

Statistics Formula Compendium

Rune Skov Christensen

Sebastian Hvid Olavsgaard

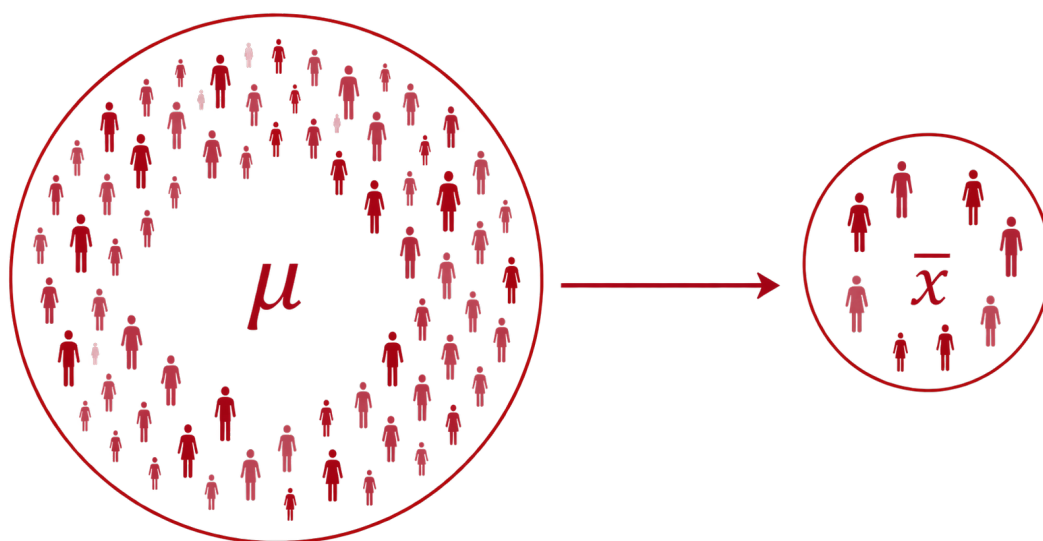


Table of Contents

1	Descriptive statistics	1
2	Probability distributions	3
	Discrete distributions	3
	General discrete random variable formulas	3
	Continuous distributions	6
	General continuous random variable formulas	6
3	Linear functions and relationships of random variables	11
	Rules for linear functions and combinations of random variables	11
	Independence and properties of random variables	12
	Measures of linear relationship between two random variables	12
4	Classic sample inference	13
	Hypothesis tests and confidence intervals	13
	One-sample	13
	Two-samples	15
	Paired	15
	Welch	16
	Pooled	18
	Statistical errors	20
	Type I error	20
	Type II error	20
5	Simulation-based inference and error propagation	21
	Python commands for simulation	21
6	Linear regression	23
	Matrix formulation for linear regression	27
7	Inference for proportions	29
	One-proportion	29
	Two-proportions	31
	Contingency-tables	33
8	One-way ANOVA	35
	Post-hoc pairwise tests	37
9	Two-way ANOVA	39
	Group hypothesis testing	41
	Block hypothesis testing	41
	Post-hoc pairwise tests	42
10	Model diagnostics and residual analysis	43
11	Notation and terminology	45
12	Python libraries	47

1 Descriptive statistics

Where n denotes the number of observations in the data sample, and x_i denotes the i th observation.

Collect data in Python

```
x = np.array([...])
```

	Description	Formula	Python
	Sample average:	$\bar{x} = \frac{1}{n} \sum_{i=1}^n x_i$	<code>np.mean(x)</code>
	Sample variance:	$s^2 = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2$	<code>np.var(x, ddof=1)</code>
	Sample standard deviation:	$s = \sqrt{s^2} = \sqrt{\frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2}$	<code>np.std(x, ddof=1)</code>
	Coefficient of variation:	$CV = \frac{s}{\bar{x}}$	<code>np.std(x, ddof=1) / np.mean(x)</code>
	Median:	$Q_2 = \begin{cases} x_{(\frac{n+1}{2})}, & \text{if } n \text{ is odd} \\ \frac{x_{(\frac{n}{2})} + x_{(\frac{n}{2}+1)}}{2}, & \text{if } n \text{ is even} \end{cases}$	<code>np.median(x)</code>
	Quantile/percentile:	$q_p = \begin{cases} \frac{x_{(np)} + x_{(np+1)}}{2}, & \text{if } np \text{ is integer} \\ x_{(\lceil np \rceil)}, & \text{if } np \text{ is not integer} \end{cases}$	<code>np.quantile(x, q)</code> where $q \in [0, 1]$ Alternatively: <code>np.percentile(x, q)</code> where $q \in [0, 100]$

	Description	Formula	Python
	Quartiles:	$Q_0 = q_{0.00} = \text{minimum}$ $Q_1 = q_{0.25} = \text{lower quartile}$ $Q_2 = q_{0.50} = \text{median}$ $Q_3 = q_{0.75} = \text{upper quartile}$ $Q_4 = q_{1.00} = \text{maximum}$	<pre>np.quantile(x, 0.00) np.quantile(x, 0.25) np.quantile(x, 0.50) np.quantile(x, 0.75) np.quantile(x, 1.00)</pre>
	Interquartile range:	$IQR = Q_3 - Q_1$	<pre>stats.iqr(x)</pre>
	Range:	$\text{Range} = \max - \min = Q_4 - Q_0$	<pre>np.ptp(x)</pre>
	Sample covariance: Also denoted: $Cov(x, y)$	$s_{xy} = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y})$	<pre>y = np.array([...]) np.cov(x,y,ddof=1)[0,1]</pre>
	Sample correlation coefficient (r): Note: $r = \hat{\rho}(x, y) = \hat{\rho}_{xy}$ $r \in [-1, 1]$	$r = \frac{1}{n-1} \sum_{i=1}^n \frac{(x_i - \bar{x})}{s_x} \frac{(y_i - \bar{y})}{s_y} = \frac{s_{xy}}{s_x s_y}$	<pre>np.corrcoef(x, y)[0,1]</pre>

2 Probability distributions

Discrete distributions

The random variable takes values in a finite or countably infinite set of possible outcomes.

General discrete random variable formulas

	Description	Formula	Python
	Probability density function (pdf): Also known as the <i>probability mass function</i> (pmf).	$f(x) = P(X = x)$ Note: $f(x) \geq 0$ and $\sum_{\text{all } x} f(x) = 1$	<code>stats.binom.pmf()</code> <code>stats.hypergeom.pmf()</code> <code>stats.poisson.pmf()</code>
	Cumulative distribution function (cdf):	$F(x) = P(X \leq x) = \sum_{t \leq x} P(X = t)$	<code>stats.binom.cdf()</code> <code>stats.hypergeom.cdf()</code> <code>stats.poisson.cdf()</code>
	Mean of a discrete random variable	$\mu = E[X] = \sum_{\text{all } x} x f(x)$	<code>stats.binom.mean()</code> <code>stats.hypergeom.mean()</code> <code>stats.poisson.mean()</code>
	Variance of a discrete random variable X:	$\sigma^2 = V[X] = E[(X - \mu)^2] = \sum_{\text{all } x} (x - \mu)^2 f(x)$	<code>stats.binom.var()</code> <code>stats.hypergeom.var()</code> <code>stats.poisson.var()</code>

Combinatorics

	Description	Formula	Notes
	Binomial coefficient: Number of ways to choose b items from a total set of a items without regard to order.	$\binom{a}{b} = \frac{a!}{b!(a-b)!}$	

Binomial distribution $X \sim B(n, p)$

Describes the number of successes X in n independent trials, where each trial has two possible outcomes: success with probability p and failure with probability $1 - p$. The probability p is the same for every trial.

$f(x)$	$F(x)$	$E[X]$	$V[X]$
$\binom{n}{x} p^x (1-p)^{n-x}$	$\sum_{k=0}^x \binom{n}{k} p^k (1-p)^{n-k}$	np	$np(1-p)$

Python functions

```
stats.binom
.pmf(k, n, p)
.cdf(k, n, p)
.mean(n, p)
.var(n, p)
.std(n, p)
.ppf(q, n, p)
.rvs(n, p, size)
```

Notation

Book	Python	Meaning
x	k	number of successes
n	n	number of trials
p	p	probability of success

Hypergeometric distribution $X \sim H(n, a, N)$

Describes the probability of obtaining a certain number of successes X when sampling without replacement from a finite population N that contains a fixed number of successes a and failures $N - a$.

$f(x)$	$F(x)$	$E[X]$	$V[X]$
$\frac{\binom{a}{x} \binom{N-a}{n-x}}{\binom{N}{n}}$	$\sum_{k=0}^{\lfloor x \rfloor} \frac{\binom{a}{k} \binom{N-a}{n-k}}{\binom{N}{n}}$	$n \frac{a}{N}$	$n \frac{a}{N} \frac{N-a}{N} \frac{N-n}{N-1}$

Python functions

```
stats.hypergeom
.pmf(k, M, n, N)
.cdf(k, M, n, N)
.mean(M, n, N)
.var(M, n, N)
.std(M, n, N)
.ppf(q, M, n, N)
.rvs(M, n, N, size)
```

Notation

Book	Python	Meaning
x	k	number of successes in sample
N	M	population size
a	n	number of successes in pop.
n	N	number of draws

Poisson distribution $X \sim Po(\lambda)$

Describes the number of events X that occur in a given time period or region where the events occur independently and with a constant average rate λ . It can be interpreted as events per unit of time or as a count per interval.

$f(x)$	$F(x)$	$E[X]$	$V[X]$
$\frac{\lambda^x}{x!} e^{-\lambda}$	$\sum_{k=0}^{\lfloor x \rfloor} \frac{\lambda^k}{k!} e^{-\lambda}$	λ	λ

Python functions

```
stats.poisson
.pmf(k, mu)
.cdf(k, mu)
.mean(mu)
.var(mu)
.std(mu)
.ppf(q, mu)
.rvs(mu, size)
```

Notation

Book	Python	Meaning
x	k	number of occurrences
λ	mu	average rate

Continuous distributions

The random variable can take values over a continuum of possible outcomes.

