

Die t-Verteilung und die Prüfstatistik

Jonathan Harrington

Standard error of the mean (SE)

ist die Standardabweichung von Mittelwerten

Ich werfe 5 Würfel und berechne den Mittelwert der Zahlen

$\mu = 3.5$ der wahrscheinlichste Wert

$$SE = \frac{\sigma}{\sqrt{5}}$$

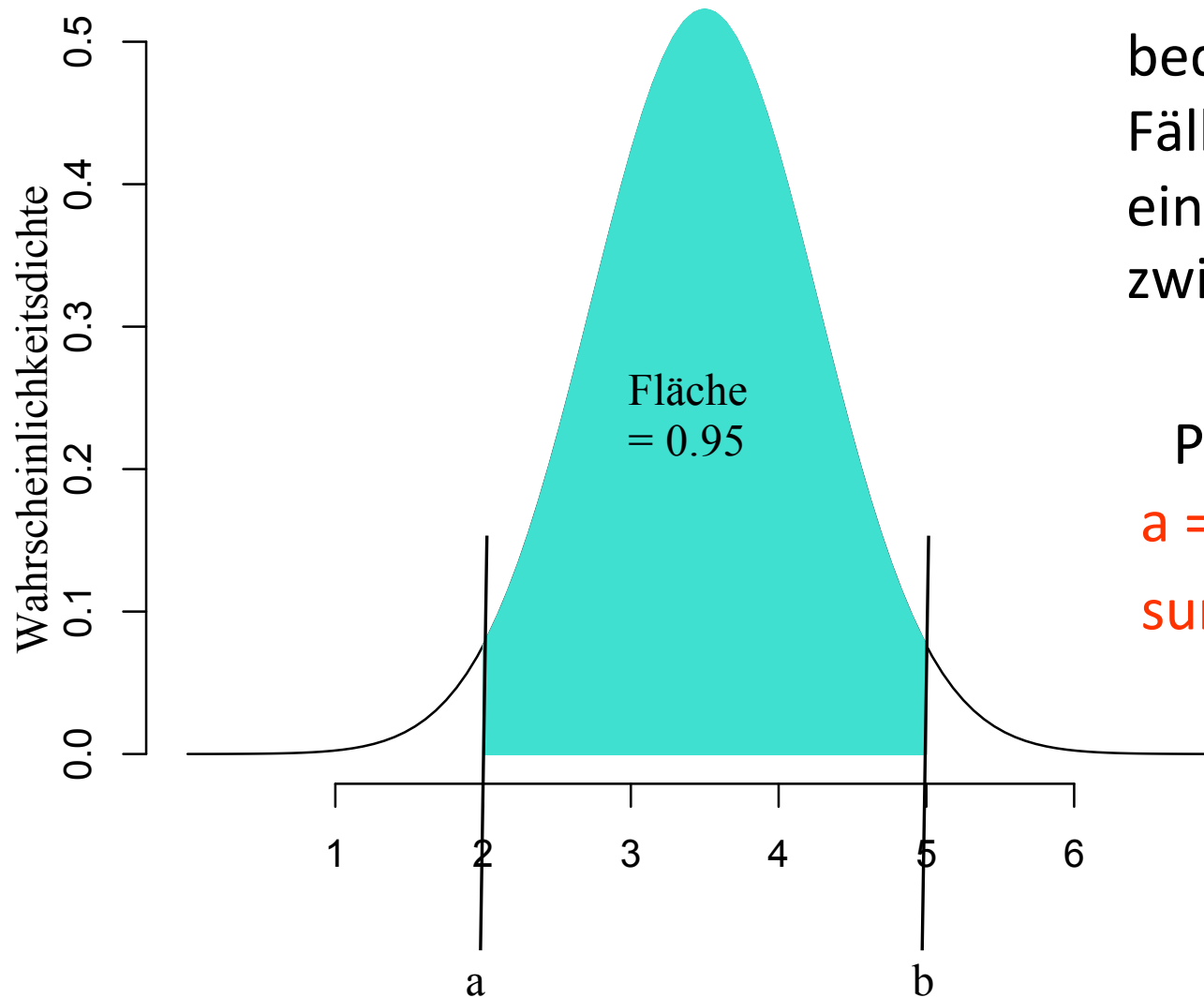
Die Verteilung der Mittelwerte.

