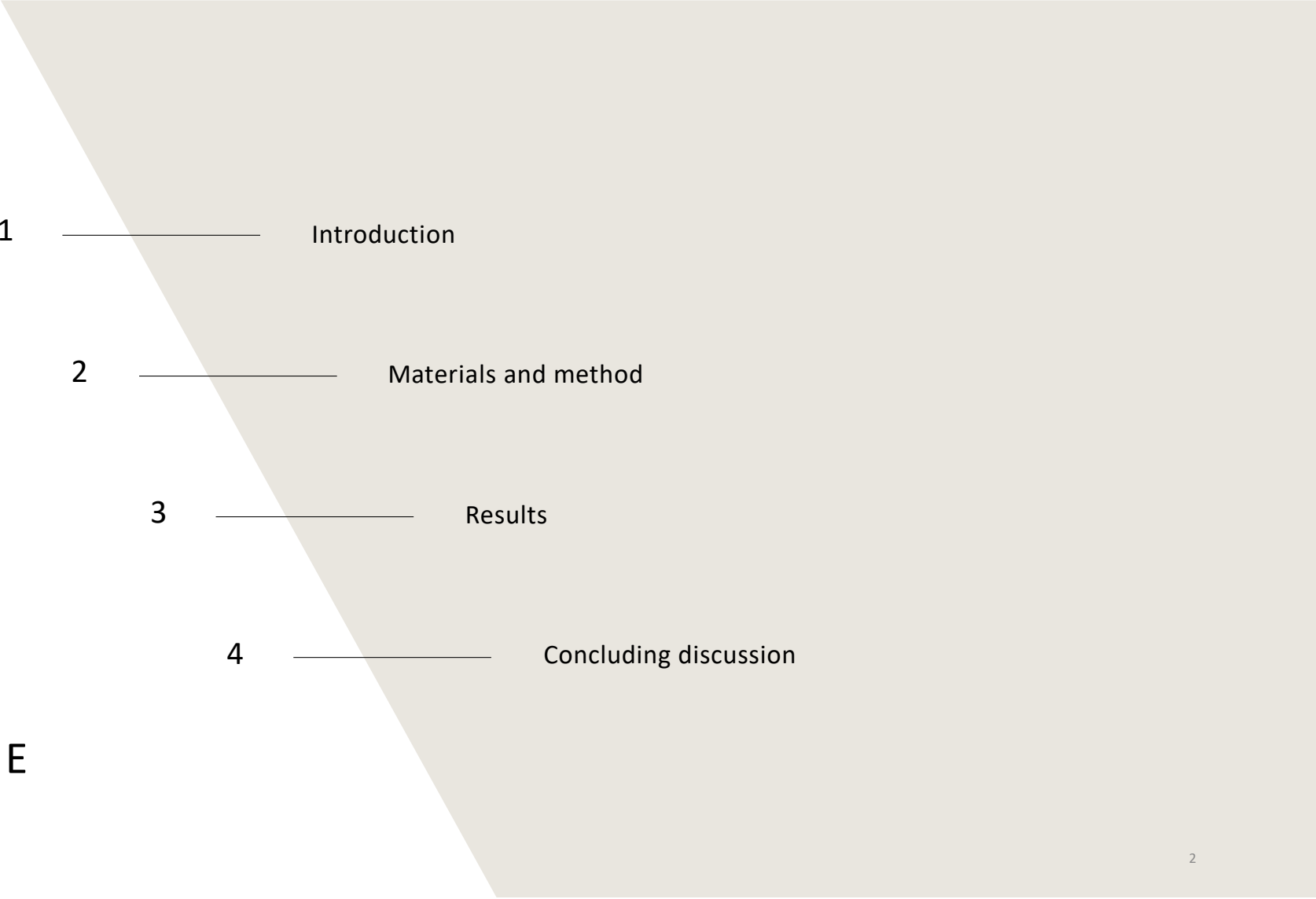
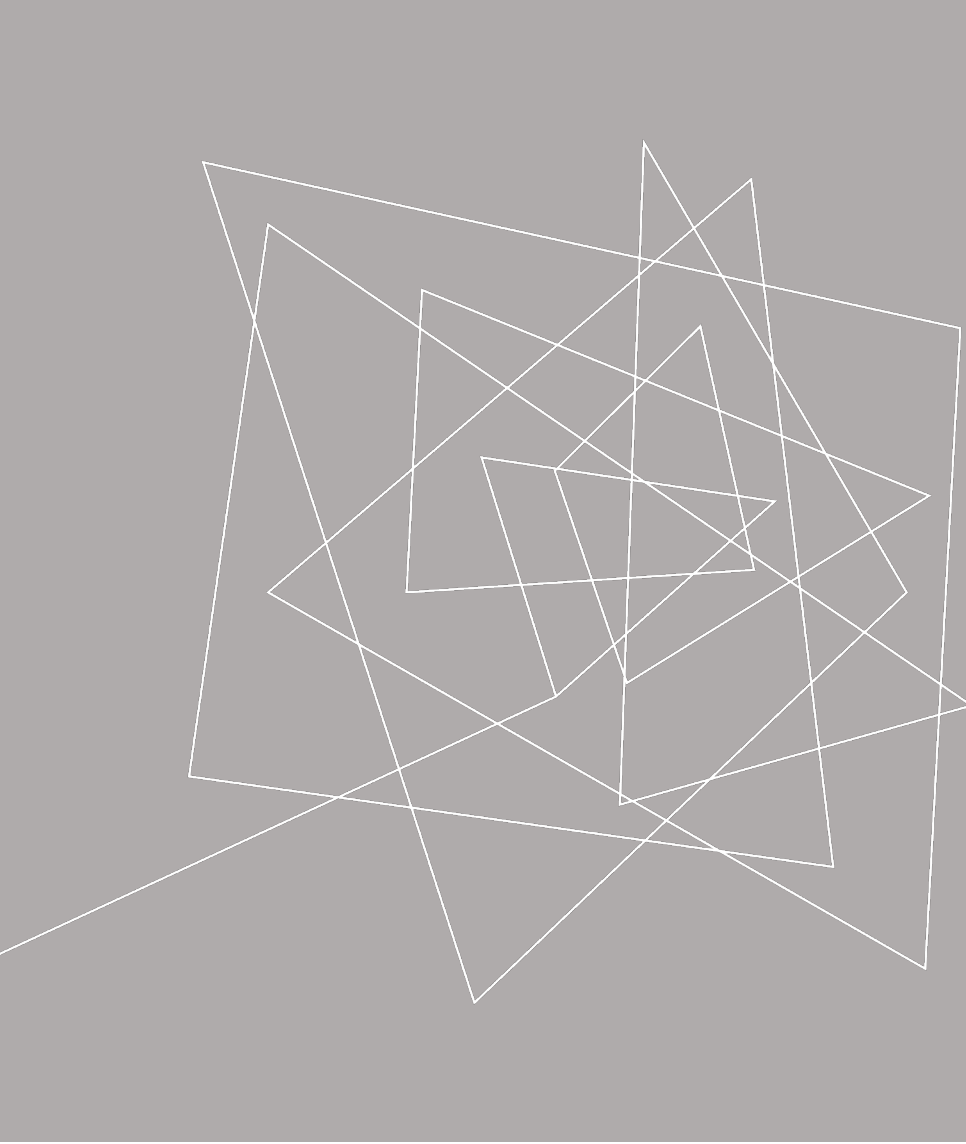


