

Introduction to Data Analysis and Statistics

Dario Paape

Winter 2024/25

Universität Osnabrück

November 14, 2024

- Mondays, 12–16 (including 30 minutes break)
- Bring a laptop
- Additional online QA (30 minutes, time to be set)
- Tutorials in small groups (times to be set)
- Two exams (30% each), four homework assignments (10% each)

Syllabus

1 Why do we need inferential statistics?

- The world is less ordered than you think: Descriptive statistics, randomness, and cognitive biases
- The quest for the data-generating process: Estimating parameters of distributions

Syllabus

2 Intro to R

- Installing R and RStudio
- Basics: Working with vectors and data frames, reading and writing data
- More data wrangling: Long and wide formats, merging
- Descriptive statistics: Calculating measures of central tendency and dispersion
- Visualization with the dplyr and ggplot2 packages
- Generating documents with R Markdown

Syllabus

- 3** The basic logic of inferential statistics: Random or not?
- The t-test: Sampling distribution of the sample mean, central limit theorem, t-value, standard error, confidence interval
 - Paired vs. unpaired tests, one-sided vs. two-sided tests
 - Alpha level, Type I/II errors, power

Syllabus

- 4** The Swiss army knife of stats: (Generalized) linear (mixed) models
 - A basic linear model
 - Categorical and continuous predictors, contrast coding
 - Interactions between predictors
 - Applying the logic to different data-generating distributions
 - Random effects: Taking sampling of data sources into account

Syllabus

5 Bonus

- Publication bias, multiple comparisons, p-hacking, replication crisis
- The Bayesian revolution: Taking prior information into account
- Fit to observed data vs. predictive fit, cross-validation

Who am I

- B.Sc. and M.Sc. in Linguistics
- Dr. phil. in Cognitive Science
- PI/Postdoc in Potsdam (Vasishth Lab)
- Substitute Professor of Cognitive Modeling in Osnabrück
- Psycholinguistic experiments, computational modeling of sentence processing (reading), stats

Let's face it: Many people hate stats

"Facts are stubborn things, but statistics are pliable"

(Attributed to Mark Twain)

"The only statistics you can trust are those you falsified yourself"

(Attributed to Winston Churchill)

- Implies that statistics have too many degrees of freedom, are manipulative, untrustworthy, and distort the truth
- Reveals a simplistic understanding of what a "fact" is

TYPES OF HEADACHES



MIGRAINE



TENSION



STRESS



STATISTICS



Dr. J. J. J.

Folgen



Doing statistics should be like going to the bathroom. Yes, you have to do it. Yes, when you do it, you want to do it right. But don't make a big deal out of it, be careful about telling other people how to do it, and if your whole life is centered on it, there's something wrong.

Statistics (from German: Statistik, orig. "description of a state, a country") is the discipline that concerns the **collection, organization, analysis, interpretation, and presentation of data**. In applying statistics to a scientific, industrial, or social problem, it is conventional to begin with a **statistical population or a statistical model** to be studied. Populations can be diverse groups of people or objects such as "all people living in a country" or "every atom composing a crystal". Statistics deals with every aspect of data, including the planning of data collection in terms of the design of surveys and experiments.

- Importance of **generalizing** from a sample to a population ("Grundgesamtheit")
- **Descriptive statistics** based on the sample (plots, means, ...) versus **inferential statistics** (inferring population characteristics from the sample)

Common sources of confusion: Statistical mismatches

- **Type 1 mismatch:**

Casual observation (plus maybe intuitive plausibility) suggests one thing, descriptive/inferential statistics/logical reasoning says otherwise

- **Type 2 mismatch:**

Descriptive statistics (plus maybe intuitive plausibility) suggest one thing, inferential statistics say otherwise ("not significant")

- **Type 3 mismatch:**

Inferential statistics (plus maybe intuitive plausibility) say one thing, repeated sampling says otherwise ("did not replicate")

→ Unholy trinity of why we can't have nice things

Bone find proves early man lived horizontally beneath the earth



Leicester (dpo) - Our ancestors climbed down from the trees – or did they? When British Archaeologists uncovered the approximately 20,000-year-old body of a man this weekend, they discovered some astonishing news. The anatomy and position of the body lead them to the conclusion that people back then actually lived beneath the earth, moved horizontally and looked like skeletons. Many textbooks will now have to be re-written.

(Der Postillon)

Type 1 mismatch

"I took the homeopathic/herbal/... remedy and felt better afterwards"

"I ate X and had a migraine attack a few hours later"

"We performed the rain dance and on the next day it rained"

"The economy is much worse since X came into power"

"The company's sales have gone up since X became CEO"

"The team is winning again after we replaced coach X"



Canadian Hockey League

<https://chl.ca> › [ohl-storm](#) › [stor...](#) · [Diese Seite übersetzen](#) ⋮

Storm hires new coach; snaps 15 game losing streak

Storm hires **new coach**; snaps 15 game **losing streak** · [Post navigation](#).