Bedeutung: ich werde nicht jedes Mal einen Mittelwert $m = 3.5$ bekommen, sondern davon abweichende

Mittelwerte. Der SE ist eine numerische Verschlüsselung dieser Abweichung.

$$SE = \text{sigma()}/\text{sqrt}(5)$$

95% Konfidenzintervall



bedeutet: in 95/100
Fälle erwarte ich
einen Mittelwert
zwischen 2 und 5.

Probieren

```
a = proben(1, 6, 5, 100)  
sum(a < 2 | a > 5)
```

```
a = qnorm(0.025, mu, SE)
```

2.003053

```
b = qnorm(0.975, mu, SE)
```

4.996947

Berechnungen wenn μ , σ unbekannt ist

1. μ , SE werden eingeschätzt
2. Verwendung der t-Verteilung statt der Normalverteilung

μ, σ ist unbekannt

Lenneberg behauptet, dass wir im Durchschnitt mit einer Geschwindigkeit von 6 Silben pro Sekunde sprechen. Hier sind 12 Werte (Silben/Sekunde) von einem Sprecher.

werte

[1] 6 5 6 9 6 5 6 8 5 6 10 9

Frage: sind diese Werte konsistent mit Lennebergs Hypothese?

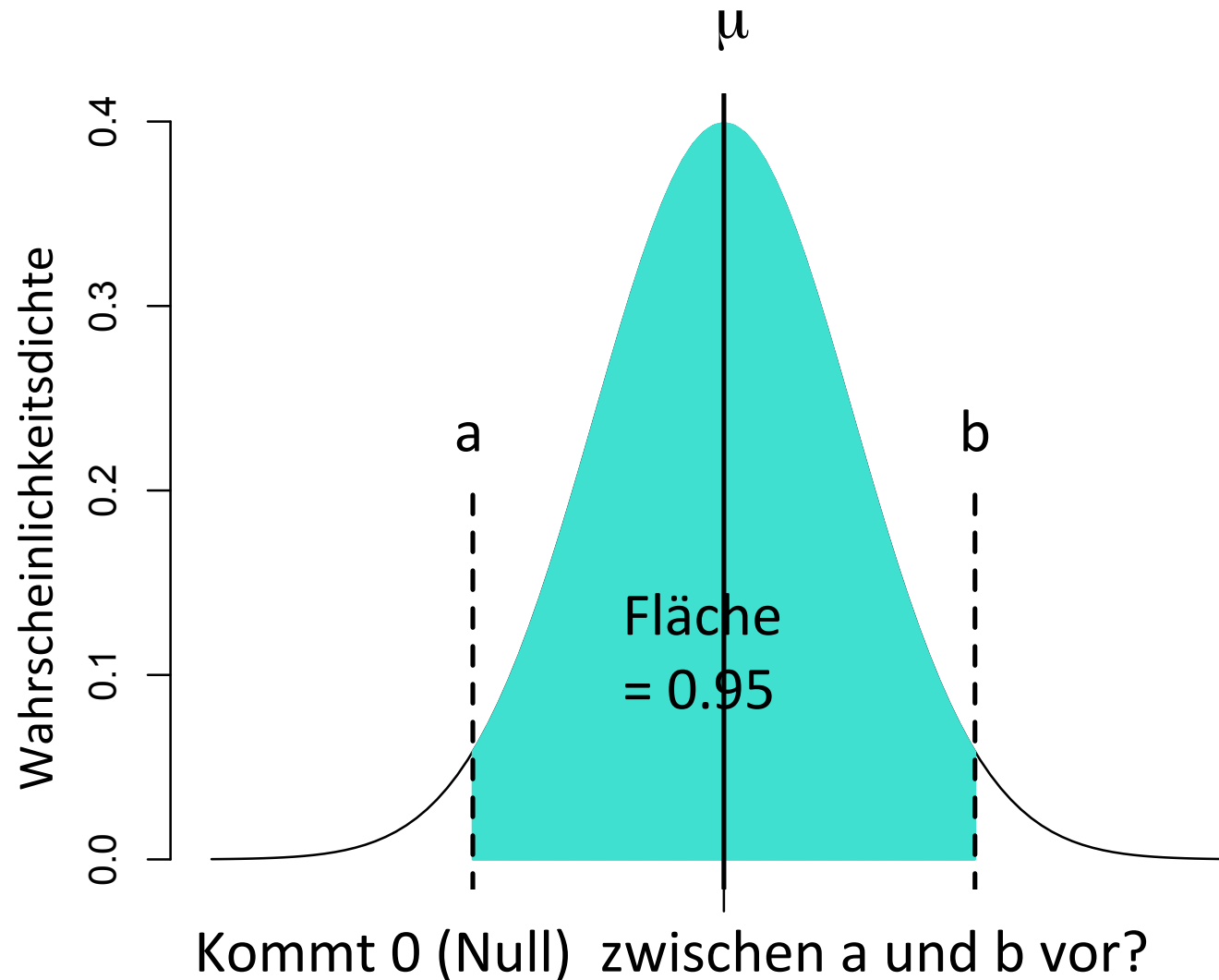
Vorgang: was ist die Wahrscheinlichkeit, dass der Unterschied zwischen dem Stichprobenmittelwert und 6 von 0 (Null) abweicht?

