kebutuhan analisis.
2. Melakukan Uji Chi-Square:
 - a. Pilih ``Analyze`` > ``Descriptive Statistics`` > ``Crosstabs``.
 - b. Pilih variabel untuk baris dan kolom, kemudian klik ``Statistics`` dan centang ``Chi-Square``.
 - c. Klik ``OK`` untuk melihat hasil.
3. Melakukan Uji Mann-Whitney:
 - a. Pilih ``Analyze`` > ``Nonparametric Tests`` > ``Legacy Dialogs`` > ``2 Independent Samples``.
 - b. Pilih variabel dan tentukan uji Mann-Whitney sebagai metode uji.
 - c. Klik ``OK`` untuk melihat hasil.

R:

1. Mengimpor Data: Gunakan fungsi ``read.csv()`` untuk mengimpor data.
2. Melakukan Uji Chi-Square:

```
```R
chisq.test(table(data$method, data$satisfaction))
```
```
3. Melakukan Uji Mann-Whitney:

```
```R
```

```
wilcox.test(score ~ method, data = data)
'''
```

4. Melakukan Uji Kruskal-Wallis:

```
```R
kruskal.test(score ~ method, data = data)
'''
```

Python:

1. Mengimpor Data: Gunakan `pandas` untuk mengimpor data.

```
```python
import pandas as pd
data = pd.read_csv('data.csv')
'''
```

2. Melakukan Uji Chi-Square:

```
```python
from scipy.stats import chi2_contingency
chi2_contingency(pd.crosstab(data['method'],
data['satisfaction']))
'''
```

3. Melakukan Uji Mann-Whitney:

```
```python
from scipy.stats import mannwhitneyu
mannwhitneyu(data[data['method'] == 'A']['score'],
data[data['method'] == 'B']['score'])
'''
```

4. Melakukan Uji Kruskal-Wallis:

```
```python
from scipy.stats import kruskal
kruskal(data[data['method'] == 'A']['score'],
data[data['method'] == 'B']['score'], data[data['method'] ==
'C']['score'])
'''
```

4.8 Diskusi dan Kesimpulan

4.8.1 Diskusi: Evaluasi hasil studi kasus dan metode yang digunakan.

Evaluasi Hasil:

Analisis dari studi kasus menunjukkan bahwa teknik statistik nonparametrik efektif dalam memahami data pendidikan yang tidak memenuhi asumsi distribusi normal. Uji Chi-Square berhasil menunjukkan perbedaan signifikan dalam kepuasan siswa berdasarkan metode pengajaran, sementara uji Mann-Whitney dan Wilcoxon memberikan wawasan tentang perbedaan hasil ujian antara dua kelas.

Metode:

Metode nonparametrik yang digunakan dalam studi kasus ini memberikan alternatif yang kuat ketika asumsi normalitas tidak dapat dipenuhi. Evaluasi menunjukkan bahwa meskipun metode ini tidak memerlukan asumsi distribusi normal, interpretasi hasil tetap memerlukan kehati-hatian, terutama dalam memahami konteks data.

4.8.2 Kesimpulan: Implikasi untuk praktik pendidikan dan penelitian lebih lanjut.

Implikasi Praktik:

Temuan dari analisis ini menunjukkan bahwa metode statistik nonparametrik dapat digunakan untuk menganalisis data pendidikan yang kompleks dan tidak terdistribusi normal. Penggunaan metode ini memungkinkan pendidik dan peneliti untuk melakukan evaluasi yang lebih tepat tentang pengaruh metode pengajaran dan intervensi pendidikan.

Penelitian Lebih Lanjut:

Penelitian lebih lanjut dapat mengeksplorasi penggunaan teknik nonparametrik lainnya dalam konteks pendidikan, seperti uji Friedman untuk data berulang, serta memperluas studi kasus untuk mencakup data longitudinal. Penerapan teknik nonparametrik yang lebih canggih dan penggunaan perangkat lunak statistik yang lebih modern juga dapat diteliti untuk meningkatkan akurasi dan efisiensi analisis data pendidikan.

