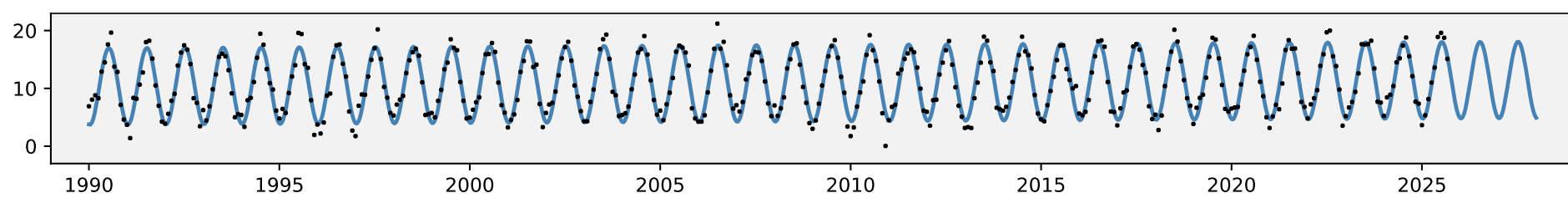


Iterative model development



At first glance, this looks like a simple periodic model will fit.

Haven't you heard of global warming 🌍?! I'll add a linear trend, say $\gamma^\circ\text{C}/\text{century}$. The mle is $\hat{\gamma}=3.0196$. 😎

Is your model really better than mine, or is it just your choice? 😏

Mine has a higher likelihood, which means it fits the data better. 💪 📈

I'm 95% confident that $1.7 \leq \hat{\gamma} \leq 4.3$, so I'm confident there's a trend. 😌
Also, when I try a quadratic, I'm not confident the extra term is non-zero. So I'll stick with linear. 🧘

OK, but maybe there's something else you haven't thought of 🧐. Why do you think your model is right? 😐

QUESTION. Why is this a silly question?

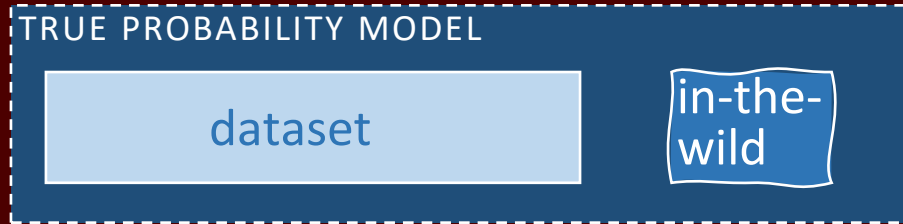


QUESTION. What's wrong with this argument?

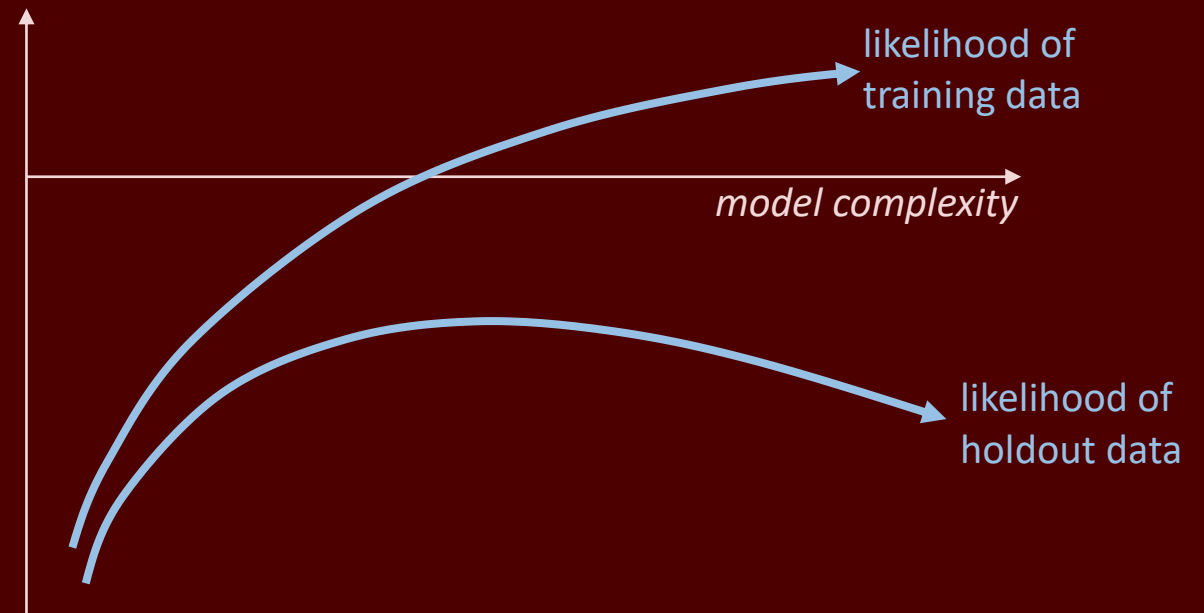
Overfitting! A more complex model always scores a higher likelihood on the training dataset. But that doesn't mean it's closer to the true distribution.
Remedies: holdout comparison, or Bayesian model choice, frequentist confidence intervals ... anything that addresses the difference between the data and the truth.

OVERFITTING AND HOLDOUT EVALUATION

A too-complex model will typically fit the dataset well, but generalize poorly because it doesn't match the true probability model.