Das Verfahren: a one-sampled t-test

Das Verfahren im t-test

μ = der Unterschied zwischen den Mittelwerten

a, b: Grenzen des 95% Konfidenzintervalls



1. Einschätzung von μ , SE

μ : 6 von jedem Wert abziehen, und den Mittelwert der Unterschiede berechnen

`werte = werte - 6`

`mu = mean(werte)`

Die beste Einschätzung von SE ist die Standardabweichung der Unterschiede (der Werte minus 6 in diesem Fall), s dividiert durch Wurzel n (Anzahl der Stichproben):

$$\hat{SE} = \frac{s}{\sqrt{n}}$$

In R:

`n = length(werte)`

`SE = sd(werte)/sqrt(n)`

2. die t-Verteilung

Wenn SE **eingeschätzt** werden muss, dann wird das Konfidenzintervall nicht mit der Normal- sondern der **t-Verteilung** mit einer gewissen Anzahl von **Freiheitsgraden** berechnet.

Die t-Verteilung ist der Normalverteilung recht ähnlich, aber die 'Glocke' und daher das Konfidenzintervall sind etwas breiter (dies berücksichtigt, die zusätzliche Unsicherheit wegen der Einschätzung von SE).

Bei diesem one-sample t-test ist die Anzahl der Freiheitsgrade, df (degrees of freedom), von der **Anzahl der Werte in der Stichprobe** abhängig: **$df = n - 1$**