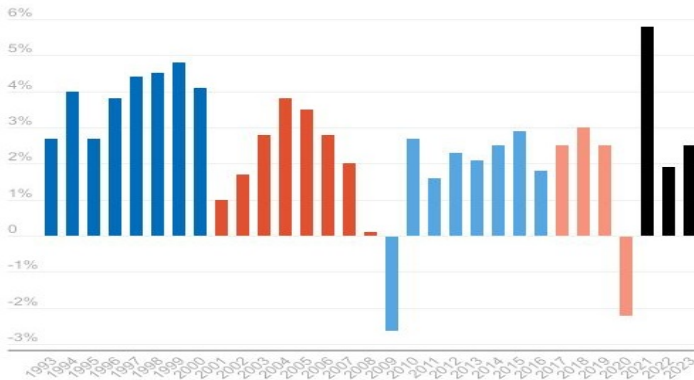
CBS News

<https://www.cbsnews.com> › [news](#) · [Diese Seite übersetzen](#) ⋮

Voters remember Trump's economy as being better than ...

06.03.2024 — In fact, almost 6 in 10 voters polled by CBS News described the U.S. **economy under Biden** as bad, even as economists' views are much more upbeat ...

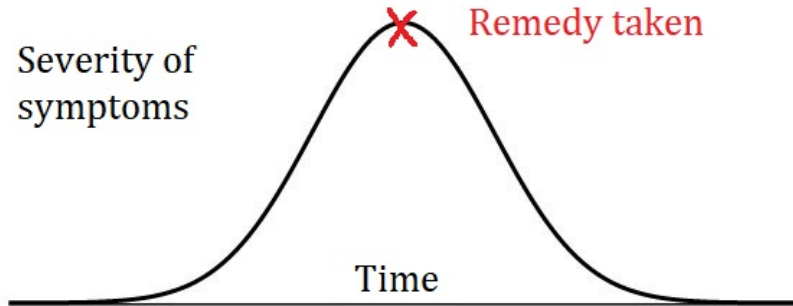
Yearly change in inflation-adjusted gross domestic product compared to previous year. Color code: Bill Clinton, George W. Bush, Barack Obama, Donald Trump, Joe Biden.



Source: Bureau of Economic Analysis

US GDP change by president

- Due to the way our minds evolved, we are **cognitively biased** to see patterns and assume causality (e.g., Matute et al., 2015)
- We often have to make decisions and build a world model based on **incomplete information**
- Many cognitive biases appear to be unrelated to intelligence, personality (Wiseman & Watt, 2006)
- There are many other cognitive biases that can interact with the causality bias: confirmation bias, sampling bias, availability bias, base rate neglect, framing/anchoring, ...



- This is an example of **regression towards the mean**
- Another example is non-causal performance change following criticism/praise

Scientific solutions to the dilemma:

- Collect **large amounts of data** (more representative of population)
- Aim for **objectivity and consensus** (standardized methods, peer review)
- Consider possible **confounds** (external factors)
- Design appropriate **controls** (e.g., placebo trials)

TABLE 2 | Contingency matrix showing a situation in which the outcome occurs with high probability but with no contingency.

	Outcome present	Outcome absent
Cause present	80	20
Cause absent	80	20

For instance, 80% of the patients who took a pill recovered from a disease, but 80% of patients who did not take the pill recovered just as well.



Plato's allegory of the cave

(4edges, Wikipedia)

Type 2 mismatch (descriptive $\neq$ inferential)

- Thought experiment:

Coach Bernard has led the men's ice hockey team to two consecutive Harris Cup wins. The team currently stands at 50 wins to 32 losses.

How many games would the team have to lose for you to conclude that Bernard has "lost his touch" and needs to be replaced?

- We intuitively recognize that many patterns we observe contain an **element of chance or randomness**
- There is a **threshold** beyond which we tend to assume that the pattern is "real"

- Hypothetical drug trial: Pain rated on 1-10 scale before and after treatment
- Hypothetical study has only one participant in each “arm”

Scenario A			Scenario B		
	Pain before	Pain after		Pain before	Pain after
Drug	9.8	7.9	Drug	9.9	7.8
Placebo	9.9	8.0	Placebo	9.8	8.0
Scenario C			Scenario D		
	Pain before	Pain after		Pain before	Pain after
Drug	9.8	7.5	Drug	9.8	6.1
Placebo	10	8.0	Placebo	10	8

- Would your conclusions change with more participants per group? If yes, why?
- People are a source of randomness; bigger samples decrease possible impact of random variability

— Question: Are the following sequences of 10 coin tosses (H = heads, T = tails) random or predetermined? (Van Prooijen et al., 2018)