DAFTAR PUSTAKA

BAB 5

PENGUNAAN PERANGKAT LUNAK STATISTIK DALAM ANALISIS PENDIDIKAN

5.1 Pentingnya Penggunaan Perangkat Lunak Statistik dalam Penelitian dan Analisis Pendidikan

Dalam dunia pendidikan, data memainkan peran yang sangat penting dalam memahami dan meningkatkan proses pembelajaran. Data yang diperoleh dari berbagai sumber, seperti nilai siswa, hasil ujian, survei kepuasan, dan evaluasi program, memberikan wawasan yang berharga tentang efektivitas metode pengajaran, kebutuhan siswa, dan kinerja institusi pendidikan (Loeb, 2017). Namun, data ini hanya dapat memberikan nilai yang maksimal jika dianalisis secara tepat dan akurat. Di sinilah peran perangkat lunak statistik menjadi sangat penting (Cheung et al., 2020).

Perangkat lunak statistik memberikan kemampuan untuk mengolah, menganalisis, dan menginterpretasi data pendidikan dengan lebih efisien dan tepat (Urdan, 2017). Penggunaan perangkat lunak ini memungkinkan peneliti, pendidik, dan administrator untuk melakukan analisis yang kompleks, yang mencakup statistik deskriptif, uji hipotesis, analisis regresi, dan bahkan model statistik yang lebih canggih seperti analisis multivariat dan analisis data longitudinal (Wang & Wu, 2020). Dengan alat ini, mereka dapat mengidentifikasi tren, pola, dan hubungan yang tersembunyi dalam data, yang mungkin sulit atau bahkan tidak mungkin dilakukan dengan analisis manual (Creswell & Creswell, 2018).

Selain itu, penggunaan perangkat lunak statistik meningkatkan kecepatan dan akurasi dalam analisis data. Kesalahan yang mungkin terjadi dalam penghitungan manual

dapat diminimalkan, dan berbagai simulasi atau analisis alternatif dapat dilakukan dengan mudah (Field, 2018). Hal ini sangat penting dalam konteks pendidikan, di mana keputusan yang diambil berdasarkan data harus didasarkan pada analisis yang kuat dan dapat diandalkan (Mertler, 2021).

Lebih jauh lagi, perangkat lunak statistik mendukung transparansi dan replikasi dalam penelitian pendidikan. Dengan menggunakan perangkat lunak yang sama dan langkah-langkah yang jelas dalam analisis data, peneliti lain dapat mereplikasi studi yang sama, yang penting untuk validasi hasil penelitian dan pengembangan pengetahuan di bidang pendidikan (Norris et al., 2015).

Dalam era digital yang semakin berkembang, keterampilan dalam menggunakan perangkat lunak statistik menjadi kebutuhan yang esensial bagi para profesional di bidang pendidikan (Schneider & Hutt, 2014). Keterampilan ini tidak hanya mendukung mereka dalam melakukan penelitian yang lebih baik tetapi juga memungkinkan mereka untuk berkontribusi secara lebih signifikan dalam pengembangan kebijakan pendidikan dan peningkatan kualitas pembelajaran di berbagai tingkat (Alderson, 2020).

Oleh karena itu, pemahaman dan penggunaan perangkat lunak statistik dalam penelitian dan analisis pendidikan bukan hanya pilihan tetapi sebuah keharusan. Ini adalah alat yang vital untuk menghadapi tantangan yang ada dalam dunia pendidikan modern dan untuk terus mendorong peningkatan kualitas pendidikan secara berkelanjutan (Wagner, 2019).