We can measure this by using holdout data to approximate the true distribution.



```
# Periodic model
```

```
model0 = sklearn.linear_model.LinearRegression()
```

```
def X0(t): return np.column_stack([np.sin(2* $\pi$ *t), np.cos(2* $\pi$ *t)])
```

```
model0.fit(X0(df.t), df.temp)
```

```
# Model with linear trend
```

```
model1 = sklearn.linear_model.LinearRegression()
```

```
def X1(t): return np.column_stack([np.sin(2* $\pi$ *t), np.cos(2* $\pi$ *t), (t-2000)/100])
```

```
model1.fit(X1(df.t), df.temp)
```

```
model1.coef_[-1]
```

```
# Compare Log Likelihoods
```

```
def loglik( $\epsilon$ ):
```

```
     $\sigma$  = np.sqrt(np.mean( $\epsilon$ **2))
```

```
    return - 1/2*np.sqrt(2* $\pi$ * $\sigma$ **2) - 1/(2* $\sigma$ **2)*np.mean( $\epsilon$ **2)
```

```
loglik(df.temp - model0.predict(X0(df.t))), loglik(df.temp - model1.predict(X1(df.t)))
```

```
# Periodic model
model0 = sklearn.linear_model.LinearRegression()
def X0(t): return np.column_stack([np.sin(2* $\pi$ *t), np.cos(2* $\pi$ *t)])
model0.fit(X0(df.t), df.temp)
```

```
# Model with linear trend
model1 = sklearn.linear_model.LinearRegression()
def X1(t): return np.column_stack([np.sin(2* $\pi$ *t), np.cos(2* $\pi$ *t), (t-2000)/100])
model1.fit(X1(df.t), df.temp)
model1.coef_[-1]
```

```
# Confidence interval for the trend term
```

```
# 1. Define the readout statistic, yhat
m = sklearn.linear_model.LinearRegression()
x = X1(df.t)
def trend(temp): return m.fit(x,temp).coef_[-1]
```

```
# 2. Fit the model, and generate resampled data based on the fitted model
pred1 = model1.predict(X1(df.t))
 $\epsilon$  = df.temp - model1.predict(X1(df.t))
 $\sigma$  = np.sqrt(np.mean( $\epsilon$ **2))
def tempstar(): return np.random.normal(loc=pred1, scale= $\sigma$ )
```

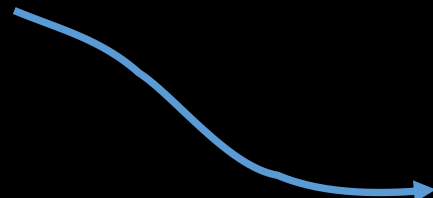
```
# 3. Find the spread of the readout statistic
y_ = [trend(tempstar()) for _ in range(20000)]
lo,hi = np.quantile(y_, [.025,.975])
lo,hi
```

Once we've tried all the models we can think of and chosen the best, what next?

1. Eyeball our model's fit, to see if we can spot any specific improvements

- Compute the residuals $\varepsilon_i = y_i - \text{pred}(x_i)$ and look for patterns
- More generally, compute the likelihood of each datapoint, and investigate datapoints that our model thinks are unlikely

2. Ask: is it plausible that our model might have generated the dataset?



this is Hypothesis Testing

RESEARCH QUESTION

Can you taste the difference between milk-first versus tea-first?

EXPERIMENT

Make 4 cups of each style, shuffle them, and ask taster to label them



RESEARCH QUESTION

Can you taste the difference between milk-first versus tea-first?

EXPERIMENT

Make 4 cups of each style, shuffle them, and ask the taster to label them

DATASET

8 pairs of (truth,label)

HYPOTHESIS / MODEL

There's no difference. Thus the dataset arose by chance, from purely random choice