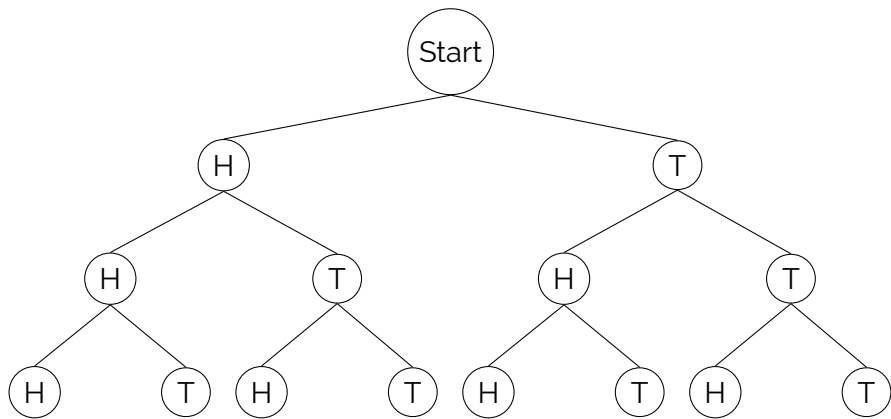
(1) HTTHHTTTHT

(2) HHHHHHHHHH

(3) HHHHHTTTTT

(4) HTHTHTHTHT

→ "Heads only" is **as likely as any other outcome** for a fair coin



$$P(TTT) = P(HTT) = P(HTH) = 0.5 * 0.5 * 0.5$$

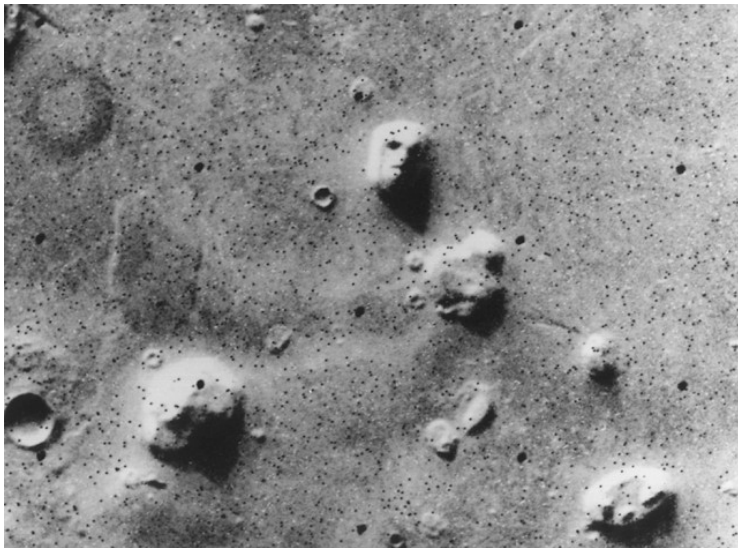
It has frequently been suggested that irrational beliefs are rooted in **pattern perception**, that is, the automatic tendency to make sense of the world by identifying meaningful relationships between stimuli (Zhao, Hahn, & Osherson, 2014). This is a functional process, as it enables people to recognize basic patterns that are real, and that are important to internalize (e.g., a red traffic light signals danger; drinking water quenches one's thirst; being unfriendly to a stranger may elicit an unfriendly response). [P]eople often have difficulty recognizing when stimuli do or do not occur through a random process. For instance, **truly random sequences typically display less variation — and hence form more clusters — than people intuitively expect**, creating the feeling of meaningful patterns that in fact occurred at random (Falk & Konold, 1997). Put differently, **a random process often generates sequences that appear nonrandom to the human mind**, and that may even contain occasional symmetries or aesthetic regularities.

(Van Prooijen et al., 2018, p. 321)

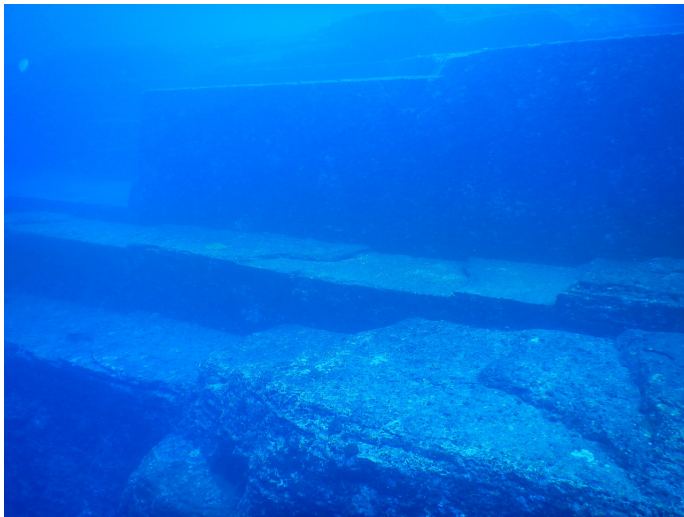
→ In their study, illusory pattern perception was correlated with belief in conspiracy theories and "the supernatural"



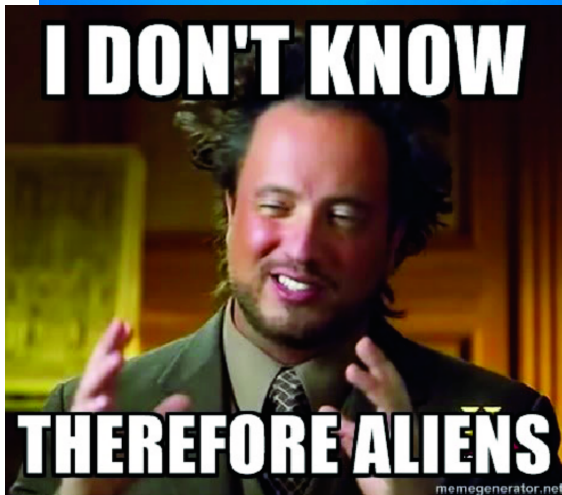
Fairy circles, Namib desert
(Stephan Getzin)