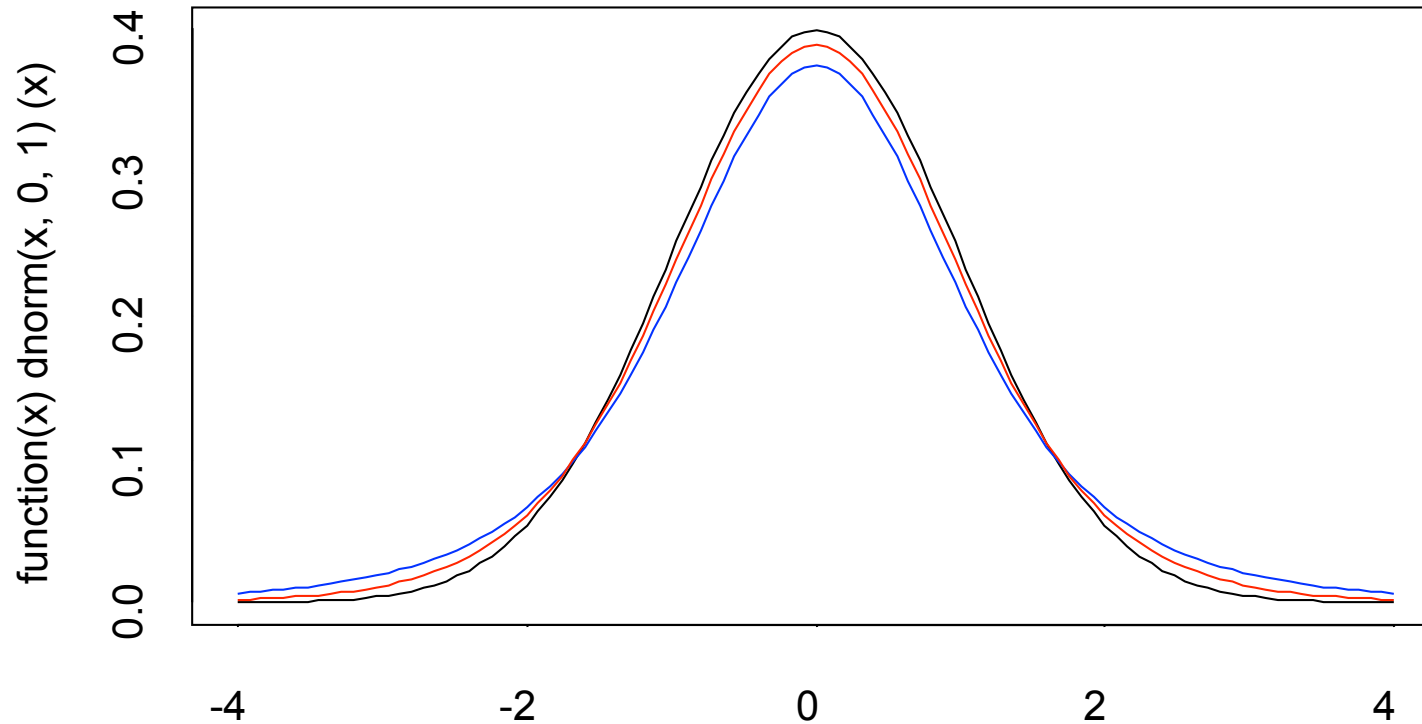
Je höher df, umso sicherer können wir sein, dass $\hat{SE} = SE$ und umso mehr nähert sich die t-Verteilung der Normalverteilung

Normalverteilung, $\mu = 0$, $SE = 1$.