General continuous random variable formulas

	Description	Formula	Python
	Probability density function (pdf):	$\int_a^b f(x) dx = P(a < X \leq b)$ Note: $f(x) \geq 0$ and $\int_{-\infty}^{\infty} f(x) dx = 1$	<pre>stats.norm.pdf() stats.uniform.pdf() stats.expon.pdf() stats.lognorm.pdf() stats.chi2.pdf() stats.t.pdf() stats.f.pdf()</pre>
	Cumulative distribution function (cdf):	$F(x) = P(X \leq x) = \int_{-\infty}^x f(u) du$	<pre>stats.norm.cdf() stats.uniform.cdf() stats.expon.cdf() stats.lognorm.cdf() stats.chi2.cdf() stats.t.cdf() stats.f.cdf()</pre>
	Mean of a continuous random variable X:	$\mu = E[X] = \int_{-\infty}^{\infty} x f(x) dx$	<pre>stats.norm.mean() stats.uniform.mean() stats.expon.mean() stats.lognorm.mean() stats.chi2.mean() stats.t.mean() stats.f.mean()</pre>
	Variance of a continuous random variable X:	$\sigma^2 = E[(X - \mu)^2] = \int_{-\infty}^{\infty} (x - \mu)^2 f(x) dx$	<pre>stats.norm.var() stats.uniform.var() stats.expon.var() stats.lognorm.var() stats.chi2.var() stats.t.var() stats.f.var()</pre>

Normal distribution $X \sim \mathcal{N}(\mu, \sigma^2)$

Describes a random variable X that follows a characteristic bell-shaped distribution. The distribution is symmetric about the mean and has spread σ .

$f(x)$	$F(x)$	$E[X]$	$V[X]$
$\frac{1}{\sqrt{2\pi\sigma^2}} \exp\left(-\frac{(x-\mu)^2}{2\sigma^2}\right)$	$\int_{-\infty}^x f(u) du$	μ	σ^2

Python functions

```
stats.norm
.pdf(x, loc, scale)
.cdf(x, loc, scale)
.mean(loc, scale)
.var(loc, scale)
.std(loc, scale)
.ppf(q, loc, scale)
.rvs(loc, scale, size)
```

Notation

Book	Python	Meaning
x	<code>x</code>	observed value
μ	<code>loc</code>	mean
σ	<code>scale</code>	standard deviation

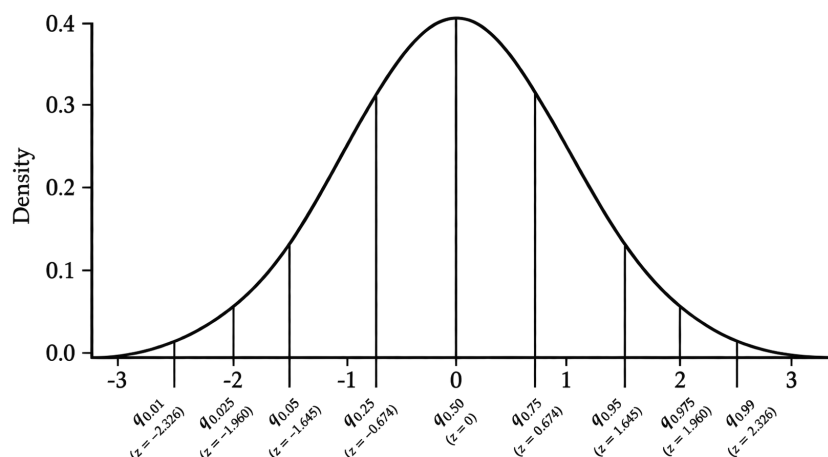
Standard normal distribution

A stochastic variable X from a given distribution can be transformed (standardized) into a standard normal distribution variable via:

$$Z = \frac{X - \mu}{\sigma} \sim N(0, 1^2)$$

Important percentiles and quantiles of the standard normal distribution

Percentile	1%	2.275%	2.5%	5%	15.866%	25%	50%	75%	84.134%	95%	97.5%	97.725%	99%
Quantile	0.01	0.0228	0.025	0.05	0.1587	0.25	0.50	0.75	0.8413	0.95	0.975	0.9772	0.99
Z-score	-2.326	-2.000	-1.960	-1.645	-1.000	-0.674	0.000	0.674	1.000	1.645	1.960	2.000	2.326



Notable quantiles of the standard normal distribution

Lognormal distribution $X \sim \text{LN}(\alpha, \beta^2)$

Describes a positive random variable X whose natural logarithm follows a normal distribution. The distribution is asymmetric and right-skewed.

$f(x)$	$F(x)$	$E[X]$	$V[X]$
$\frac{1}{x\beta\sqrt{2\pi}} \exp\left(-\frac{(\ln(x) - \alpha)^2}{2\beta^2}\right), x > 0$	$\int_0^x f(u) du, x > 0$	$e^{\alpha + \frac{\beta^2}{2}}$	$e^{2\alpha + \beta^2} (e^{\beta^2} - 1)$

Python functions

```
stats.lognorm
.pdf(x, s, loc, scale)
.cdf(x, s, loc, scale)
.mean(s, loc, scale)
.var(s, loc, scale)
.std(s, loc, scale)
.ppf(q, s, loc, scale)
.rvs(s, loc, scale, size)
```

Notation

Book	Python	Meaning
x	x	observed value
e^α	scale	scale parameter
β	s	std of $\ln(X)$
α		mean of $\ln(X)$

Transformation from Lognormal to Normal

A lognormal distributed variable $X \sim \text{LN}(\alpha, \beta^2)$ can be transformed to a normal distributed random variable via: $Y = \ln(X) \sim N(\alpha, \beta^2)$

Exponential distribution $X \sim \text{Exp}(\lambda)$

Describes a positive random variable X representing waiting times or lifetimes. It is commonly used to model the time between events in a Poisson process.

$f(x)$	$F(x)$	$E[X]$	$V[X]$
$\lambda e^{-\lambda x}, x \geq 0$	$1 - e^{-\lambda x}, x \geq 0$	$\frac{1}{\lambda}$	$\frac{1}{\lambda^2}$

Python functions

```
stats.expon
.pdf(x, loc, scale)
.cdf(x, loc, scale)
.mean(loc, scale)
.var(loc, scale)
.std(loc, scale)
.ppf(q, loc, scale)
.rvs(loc, scale, size)
```

Notation

Book	Python	Meaning
x	x	waiting time
$1/\lambda$	scale	mean waiting time
λ		rate parameter

Uniform distribution $X \sim U(\alpha, \beta)$

Describes a random variable X whose values are equally likely within the interval $[\alpha, \beta]$.

$f(x)$	$F(x)$	$E[X]$	$V[X]$
$\begin{cases} 0, & x < \alpha \\ \frac{1}{\beta - \alpha}, & \alpha \leq x \leq \beta \\ 0, & x > \beta \end{cases}$	$\begin{cases} 0, & x < \alpha \\ \frac{x - \alpha}{\beta - \alpha}, & \alpha \leq x \leq \beta \\ 1, & x > \beta \end{cases}$	$\frac{\alpha + \beta}{2}$	$\frac{(\beta - \alpha)^2}{12}$

Python functions

```
stats.uniform
.rvs(loc, scale, size)
.pdf(x, loc, scale)
.cdf(x, loc, scale)
.mean(loc, scale)
.var(loc, scale)
.std(loc, scale)
.ppf(q, loc, scale)
```

Notation

Book	Python	Meaning
x	x	observed value
α	loc	lower limit
β	loc + scale	upper limit

Chi-squared distribution $X \sim \chi^2(\nu)$

Describes a random variable that can be expressed as $X = \sum_{i=1}^{\nu} Z_i^2$, the sum of ν squared independent standard normally distributed variables. In the context of a normally distributed sample, it also arises as $X = \frac{(n-1)S^2}{\sigma^2} \sim \chi^2(n-1)$, where S^2 is the sample variance and σ^2 is the population variance. It is only defined for non-negative values and is *right-skewed*, with its shape depending on the degrees of freedom ν .

$f(x)$	$F(x)$	$E[X]$	$V[X]$
$\frac{1}{2^{\nu/2}\Gamma(\nu/2)} x^{\frac{\nu}{2}-1} e^{-x/2}, x \geq 0$	$\int_0^x f(u) du, \quad x \geq 0$	ν	2ν

Python functions

```
stats.chi2
.pdf(x, df, loc, scale)
.cdf(x, df, loc, scale)
.mean(df, loc, scale)
.var(df, loc, scale)
.std(df, loc, scale)
.ppf(q, df, loc, scale)
.rvs(df, loc, scale, size)
```

Notation

Book	Python	Meaning
x	x	observed value
ν	df	degrees of freedom
$\Gamma(I)$		gamma function

Student's t -distribution $X \sim t(\nu)$

Describes a random variable $X = \frac{Z}{\sqrt{Y/\nu}}$, where $Z \sim \mathcal{N}(0, 1)$ and $Y \sim \chi^2(\nu)$ are independent. In statistical inference, it arises as $X = \frac{\bar{X} - \mu}{S/\sqrt{n}} \sim t(n-1)$ for an i.i.d. normal sample. It is symmetric about 0 with heavier tails than the standard normal, with shape depending on ν .

$f(x)$	$F(x)$	$E[X]$	$V[X]$
$\frac{\Gamma(\frac{\nu+1}{2})}{\sqrt{\nu\pi}\Gamma(\frac{\nu}{2})} \left(1 + \frac{x^2}{\nu}\right)^{-\frac{\nu+1}{2}}$	$\int_{-\infty}^x f(u) du$	0, $\nu > 1$	$\frac{\nu}{\nu-2}, \nu > 2$

Python functions

```
stats.t
.pdf(x, df, loc, scale)
.cdf(x, df, loc, scale)
.mean(df, loc, scale)
.var(df, loc, scale)
.std(df, loc, scale)
.ppf(q, df, loc, scale)
.rvs(df, loc, scale, size)
```

Notation

Book	Python	Meaning
x	x	observed value
ν	df	degrees of freedom
$\Gamma(I)$		gamma function

F-distribution $X \sim F(\nu_1, \nu_2)$

Describes a random variable $X = \frac{S_1^2/\sigma_1^2}{S_2^2/\sigma_2^2} \sim F(n_1-1, n_2-1)$ for two independent normal samples. More generally, $X = \frac{U/\nu_1}{V/\nu_2}$, where $U \sim \chi^2(\nu_1)$ and $V \sim \chi^2(\nu_2)$ are independent. It is non-negative and right-skewed, with shape depending on ν_1, ν_2 .