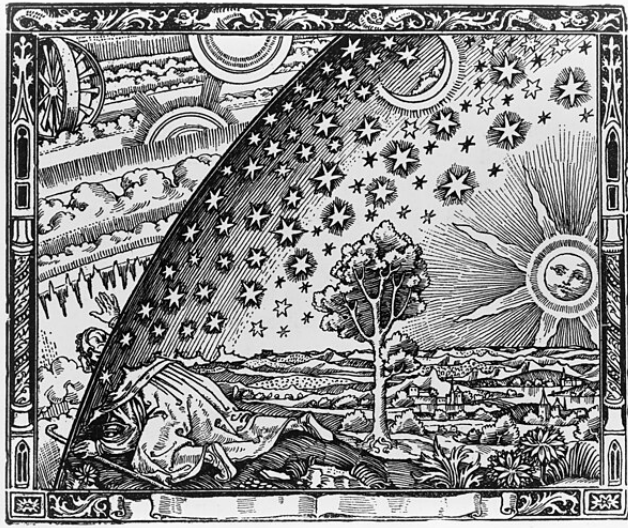


Cydonia Mensae, Mars









"Get me out of here!"

JACOBI BERNOULLI,
Profess. Basil. & utriusque Societ. Reg. Scientiar.
Gall. & Pruss. Sodal.
MATHEMATICI CELEBERRIMI,
ARS CONJECTANDI,
OPUS POSTHUMUM.

Accedit
TRACTATUS
DE SERIEBUS INFINITIS,

Et EPISTOLA Gallicè scripta
DE LUDO PILÆ
RETICULARIS.



BASILEÆ

Back to coin tosses

- What we would like is to be able to quantify **how likely an observed pattern is to be due to “chance” (random variability)**
- “How likely are these data assuming a fair coin?” $P(T) = P(H) = 0.5$

(1) HTTHHHTTHT

(2) HHHHHHHHHH

(3) HHHHHTTTTT

(4) HTHTHTHTHT

→ “Heads only” is **as likely as any other specific outcome** for a fair coin

→ ... but it is **more** likely for a heads-only coin!

- Turns out the number of “successes” in n independent “experiments” is described by a **probability distribution**, namely the **binomial distribution**:

In general, if the **random variable** X follows the binomial distribution with parameters $n \in \mathbb{N}$ and $p \in [0, 1]$, we write $X \sim B(n, p)$. The probability of getting exactly k successes in n independent Bernoulli trials (with the same rate p) is given by the **probability mass function**:

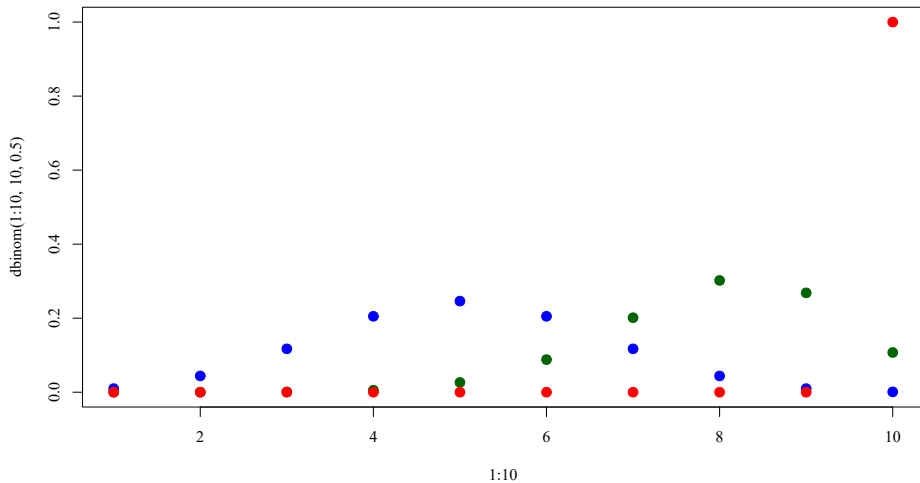
$$f(k, n, p) = \Pr(X = k) = \binom{n}{k} p^k (1 - p)^{n-k}$$

for $k = 0, 1, 2, \dots, n$, where

$$\binom{n}{k} = \frac{n!}{k!(n-k)!}$$

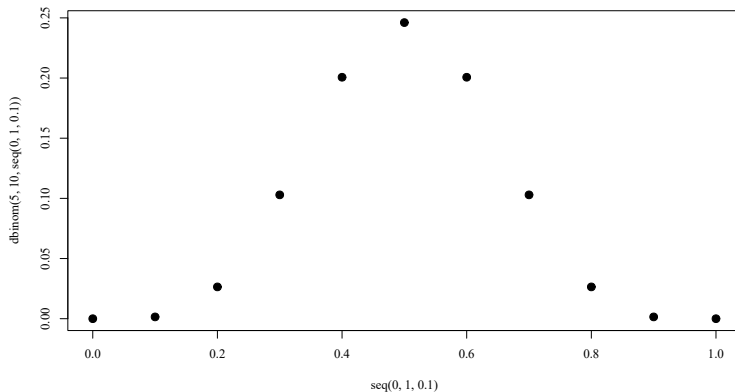
(Wikipedia)