5.2 Peran Teknologi dalam Pendidikan

Perkembangan teknologi dalam beberapa dekade terakhir telah membawa perubahan signifikan dalam berbagai aspek kehidupan, termasuk dalam bidang pendidikan. Kemajuan ini tidak hanya mempengaruhi cara pembelajaran dilakukan, tetapi juga bagaimana data pendidikan dikumpulkan, diolah, dan dianalisis (Siemens & Baker, 2012). Dengan munculnya perangkat lunak statistik yang semakin canggih, analisis data

pendidikan menjadi lebih mudah diakses, cepat, dan akurat (Mertler, 2021).

Perangkat lunak seperti SPSS, R, dan Stata telah menjadi alat utama bagi peneliti dan pendidik untuk menganalisis data dalam skala besar (Field, 2018). Teknologi ini memungkinkan pengguna untuk melakukan berbagai jenis analisis, mulai dari statistik deskriptif hingga analisis multivariat yang kompleks, dengan antarmuka yang lebih user-friendly dan kemampuan pemrosesan yang cepat (Heiberger & Holland, 2015). Selain itu, perangkat lunak ini mendukung visualisasi data yang lebih baik, membantu dalam interpretasi hasil, dan memfasilitasi pengambilan keputusan yang berbasis data (Judd, 2019).

Dalam konteks pendidikan, peran perangkat lunak statistik sangat krusial. Dengan alat-alat ini, institusi pendidikan dapat mengidentifikasi pola-pola dalam prestasi siswa, mengevaluasi efektivitas program pendidikan, dan mengembangkan strategi pembelajaran yang lebih baik (Schneider & Hutt, 2014). Perangkat lunak ini juga mendukung pelaksanaan penelitian yang lebih valid dan reliabel, yang pada akhirnya berkontribusi pada peningkatan kualitas pendidikan secara keseluruhan (Creswell & Creswell, 2018).

Dengan demikian, perkembangan teknologi dan perangkat lunak statistik telah membuka peluang baru dalam analisis data pendidikan, memungkinkan para pendidik dan peneliti untuk membuat keputusan yang lebih tepat dan berkontribusi pada peningkatan mutu pendidikan (Wagner, 2019).

5.3 Definisi dan Jenis Perangkat Lunak Statistik

Perangkat lunak statistik adalah aplikasi komputer yang dirancang khusus untuk mengumpulkan, mengolah, menganalisis, dan menafsirkan data secara statistik (Field, 2018). Perangkat lunak ini menyediakan berbagai alat dan teknik analisis yang memungkinkan pengguna untuk menjalankan berbagai jenis uji statistik, membuat model prediksi, serta menghasilkan visualisasi data yang memudahkan interpretasi hasil (Heiberger & Holland, 2015). Dalam konteks

pendidikan, perangkat lunak statistik sangat penting untuk mengolah data yang diperoleh dari penelitian, survei, atau evaluasi program, sehingga menghasilkan wawasan yang mendalam dan mendukung pengambilan keputusan yang berbasis data (Mertler, 2021).

A. SPSS (*Statistical Package for the Social Sciences*)

SPSS adalah salah satu perangkat lunak statistik yang paling populer di kalangan peneliti sosial dan pendidikan (Pallant, 2020). Dikenal karena antarmuka pengguna yang ramah, SPSS memungkinkan pengguna untuk melakukan berbagai analisis statistik tanpa perlu pengetahuan mendalam tentang pemrograman (Field, 2018).

SPSS menyediakan fitur untuk analisis deskriptif, uji hipotesis, regresi, analisis faktor, dan banyak lagi (Green & Salkind, 2017). Selain itu, SPSS juga mendukung visualisasi data seperti grafik batang, histogram, dan plot scatter (Wagner, 2019). Kelebihan dari program SPSS ini adalah penggunaannya yang mudah bagi pemula, mendukung berbagai format data, dan memiliki dokumentasi serta dukungan komunitas yang luas (Pallant, 2020).

B. R

R adalah bahasa pemrograman dan lingkungan perangkat lunak yang digunakan untuk analisis statistik dan grafis (R Core Team, 2020). R sangat populer di kalangan akademisi dan peneliti yang membutuhkan fleksibilitas dan kemampuan untuk melakukan analisis statistik yang lebih kompleks (Wickham, 2016).