$f(x)$	$F(x)$	$E[X]$	$V[X]$
$\frac{1}{B(\frac{\nu_1}{2}, \frac{\nu_2}{2})} \left(\frac{\nu_1}{\nu_2}\right)^{\nu_1/2} \times x^{\frac{\nu_1}{2}-1} \left(1 + \frac{\nu_1}{\nu_2}x\right)^{-\frac{\nu_1+\nu_2}{2}},$	$\int_0^x f(u) du, x \geq 0$	$\frac{\nu_2}{\nu_2-2}, \nu_2 > 2$	$\frac{2\nu_2^2(\nu_1+\nu_2-2)}{\nu_1(\nu_2-2)^2(\nu_2-4)}, \nu_2 > 4$
$x \geq 0, \nu_1, \nu_2 > 0$			

Python functions

```
stats.f
.pdf(x, dfn, dfd, loc, scale)
.cdf(x, dfn, dfd, loc, scale)
.mean(dfn, dfd, loc, scale)
.var(dfn, dfd, loc, scale)
.std(dfn, dfd, loc, scale)
.ppf(q, dfn, dfd, loc, scale)
.rvs(dfn, dfd, loc, scale, size)
```

Notation

Book	Python	Meaning
x	x	observed value
ν_1	dfn	df numerator
ν_2	dfd	df denominator
$B(I_1, I_2)$		beta function

3 Linear functions and relationships of random variables

Rules for linear functions and combinations of random variables

The following rules apply to both discrete and continuous random variables. Let X and Y denote random variables.

	Description	Formula	Notes
	Mean of linear functions: $Y = aX + b.$	$E(Y) = E(aX + b) = aE(X) + b$	
	Variance of linear functions: $Y = aX + b.$	$V(Y) = V(aX + b) = a^2V(X)$ $V(X) = \frac{V(Y)}{a^2}$	
	Standard deviation of linear functions: $Y = aX + b.$	$\sigma_Y = \sqrt{V(Y)} = a \sigma_X$ $\sigma_X = \sqrt{V(X)} = \frac{\sigma_Y}{ a }$	
	Mean of linear combinations: $Y = a_1X_1 + \cdots + a_nX_n.$ for independent random variables.	$E(Y) = E(a_1X_1 + \cdots + a_nX_n)$ $= a_1E(X_1) + a_2E(X_2) + \cdots + a_nE(X_n)$	
	Variance of a linear combination: $Y = a_1X_1 + \cdots + a_nX_n.$ for independent random variables.	$V(Y) = V(a_1X_1 + a_2X_2 + \cdots + a_nX_n)$ $= a_1^2V(X_1) + a_2^2V(X_2) + \cdots + a_n^2V(X_n)$	

Independence and properties of random variables

	Description	Formula	Notes
	Independence of discrete random variables: Proof of independence between two random variables.	$P(X = x, Y = y) = P(X = x)P(Y = y)$	
	Independence of continuous random variables: Two continuous random variables X and Y are said to be independent if and only if	$f(x, y) = f(x)f(y)$	
	Properties of independent random variables: Requirement: X and Y are independent.	$E(XY) = E(X)E(Y)$ $Cov(X, Y) = 0$	

Measures of linear relationship between two random variables

	Description	Formula	Notes
	Covariance: Also denoted σ_{XY} . The covariance between X and Y .	$Cov(X, Y) = E[(X - E(X))(Y - E(Y))]$	
	Correlation coefficient: Measures the strength of the linear relationship between two stochastic variables X and Y .	$\rho(X, Y) = \frac{Cov(X, Y)}{\sigma_X \sigma_Y} \in [-1, 1]$	

4 Classic sample inference

Given a sample $\{x_1, x_2, \dots, x_n\}$ from a population with mean μ and variance σ^2 , sample statistics can be used to estimate population parameters. In particular, the Central Limit Theorem (CLT) states that, for sufficiently large $n \geq 30$, the sample mean is approximately normally distributed $\bar{X} = \frac{1}{n} \sum_{i=1}^n X_i \sim \mathcal{N}\left(\mu, \frac{\sigma^2}{n}\right)$. For already normally distributed variables $n \geq 30$ is not required. The CLT therefore provides the theoretical basis for the statistical inference methods introduced in the following section.

Hypothesis tests and confidence intervals

One-sample

Collect data in Python

```
x = np.array([...])
```

	Description	Formula	Python
	Variation of the mean ($\sigma_{\bar{x}}^2$):	$\hat{\sigma}_{\bar{x}}^2 = \frac{s^2}{n}$	
	Standard error of the mean ($\sigma_{\bar{x}}$):	$\hat{\sigma}_{\bar{x}} = \sqrt{\frac{s^2}{n}} = \frac{s}{\sqrt{n}}$	SEM = stats.sem(x)
	Confidence interval of the mean:	$\bar{x} \pm t_{1-\alpha/2, (n-1)} \frac{s}{\sqrt{n}}$	CI = stats.ttest_1samp(x, popmean= μ_0).confidence_interval(1- α)
	Margin of error:	$ME = t_{1-\alpha/2, (n-1)} \frac{s}{\sqrt{n}}$	
	Confidence interval for the variance and standard deviation: Rule of thumb: Only applies to normally distributed data; unlike mean CIs, the CLT does not extend to variance CIs.	$\sigma^2 : \left[\frac{(n-1)s^2}{\chi_{1-\alpha/2, (n-1)}^2}, \frac{(n-1)s^2}{\chi_{\alpha/2, (n-1)}^2} \right]$ $\sigma : \left[\sqrt{\frac{(n-1)s^2}{\chi_{1-\alpha/2, (n-1)}^2}}, \sqrt{\frac{(n-1)s^2}{\chi_{\alpha/2, (n-1)}^2}} \right]$	lower quantile: stats.chi2.ppf(1- $\alpha/2$, df=n-1) upper quantile: stats.chi2.ppf($\alpha/2$, df=n-1)

	Description	Formula	Python
	t-test statistic (t_{obs}): $H_0 : \mu = \mu_0$ $H_1 : \mu \neq \mu_0$	$t_{obs} = \frac{\bar{x} - \mu_0}{s/\sqrt{n}}$	<pre>t_obs, pval = stats.ttest_1samp(x, popmean=μ_0)</pre>
	Critical value (t_{crit}):	$t_{crit} = t_{1-\alpha/2, (n-1)}$ if $t_{obs} \notin [-t_{crit}, t_{crit}] \Rightarrow \text{reject } H_0$ else $\Rightarrow \text{accept } H_0$	<pre>crit = stats.t.ppf(1-$\alpha/2$, df=n-1)</pre>
	P-value:	$pvalue = 2 \cdot P(T > t_{obs})$ if $pvalue < \alpha \Rightarrow \text{reject } H_0$ else $\Rightarrow \text{accept } H_0$	<pre>pval = 2 * (1 - stats.t.cdf(abs(t_obs), df=n-1))</pre>

Sample size n for one sample

	Description	Formula	Python
	Sample size n for specified confidence interval:	$n = \left\lceil \left(\frac{z_{1-\alpha/2} \cdot \sigma}{ME} \right)^2 \right\rceil$	<pre>z = stats.norm.ppf(1-$\alpha/2$) n = np.ceil((z*sigma/ME)**2)</pre>
	Sample size n for given power $(1 - \beta)$:	$n = \left\lceil \left(\sigma \frac{z_{1-\beta} + z_{1-\alpha/2}}{\mu_0 - \mu_1} \right)^2 \right\rceil$	<pre>n = smp.TTestPower(). solve_power(effect_ size=($\mu_0 - \mu_1$)/σ, alpha=α, power=1-β)</pre>

Two-samples

Paired

When two “paired” samples are available, they can be transformed into a single sample of measured differences. The data is therefore analyzed using a standard one-sample t -test.

Collect data in Python

```
x1 = np.array([...]) # e.g. patient results before clinical trial
x2 = np.array([...]) # patient results after clinical trial
diff_data = x1 - x2
```

Hypothesis test

For a paired t -test the proposed hypotheses are usually

$$H_0 : \mu = 0,$$

$$H_1 : \mu \neq 0$$

since we want to test if the mean difference is statistically significant from 0.

Method 1

```
t_obs, p_val = stats.ttest_1samp(diff_data, popmean= $\mu_0$ )
```

Method 2

```
t_obs, p_val = stats.ttest_rel(x1, x2) # note this method only computes for  $\mu_0 = 0$ 
```

Confidence interval

For a paired confidence interval, we usually check to see if 0 is included. If 0 is included in the interval, there is not enough evidence to suggest that the two population means are statistically significant from one another.

Method 1

```
CI = stats.ttest_1samp(diff_data, popmean= $\mu_0$ ).confidence_interval(1- $\alpha$ )
```

Method 2

```
CI = stats.ttest_rel(x1, x2).confidence_interval(1- $\alpha$ )
```

Welch

The Welch t-test is used when we assume unequal variances in each distribution.