milk first



tea first



milk first



tea first



tea first



milk first



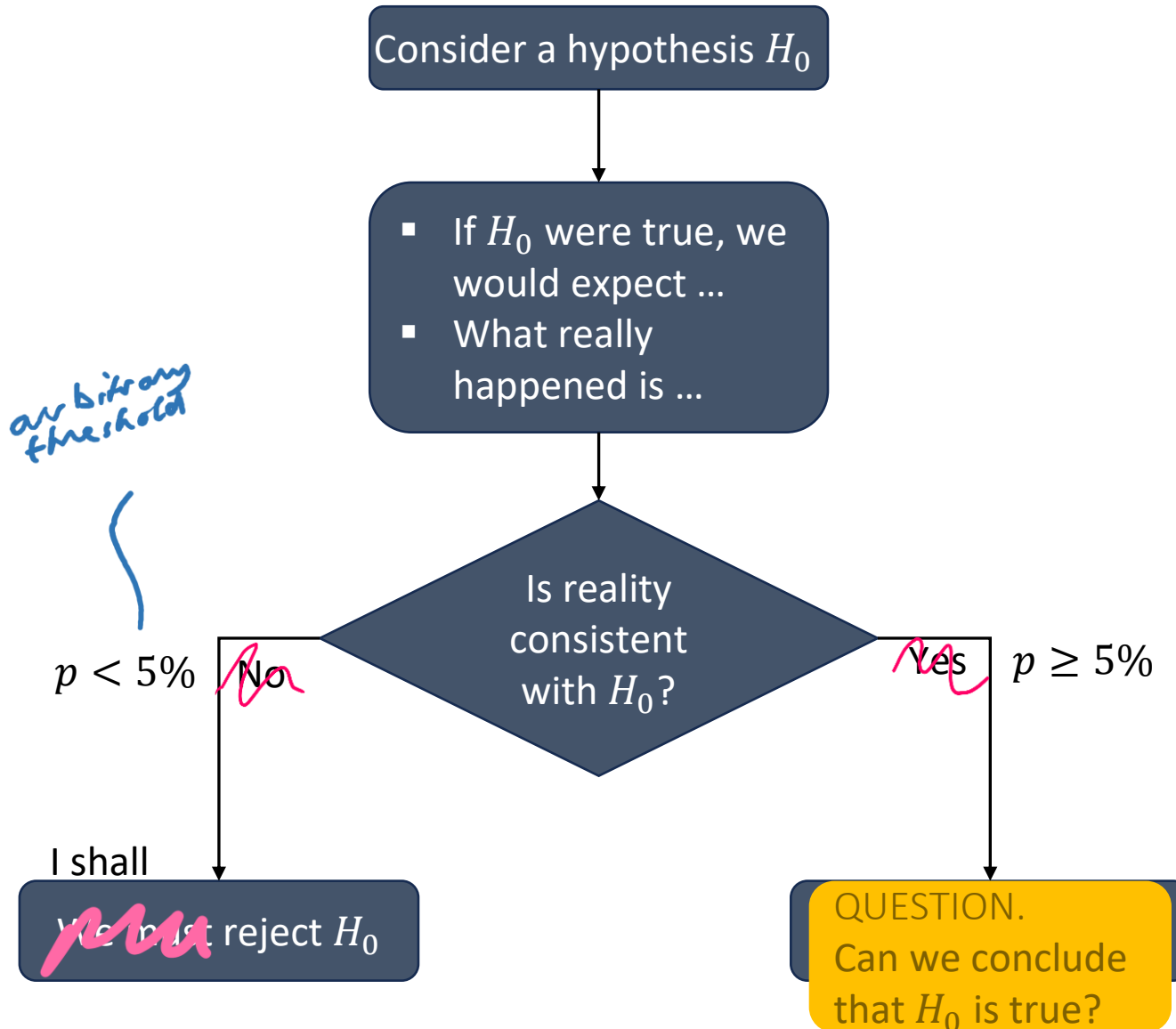
tea first



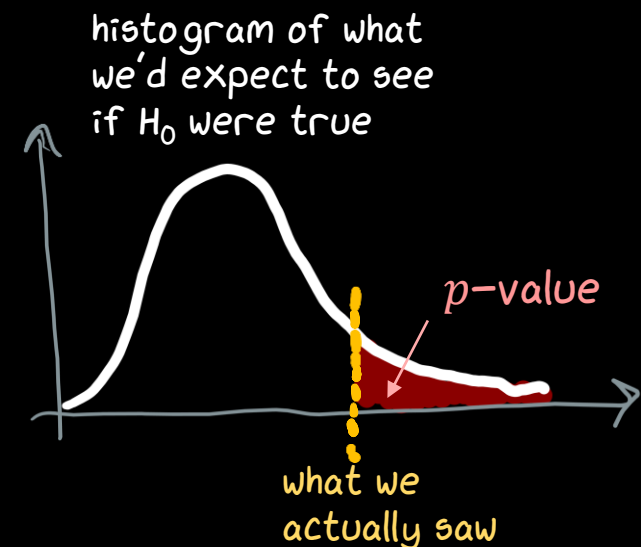
tea first

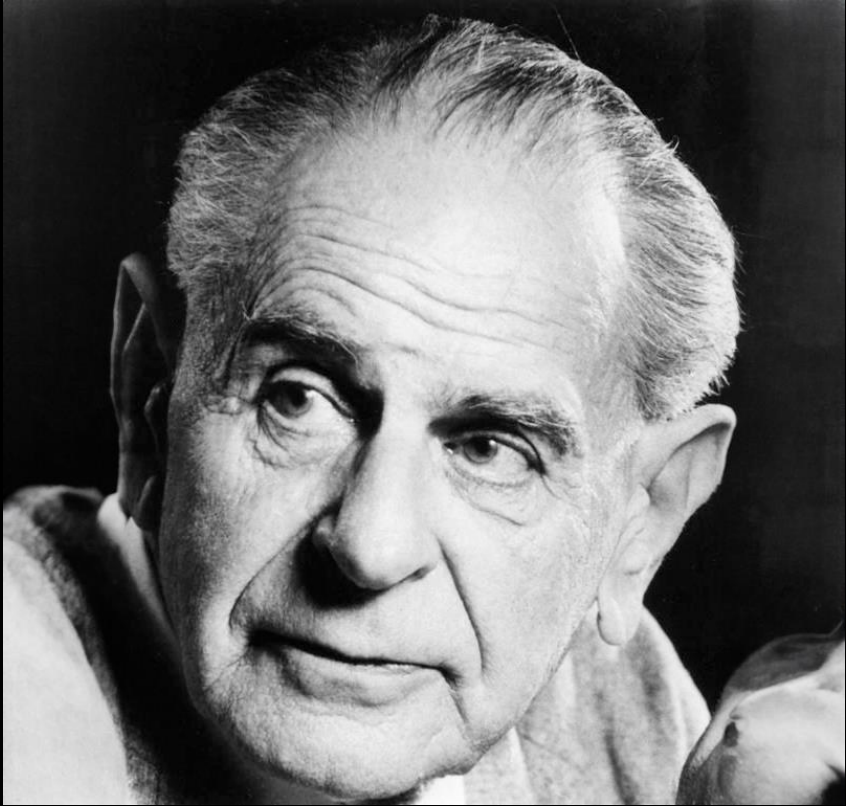


Hypothesis testing asks whether a proposed probability model H_0 is consistent with the dataset.



