

SDsanitycheck

2025-06-15

This file contains analyses (not on the entire data set) that would confirm the phenomenon of stress deafness (by comparing voicing to stress). There is not really enough data to run these.

```
table = read.delim ("data/SDsanitycheck.txt", stringsAsFactors=TRUE)
options(width = 200)
head(table)
```

##	participant	training	LexTale	trialIndex	item	correct	Xtype	contrast	pair	response	accuracy
## 1	p929263	stress	71.25	1	i44	A	irrelevant	voicing	faces	B	0
## 2	p929263	stress	71.25	2	i5	A	extreme	stress	extract	A	1
## 3	p929263	stress	71.25	3	i48	A	irrelevant	voicing	dose	B	0
## 4	p929263	stress	71.25	4	i2	B	extreme	stress	refund	B	1
## 5	p929263	stress	71.25	5	i46	A	irrelevant	voicing	faces	A	1
## 6	p929263	stress	71.25	6	i47	B	irrelevant	voicing	dose	A	0

```
tail(table)
```

##	participant	training	LexTale	trialIndex	item	correct	Xtype	contrast	pair	response	accuracy
## 363	p162417	stress	76.25	41	i33	A	irrelevant	voicing	device	B	
## 364	p162417	stress	76.25	42	i24	A	ambiguous	stress	conduct	B	
## 365	p162417	stress	76.25	43	i42	B	irrelevant	voicing	device	A	
## 366	p162417	stress	76.25	44	i22	B	ambiguous	stress	conduct	B	
## 367	p162417	stress	76.25	45	i20	A	ambiguous	stress	extract	B	
## 368	p162417	stress	76.25	48	i8	A	extreme	stress	extract	A	

set contrast for contrast type Voicing is expected to be easiest, thus coded positively:

```
levels (table$contrast)
```

```
## [1] "stress" "voicing"
```

```
contrast <- cbind (c(-0.5, 0.5))
colnames (contrast) <- c("-S+V")
contrasts (table$contrast) <- contrast
contrasts (table$contrast)
```

```
##      -S+V
## stress -0.5
## voicing 0.5
```

Is accuracy in general worse on stress contrast than on voicing?

```
library(lme4)
```

```
## Loading required package: Matrix
```

```
contrastModel = glmer(formula = accuracy ~ contrast + (contrast | participant) + (1 | item), data = table)
```

```
## boundary (singular) fit: see help('isSingular')
```

```
summary(contrastModel)
```

```
## Generalized linear mixed model fit by maximum likelihood (Laplace Approximation) ['glmerMod']
## Family: binomial ( logit )
## Formula: accuracy ~ contrast + (contrast | participant) + (1 | item)
## Data: table
##
##      AIC      BIC   logLik deviance df.resid
##    471.2    494.6   -229.6    459.2     362
##
## Scaled residuals:
##      Min       1Q   Median       3Q      Max
## -2.7477 -1.1191  0.5747  0.7046  1.0919
##
## Random effects:
## Groups      Name                Variance Std.Dev. Corr
## item        (Intercept)         0.2243   0.4736
## participant (Intercept)         0.1852   0.4304
##              contrast-S+V       0.3165   0.5626   1.00
## Number of obs: 368, groups:  item, 40; participant, 10
##
## Fixed effects:
##              Estimate Std. Error z value Pr(>|z|)
## (Intercept)    0.8175     0.2028   4.031 5.55e-05 ***
## contrast-S+V   0.5499     0.3443   1.597    0.11
## ---
## Signif. codes:  0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1
##
## Correlation of Fixed Effects:
##              (Intr)
## contrst-S+V  0.552
## optimizer (Nelder_Mead) convergence code: 0 (OK)
## boundary (singular) fit: see help('isSingular')
```

adding a contrast for training

This is post-hoc: we've seen those trained in the stress-ignoring lesson more so have the expected voicing over stress advantage, thus "nonStress" is coded positively and "stress" training is coded negatively.

```
levels (table$training)
```

```
## [1] "nonStress" "stress"
```

```
traincontrast <- cbind (c(0.5, -0.5))
colnames (traincontrast) <- c("-stress+nonStress")
contrasts (table$training) <- traincontrast
contrasts (table$training)
```

```
##              -stress+nonStress
## nonStress              0.5
## stress                 -0.5
```

```
library(lme4)
```

```
contrastModel = glmer(formula = accuracy ~ contrast * training + (contrast | participant) + (training |
```

```
## Warning in checkConv(attr(opt, "derivs"), opt$par, ctrl = control$checkConv, : Model failed to converge
```

```
summary(contrastModel)
```

```
## Generalized linear mixed model fit by maximum likelihood (Laplace Approximation) ['glmerMod']
## Family: binomial ( logit )
## Formula: accuracy ~ contrast * training + (contrast | participant) + (training | item)
## Data: table
##
##      AIC      BIC   logLik deviance df.resid
##    467.0    506.1   -223.5    447.0     358
##
## Scaled residuals:
##      Min       1Q   Median       3Q      Max
## -2.9442 -1.0857  0.4654  0.7564  1.1063
##
## Random effects:
##   Groups                Name                Variance Std.Dev. Corr
##   item                (Intercept)            0.3345241 0.57838
##                training-stress+nonStress 0.3945165 0.62811  1.00
##   participant (Intercept)            0.0461587 0.21485
##                contrast-S+V            0.0004719 0.02172 -1.00
## Number of obs: 368, groups:  item, 40; participant, 10
##
## Fixed effects:
##                                Estimate Std. Error z value Pr(>|z|)
## (Intercept)                   0.8847     0.1826   4.845 1.27e-06 ***
## contrast-S+V                   0.6332     0.3257   1.944  0.05184 .
## training-stress+nonStress      0.9075     0.3321   2.733  0.00628 **
## contrast-S+V:training-stress+nonStress 1.5233     0.5710   2.668  0.00764 **
## ---
## Signif. codes:  0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1
##
## Correlation of Fixed Effects:
##              (Intr) cn-S+V trn-+S
## contrst-S+V 0.310
## trnng-str+S 0.402  0.237
## cnt-S+V:-+S 0.242  0.430  0.343
## optimizer (Nelder_Mead) convergence code: 0 (OK)
## Model failed to converge with max|grad| = 0.00271705 (tol = 0.002, component 1)
```

This converges; interaction p-value isn't valid when removing slopes for training/contrast so won't simplify