R memiliki ribuan paket (libraries) yang dapat digunakan untuk hampir semua jenis analisis statistik, termasuk analisis regresi, analisis multivariat, analisis data waktu, dan machine learning (James et al., 2013). R juga dikenal karena kemampuannya dalam membuat visualisasi data yang sangat baik, seperti dengan menggunakan paket ggplot2 (Wickham, 2016). R ini sangat fleksibel, open-source (gratis), dan terus berkembang dengan kontribusi dari komunitas global (Peng, 2021).

C. SAS (*Statistical Analysis System*)

SAS adalah perangkat lunak statistik yang banyak digunakan di industri, pemerintah, dan akademisi untuk analisis data yang kompleks, manajemen data, dan pelaporan (SAS Institute Inc., 2020). SAS menawarkan berbagai modul untuk analisis statistik, data mining, analisis prediktif, serta manajemen data dalam jumlah besar (Cody, 2018). SAS juga menyediakan antarmuka pengguna yang berbasis GUI (Graphical User Interface) dan berbasis kode (Dilorio, 2016). SAS ini sangat kuat dalam manajemen data, memiliki dukungan pelanggan yang solid, dan sering digunakan dalam industri yang membutuhkan analisis data skala besar (Lora, 2019).

D. Stata

Stata adalah perangkat lunak statistik yang sering digunakan dalam bidang ekonomi, sosiologi, dan kesehatan masyarakat (Acock, 2018). Stata dikenal karena kemampuannya dalam menangani data panel dan analisis ekonometrika (Rabe-Hesketh & Skrondal, 2022).

Stata menawarkan alat untuk analisis data lintas-seksional, data panel, regresi, analisis survival, dan berbagai uji statistik lainnya (Hamilton, 2013). Stata juga memiliki antarmuka pengguna yang intuitif dan mendukung skrip otomatisasi (Cameron & Trivedi, 2010). Stata ini juga dikenal untuk analisis data panel, dokumentasi yang sangat baik, dan kinerja yang cepat bahkan dengan dataset yang besar (Long & Freese, 2014).

Setiap perangkat lunak statistik memiliki kelebihan dan kekurangan masing-masing, serta ditujukan untuk kebutuhan dan tingkat keahlian yang berbeda-beda. Pemilihan perangkat lunak yang tepat sangat bergantung pada jenis analisis yang akan dilakukan, volume data, serta preferensi dan keterampilan pengguna (Acock, 2018).

5.4 Kriteria Pemilihan Perangkat Lunak

Saat memilih perangkat lunak statistik untuk analisis pendidikan, beberapa faktor penting perlu dipertimbangkan untuk memastikan alat yang dipilih sesuai dengan kebutuhan dan konteks analisis. Berikut adalah faktor-faktor yang harus diperhatikan:

1. Tujuan dan Jenis Analisis. Pertimbangkan jenis analisis yang akan dilakukan, seperti analisis deskriptif, inferensial, regresi, atau analisis multivariat. Beberapa perangkat lunak lebih cocok untuk analisis tertentu; misalnya, Stata unggul dalam analisis data panel, sementara R lebih fleksibel untuk berbagai jenis analisis.
2. Kompleksitas Analisis. Perangkat lunak tertentu mungkin lebih cocok untuk analisis sederhana atau lanjutan. Jika analisis yang diperlukan kompleks, seperti model statistik canggih atau analisis data besar, perangkat lunak seperti SAS atau R mungkin lebih sesuai.
3. Kemudahan Penggunaan. Pilih perangkat lunak dengan antarmuka yang sesuai dengan tingkat kenyamanan pengguna. SPSS, misalnya, terkenal dengan antarmuka grafis yang user-friendly, sementara R membutuhkan pemahaman pemrograman.
4. Kurva Pembelajaran. Pertimbangkan tingkat kesulitan dalam mempelajari perangkat lunak. Bagi pengguna yang tidak berpengalaman, perangkat lunak dengan kurva pembelajaran yang rendah, seperti SPSS atau Minitab, mungkin lebih tepat.
5. Dukungan dan Sumber Daya. Perangkat lunak dengan komunitas pengguna yang aktif dan sumber daya pendukung yang banyak (tutorial, forum, dokumentasi) sangat membantu, terutama bagi pemula. R dan Python, misalnya, memiliki komunitas besar yang terus mengembangkan paket dan memberikan dukungan.
6. Ketersediaan Pelatihan. Pastikan ada cukup sumber daya pelatihan, baik online maupun offline, untuk perangkat lunak yang dipilih, terutama jika organisasi memerlukan pelatihan bagi staf atau peneliti.