- This is our model of the **data-generating process**
- Note that the meaning of “random” here is a bit different from its everyday usage (“**random within bounds**”)



$p = 0.5$ — $p = 0.8$ — $p = 1$

- Turning the logic on its head: Which value of p is the one that **maximizes the PMF** for a given k and n ?
- This is called the **likelihood function**
- The **maximum likelihood estimate** (MLE) of p can be easily identified for $k = 5, n = 10$:



- This is a step forward, but the MLE itself is still descriptive – can we quantify how **sure** we are about the underlying p ?
- (Turns out that the MLE also corresponds to the empirical proportion, which can be directly observed in the sample)
- In frequentist statistics, we always need to think about **repeated sampling** – what would happen if we repeated the same procedure N times?
 - Generalizing from the sample to the “population”
- Turns out there are ways to “simulate”/**estimate** this based **on one sample alone**
- Basic logic: If we repeated the procedure N times, what would be the **average** and the **spread** of the samples?
 - Sampling distribution converges to true population parameter(s)

- **(Arithmetic) mean** and **standard error** are usually used to index averages and (expected) spread:

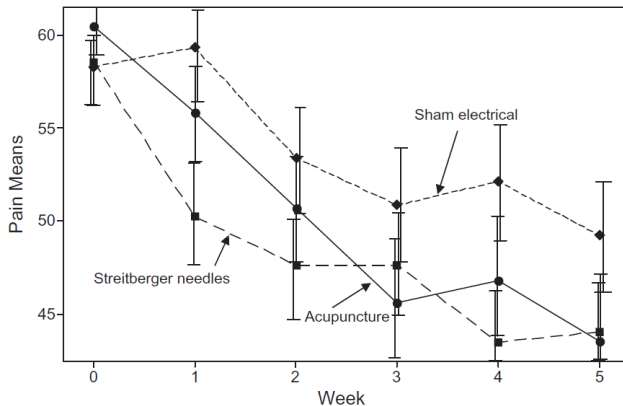
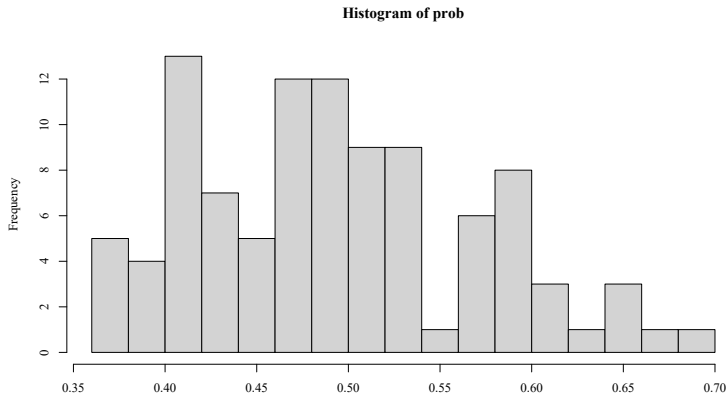


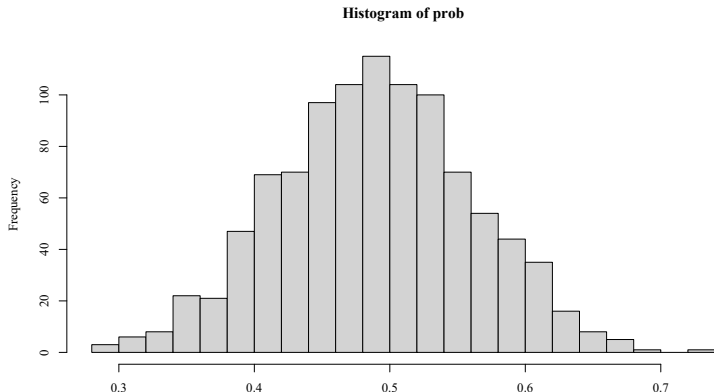
Fig. 2. Mean pain scores from pretreatment to week 5 by treatment group. NB: Error bars are based on the means plus and minus the standard error of the mean (they are not confidence intervals)

```
nrep <- 100 # Number of repeated "experiments"
eachexp <- 50 # Number of coin tosses in each "experiment"
rep <- rbinom(nrep,eachexp,0.5) # Simulate with p = 0.5
prob <- rep/eachexp # Convert no. of successes into proportion

hist(prob,breaks=nrep/5)
```

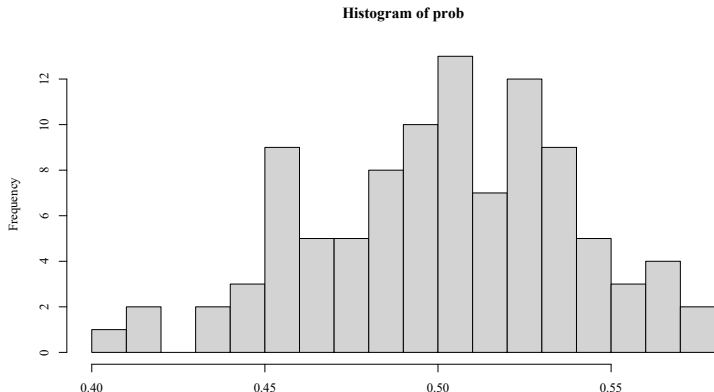


```
nrep <- 1000 # Number of repeated "experiments"  
eachexp <- 50 # Number of coin tosses in each "experiment"  
rep <- rbinom(nrep,eachexp,0.5) # Simulate with  $p = 0.5$   
prob <- rep/eachexp # Convert no. of successes into proportion  
  
hist(prob,breaks=nrep/5)
```