Collect data in Python

```
x1 = np.array([...])
```

```
x2 = np.array([...])
```

	Description	Formula	Python
	Difference between two samples (d):	$d = \bar{x}_1 - \bar{x}_2$	<code>d = np.mean(x1) - np.mean(x2)</code>
	Variance of the difference (σ_d^2):	$\hat{\sigma}_d^2 = \frac{s_1^2}{n_1} + \frac{s_2^2}{n_2} = (\hat{\sigma}_{\bar{x}_1})^2 + (\hat{\sigma}_{\bar{x}_2})^2$	
	Standard error of the difference (σ_d):	$\hat{\sigma}_d = \sqrt{\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}} = \sqrt{(\hat{\sigma}_{\bar{x}_1})^2 + (\hat{\sigma}_{\bar{x}_2})^2}$	<code>SED = np.sqrt(stats.sem(x1)**2 + stats.sem(x2)**2)</code>
	Welch's degrees of freedom (ν):	$\nu = \frac{\left(\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}\right)^2}{\frac{(s_1^2/n_1)^2}{n_1 - 1} + \frac{(s_2^2/n_2)^2}{n_2 - 1}} = \frac{(\hat{\sigma}_{\bar{x}_1} + \hat{\sigma}_{\bar{x}_2})^4}{\frac{(\hat{\sigma}_{\bar{x}_1})^4}{n_1 - 1} + \frac{(\hat{\sigma}_{\bar{x}_2})^4}{n_2 - 1}}$	
	Confidence interval of difference:	$\bar{x}_1 - \bar{x}_2 \pm t_{1-\alpha/2, \nu} \sqrt{\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}}$	<code>CI = stats.ttest_ind(x1, x2, equal_var=False).confidence_interval(1-alpha)</code>
	Margin of error:	$ME = t_{1-\alpha/2, \nu} \sqrt{\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}}$	

	Description	Formula	Python
	Welch's t-test statistic (t_{obs}): $H_0: \mu_1 - \mu_2 = 0$ $H_1: \mu_1 - \mu_2 \neq 0$	$t_{obs} = \frac{(\bar{x}_1 - \bar{x}_2) - 0}{\sqrt{\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}}} = \frac{d - 0}{\hat{\sigma}_d}$	<pre>t_obs, pval = stats.ttest_ind(x1, x2, equal_var=False)</pre>
	Critical value (t_{crit}):	$t_{crit} = t_{1-\alpha/2, \nu}$ if $t_{obs} \notin [-t_{crit}, t_{crit}] \Rightarrow$ reject else $\Rightarrow$ accept H_0	<pre>crit = stats.t.ppf(1-alpha/2, df=nu)</pre>
	P-value:	$p_{value} = 2 \cdot P(T > t_{obs})$ if $p_{value} < \alpha \Rightarrow$ reject H_0 else $\Rightarrow$ accept H_0	<pre>pval = 2 * (1 - stats.t.cdf(abs(t_obs), df=nu))</pre>

Pooled

The pooled t-test is used when we assume equal variances in each population.

Collect data in Python

```
x1 = np.array([...])
```

```
x2 = np.array([...])
```

	Description	Formula	Python
	Difference between two samples (d):	$d = \bar{x}_1 - \bar{x}_2$	<code>d = np.mean(x1) - np.mean(x2)</code>
	Pooled variance (s_p^2):	$s_p^2 = \frac{(n_1 - 1)s_1^2 + (n_2 - 1)s_2^2}{n_1 + n_2 - 2}$	
	Pooled standard deviation (s_p):	$s_p = \sqrt{s_p^2}$	
	Variation of the difference (σ_d^2):	$\hat{\sigma}_d^2 = \frac{s_p^2}{n_1} + \frac{s_p^2}{n_2} = s_p^2 \left(\frac{1}{n_1} + \frac{1}{n_2} \right)$	
	Standard error of the difference (σ_d):	$\hat{\sigma}_d = \sqrt{\frac{s_p^2}{n_1} + \frac{s_p^2}{n_2}} = s_p \sqrt{\frac{1}{n_1} + \frac{1}{n_2}}$	
	Confidence interval of the difference:	$\bar{x}_1 - \bar{x}_2 \pm t_{1-\alpha/2, (n_1+n_2-2)} \sqrt{\frac{\sigma_p^2}{n_1} + \frac{\sigma_p^2}{n_2}}$	<code>CI = stats.ttest_ind(x1, x2, equal_var=True).confidence_interval(1-alpha)</code>
	Margin of error:	$ME = t_{1-\alpha/2, (n_1+n_2-2)} \sqrt{\frac{\sigma_p^2}{n_1} + \frac{\sigma_p^2}{n_2}}$	

	Description	Formula	Python
	Pooled t-test statistic (t_{obs}): $H_0: \mu_1 - \mu_2 = 0$ $H_1: \mu_1 - \mu_2 \neq 0$	$t_{obs} = \frac{(\bar{x}_1 - \bar{x}_2) - 0}{\sqrt{\frac{s_p^2}{n_1} + \frac{s_p^2}{n_2}}} = \frac{d - 0}{\hat{\sigma}_d}$	<pre>t_obs, pval = stats.ttest_ind(x1, x2, equal_var=True)</pre>
	Critical value (t_{crit}):	$t_{crit} = t_{1-\alpha/2, (n_1+n_2-2)}$ if $t_{obs} \notin [-t_{crit}, t_{crit}] \Rightarrow \text{reject } H_0$ else $\Rightarrow \text{accept } H_0$	<pre>crit = stats.t.ppf(1-alpha/2, df=n1+n2-2)</pre>
	P-value:	$p_{value} = 2 \cdot P(T > t_{obs})$ if $p_{value} < \alpha \Rightarrow \text{reject } H_0$ else $\Rightarrow \text{accept } H_0$	<pre>pval = 2 * (1 - stats.t.cdf(abs(t_obs), df=n1+n2-2))</pre>

Sample size n for pooled sample

	Description	Python	Notes
	Sample size n for given power $(1 - \beta)$: Where σ is estimated with s_p and where $k = n_1/n_2$, that is the ratio of the two sample sizes. Typically $k = 1$.	<pre>n = smp.TTestIndPower(). solve_power(effect_size = $\mu_0 - \mu_1$) / σ, alpha=α, power=$1 - \beta$, ratio=k)</pre>	

Statistical errors

When carrying out hypothesis tests there are two types of errors you can commit:

Type I error

Also known as a false positive. It's the rejection of the null hypothesis, H_0 , when H_0 is true.

Probability of a Type I error: $P(\text{Type I error}) = \alpha$.

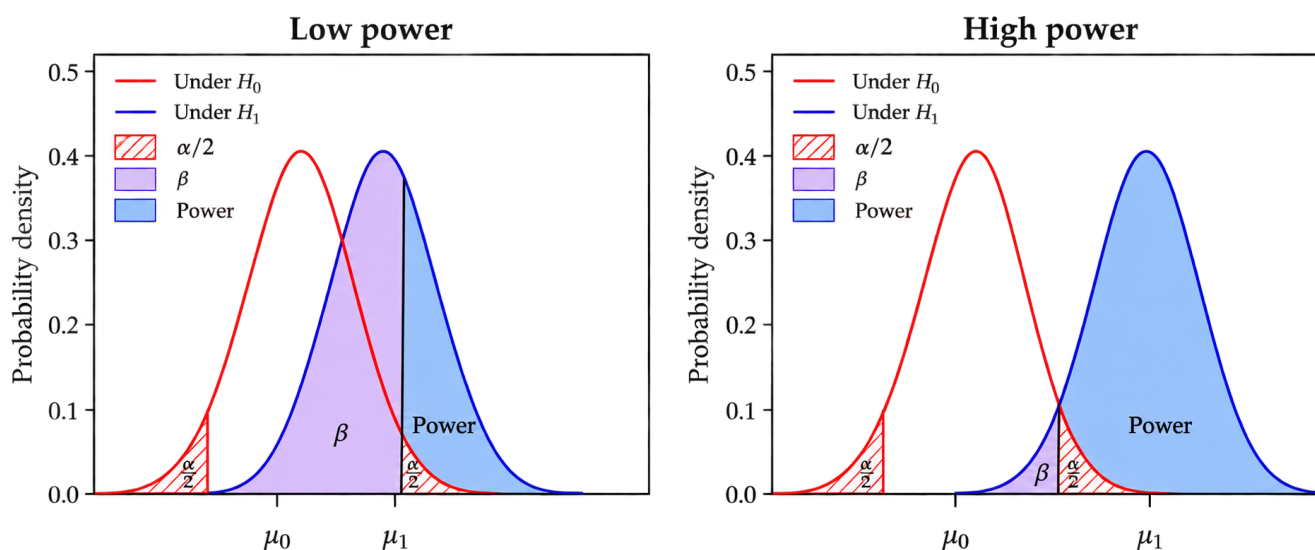
Confidence is therefore defined: $P(\text{reject } H_1 | H_0 \text{ true}) = 1 - \alpha$.

Type II error

Also known as a false negative. It's the acceptance of the null hypothesis, H_0 , when H_1 is true.

Probability of a Type II error: $P(\text{Type II error}) = \beta$.

Power is therefore defined: $P(\text{reject } H_0 | H_1 \text{ true}) = 1 - \beta$.



Red striped area (α): probability of type I error. Lavender area (β): probability of type II error.

Note: When the sample size and effect size are fixed α and β share an inverse relationship, meaning if you decrease one the other increases. A larger sample size is therefore necessary in order to increase power without decreasing confidence.

5 Simulation-based inference and error propagation

Python commands for simulation

1. Simulate from a specified discrete sample space

To obtain the same simulated results each time the code is run, set a seed:

```
np.random.seed(#)
```

Use `np.random.choice()` when the possible outcomes and their probabilities are specified directly:

```
np.random.choice(a=, size=, replace=, p=)
```

a — List containing the possible outcomes (sample space).

size — Number (or shape) of observations to simulate.

replace — Whether sampling is with replacement (**True**) or without replacement (**False**).

p — List containing the probability of each outcome. If **None**, all outcomes are equally likely.

2. Simulate from a probability distribution

If the random variable follows a known probability distribution, observations can be simulated directly using:

```
stats.<distribution>.rvs(..., size=, random_state=)
```

size — specifies the number (or shape) of simulated observations.

random_state — sets the random seed for reproducible simulation results; **None**, no fixed seed is used.

Discrete distributions:

- Binomial: `stats.binom.rvs(n=, p=, loc=0, size=, random_state=)`
- Hypergeometric: `stats.hypergeom.rvs(M=, n=, N=, loc=0, size=, random_state=)`
- Poisson: `stats.poisson.rvs(mu=, loc=0, size=, random_state=)`

Continuous distributions:

- Normal: `stats.norm.rvs(loc= $\mu$ , scale= $\sigma$ , size=, random_state=)`
- Uniform: `stats.uniform.rvs(loc= $\alpha$ , scale= $\beta - \alpha$ , size=, random_state=)`
- Exponential: `stats.expon.rvs(loc=0, scale= $1/\lambda$ , size=, random_state=)`
- Log-normal: `stats.lognorm.rvs(s= $\sigma$ , loc=0, scale= $\exp(\mu)$ , size=, random_state=)`
- Chi-squared: `stats.chi2.rvs(df=, loc=0, scale=1, size=, random_state=)`
- Student's *t*: `stats.t.rvs(df=, loc=0, scale=1, size=, random_state=)`
- *F*: `stats.f.rvs(dfn=, dfd=, loc=0, scale=1, size=, random_state=)`

	Description	Formula	Notes
	Error propagation: The variance of a function $f(X_1, \dots, X_n)$ of random variables.	$\sigma_{f(X_1, \dots, X_n)}^2 = \sum_{i=1}^n \left(\frac{\partial f}{\partial x_i} \sigma_i \right)^2$	Where it is assumed x_i and x_j are independent for all i 's and j 's.