7. Kompatibilitas dan Integrasi. Pastikan perangkat lunak dapat dijalankan pada sistem operasi yang digunakan (Windows, Mac, Linux). Beberapa perangkat lunak, seperti SAS, mungkin lebih kompatibel dengan sistem tertentu.
8. Integrasi dengan Perangkat Lain. Pertimbangkan apakah perangkat lunak dapat dengan mudah diintegrasikan dengan alat atau platform lain yang digunakan untuk pengumpulan dan pengolahan data, seperti Excel atau database institusi.
9. Harga Perangkat Lunak. Biaya lisensi perangkat lunak dapat menjadi faktor penting, terutama dalam konteks pendidikan di mana anggaran mungkin terbatas. Perangkat lunak open-source seperti R atau Python adalah pilihan populer karena gratis, sementara SPSS dan SAS memerlukan biaya lisensi yang signifikan.
10. Biaya Pelatihan dan Dukungan. Selain biaya lisensi, pertimbangkan biaya tambahan yang mungkin diperlukan untuk pelatihan, dukungan teknis, atau pemeliharaan perangkat lunak.
11. Keandalan dan Reputasi. Pilih perangkat lunak yang telah terbukti andal dan memiliki reputasi baik di kalangan peneliti pendidikan. Perangkat lunak yang telah lama digunakan dan diakui di bidang akademik atau industri biasanya lebih terpercaya.
12. Kemampuan Penanganan Data Besar. Jika analisis melibatkan data dalam jumlah besar, pastikan perangkat lunak dapat menangani volume data tersebut tanpa masalah kinerja.
13. Fleksibilitas dalam Analisis. Pilih perangkat lunak yang fleksibel dan dapat disesuaikan dengan berbagai jenis analisis dan dataset. R, misalnya, sangat fleksibel dengan ribuan paket tambahan yang dapat disesuaikan dengan kebutuhan spesifik.
14. Kemampuan untuk Mengembangkan dan Memodifikasi. Jika ada kebutuhan untuk mengembangkan analisis khusus atau menyesuaikan algoritma, perangkat lunak seperti R atau Python mungkin lebih cocok karena mendukung scripting dan pemrograman kustom.

Dengan mempertimbangkan faktor-faktor ini, pengguna dapat memilih perangkat lunak statistik yang paling sesuai dengan kebutuhan analisis pendidikan mereka, memastikan hasil yang lebih akurat dan relevan.

5.5 Tantangan dan Solusi dalam Penggunaan Perangkat Lunak Statistik

Tantangan teknis dalam penggunaan perangkat lunak statistik dalam pendidikan sering kali berkaitan dengan kesulitan dalam instalasi, kompatibilitas sistem, dan kompleksitas dalam penggunaannya, terutama bagi pengguna yang tidak memiliki latar belakang teknis yang kuat (Field, 2018). Beberapa perangkat lunak statistik memiliki antarmuka yang rumit atau memerlukan pemahaman tentang pemrograman, yang dapat menjadi penghalang bagi pendidik dan peneliti yang tidak terbiasa dengan teknologi ini (Pallant, 2020). Selain itu, kendala teknis lain seperti keterbatasan perangkat keras atau masalah lisensi juga bisa menjadi tantangan yang signifikan (Wickham, 2016).