- ❖ Because of noise, this isn't a simple yes/no decision.
- ❖ It's about *degree of consistency*, which we measure by the p -value. Instead of yes/no, we can go by $p < 5\%$ / $p \geq 5\%$.
- ❖ The p -value is the probability, according to the model H_0 , of seeing something at least as extreme as what we actually saw.





“Every genuine scientific theory must be falsifiable. It is easy to obtain evidence in support of virtually any theory; the evidence only counts if it is the positive result of a genuinely risky prediction.”

Karl Popper (1902–1994)

Fisher's hypothesis testing

Let x be the dataset.

State a null hypothesis H_0 , i.e. a probability model for the dataset

1. Choose a test statistic

$$t : \text{dataset} \mapsto \mathbb{R}$$

2. Define a random synthetic dataset X^* , what we might see if H_0 were true.

3. Look at the histogram of $t(X^*)$.
Let p be the probability of seeing a value as extreme or more so than the observed $t(x)$.

A low p -value is a sign that H_0 should be rejected.

cup id	truth	taster
1	milk	milk
2	tea	tea
$\vdots$		

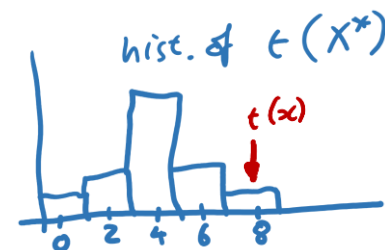
x = taster's assignment of labels

H_0 : the taster can't tell the difference and just assigned labels randomly

$$t(x) = \text{\#correct}$$

def $X^*()$: return random permutation of {t,t,t,t,m,m,m,m}

What would be the distribution of the test statistic if H_0 were true?



$$p = P(t(X^*) \geq t(x)) = 1.4\%$$

$p < 5\%$ so we shall reject H_0 .

"The evidence against H_0 is statistically significant ($p=1.4\%$)"

Example 9.6.2.

I have a dataset with readings from two groups, $x = [x_1, \dots, x_m]$ and $y = [y_1, \dots, y_n]$. Test whether the two groups are significantly different, using the test statistic $\bar{y} - \bar{x}$.

Dataset: (x, y)

H_0 : x and y are samples from the same distribution.

Test statistic: given in the question.

```
1  # 1. Define the test statistic
2  def t(x,y): return np.mean(y) - np.mean(x)

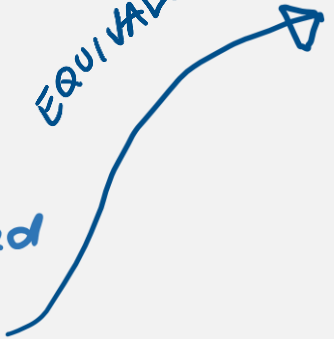
3  # 2. To generate a synthetic dataset, assuming  $H_0$ , ...
4  xy = np.concatenate([x,y])
5  def rxy_star():
6      return (np.random.choice(xy, size=len(x)),
7              np.random.choice(xy, size=len(y)))

8  # 3. Sample the test statistic under  $H_0$ ; find p-value for observed data
9  t_ = np.array([t(*rxy_star()) for _ in range(10000)])
10 p = ...
```

Example 9.3.1.

I have a dataset with readings from two groups, $x = [x_1, \dots, x_m]$ and $y = [y_1, \dots, y_n]$. Test whether the two groups are significantly different, using the test statistic $\bar{y} - \bar{x}$.

H_0 : x, y are both sampled from $N(\mu, \sigma^2)$ (μ, σ unknown).

EQUIVALENTLY 

General model:

X sampled from $N(\mu, \sigma^2)$

Y sampled from $N(\mu + \delta, \sigma^2)$

$H_0: \delta = 0.$

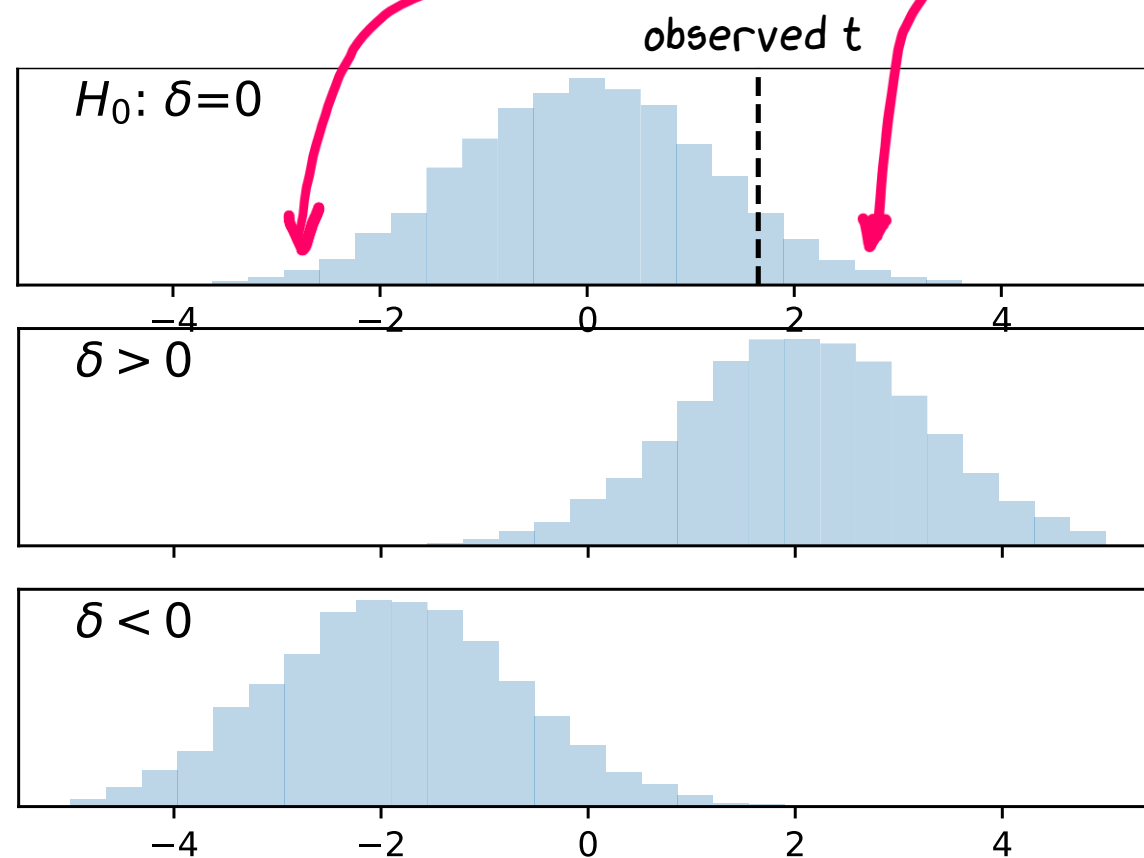
```
1 # 1. Define the test statistic
2 def t(x,y): return np.mean(y) - np.mean(x)

3 # 2. To generate a synthetic dataset, assuming  $H_0$ , ...
4 xy = np.concatenate([x,y])
5  $\hat{\mu}$  = np.mean(xy)
6  $\hat{\sigma}$  = np.sqrt(np.mean((xy -  $\hat{\mu}$ )**2))
7 def rxy_star():
8     return (np.random.normal(loc= $\hat{\mu}$ , scale= $\hat{\sigma}$ , size=len(x)),
9            np.random.normal(loc= $\hat{\mu}$ , scale= $\hat{\sigma}$ , size=len(y)))

10 # 3. Sample the test statistic under  $H_0$ ; find p-value for observed data
11 t_ = np.array([t(*rxy_star()) for _ in range(10000)])
12 p = 2 * min(np.mean(t_ >= t(x,y)), np.mean(t_ <= t(x,y)))
```

What counts as 'more extreme'?

- Plot the histogram for $t(X^*)$, assuming H_0 is true
- Also plot the histogram for some scenarios where H_0 is false
- Do the alternatives push $t(X^*)$ bigger, or smaller, or either? This determines what 'more extreme' means — either one-tailed or two-tailed.

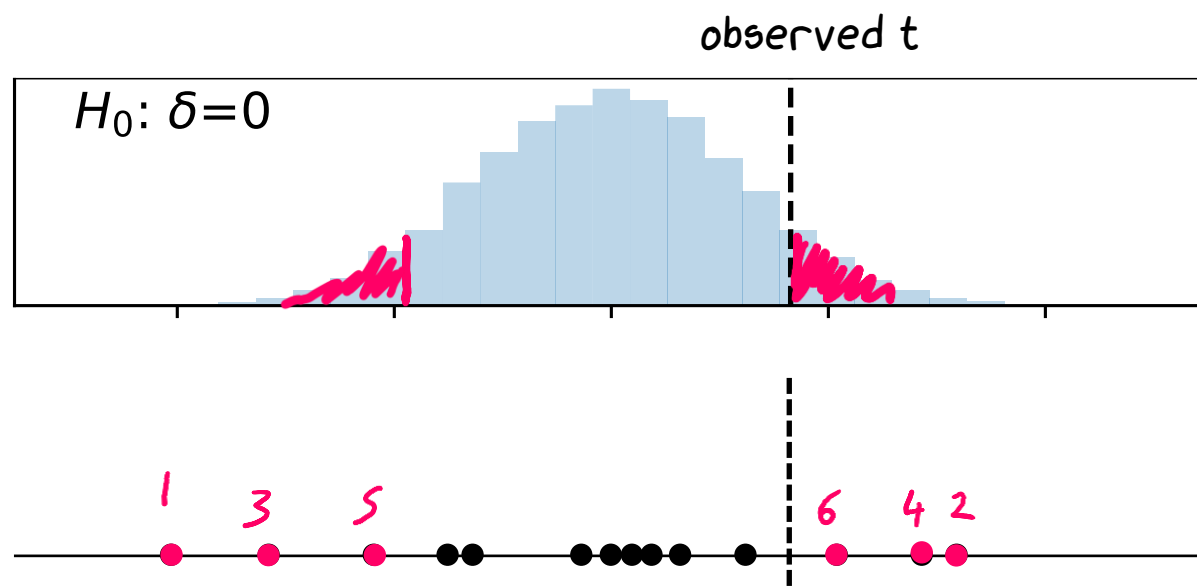


If the observed t lies at either extreme, it's evidence against $H_0: \delta=0$

How do we compute p for a two-tailed test?

The p -value is

$$\mathbb{P} \left(t(X^*) \text{ at least as extreme as } t(x) \mid H_0 \text{ is true} \right)$$



"6 of my samples of $t(X^*, Y^*)$ are more extreme than $t(x, y)$."

$$p = 2 * \min(\text{np.mean}(\mathbf{t_} \geq t(x, y)), \text{np.mean}(\mathbf{t_} \leq t(x, y)))$$

Where do test statistics come from?

Indeed, why do we even need test statistics?

We want to know if the dataset is plausible according to H_0 ,
so why not simply measure its likelihood under H_0 ?