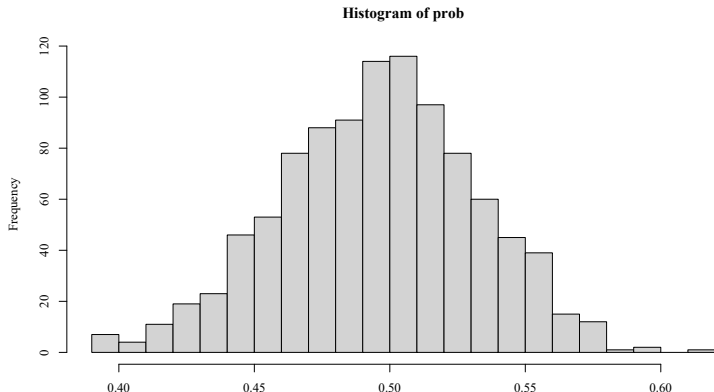
```
nrep <- 100 # Number of repeated "experiments"
eachexp <- 200 # Number of coin tosses in each "experiment"
rep <- rbinom(nrep,eachexp,0.5) # Simulate with  $p = 0.5$ 
prob <- rep/eachexp # Convert no. of successes into proportion

hist(prob,breaks=nrep/5)
```



```
nrep <- 1000 # Number of repeated "experiments"
eachexp <- 200 # Number of coin tosses in each "experiment"
rep <- rbinom(nrep,eachexp,0.5) # Simulate with p = 0.5
prob <- rep/eachexp # Convert no. of successes into proportion

hist(prob,breaks=nrep/5)
```



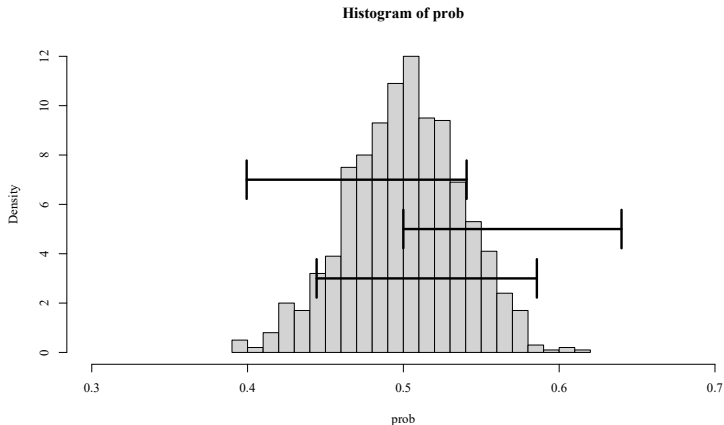
- Increasing the number of “experiments” doesn’t affect the true shape of the distribution, but more experiments are more representative
- Increasing the number of tosses per “experiment” makes the distribution **tighter** – certainty about the underlying true probability increases!
- The spread of the sampling distribution is called the **standard error**, which is equal to the **standard deviation** (root of variance) of the sampling distribution
- The binomial standard error can be estimated from a single sample:

$$\sqrt{\frac{\hat{p} \cdot (1 - \hat{p})}{n}}$$

```
> sqrt(prob[1]*(1-prob[1])/eachexp)
[1] 0.03534827
> sd(prob)
[1] 0.03629616
```

- The standard error is the basis of the **95% confidence interval**: for a normal distribution, 95% of values will be within ± 2 SD of the mean
- We can take the **mean of our sample**, which is our best estimate of the mean of the sampling distribution, and compute ± 2 standard errors from there to cover 95% of the sampling distribution
- Informally, we can say that we are 95% sure that the **true parameter** lies within this range
- (More formally, if we repeat this procedure 100 times, we expect 95 confidence intervals to contain the true value)

```
mn <- prob[1]
high <- min(mn+2*sqrt(mn*(1-mn)/eachexp),1)
low <- max(mn-2*sqrt(mn*(1-mn)/eachexp),0)
hist(prob,breaks=20,right=FALSE,freq=FALSE)
arrows(x0=low,x1=high,y0=7,y1=7,code=3,angle=90,lwd=3)
```



Key insights so far

- Randomness exists in abundance
- Products of random fluctuation can appear non-random
- To decide between (probably) random and (probably) non-random outcomes, we need to know the **bounds** of randomness
- Uncertainty **always remains**

- In principle, the statistical pipeline **works** and offers relatively little room for manipulation
- A single study can nevertheless give a **false positive** if the sample happens to be “strange” (non-representative)

But there are lots of problems in reality (Type 3 mismatch):

- Assuming the wrong generating distribution, other errors in the analysis
- Incentives to publish “significant” results (null hypothesis value outside 95% interval) and put non-significant ones in the “file drawer”
- Manipulated or straight-out fake data
- Testing too many things at the same time (increases risk of false positives; hello ML!)
- Overinterpretation of results (either in the paper or in secondary media)



Having an older sibling increases the odds of being gay by a fifth, major study finds

Being a younger sibling increases your chances of being gay, according to a landmark study of Brits.