Tantangan non-teknis mencakup hambatan seperti kurangnya keterampilan statistik di kalangan pendidik atau peneliti, resistensi terhadap perubahan teknologi, dan keterbatasan waktu untuk mempelajari perangkat lunak baru (Heiberger & Holland, 2015). Kurangnya dukungan atau sumber daya untuk pelatihan juga dapat memperburuk situasi ini, menghambat penerapan perangkat lunak statistik dalam pendidikan (Mertler, 2021).

Solusi untuk mengatasi tantangan ini melibatkan pelatihan yang tepat dan berkelanjutan, yang dapat membantu pengguna memahami dan menguasai perangkat lunak statistik yang mereka gunakan (Pallant, 2020). Selain itu, dukungan teknis yang kuat, baik dari tim IT internal maupun vendor perangkat lunak, sangat penting untuk membantu pengguna mengatasi masalah teknis yang mereka hadapi (Heiberger & Holland, 2015). Bergabung dengan komunitas pengguna perangkat lunak, baik melalui forum online atau kelompok studi lokal, juga dapat memberikan dukungan tambahan, termasuk

berbagi pengetahuan dan pengalaman praktis (Wickham, 2016). Dengan adanya dukungan ini, tantangan teknis dan non-teknis dapat diatasi, memungkinkan penggunaan perangkat lunak statistik yang lebih efektif dalam pendidikan.

5.6 Kesimpulan

Perangkat lunak statistik memainkan peran krusial dalam penelitian dan analisis pendidikan dengan memungkinkan pengolahan dan interpretasi data yang lebih efisien dan akurat. Alat ini mendukung analisis kompleks seperti statistik deskriptif, uji hipotesis, dan model statistik canggih yang penting untuk mengevaluasi metode pengajaran dan kinerja institusi pendidikan (Loeb, 2017; Cheung et al., 2020). Penggunaan perangkat lunak statistik tidak hanya meningkatkan kecepatan dan akurasi analisis tetapi juga mendukung transparansi dan replikasi penelitian (Field, 2018; Mertler, 2021). Dalam era digital, keterampilan dalam perangkat lunak statistik menjadi kebutuhan esensial bagi profesional pendidikan untuk mendorong peningkatan kualitas pembelajaran (Schneider & Hutt, 2014; Alderson, 2020).

Perangkat lunak statistik akan terus berkembang dengan kemajuan teknologi, menawarkan fitur yang lebih canggih dan antarmuka yang lebih user-friendly (Mertler, 2021). Arah perkembangan teknologi yang mendukung analisis data pendidikan termasuk peningkatan kemampuan untuk menangani data besar, integrasi dengan alat pembelajaran lainnya, dan pengembangan teknik analisis berbasis kecerdasan buatan (AI) dan machine learning (Siemens & Baker, 2012; Judd, 2019). Selain itu, ada tren yang meningkat dalam penyediaan solusi berbasis cloud yang memudahkan kolaborasi dan akses data secara real-time (Wagner, 2019).

Kemajuan ini akan memungkinkan pendidik dan peneliti untuk melakukan analisis yang lebih mendalam dan relevan, mendukung pengambilan keputusan yang lebih tepat dan berbasis data, serta berkontribusi pada peningkatan kualitas pendidikan secara keseluruhan (Creswell & Creswell, 2018; Schneider & Hutt, 2014). Dengan pelatihan yang tepat dan

dukungan teknis yang memadai, tantangan dalam penggunaan perangkat lunak statistik dapat diatasi, memaksimalkan potensi alat ini dalam dunia pendidikan.