6 Linear regression

$$Y_i = \beta_0 + \beta_1 X_{1,i} + \cdots + \beta_p X_{p,i} + \varepsilon_i$$

Description of components:

- Y_i — the response variable (observed outcome).
- $X_{1,i}, X_{2,i}, \dots, X_{p,i}$ — the explanatory variables. There are p explanatory variables (features).
- ε_i — the random error and the model assumption is $\varepsilon_i \sim \mathcal{N}(0, \sigma^2)$, $\varepsilon_1, \dots, \varepsilon_n$ i.i.d.

Fitting the model in Python

```
data = pd.DataFrame({'y': [...], 'x1': [...], ..., 'xp': [...]})
model = smf.ols(formula='y ~ x1 + x2 + ... + xp', data=data).fit()
print(model.summary(slim=True, alpha=alpha))
```

	Description	Formula	Python
	Sample estimated linear regression model:	$\hat{y}_i = \hat{\beta}_0 + \hat{\beta}_1 x_{1,i} + \cdots + \hat{\beta}_p x_{p,i}$	<code>fittedvalues = model.fittedvalues</code>
	Residual (e_i):	$e_i = y_i - \hat{\beta}_0 - \hat{\beta}_1 x_{1,i} - \cdots - \hat{\beta}_p x_{p,i}$	<code>residuals = model.resid</code>
	Residual sum of squares (RSS):	$RSS = \sum_{i=1}^n (y_i - \hat{y}_i)^2 = \sum_{i=1}^n e_i^2$	<code>RSS = model.ssr</code>
	Estimate of variance (σ^2):	$\hat{\sigma}^2 = \frac{RSS}{n - (p + 1)}$	<code>var = model.scale</code> Alternatively: <code>var = model.mse_resid</code>
	Estimate of standard deviation (σ):	$\hat{\sigma} = \sqrt{\hat{\sigma}^2} = \sqrt{\frac{RSS}{n - (p + 1)}}$	<code>sigma = np.sqrt(model.scale)</code>

Simple linear regression-specific formulas ($Y_i = \beta_0 + \beta_1 X_{1,i} + \varepsilon_i$)

	Description	Formula	Python
	Spread of x-values:	$S_{xx} = \sum_{i=1}^n (x_i - \bar{x})^2$	<pre>S_xx = sum((x-np.mean(x))**2)</pre>
	Sum of cross-products:	$S_{xy} = \sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y})$	<pre>S_xy = sum((x-np.mean(x)) *(y-np.mean(y)))</pre>
	Slope (β_1) and intercept (β_0):	$\hat{\beta}_1 = \frac{\sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y})}{S_{xx}} = \frac{S_{xy}}{S_{xx}}$ $\hat{\beta}_0 = \bar{y} - \hat{\beta}_1 \bar{x}$	<pre>print(model.params)</pre>
	Covariance of β_0 and β_1:	$Cov(\hat{\beta}_0, \hat{\beta}_1) = -\frac{\bar{x}\sigma^2}{S_{xx}}$	<pre>print(model.cov_ params().iloc[0,1])</pre>
	Variance of β_0 ($\sigma_{\hat{\beta}_0}^2$): Also denoted $V[\hat{\beta}_0]$	$\hat{\sigma}_{\hat{\beta}_0}^2 = \frac{\sigma^2}{n} + \frac{\bar{x}^2 \sigma^2}{S_{xx}}$	<pre>print(model.bse**2)</pre>
	Variance of β_1 ($\sigma_{\hat{\beta}_1}^2$): Also denoted $V[\hat{\beta}_1]$	$\hat{\sigma}_{\hat{\beta}_1}^2 = \frac{\sigma^2}{S_{xx}}$	<pre>print(model.bse**2)</pre>
	Standard error of β_0 (σ_{β_0}):	$\hat{\sigma}_{\hat{\beta}_0} = \sqrt{\hat{\sigma}_{\hat{\beta}_0}^2} = \sqrt{\frac{\sigma^2}{n} + \frac{\bar{x}^2 \sigma^2}{S_{xx}}} = \hat{\sigma} \sqrt{\frac{1}{n} + \frac{\bar{x}^2}{S_{xx}}}$	<pre>print(model.bse)</pre>
	Standard error of β_1 (σ_{β_1}):	$\hat{\sigma}_{\hat{\beta}_1} = \sqrt{\hat{\sigma}_{\hat{\beta}_1}^2} = \frac{\sigma}{\sqrt{S_{xx}}} = \hat{\sigma} \sqrt{\frac{1}{S_{xx}}}$	<pre>print(model.bse)</pre>

General linear regression continued

	Description	Formula	Python
	Spread of y-values:	$S_{yy} = \sum_{i=1}^n (y_i - \bar{y})^2$	<pre>S_yy = sum((y-np.mean(y))**2)</pre>
	t-test for parameters: $H_0 : \beta_i = \beta_{0,i}$ $H_1 : \beta_i \neq \beta_{0,i}$	$t_{obs,\beta_i} = \frac{\hat{\beta}_i - \beta_{0,i}}{\hat{\sigma}_{\hat{\beta}_i}}$	<pre>print(model.tvalues)</pre>
	Critical value (t_{crit}):	$t_{crit} = t_{1-\alpha/2, n-(p+1)}$ if $t_{obs,\beta_i} \notin [-t_{crit}, t_{crit}] \Rightarrow$ reject H_0 else $\Rightarrow$ accept H_0	<pre>crit = stats.t.ppf(1-α/2, df=n-(p+1))</pre>
	P-value:	$p_{value} = 2 \cdot P(T > t_{obs,\beta_i})$ if $p_{value} < \alpha \Rightarrow$ reject H_0 else $\Rightarrow$ accept H_0	<pre>pvals = model.pvalues Alternatively: pval = 2 * (1 - stats.t.cdf(abs(t_obs,β_i), df=n-(p+1)))</pre>
	Confidence interval for parameters: Note: $\hat{\beta}_i$ and $\hat{\sigma}_{\hat{\beta}_i}$ are computed in Python when $p > 1$.	$\hat{\beta}_i \pm t_{1-\alpha/2, n-(p+1)} \hat{\sigma}_{\hat{\beta}_i}$	<pre>confint = model. conf_int(alpha=α)</pre>
	Confidence interval for the regression function for a given point $\mathbf{x}_0$: Fitted value: $\hat{y}_0 = \hat{\beta}_0 + \hat{\beta}_1 x_{1,0} + \dots + \hat{\beta}_p x_{p,0}$	$\hat{y}_0 \pm t_{1-\alpha/2, n-p-1} \sigma_{\hat{y}_0}$ Simple linear regression: $\sigma_{\hat{y}_0} = \hat{\sigma} \sqrt{\frac{1}{n} + \frac{(x_0 - \bar{x})^2}{S_{xx}}}$	<pre>x_0 = pd.DataFrame({"x1": [...], ..., "xp": [...]}) pred = model.get_ prediction(x_0).sum mary_frame(alpha=α) print(pred.head())</pre>

	Description	Formula	Python
	Prediction interval for a given point $\mathbf{x}_0$: Fitted value: $\hat{y}_0 = \hat{\beta}_0 + \hat{\beta}_1 x_{1,0} + \dots + \hat{\beta}_p x_{p,0}$	$\hat{y}_0 \pm t_{1-\alpha/2, n-p-1} \sigma_{\tilde{y}_0}$ Simple linear regression: $\sigma_{\tilde{y}_0} = \hat{\sigma} \sqrt{1 + \frac{1}{n} + \frac{(x_0 - \bar{x})^2}{S_{xx}}}$	<pre>pred = model.get_prediction(x_0).summary_frame(alpha=alpha) print(pred.head())</pre>
	Coefficient of determination R^2:	$R^2 = 1 - \frac{RSS}{S_{yy}} = 1 - \frac{\sum_{i=1}^n (y_i - \hat{y})^2}{\sum_{i=1}^n (y_i - \bar{y})^2}$	<pre>print(model.rsquared) adjusted R-squared: print(model.rsquared_adj)</pre>
	Correlation coefficient r:	Simple linear regression: $r = \hat{\rho} = \text{sign}(\hat{\beta}_1) \sqrt{R^2}$	

Backward selection

Start with the full model and iteratively remove the least significant variable, refitting the model and testing the remaining parameters after each removal. Continue until all remaining explanatory variables are significant.

Matrix formulation for linear regression

	Description	Formula	Python
	Linear regression model:	$\mathbf{Y} = \mathbf{X}\boldsymbol{\beta} + \boldsymbol{\varepsilon}, \varepsilon_i \sim N(0, \sigma^2)$ $\mathbf{Y} = \begin{bmatrix} Y_1 \\ Y_2 \\ \vdots \\ Y_n \end{bmatrix} = \begin{bmatrix} 1 & x_{1,1} & \dots & x_{1,p} \\ 1 & x_{2,1} & \dots & x_{2,p} \\ \vdots & \vdots & \ddots & \vdots \\ 1 & x_{n,1} & \dots & x_{n,p} \end{bmatrix} \begin{bmatrix} \beta_0 \\ \beta_1 \\ \vdots \\ \beta_p \end{bmatrix} + \begin{bmatrix} \varepsilon_1 \\ \varepsilon_2 \\ \vdots \\ \varepsilon_n \end{bmatrix}$	
	Residual e:	$\mathbf{e} = \begin{bmatrix} e_1 \\ e_2 \\ \vdots \\ e_n \end{bmatrix}$	<code>e = np.array([...])</code>
	Residual sum of squares (RSS):	$RSS = \mathbf{e}^T \mathbf{e}$	<code>RSS = e @ e.T</code>
	Estimated $\boldsymbol{\beta}$:	$\hat{\boldsymbol{\beta}} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{Y} = \begin{bmatrix} \hat{\beta}_0 \\ \hat{\beta}_1 \\ \vdots \\ \hat{\beta}_p \end{bmatrix}$	
	Sample estimated linear regression model: The fitted value $\hat{Y}_i$ for vector $\mathbf{x}_i$	$\mathbf{x}_i = [1, x_{1,i}, \dots, x_{p,i}]$ $\hat{Y}_i = \mathbf{x}_i \hat{\boldsymbol{\beta}}$	