`curve(dnorm(x, 0, 1), -4, 4)`

t-Verteilung, $\mu = 0$, $SE = 1$, $df = 3$

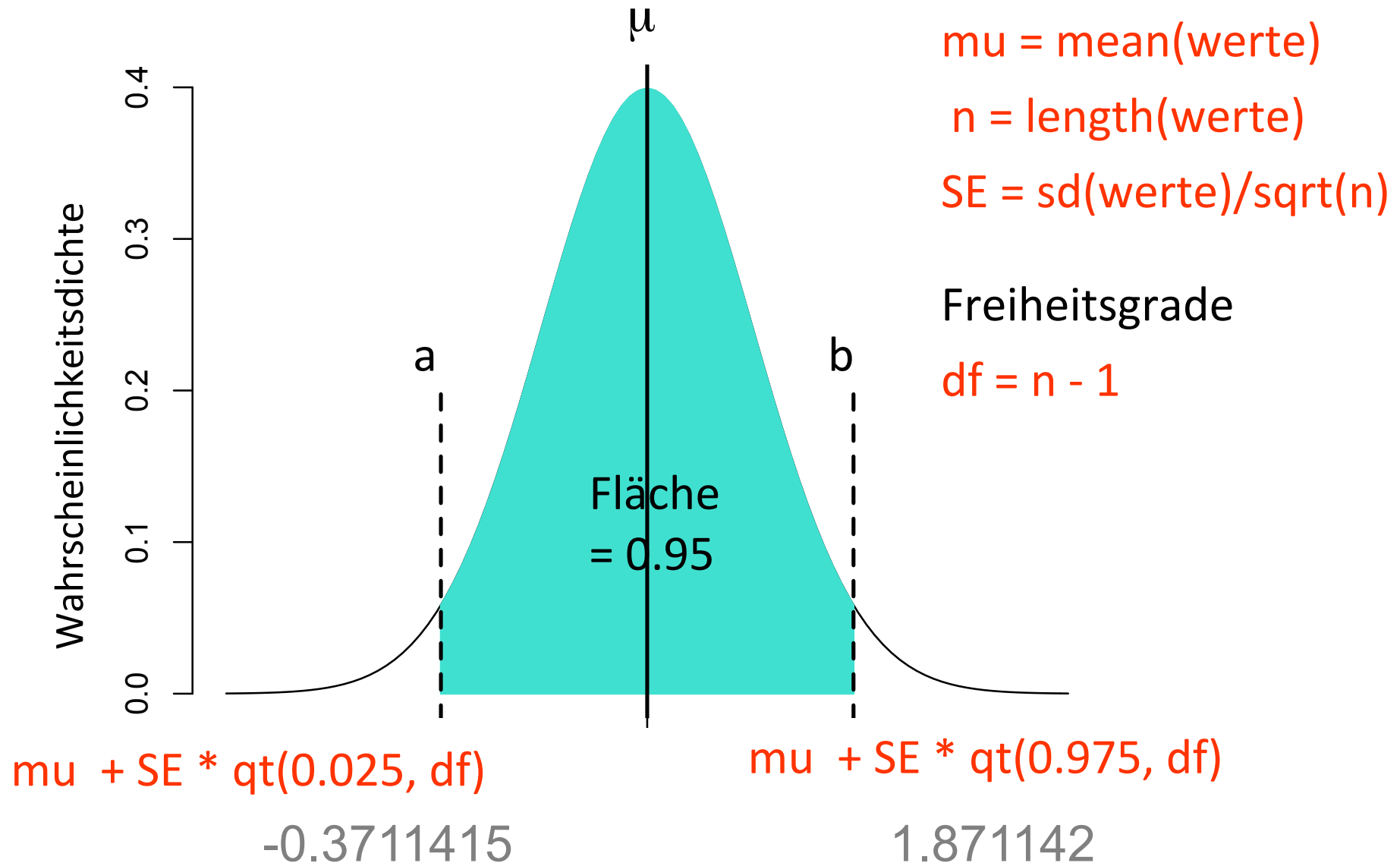
`curve(dt(x, 3), add=T, col="blue")`



`curve(dt(x, 10), add=T, col="red")`

Fällt 0 außerhalb des 95% Konfidenzintervalls von μ ?

= kommt 0 zwischen a und b vor?



Auf der Basis dieser Stichprobe liegt μ (der Unterschied zwischen den Mittelwerten) zwischen **-0.3711415** und **1.871142** mit einer Wahrscheinlichkeit von 95%.

Frage: sind diese Werte konsistent mit Lennebergs Hypothese?

Ja.

Ein einseitiger t-Test in der Phonetik

wird eingesetzt, wenn der Mittelwert aus Unterschieden **pro Versuchsperson** berechnet wird (auch ein gepaarter t-test).

12 Sprecher produzierten /i/ in einer betonten und unbetonten Silbe. Hat die Betonung (=Faktor) einen Einfluss auf F2 (= die abhängige Variable)?

```
F2 = read.table(file.path(pfadu, "bet.txt"))
```

Was ist die Wahrscheinlichkeit, dass der Unterschied zwischen den Mittelwerten (für betont und unbetont) 0 sein könnte (= kommt 0 innerhalb des 95% Konfidenzintervalls vor)?

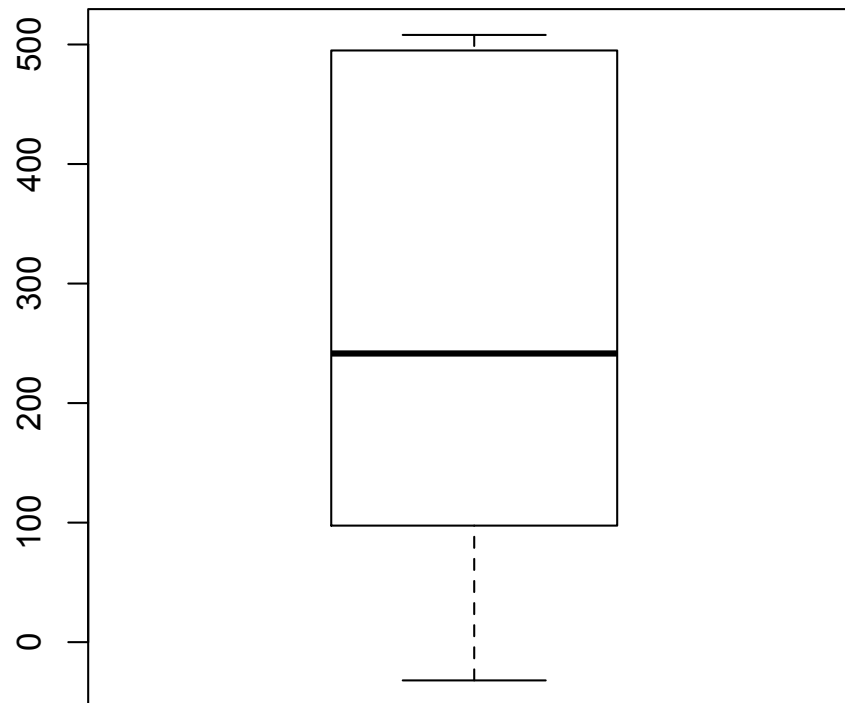
μ , SE der Mittelwert-*Unterschiede* einschätzen.