DAFTAR PUSTAKA

- Acock, A. C. (2018). *A gentle introduction to Stata* (6th ed.). Stata Press.
- Alderson, J. C. (2020). *Assessing the quality of language teaching: An overview*. Routledge.
- Cameron, A. C., & Trivedi, P. K. (2010). *Microeconometrics using Stata*. Stata Press.
- Cheung, A. C. K., Chan, A. K. W., & Tsang, A. Y. H. (2020). *Introduction to statistical methods for education and psychology*. Wiley.
- Cody, R. P. (2018). *SAS statistics by example*. SAS Institute Inc.
- Creswell, J. W., & Creswell, J. D. (2018). *Research design: Qualitative, quantitative, and mixed methods approaches* (5th ed.). SAGE Publications.
- Dilorio, R. (2016). *SAS for data analysis and visualization*. Wiley.
- Field, A. (2018). *Discovering statistics using IBM SPSS Statistics* (5th ed.). SAGE Publications.
- Green, S. B., & Salkind, N. J. (2017). *Using SPSS for Windows and Macintosh: Analyzing and understanding data* (8th ed.). Pearson.
- Hamilton, L. C. (2013). *Statistics with STATA: Version 12*. Cengage Learning.
- Heiberger, R. M., & Holland, B. (2015). *Statistical analysis and data display: An intermediate course with examples in R*. Springer.
- James, G., Witten, D., Hastie, T., & Tibshirani, R. (2013). *An introduction to statistical learning: With applications in R*. Springer.
- Judd, C. M. (2019). *Introduction to data analysis in research*. SAGE Publications.
- Loeb, S. (2017). *Creating a culture of continuous improvement in education*. Stanford University Press.
- Long, J. S., & Freese, J. (2014). *Regression models for categorical dependent variables using Stata* (3rd ed.). Stata Press.
- Lora, J. (2019). *SAS for big data: Managing large-scale datasets with SAS*. Wiley.

- Mertler, C. A. (2021). *Introduction to educational research*. SAGE Publications.
- Norris, D. M., Mason, M., & McPherson, M. (2015). *Validating research findings in education*. Routledge.
- Pallant, J. (2020). *SPSS survival manual: A step by step guide to data analysis using IBM SPSS (7th ed.)*. Open University Press.
- Peng, R. D. (2021). *R for everyone: Advanced analytics and graphics*. Pearson.
- R Core Team. (2020). *R: A language and environment for statistical computing*. R Foundation for Statistical Computing.
- Rabe-Hesketh, S., & Skrondal, A. (2022). *Multilevel and longitudinal modeling using Stata (3rd ed.)*. Stata Press.
- SAS Institute Inc. (2020). *SAS system for statistical analysis*. SAS Institute Inc.
- Schneider, J., & Hutt, E. (2014). *Data-driven decision making in education*. Routledge.
- Siemens, G., & Baker, R. S. J. D. (2012). *Learning analytics and educational data mining: Towards communication and collaboration*. *Proceedings of the 2nd International Conference on Learning Analytics and Knowledge*.
- Urdu, T. C. (2017). *Statistics in education research: An introduction*. Routledge.
- Wagner, D. (2019). *The future of educational technology: Data and decision-making*. Routledge.
- Wang, X., & Wu, T. (2020). *Applied statistical methods in education*. Wiley.
- Wickham, H. (2016). *ggplot2: Elegant graphics for data analysis (2nd ed.)*. Springer.

BIODATA PENULIS



Hetty Elfina, S.Pd., M.Pd

Dosen Program Studi Sistem Informasi
Fakultas Teknik dan Ilmu Komputer
Universitas Pembinaan Masyarakat Indonesia
Medan

Penulis lahir di kota Sibolga pada tanggal 07 Januari 1990. Penulis merupakan dosen tetap pada Program Studi Sistem Informasi, Fakultas Teknik dan Ilmu Komputer, Universitas Pembinaan Masyarakat Indonesia. Penulis menyelesaikan pendidikan S1 pada Jurusan Pendidikan Matematika di Universitas Negeri Medan tahun 2013 dan melanjutkan S2 pada tahun yang sama di Jurusan Pendidikan Matematika di Universitas Negeri Medan dan lulus tahun 2015. Saat ini penulis sedang menempuh kuliah S3 Prodi Pendidikan Matematika di Universitas Negeri Medan. Penulis telah menulis 6 buku, yaitu Kalkulus, Statistika Pendidikan, Teori Bilangan, Metodologi Penelitian, Statistik, dan Bagian-Bagian Matematika Dasar.