	Description	Formula	Python
	Estimated variance and covariance of β:	$\mathbf{V}[\hat{\beta}] = \sigma^2 (\mathbf{X}^T \mathbf{X})^{-1}$ $= \begin{bmatrix} \mathbf{V}[\hat{\beta}_0] & \text{Cov}(\hat{\beta}_0, \hat{\beta}_1) & \cdots & \text{Cov}(\hat{\beta}_0, \hat{\beta}_p) \\ \text{Cov}(\hat{\beta}_1, \hat{\beta}_0) & \mathbf{V}[\hat{\beta}_1] & \cdots & \text{Cov}(\hat{\beta}_1, \hat{\beta}_p) \\ \vdots & \vdots & \ddots & \vdots \\ \text{Cov}(\hat{\beta}_p, \hat{\beta}_0) & \text{Cov}(\hat{\beta}_p, \hat{\beta}_1) & \cdots & \mathbf{V}[\hat{\beta}_p] \end{bmatrix}$	
	Confidence interval for the regression line:	$\mathbf{x}_0 = [1, x_{1,0}, \dots, x_{p,0}]$ $V(\hat{Y}_0) = V(\mathbf{x}_0 \hat{\beta})$ $= \mathbf{x}_0 \mathbf{V}[\hat{\beta}] \mathbf{x}_0^T$ $= \sigma^2 \mathbf{x}_0 (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{x}_0^T$ $\hat{Y}_0 = \mathbf{x}_0 \hat{\beta},$ $\hat{Y}_0 \pm t_{1-\alpha/2, n-p-1} \sqrt{V(\hat{Y}_0)}$	
	Prediction interval for the regression line:	$\mathbf{x}_0 = [1, x_{1,0}, \dots, x_{p,0}]$ $V(\tilde{Y}_0) = V(\mathbf{x}_0 \hat{\beta} + \varepsilon_0)$ $= \mathbf{x}_0 \mathbf{V}[\hat{\beta}] \mathbf{x}_0^T + \sigma^2$ $= \sigma^2 (1 + \mathbf{x}_0 (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{x}_0^T)$ $\hat{Y}_0 = \mathbf{x}_0 \hat{\beta},$ $\hat{Y}_0 \pm t_{1-\alpha/2, n-p-1} \sqrt{V(\tilde{Y}_0)}$	

7 Inference for proportions

Where x denotes the number of successes out of n independent trials with probability of success $\hat{p}$ and probability of failure $1 - \hat{p}$.

One-proportion

Collect data in Python

```
data = np.array([0, 0, 1, 0, 1, 1, ...])
x = sum(data)
n = len(data)
```

	Description	Formula	Python
	Sample proportion $\hat{p}$:	$\hat{p} = \frac{x}{n}$	
	Variance of a proportion ($\sigma_{\hat{p}}^2$): Also denoted $V(\hat{p})$	$\hat{\sigma}_{\hat{p}}^2 = \frac{\hat{p}(1 - \hat{p})}{n}$	
	Standard error of p ($\sigma_{\hat{p}}$):	$\hat{\sigma}_{\hat{p}} = \sqrt{\hat{\sigma}_{\hat{p}}^2} = \sqrt{\frac{\hat{p}(1 - \hat{p})}{n}}$	
	Confidence interval for a proportion: Rule of thumb: $np > 15$ and $n(1 - p) > 15$	$\hat{p} \pm z_{1-\alpha/2} \sqrt{\frac{\hat{p}(1 - \hat{p})}{n}}$	<pre>confint = smprop.proportion_ confint(count=x, nobs=n, alpha=α, method="normal")</pre>
	Margin of error:	$ME = z_{1-\alpha/2} \sqrt{\frac{\hat{p}(1 - \hat{p})}{n}}$	
	Plus 2 method: Used when $np \leq 15$ and $n(1 - p) \leq 15$.	$\tilde{p} = \frac{\tilde{x}}{\tilde{n}} = \frac{x + 2}{n + 4}$ $\tilde{p} \pm z_{1-\alpha/2} \sqrt{\frac{\tilde{p}(1 - \tilde{p})}{\tilde{n}}}$	<pre>confint = smprop.proportion_ confint(count=x, nobs=n, alpha=α, method="agresti_coull")</pre>

	Description	Formula	Python
	z-test statistic (z_{obs}): $H_0 : p = p_0$ $H_1 : p \neq p_0$	$z_{obs} = \frac{x - n \cdot p_0}{\sqrt{n \cdot p_0(1 - p_0)}}$ $= \frac{\hat{p} - p_0}{\sqrt{p_0(1 - p_0)/n}} \sim \mathcal{N}(0, 1^2)$	<pre>z_obs, pval = smprop.proportions_ ztest(count=x, nobs=n, value=p0, prop_var=0.5)</pre>
	Critical value (z_{crit}):	$z_{crit} = z_{1-\alpha/2}$ if $z_{obs} \notin [-z_{crit}, z_{crit}] \Rightarrow \text{reject } H_0$ else $\Rightarrow \text{accept } H_0$	<pre>crit = stats.norm.ppf(1-alpha/2)</pre>
	P-value:	$p_{value} = 2 \cdot P(Z > z_{obs})$ if $p_{value} < \alpha \Rightarrow \text{reject } H_0$ else $\Rightarrow \text{accept } H_0$	<pre>pval = 2 * (1 - stats.norm.cdf(abs(z_obs)))</pre>

Sample size n for one sample

	Description	Formula	Python
	Sample size n for specified confidence interval: For unknown p (worst case), use $p = 0.5$.	$n = \left\lceil p(1 - p) \left(\frac{z_{1-\alpha/2}}{ME} \right)^2 \right\rceil$	
	Sample size n for given power $(1 - \beta)$:	$n = \left\lceil \frac{[z_{1-\alpha/2} \sqrt{p_0(1 - p_0)} + z_{1-\beta} \sqrt{p_1(1 - p_1)}]^2}{(p_1 - p_0)^2} \right\rceil$	

Two-proportions

Collect data in Python

```
A = np.array([0, 0, 1, 0, 1, 1, ...])
B = np.array([0, 1, 0, 1, 1, 1, ...])
x1, x2 = sum(A), sum(B)
n1, n2 = len(A), len(B)
```

	Description	Formula	Python
	Difference between two proportions ($\hat{d}$):	$\hat{d} = \hat{p}_1 - \hat{p}_2$	
	Variance of the difference $\sigma_{\hat{p}_1 - \hat{p}_2}^2$: Also denoted $V(\hat{p}_1 - \hat{p}_2)$	$\hat{\sigma}_{\hat{p}_1 - \hat{p}_2}^2 = \hat{\sigma}_{\hat{p}_1}^2 + \hat{\sigma}_{\hat{p}_2}^2$	
	Standard error of the difference ($\sigma_{\hat{d}}$): Rule of thumb: $n_i p_i \geq 10$ and $n_i(1 - p_i) \geq 10$ for $i = 1, 2$.	$\begin{aligned}\hat{\sigma}_{\hat{d}} &= \sqrt{\frac{\hat{p}_1(1 - \hat{p}_1)}{n_1} + \frac{\hat{p}_2(1 - \hat{p}_2)}{n_2}} \\ &= \sqrt{(\hat{\sigma}_{\hat{p}_1})^2 + (\hat{\sigma}_{\hat{p}_2})^2}\end{aligned}$	
	Confidence interval for the difference: Rule of thumb: $n_i p_i \geq 10$ and $n_i(1 - p_i) \geq 10$ for $i = 1, 2$.	$(\hat{p}_1 - \hat{p}_2) \pm z_{1-\alpha/2} \sqrt{\frac{\hat{p}_1(1 - \hat{p}_1)}{n_1} + \frac{\hat{p}_2(1 - \hat{p}_2)}{n_2}}$	<pre>confint = confint_proportions_ 2indep(count1=x1, nobs1=n1, count2=x2, nobs2=n2, alpha=α, method="normal")</pre>
	Margin of Error:	$ME = z_{1-\alpha/2} \sqrt{\frac{\hat{p}_1(1 - \hat{p}_1)}{n_1} + \frac{\hat{p}_2(1 - \hat{p}_2)}{n_2}}$	

	Description	Formula	Python
	z-test statistic (z_{obs}): $H_0 : p_1 - p_2 = 0$ $H_1 : p_1 - p_2 \neq 0$	$\hat{p} = \frac{x_1 + x_2}{n_1 + n_2}$ $z_{obs} = \frac{\hat{p}_1 - \hat{p}_2}{\sqrt{\hat{p}(1 - \hat{p}) \left(\frac{1}{n_1} + \frac{1}{n_2} \right)}}$	<pre>z_obs, pval = proportions_ztest(count=[x1, x2], nobs=[n1, n2])</pre>
	Critical value (z_{crit}):	$z_{crit} = z_{1-\alpha/2}$ if $z_{obs} \notin [-z_{crit}, z_{crit}] \Rightarrow \text{reject } H_0$ else $\Rightarrow \text{accept } H_0$	<pre>crit = stats.norm.ppf(1-alpha/2)</pre>
	P-value:	$pvalue = 2 \cdot P(Z > z_{obs})$ if $pvalue < \alpha \Rightarrow \text{reject } H_0$ else $\Rightarrow \text{accept } H_0$	<pre>pval = 2 * (1 - stats.norm.cdf(abs(z_obs)))</pre>

Sample size n for two samples

	Description	Formula	Python
	Sample size n for specified confidence interval: $n = n_1 = n_2$ If p_1 or p_2 are unknown maximize variance by setting $p_1 = p_2 = 0.5$	$n = \left\lceil \frac{(z_{1-\alpha/2})^2 [p_1(1-p_1) + p_2(1-p_2)]}{ME^2} \right\rceil$	
	Sample size n for given power $(1 - \beta)$: $n = n_1 = n_2$	$\bar{p} = \frac{p_1 + p_2}{2}$ $n = \left\lceil \frac{\left(z_{1-\alpha/2} \sqrt{2\bar{p}(1-\bar{p})} + z_{1-\beta} \sqrt{p_1(1-p_1) + p_2(1-p_2)} \right)^2}{(p_1 - p_2)^2} \right\rceil$	

Contingency-tables

An $r \times c$ contingency table summarizes the joint distribution of two categorical variables by recording observed cell counts.