Die Unterschiede (betont-unbetont) pro Sprecher

```
F2unt = F2$betont - F2$unbetont
```

Zuerst eine Abbildung

`boxplot(F2unt)`



Wir werden die Wahrscheinlichkeit prüfen, dass der Mittelwert dieser Verteilung von 0 abweicht.

Berechnungen

μ

$\mu = \text{mean}(F2\text{unt})$

SE

$n = \text{length}(F2\text{unt})$

$SE = \text{sd}(F2\text{unt})/\text{sqrt}(n)$

df

$df = n - 1$

Konfidenzintervall

$\text{unten} = \mu + SE * qt(0.025, df)$

134.0163

$\text{oben} = \mu + SE * qt(0.725, df)$

309.3802

Kommt 0 innerhalb des Konfidenzintervalls vor?

unten = $\mu + SE * qt(0.025, df)$

134.0163

oben = $\mu + SE * qt(0.975, df)$

407.9837

Auf der Basis dieser Stichprobe liegt μ (der Unterschied zwischen den Mittelwerten) zwischen 134.0163 und 407.9837 mit einer Wahrscheinlichkeit von 95%.

0 kommt innerhalb dieses Konfidenzintervalls nicht vor, daher:

Betonung beeinflusst F2 ($p < 0.05$)

(= die Wahrscheinlichkeit, dass Betonung keinen Einfluss auf F2 hat, liegt unter 0.05).

`t.test(F2unt)`

Die Wahrscheinlichkeit, dass der Unterschied zwischen den Mittelwerten 0 sein könnte.

Die t-Statistik oder critical-ratio: μ/SE

= wieviele Standard-Errors μ und 0 voneinander entfernt sind

Freiheitsgrade

```
data:  F2unt  
t = 4.3543, df = 11, p-value = 0.001147  
alternative hypothesis: true mean is not equal  
to 0  
95 percent confidence interval:  
134.0163 407.9837 ← Konfidenzintervall  
sample estimates:  
mean of x  
271 ←  $\mu$  (der Mittelwert der Unterschiede)
```

Betonung hat einen signifikanten Einfluss auf F2 oder F2 wird signifikant von der Betonung beeinflusst ($t[11] = 4.4, p < 0.01$)

In der Phonetik wird ein solcher one-sample t-test eingesetzt, wenn gepaarte Werte **für die selbe Versuchsperson** vorliegen – wie im vorigen Fall: es gab einen betonten F2-Wert und einen unbetonten F2-Wert **pro Sprecher**, und der Unterschied wurde **pro Sprecher berechnet** (und dann der Mittelwert dieser Unterschiede).

Der two-sample t-test wird dagegen in der Phonetik dann eingesetzt, wenn sich **die Versuchspersonen unterscheiden**: z.B. wir wollen F2 in Männern und Frauen vergleichen; die Grundfrequenz von deutschen vs. französischen Sprechern usw.