BIODATA PENULIS



Dr.Ir. Patrich Phill Edrich Papilaya M.Sc.Forest.Trop

Dosen Program Pascasarjana Program Studi Manajemen Hutan
Universitas Pattimura

Penulis lahir di Ambon Tanggal 14 Pebruari 1967. Penulis adalah dosen tetap pada Program Studi Manajemen Hutan Pascasarjana. Universitas Pattimura. Ambon. Menyelesaikan pendidikan S1 pada Jurusan Kehutanan Fakultas Pertanian Universitas Pattimura. menyelesaikan Pendidikan S2 pada Georg-August Universitas. Faculty of Forestry. Göttingen. dalam bidang Remote sensing dan GIS. Pendidikan S3 diselesaikan pada IPB University. Bogor

Penulis menjadi bagian dalam penulisan beberapa Book Chapter; antara lain : Metodologi Penelitian Eksperimen tahun 2023; Memahami Konsep Dan Teori Ekologi Dan Lingkungan tahun 2024; Metodologi Penelitian tahun 2024 dan Dasar Biostatistika Untuk Peneliti tahun 2024.

BIODATA PENULIS



Muh. Khaedir Lutfi, S.Pd., Gr., M.Pd.

Dosen Program Studi Pendidikan Matematika
Fakultas Keguruan dan Ilmu Pendidikan
Universitas Tangerang Raya

Penulis lahir di Ujung Pandang, 28 April 1991. Penulis adalah dosen tetap pada Program Studi Pendidikan Matematika Fakultas Keguruan dan Ilmu Pendidikan, Universitas Tangerang Raya. Penulis menyelesaikan pendidikan S1 pada Program Studi Pendidikan Matematika di Universitas Muhammadiyah Makassar tahun 2014 dan PPG di Universitas Islam Nusantara pada tahun 2017. Penulis menyelesaikan S2 di Universitas Pendidikan Indonesia pada Program Studi Pendidikan Matematika. Salah satu matakuliah yang diampu oleh penulis adalah matakuliah Statistika. Penulis juga aktif dalam mempublikasikan karya ilmiah yang beberapa di antaranya menggunakan uji statistik.

BIODATA PENULIS

Fitri Anisa Kusumastuti, M.Pd

Dosen Program Studi Pendidikan Matematika
Fakultas Keguruan dan Ilmu Pendidikan Universitas Tangerang
Raya

Penulis lahir di Wonogiri tanggal 29 Maret 1995. Penulis adalah dosen tetap pada Program Studi Pendidikan Matematika Fakultas Keguruan dan Ilmu Pendidikan, Universitas Tangerang Raya. Menyelesaikan pendidikan S1 pada Jurusan Pendidikan Matematika di Universitas Negeri Surabaya dan melanjutkan S2 pada Jurusan Pendidikan Matematika Universitas Pendidikan Indonesia. Penulis menekuni bidang literasi-numerasi.

BIODATA PENULIS



Dr. Isry Laila Syathroh, M.Pd.

Dosen Program Studi Bahasa Inggris
Fakultas Pendidikan Bahasa – IKIP Siliwangi

Penulis memperoleh gelar Doktor Pendidikan Bahasa Inggris dari Universitas Pendidikan Indonesia pada tahun 2021. Belasan buku antologi fiksi pernah beliau tulis di komunitas Rumah Antologi Indonesia. Beliau juga adalah salah satu penulis buku *“Life Today”*, buku ajar bahasa Inggris nasional fase F, berdasarkan Kurikulum Merdeka, yang diterbitkan pada 2022 oleh Kementerian Pendidikan dan Kebudayaan Republik Indonesia. Penulis dapat dihubungi melalui e-mail: islaisya@ikipsiliwangi.ac.id