—*this is answered in §9.4 of lecture notes.*

Example 9.3.1.

I have a dataset with readings from two groups, $x = [x_1, \dots, x_m]$ and $y = [y_1, \dots, y_n]$. Test whether the two groups are significantly different, using the test statistic $\bar{y} - \bar{x}$.

In some problems, it's natural to express H_0 as a *constraint* on the parameters of a more general model H_1 .

$$H_1: \begin{aligned} x &\sim N(\mu, \sigma^2) \\ y &\sim N(\nu, \sigma^2) \end{aligned}$$

$$H_0: \mu = \nu$$

❖ For the test statistic:

estimate the parameters under H_1 , and let t measure how close they are to meeting the constraint.

$$\begin{aligned} \hat{\mu} &= \bar{x}, \quad \hat{\nu} = \bar{y} \\ t &= |\hat{\nu} - \hat{\mu}| = |\bar{y} - \bar{x}| \quad \text{one-tailed test} \\ &\quad \text{reject } H_0 \text{ if } t \text{ large} \\ t &= \bar{y} - \bar{x}, \quad \text{two-tailed test.} \end{aligned}$$

❖ For resampling:

Fit H_0 by maximizing the likelihood of the parameters under the constraint. Then, sample from this fitted model.

$$\begin{aligned} \max_{\mu, \nu, \sigma: \mu = \nu} \log \text{Pr}(\text{data}; \mu, \nu, \sigma) \\ \text{equiv. to} \quad \max_{\lambda, \sigma} \log \text{Pr}(\text{data}; \lambda, \lambda, \sigma) \end{aligned}$$

Exercise 9.3.2 (Equality of group means).

We are given three groups of observations from three different systems

$$x = [7.2, 7.3, 7.8, 8.2, 8.8, 9.5]$$

$$y = [8.3, 8.5, 9.2]$$

$$z = [7.4, 8.5, 9.0]$$

Do all three groups have the same mean?

$$H_1: X \sim N(a, \sigma^2) \quad Y \sim N(b, \sigma^2) \quad Z \sim N(c, \sigma^2)$$

$$H_0: a = b = c$$

$$\hat{a} = \bar{x} \quad \hat{b} = \bar{y} \quad \hat{c} = \bar{z}$$

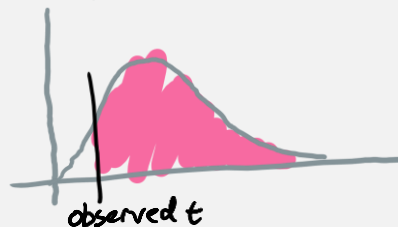
$$t = (\hat{a} - \hat{b})^2 + (\hat{b} - \hat{c})^2 + (\hat{c} - \hat{a})^2$$

$$\text{OR } t = (\hat{a} - \hat{\mu})^2 + (\hat{b} - \hat{\mu})^2 + (\hat{c} - \hat{\mu})^2$$

where $\hat{\mu}$ is MLE of

$$H_0: \text{all } N(\mu, \sigma^2)$$

hist. of sampled t



QUESTION. Why did I choose a common variance?

ANSWER: I didn't have to, and it'd be wiser to propose an H_1 that has different variances.