4 days ago



 Daily Mail

Having an older sibling increases the odds of being gay by a fifth



 Times Now

Being a Stanford University Study Finds That Going To Bed After This Specific Time Could Be Harmful



4 days The study conducted at Stanford seems to find that those people who wake up early and fall asleep early consistently face fewer mental...

07.07.2024



Daily Mail



H
fi

W The Witness | Your compass in the community

Be
st

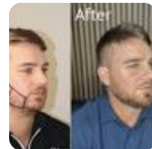
Beard transplants on the rise as study finds that men with beards are more attractive

4

A recent study by dating.com found that men with beards are more attractive to singles, with 86% of respondents preferring facial hair.

30.06.2024

07.07.2024



666 0 11 11

S Surrey Live

Finger length could be linked to how ill Covid will make you, study suggests

The length of your fingers could indicate how ill you will get if you contract Covid, according to a new study from the University of...

7 Apr 2022

07.07.2024





MinnPost

It may be counterintuitive, but we're actually better at remembering names than faces, study finds

Think you're bad at remembering names? You're probably worse at remembering faces. At least, that's what a new British study suggests.

7 Apr 2022

07.07.2024



Ψ PsyPost

Study: 'Bad trips' from magic mushrooms often result in an improved sense of personal well-being

New research suggests a bad trip isn't always bad. About 84 percent of drug users who have experienced a "bad trip" from hallucinogenic...

30 Aug 2016



07.07.2024

Study offers new evidence that drinking coffee can protect against Alzheimer's disease

Drinking coffee may reduce your risk of Alzheimer's disease, according to findings published in the journal *Frontiers in Aging Neuroscience*.

20 Apr 2022

ε ' 30 Aug 2016

7 Apr. ----

07.07.2024



Ψ PsyPost

hub. Newshub

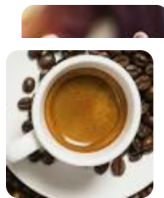
Why men should drink filtered coffee and women should drink espressos

A new development has emerged in the age-old debate as to whether the humble cup of joe is good for your heart.

13 May 2022

/ P. ---

07.07.2024



Humans

The replication crisis has spread through science – can it be fixed?

It started in psychology, but now findings in many scientific fields are proving impossible to replicate. Here's what researchers are doing to restore science's reputation

Why Most Published Research Findings Are False

John P. A. Ioannidis

Summary

There is increasing concern that most current published research findings are false. The probability that a research claim is true may depend on study power and bias, the number of other studies on the same question, and, importantly, the ratio of true to no relationships among the relationships probed in each scientific field. In this framework, a research finding is less likely to be true when the studies conducted in a field are smaller; when effect sizes are smaller; when there is a greater number and lesser preselection of tested relationships; where there is greater flexibility in designs, definitions, outcomes, and analytical modes; when there is greater financial and other interest and prejudice; and when more teams are involved in a scientific field in chase of statistical significance. Simulations show that for most study designs and settings, it is more likely for a research claim to be false than true. Moreover, for many current scientific fields, claimed research findings may

factors that influence this problem and some corollaries thereof.

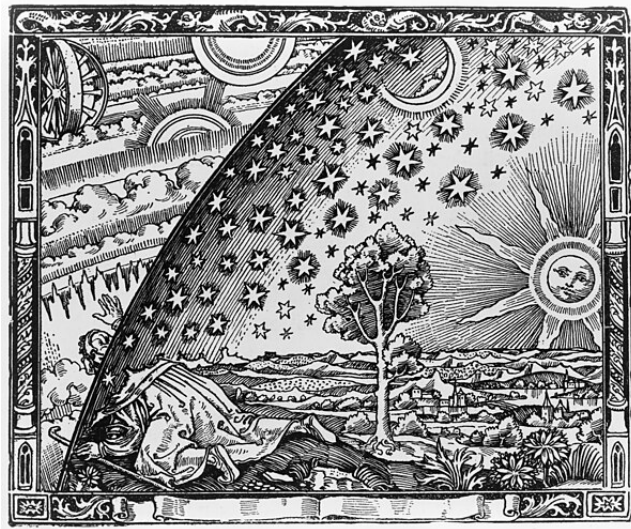
Modeling the Framework for False Positive Findings

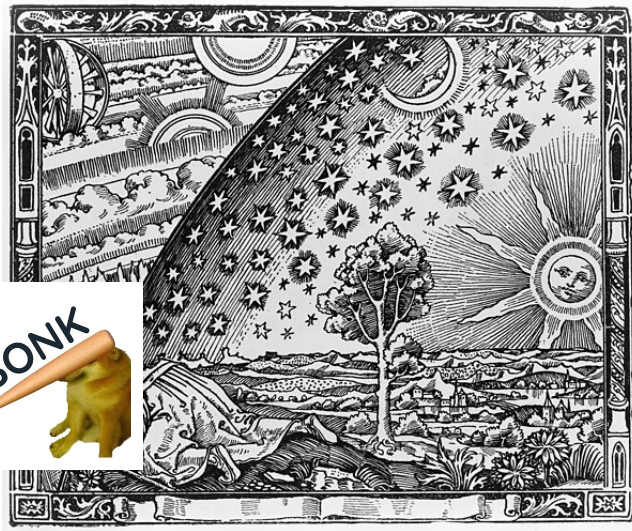
Several methodologists have pointed out [9–11] that the high rate of nonreplication (lack of confirmation) of research discoveries is a consequence of the convenient, yet ill-founded strategy of claiming conclusive research findings solely on the basis of a single study assessed by formal statistical significance, typically for a p -value less than 0.05. Research is not most appropriately represented and summarized by p -values, but, unfortunately, there is a widespread notion that medical research articles