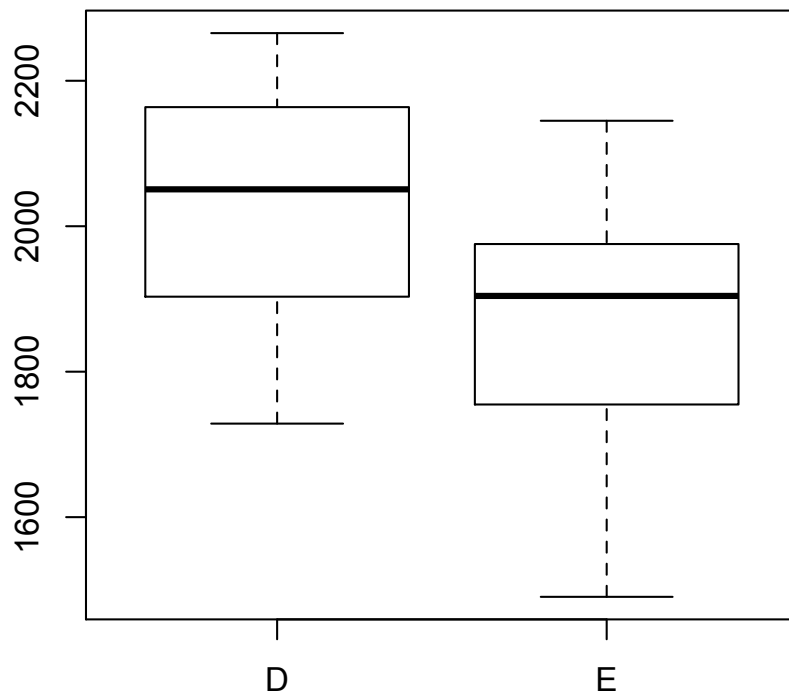
Die Frage ist genau dieselbe, aber diesmal für 2 Gruppen: was ist die Wahrscheinlichkeit, dass der Mittelwert-Unterschied (zwischen den Gruppen) 0 (Null) sein könnte?

Unterscheiden sich Deutsche und Engländer in F2 von /e/?

```
e.df = read.table(file.path(pfadu, "e.txt"))
```

```
head(e.df)
```

```
boxplot(F2 ~ Sprache, data = e.df)
```



= Was ist die Wahrscheinlichkeit, dass der Unterschied zwischen den Mittelwerten der Gruppen von 0 (Null abweicht)?

```
t.test(F2 ~ Sprache, data = e.df)
```

```
t.test(F2 ~ Sprache, data = e.df)
```

```
data:  F2 by Sprache  
t = 2.2613, df = 21.101, p-value =  
0.03443  
alternative hypothesis: true  
difference in means is not equal to 0  
95 percent confidence interval:  
 13.46719 320.73097  
sample estimates:  
mean in group D mean in group E  
    2031.672      1864.573
```

Unterschiede zwischen den Mittelwerten: 167.099

Der Unterschied zwischen den Mittelwerten liegt zwischen 13.46719 und 320.73097 mit einer Wahrscheinlichkeit von 95%.

Die Wahrscheinlichkeit, dass die Mittelwert-Unterschiede 0 (Null) sein könnte = 0.03443

```
data:  F2 by Sprache
t = 2.2613, df = 21.101, p-value =
0.03443
alternative hypothesis: true
difference in means is not equal to 0
95 percent confidence interval:
 13.46719 320.73097
sample estimates:
mean in group D mean in group E
      2031.672      1864.573
```

Die Sprache hat einen signifikanten Einfluss auf F2 ($t[21.1] = 2.3, p < 0.05$)

oder

F2 wurde signifikant von der Sprache beeinflusst ($t[21.1] = 2.3, p < 0.05$)

Die Reaktionszeiten wurde in 15 Versuchspersonen gemessen, um Wörter zu identifizieren, wenn sie akzentuiert oder unakzentuiert waren. Hat Akzentuierung einen Einfluss auf die Reaktionszeit?

RT-akzentuiert, Sprecher 1-15

rtaz = c(56, 49, 50, 39, 49, 60, 51, 39, 67, 49, 60, 46, 55, 54, 52)

RT-unakzentuiert, in denselben Sprechern 1-15

rtun = c(95, 94, 121, 48, 135, 87, 94, 135, 98, 125, 92, 115, 80, 98, 108)

Boxplot

Test

Schlussfolgerung

Prüfen Sie für den Data-Frame owl, ob die Sprache (Lang) einen Einfluss auf die Reaktionszeit hatte (rt), um dieses Wort zu identifizieren.

```
owl = read.table(file.path(pfadu, "owl.txt"))
```

Boxplot

Test

Schlussfolgerung