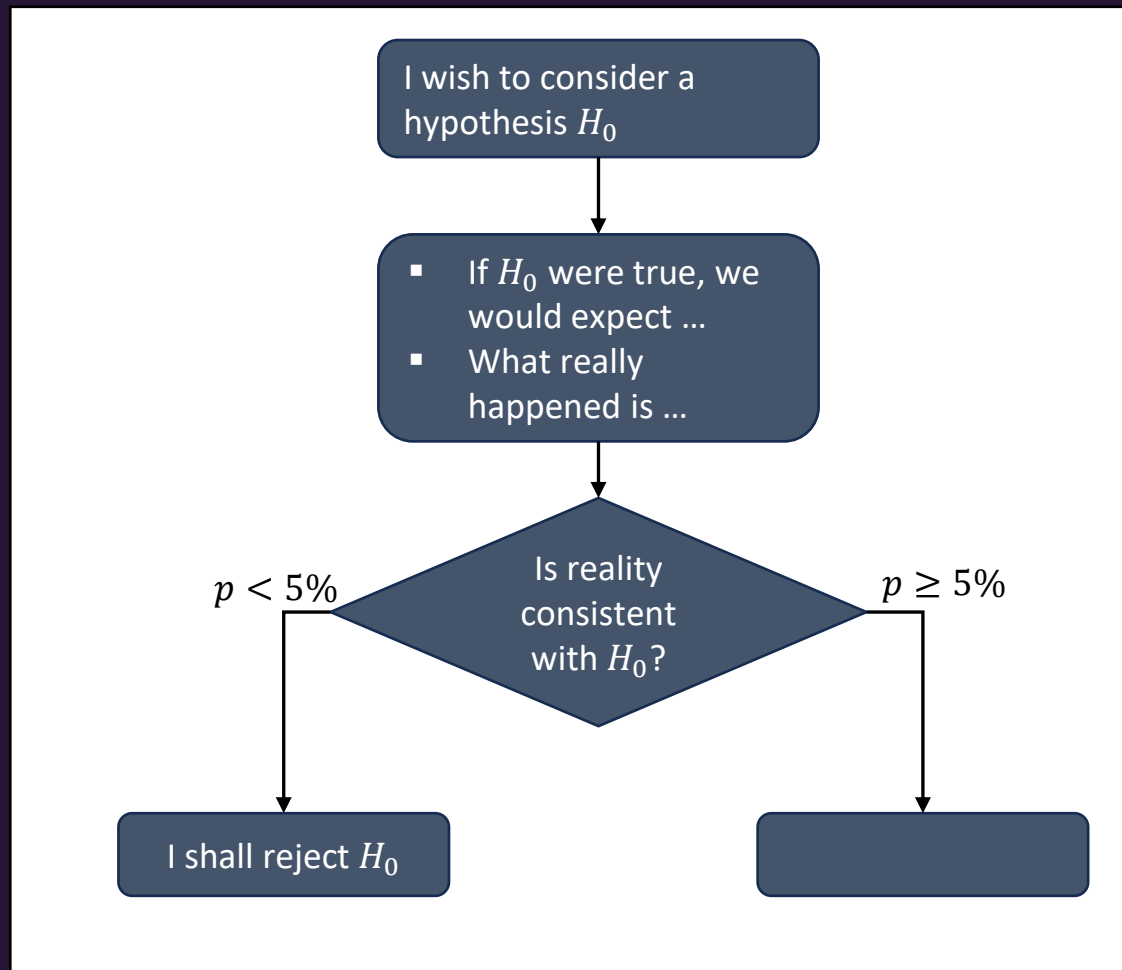
(Though, with such a tiny dataset, it's extremely unlikely we'll see evidence that rejects the hypothesis of equal variance.)

```
1 # 1. Define test statistic
2 def t(x,y,z):
3     mu = np.mean(np.concatenate([x,y,z]))
4     a,b,c = [np.mean(v) for v in [x,y,z]]
5     return (a-mu)**2 + (b-mu)**2 + (c-mu)**2

6 # 2. To generate a synthetic dataset, assuming  $H_0$  ...
7 xyz = np.concatenate([x,y,z])
8 mu_hat = np.mean(xyz)
9 sigma_hat = np.sqrt(np.mean((xyz-mu_hat)**2))
10 def rxyz_star():
11     return (np.random.normal(size=len(x), loc=mu_hat, scale=sigma_hat),
12             np.random.normal(size=len(y), loc=mu_hat, scale=sigma_hat),
13             np.random.normal(size=len(z), loc=mu_hat, scale=sigma_hat))

14 # 3. Sample the test statistic, find the p-value
15 t_ = np.array([t(*rxyz_star()) for _ in range(10000)])
16 p = np.mean(t_>=t(x,y,z))
```

“Data science is
quantitative rhetoric”



We only get a definite publishable
conclusion if we reject H_0 .

Anything we want to argue, we have to
phrase it as “reject H_0 ” for a suitable H_0 .

EXERCISE.

Here are marks for IA Algorithms questions.

I think there might be a gender bias.

What H_0 do I choose?



H_0 : I think everyone gets pretty much the same marks, regardless of gender

gender	mark
F	17
F	14
M	18
O	11
M	17
$\vdots$	$\vdots$

Null hypothesis:

Let's introduce a richer model, H_1 : $\text{Mark} \sim \mu_{\text{gender}} + N(0, \sigma^2)$
and express H_0 as a constraint: $\mu_F = \mu_M = \mu_O$

H_1 : I think gender affects marks



Test statistic: $t = (\hat{\mu}_F - \hat{\mu})^2 + (\mu_M - \hat{\mu})^2 + (\hat{\mu}_O - \hat{\mu})^2$ where $\hat{\mu}_F, \hat{\mu}_M, \hat{\mu}_O$ are MLE under H_1
and $\hat{\mu}$ is MLE under H_0

Resampling: Fit H_0 which says $\text{Mark} \sim N(\mu, \sigma^2)$ then sample from this fitted $N(\hat{\mu}, \hat{\sigma}^2)$

Here are marks for IA Algorithms questions.
I think there might be a gender bias.
What H_0 do I choose?



Under the general H_1 model
Mark $\sim \mu_{\text{gender}} + N(0, \sigma^2)$
I propose $H_0: \mu_M = \mu_F = \mu_O$

gender	mark
F	17
F	14
M	18
O	11
M	17
$\vdots$	$\vdots$

QUESTION. Suppose we reject H_0 . Does this mean we believe that the means μ_g are not all equal?

- ❖ Our H_0 claims several things at the same time:
(means are equal) & (variances are equal) & (all are Gaussian) & (all are independent)
- ❖ If we reject H_0 , we reject it in its entirety.
We might be rejecting it because of evidence against any one or more of its subclaims.

What makes a good hypothesis test?

- ❖ Your H_0 is credible to your audience.
If you propose a non-credible H_0 and then reject it — who cares?
- ❖ Your H_0 matches the research question you want to answer, and doesn't bring in contentious subclaims.
- ❖ Your resampling method assumes H_0 and nothing more.