	Description	Formula	Python
	Expected value (e_{ij}): Also denoted E .	$e_{ij} = \frac{(\text{row total } i)(\text{column total } j)}{\text{grand total}}$	
	χ^2-test for multi-sample proportions: $H_0 : p_1 = \dots = p_c$ Rule of thumb: The test is valid if all $e_{ij} \geq 5$.	$\chi_{obs}^2 = \sum_{i=1}^2 \sum_{j=1}^c \frac{(o_{ij} - e_{ij})^2}{e_{ij}}$	<pre>x = [x1,x2,...,xn] n = [n1,n2,...,nk] chi2_obs, pval, (obs, exp) = smprop.proportions_ chisquare(x, n)</pre>
	χ^2-test for $r \times c$ frequency table: $H_0 : p_{i1} = \dots = p_{ic} = p_i, \quad i = 1, 2, \dots, r$ Test for independence between rows and columns. Rule of thumb: The test is valid if all $e_{ij} \geq 5$	$\chi_{obs}^2 = \sum_{i=1}^r \sum_{j=1}^c \frac{(o_{ij} - e_{ij})^2}{e_{ij}}$	<pre>X = [[x11, x12, ..., x1c], [x21, x22, ..., x2c], ... [xr1, xr2, ..., xrc]] chi2_obs, pval, dof, exp = stats.chi2_contingency(X, correction=False)</pre>
	Critical value (χ_{crit}^2):	$\chi_{crit}^2 = \chi_{1-\alpha, (r-1)(c-1)}^2$ if $\chi_{obs}^2 > \chi_{crit}^2 \Rightarrow \text{reject } H_0$ else $\Rightarrow \text{accept } H_0$	<pre>crit = stats.chi2.ppf(1-alpha, df=(r-1)*(c-1))</pre>
	P-value:	$pvalue = P(\chi^2 > \chi_{obs}^2)$ if $pvalue < \alpha \Rightarrow \text{reject } H_0$ else $\Rightarrow \text{accept } H_0$	<pre>pval = 1 - stats.chi2.cdf(chi2_obs, df=(r-1)*(c-1))</pre>

8 One-way ANOVA

$$Y_{ij} = \mu + \alpha_i + \varepsilon_{ij}$$

Description of components:

- Y_{ij} — the response from group i for observation j , where $j \in \{1, \dots, n_i\}$
- μ — the grand (overall) mean.
- α_i — the effect of group (treatment) i , $i \in \{1, \dots, k\}$, where there are k groups each with n_i observations.
- ε_{ij} — the random error and the model assumption is $\varepsilon_{ij} \sim \mathcal{N}(0, \sigma^2)$, ε_{ij} are i.i.d.
- $\bar{y}_i$ — mean in group i .

One-way ANOVA

Source of variation	Degrees of freedom	Sum of squares	Mean square	Test statistic F	p-value
Treatment	$k - 1$	$SS(Tr)$	$MS(Tr) = \frac{SS(Tr)}{k-1}$	$F_{obs} = \frac{MS(Tr)}{MSE}$	$P(F > F_{obs})$
Residual (Error)	$n - k$	SSE	$MSE = \hat{\sigma}^2 = \frac{SSE}{n-k}$		
Total	$n - 1$	SST			

Fitting the model in Python

```
data = pd.DataFrame('y': [...], 'group': ["A", "A", ..., "B", ..., "K"])
model = smf.ols('y ~ group', data=data).fit()
anova = sm.stats.anova_lm(model)
print(anova)
```

	Description	Formula	Python
	Grand mean (μ):	$\bar{y} = \hat{\mu} = \frac{1}{n} \sum_{i=1}^k \sum_{j=1}^{n_i} y_{ij}$	
	Mean of group i :	$\bar{y}_i = \frac{1}{n_i} \sum_{j=1}^{n_i} y_{ij}$	
	Marginal effect α_i :	$\hat{\alpha}_i = \bar{y}_i - \bar{y}$	

	Description	Formula	Python
	Total sum of squares (SST):	$SST = \sum_{i=1}^k \sum_{j=1}^{n_i} (y_{ij} - \bar{y})^2 = (n-1)\sigma_y^2$	<code>SST = model.centered_tss</code>
	Treatment sum of squares ($SS(Tr)$):	$SS(Tr) = \sum_{i=1}^k n_i (\bar{y}_i - \bar{y})^2 = \sum_{i=1}^k n_i \hat{\alpha}_i^2$	<code>SSTr = model.ess</code>
	Sum of squared errors (SSE):	$SSE = \sum_{i=1}^k \sum_{j=1}^{n_i} (y_{ij} - \bar{y}_i)^2 = \sum_{i=1}^k (n_i - 1)\sigma_i^2$	<code>SSE = model.ssr</code>
	Variance decomposition:	$SST = SS(Tr) + SSE$	
	Mean squared errors (MSE): Also denoted $\hat{\sigma}^2$	$MSE = \frac{SSE}{n-k} = \frac{(n_1-1)s_1^2 + \dots + (n_k-1)s_k^2}{n-k}$ $s_i^2 = \frac{1}{n_i-1} \sum_{j=1}^{n_i} (y_{ij} - \bar{y}_i)^2$	<code>MSE = model.mse_resid</code>
	F-test for differences in group means: $H_0 : \alpha_i = 0$ $H_1 : \alpha_i \neq 0, i = 1, \dots, k$	$F_{obs} = \frac{SS(Tr)/(k-1)}{SSE/(n-k)} = \frac{MS(Tr)}{MSE}$	<code>print(anova)</code>
	Critical value (F_{crit}):	$F_{crit} = F_{1-\alpha, k-1, n-k}$ if $F_{obs} > F_{crit} \Rightarrow$ reject H_0 else $\Rightarrow$ accept H_0	<code>crit = stats.f.ppf(1-alpha, dfn=k-1, dfd=n-k)</code>
	P-value:	$pvalue = P(F > F_{obs})$ if $pvalue < \alpha \Rightarrow$ reject H_0 else $\Rightarrow$ accept H_0	<code>pval = 1 - stats.f.cdf(F_obs, dfn=k-1, dfd=n-k)</code>

Post-hoc pairwise tests

When performing multiple hypothesis tests, the probability of a Type I error increases. Therefore, use the Bonferroni-corrected significance level $\alpha_{\text{Bonferroni}} = \alpha/M$, where M is the number of pairwise comparisons. Among k groups the total possible comparisons are found as $M = \frac{k(k-1)}{2}$.

Note: the Bonferroni correction is conservative and may increase the probability of Type II errors.

Important distinction! for post hoc pairwise tests $\bar{y}_j$ is not to be confused with the mean of block j but rather the mean of group j you are comparing against group i with.

	Description	Formula	Python
	Confidence intervals:	$\bar{y}_i - \bar{y}_j \pm t_{1-\alpha/2, n-k} \sqrt{MSE \left(\frac{1}{n_i} + \frac{1}{n_j} \right)}$	
	Least significant difference (LSD): Formula requirement: $m = n_1 = n_k$	$LSD_{\alpha} = t_{1-\alpha/2, n-k} \sqrt{2 \cdot MSE / m}$ $\bar{y}_i - \bar{y}_j > LSD_{\alpha} \Rightarrow \text{reject } H_0$ $\text{else} \Rightarrow \text{accept } H_0$	
	t-test statistic (t_{obs}): $H_0 : \mu_i = \mu_j$ $H_1 : \mu_i \neq \mu_j$	$t_{obs} = \frac{\bar{y}_i - \bar{y}_j}{\sqrt{MSE \left(\frac{1}{n_i} + \frac{1}{n_j} \right)}}$	
	Critical value (t_{crit}):	$t_{crit} = t_{1-\alpha/2, n-k}$ if $t_{obs} \notin [-t_{crit}, t_{crit}] \Rightarrow \text{reject } H_0$ else $\Rightarrow \text{accept } H_0$	$\text{crit} = \text{stats.t.ppf}(1-\alpha/2, \text{df}=n-k)$
	P-value:	$p_{value} = 2 \cdot P(T > t_{obs})$ if $p_{value} < \alpha \Rightarrow \text{reject } H_0$ else $\Rightarrow \text{accept } H_0$	$\text{pval} = 2 * (1 - \text{stats.t.cdf}(\text{abs}(t_{obs}), \text{df}=n-k))$

9 Two-way ANOVA

$$Y_{ij} = \mu + \alpha_i + \beta_j + \varepsilon_{ij}$$

Description of components:

- Y_{ij} — the response for treatment i , block j .
- μ — the grand (overall) mean.
- α_i — the effect of group (Treatment), $i \in \{1, \dots, k\}$, where there are k groups each with n_i observations.
- β_j — the effect of block, $j \in \{1, \dots, l\}$, where there are j blocks.
- ε_{ij} — the random error and the model assumption is $\varepsilon_{ij} \sim \mathcal{N}(0, \sigma^2)$, ε_{ij} are i.i.d.

Two-way ANOVA

Source of variation	Degrees of freedom	Sum of squares	Mean square	Test statistic F	p-value
Treatment	$k - 1$	$SS(Tr)$	$MS(Tr) = \frac{SS(Tr)}{k-1}$	$F_{Tr} = \frac{MS(Tr)}{MSE}$	$P(F > F_{Tr})$
Block	$l - 1$	$SS(Bl)$	$MS(Bl) = \frac{SS(Bl)}{l-1}$	$F_{Bl} = \frac{MS(Bl)}{MSE}$	$P(F > F_{Bl})$
Residual (Error)	$(l - 1)(k - 1)$	SSE	$MSE = \hat{\sigma}^2 = \frac{SSE}{(l-1)(k-1)}$		
Total	$n - 1$	SST			

Fitting the model in Python

```
data = pd.DataFrame({'y': [...], 'treatment': ["A",..., "B",..., "K" ], 'block': ["1", ..., "2", ..., "1"]})
model = smf.ols('y ~ treatment + block', data=data).fit()
anova = sm.stats.anova_lm(model)
print(anova)
```

	Description	Formula	Python
	Grand mean (μ):	$\bar{y} = \hat{\mu} = \frac{1}{k \cdot l} \sum_{i=1}^k \sum_{j=1}^l y_{ij}$	
	Mean of group i and block j :	$\bar{y}_i = \frac{1}{l} \sum_{j=1}^l y_{ij}$ $\bar{y}_j = \frac{1}{k} \sum_{i=1}^k y_{ij}$	