It can be proven that most claimed research findings are false.

should be interpreted based only on p -values. Research findings are defined

is characteristic of the field and can vary a lot depending on whether the field targets highly likely relationships or searches for only one or a few true relationships among thousands and millions of hypotheses that may be postulated. Let us also consider, for computational simplicity, circumscribed fields where either there is only one true relationship (among many that can be hypothesized) or the power is similar to find any of the several existing true relationships. The pre-study probability of a relationship being true is $R/(R + 1)$. The probability of a study finding a true relationship reflects the power $1 - \beta$ (one minus the Type II error rate). The probability of claiming a relationship when none truly exists reflects the Type I error rate, α . Assuming that c relationships are being probed in the field, the expected values of the 2×2 table are given in Table 1. After a research finding has been claimed based on achieving formal statistical significance, the post-study probability that it is true is the positive predictive value, PPV.

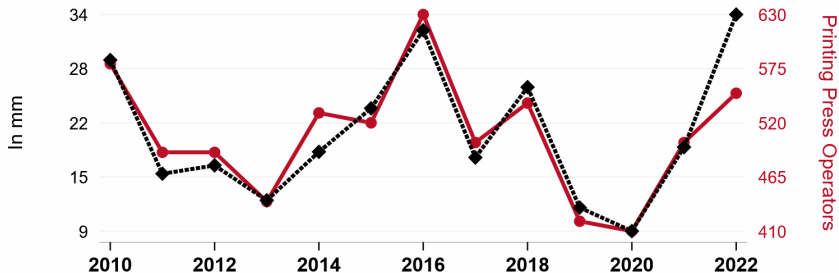




Rainfall in San Francisco

correlates with

The number of printing press operators in Rhode Island



◆-- Rainfall in San Francisco · Source: Golden Gate Weather Service

●— BLS estimate of printing press operators in Rhode Island · Source: Bureau of Labor Statistics

2010-2022, $r=0.912$, $r^2=0.832$, $p<0.01$ · tylervigen.com/spurious/correlation/2948

Number of people who drowned by falling into a pool

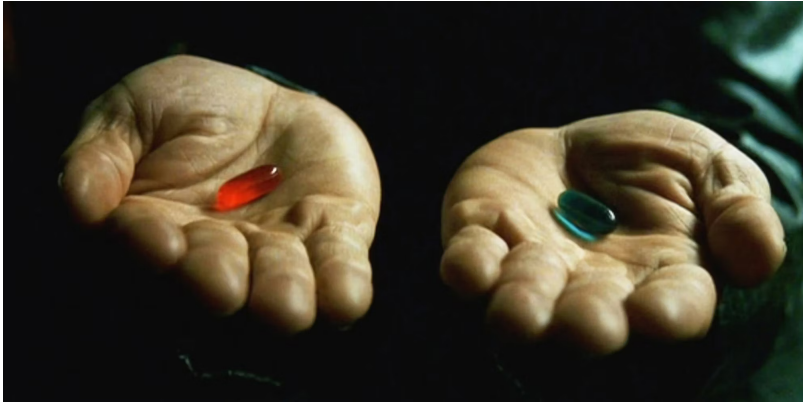
correlates with

Films Nicolas Cage appeared in

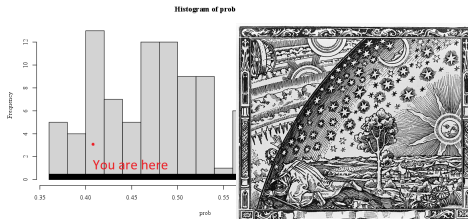
Correlation: 66.6% ($r=0.666004$)



tylervigen.com



- Stats is not your enemy
- The machinery of inferential statistics is needed to look beyond the specific sample and quantify uncertainty about the “truth”
- Machinery is mathematically justified (which can also be shown through simulation, where the truth is known)
- In reality, the truth is not known, and rarely (if ever) self-evident



- It's **really** important to cultivate doubt about published findings (without becoming anti-science)

Relationship to ML

- ML and stats use much of the same machinery
- Stats usually focuses on problems with **relatively few, human-interpretable variables** in comparatively “**small data**” environments
- A **single** variable is often of primary interest – is the difference between X and Y significantly different from zero or not?
- Different **aims**: “Which model can better distinguish between cats and dogs?” vs. “Is there even a measurable difference between cats and dogs, and if so, what is it?” (**novel discoveries**)
- The null hypothesis has a special status; **skepticism** is very important
- Fitting of models is (usually) super fast and doesn't require much computing power