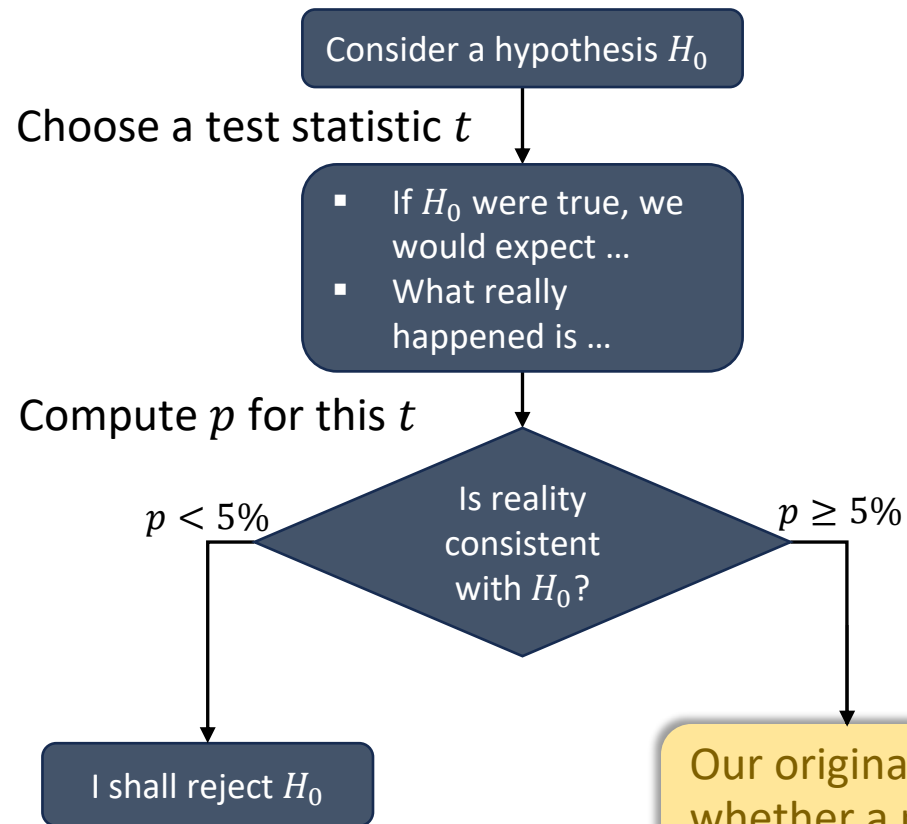
Under the general H_1 model
 $\text{Mark} \sim \mu_{\text{gender}} + N(0, \sigma^2)$
I propose $H_0: \mu_M = \mu_F = \mu_O$

H_0 : the marks have the same distribution, for each gender.

Now imagine a parallel universe where every student gets assigned a random gender (with the same totals as in our observed dataset x). Simulate this by creating a dataset X^* with a randomly permuted Gender column.

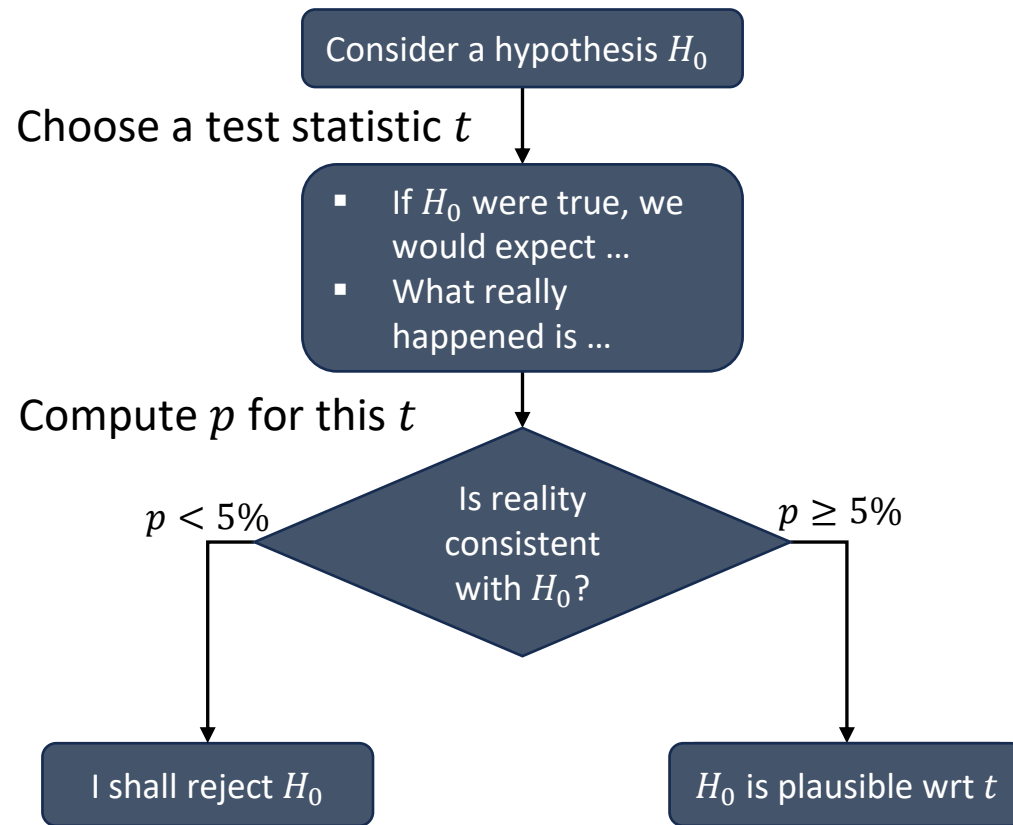
If H_0 is true, then $t(x)$ is a sample from $t(X^*)$. This lets me test H_0 .





Our original goal today was to decide whether a model is adequate.

QUESTION. Can I at least conclude that " H_0 could plausibly have generated the dataset"?



Different t are sensitive to different violations of H_0 .

If I settle for my H_0 and then someone comes up with a better model, I lose. So it's on me to test H_0 using several test statistics.

THE GREAT GENERALIZATION **SMACK DOWN** BAYESIANIST V FREQUENTIST



Climate confidence challenge

Find a 95% confidence interval for the rate of temperature increase in Cambridge from 1985 to the present, in $^{\circ}\text{C}/\text{year}$