	Description	Formula	Python
	Marginal effects α_i and β_j:	$\hat{\alpha}_i = \bar{y}_i - \bar{y}$ $\hat{\beta}_j = \bar{y}_j - \bar{y}$	
	Total sum of squares (SST):	$SST = \sum_{i=1}^k \sum_{j=1}^l (y_{ij} - \bar{y})^2$	
	Treatment sum of squares ($SS(Tr)$):	$SS(Tr) = l \sum_{i=1}^k (\bar{y}_i - \bar{y})^2 = l \sum_{i=1}^k \hat{\alpha}_i^2$	
	Block sum of squares ($SS(Bl)$):	$SS(Bl) = k \sum_{j=1}^l (\bar{y}_j - \bar{y})^2 = k \sum_{j=1}^l \hat{\beta}_j^2$	
	Sum of squared errors (SSE):	$SSE = \sum_{i=1}^k \sum_{j=1}^l \hat{\varepsilon}_{ij}^2$ $= \sum_{i=1}^k \sum_{j=1}^l (y_{ij} - \hat{y}_{ij})^2$ $= \sum_{i=1}^k \sum_{j=1}^l (y_{ij} - \bar{y} - \hat{\alpha}_i - \hat{\beta}_j)^2$ $= \sum_{i=1}^k \sum_{j=1}^l (y_{ij} - \bar{y}_i - \bar{y}_j + \bar{y})^2$	
	Variance decomposition:	$SST = SS(Tr) + SS(Bl) + SSE$	
	Mean squared errors (MSE):	$MSE = \hat{\sigma}^2 = \frac{SSE}{(l-1)(k-1)}$	<code>MSE = model.mse_resid</code>

Group hypothesis testing

	Description	Formula	Python
	F-test for group mean differences: $H_{0,Tr} : \alpha_i = 0$ $H_{1,Tr} : \alpha_i \neq 0,$ $i = 1, \dots, k$	$F_{Tr} = \frac{SS(Tr)/(k-1)}{SSE/((k-1)(l-1))} = \frac{MS(Tr)}{MSE}$	<pre>print(anova)</pre>
	Critical value $(F_{crit,Tr})$:	$F_{crit,Tr} = F_{1-\alpha, k-1, (k-1)(l-1)}$ if $F_{Tr} > F_{crit,Tr} \Rightarrow$ reject H_0 else $\Rightarrow$ accept H_0	<pre>crit = 1 - stats.f. ppf(1-alpha, dfn=k-1, dfd=(k-1)*(l-1))</pre>
	P-value:	$pvalue = P(F > F_{Tr})$ if $pvalue < \alpha \Rightarrow$ reject H_0 else $\Rightarrow$ accept H_0	<pre>pval = 1 - stats.f. cdf(F_Tr, dfn=k-1, dfd=(k-1)*(l-1))</pre>

Block hypothesis testing

	Description	Formula	Python
	F-test for block mean differences: $H_{0,Bl} : \beta_j = 0$ $H_{1,Bl} : \beta_j \neq 0$ $j = 1, \dots, l$	$F_{Bl} = \frac{SS(Bl)/(l-1)}{SSE/((k-1)(l-1))} = \frac{MS(Bl)}{MSE}$	<pre>print(anova)</pre>
	Critical value $(F_{crit,Bl})$:	$F_{crit,Bl} = F_{1-\alpha, l-1, (k-1)(l-1)}$ if $F_{Bl} > F_{crit,Bl} \Rightarrow$ reject H_0 else $\Rightarrow$ accept H_0	<pre>crit = 1 - stats.f. ppf(1-alpha, dfn=l-1, dfd=(k-1)*(l-1))</pre>
	P-value:	$pvalue = P(F > F_{Bl})$ if $pvalue < \alpha \Rightarrow$ reject H_0 else $\Rightarrow$ accept H_0	<pre>pval = 1 - stats.f. cdf(F_Bl, dfn=l-1, dfd=(k-1)*(l-1))</pre>

Post-hoc pairwise tests

When performing multiple hypothesis tests, the probability of a Type I error increases. Therefore, use the Bonferroni-corrected significance level $\alpha_{\text{Bonferroni}} = \alpha/M$, where M is the number of pairwise comparisons. Among k groups the total possible comparisons are found as $M = \frac{k(k-1)}{2}$.

Note: the Bonferroni correction is conservative and may increase the probability of Type II errors.

Important distinction! for post hoc pairwise tests $\bar{y}_j$ is not to be confused with the mean of block j but rather the mean of group j you are comparing against group i with.

	Description	Formula	Python
	Confidence intervals:	$\bar{y}_i - \bar{y}_j \pm t_{1-\alpha/2, (k-1)(l-1)} \sqrt{MSE \left(\frac{1}{n_i} + \frac{1}{n_j} \right)}$	
	Least significant difference (LSD): Formula requirement: $m = n_1 = n_k$.	$LSD_{\alpha} = t_{1-\alpha/2, (k-1)(l-1)} \sqrt{\frac{2 \cdot MSE}{m}}$ $\bar{y}_i - \bar{y}_j > LSD \Rightarrow \text{reject } H_0$ $\text{else} \Rightarrow \text{accept } H_0$	
	t-test statistic (t_{obs}): $H_0 : \mu_i = \mu_j$ $H_1 : \mu_i \neq \mu_j$	$t_{obs} = \frac{\bar{y}_i - \bar{y}_j}{\sqrt{MSE \left(\frac{1}{n_i} + \frac{1}{n_j} \right)}}$	
	Critical value (t_{crit}):	$t_{crit} = t_{1-\alpha/2, (k-1)(l-1)}$ $\text{if } t_{obs} \notin [-t_{crit}, t_{crit}] \Rightarrow \text{reject } H_0$ $\text{else} \Rightarrow \text{accept } H_0$	$\text{crit} = \text{stats.t.ppf}(1-\alpha/2, \text{df}=(k-1)*(l-1))$
	P-value:	$p_{value} = 2 \cdot P(T > t_{obs})$ $\text{if } p_{value} < \alpha \Rightarrow \text{reject } H_0$ $\text{else} \Rightarrow \text{accept } H_0$	$\text{pval} = 2 * (1 - \text{stats.t.cdf}(\text{abs}(t_{obs}), \text{df}=(k-1)*(l-1)))$

10 Model diagnostics and residual analysis

To assess whether the model assumptions are reasonable, we examine the estimated residuals from the fitted model. For each observation, the residual is defined as

$$e_i = y_i - \hat{y}_i,$$

where y_i is the observed response and $\hat{y}_i$ is the corresponding fitted value. This general diagnostic procedure applies to linear regression, one-way ANOVA, and two-way ANOVA.

Checking the normality assumption

The normality assumption is assessed using the pooled set of residuals from the fitted model. The following graphical diagnostics can be used:

1. **Q–Q plot of the residuals.**

If the residuals are approximately normally distributed, the points should lie close to the straight reference line. Small deviations are generally not problematic, whereas systematic curvature, large deviations in the tails, or extreme observations may indicate departures from normality.

2. **Histogram of the residuals.**

The histogram provides a supplementary visual check. Under approximate normality, the residual distribution should be roughly symmetric and bell-shaped around zero. In general, the Q–Q plot is more informative for assessing normality.

Other residual diagnostics

In addition to normality, the residuals can be examined to assess the other model assumptions:

- **Constant variance:** Plot the residuals against the fitted values. The residuals should have approximately the same vertical spread across the range of fitted values.
- **Independence and model fit:** The residuals should be randomly scattered around zero without any clear systematic pattern.
- **Group comparisons (ANOVA):** Boxplots of the response within each group can be used as an additional visual check for differences in spread and potential outliers.

11 Notation and terminology

Symbol	Definition
Random variables and distributions	
X	Random variable representing an outcome of any distribution.
$\bar{X}$	Random variable representing the sample mean.
S^2	Random variable representing the sample variance.
S	Random variable representing the sample standard deviation.
Z	Random variable representing an outcome from the standard normal distribution.
T	Random variable representing an outcome from the Student's t -distribution.
χ^2	Random variable representing an outcome from the χ^2 -distribution.
F	Random variable representing an outcome from the F -distribution.
Population and sample quantities	
N	Population size, the total number of elements in a population.
μ	Population mean (also denoted $E[X]$).
σ^2	Population variance (also denoted $Var(X)$ or $V[X]$).
σ	Population standard deviation.
n	Sample size, the number of observations in a sample.
$\bar{x}$	Sample mean (also denoted $\hat{\mu}$).
x_i	The i -th observation in a sample of size n .
s^2	Sample variance.
s	Sample standard deviation.
Standard errors	
$\hat{\sigma}_{\bar{x}}$	Estimated standard error of the mean.
$\hat{\sigma}_d$	Estimated standard error of the difference between two means (also denoted $\hat{\sigma}_{\bar{x}_1 - \bar{x}_2}$).
$\hat{\sigma}_{\hat{p}}$	Estimated standard error of a proportion.
$\hat{\sigma}_{\hat{d}}$	Estimated standard error of the difference between two proportions (also denoted $\hat{\sigma}_{\hat{p}_1 - \hat{p}_2}$).
Hypothesis testing	
H_0	Null-hypothesis.
H_1	Alternative hypothesis.
z_{obs}	Observed z -statistic.
t_{obs}	Observed t -statistic.
χ^2_{obs}	Observed χ^2 -statistic.
F_{obs}	Observed F -statistic.
z_{crit}	Critical quantile from the standard normal distribution.
t_{crit}	Critical quantile from the Student's t -distribution.
χ^2_{crit}	Critical quantile from the chi-squared distribution.
F_{crit}	Critical quantile from the F -distribution.
p_{value}	Probability of observing a test statistic at least as extreme as the observed one under H_0 .

Correlation and regression

s_{xy}	Sample covariance: Measures the direction of the linear relationship between two variables (unit-dependent).
r	Sample correlation coefficient (also denoted $\hat{\rho}$) it's the estimate of the population correlation ρ : Measures the strength and direction of the linear relationship between two variables, in the range $[-1, 1]$.
R^2	Coefficient of determination in linear regression. The proportion of variance in Y explained by X , in the range $[0, 1]$.

12 Python libraries

```
import numpy as np
import pandas as pd
import scipy.stats as stats
import statsmodels.api as sm
import matplotlib.pyplot as plt
import statsmodels.formula.api as smf
import statsmodels.stats.power as smp
import statsmodels.stats.proportion as smprop
```