THE ROLE OF PROSODIC ACCENTUATION ON VOWEL PRODUCTION IN ENGLISH

Gil VERBEKE (UGent) & Wout VAN PRAET (UC Louvain)

- 
- 1 — Introduction
 - 2 — Materials and method
 - 3 — Results
 - 4 — Concluding discussion

OUTLINE



INTRODUCTION

INTRODUCTION

In previous research:

suprasegmental aspects of speech production affect the production of individual speech sounds (Fougeron & Keating 1997; Byrd 1996, 2000; Breen *et al.* 2010; Byrd & Choi 2010; Krivokapić 2014), e.g.:

- syllable structure (e.g. Browman & Goldstein 1988, 2000; Gafos & Goldstein 2012; Gibson *et al.* 2019; Hermes *et al.* 2012),
 - e.g. less overlap between consonants in cluster in onset than in coda
- prosodic structure:
 - proximity to strong prosodic boundary
 - accentuation

INTRODUCTION

Prosodic strengthening of vowel:

- proximity to strong prosodic boundary (e.g. IU boundary), e.g.
 - domain-final lengthening (e.g. Edwards *et al.* 1991; Wightman *et al.* 1992)
 - domain-final vowels produced with greater articulatory magnitude → greater V-to-C displacement (e.g. Fougeron & Keating 1997; Cho 2002)
 - domain-initial strengthening, e.g. consonants in IU-initial position produced with more constriction than in middle of IU (e.g. Fougeron & Keating 1997; Cho & Keating 2009), greater intensity (Cho & Keating 2009)

INTRODUCTION

Prosodic strengthening of vowel:

- accentuation

INTRODUCTION

Two models of prosodic transcription

- ToBI (e.g. Pierrehumbert & Hirschberg 1990; Silverman *et al.* 1992)
- nuclear tone approach (e.g. Halliday 1967; Tench 1996)



MATERIALS AND METHOD

MATERIALS AND METHOD

London Lund Corpus of Spoken English:

- spontaneous conversations
- mostly educated Londoners (unaware of being recorded)
- audiomaterial and transcriptions (in nuclear tone approach)

MATERIALS AND METHOD

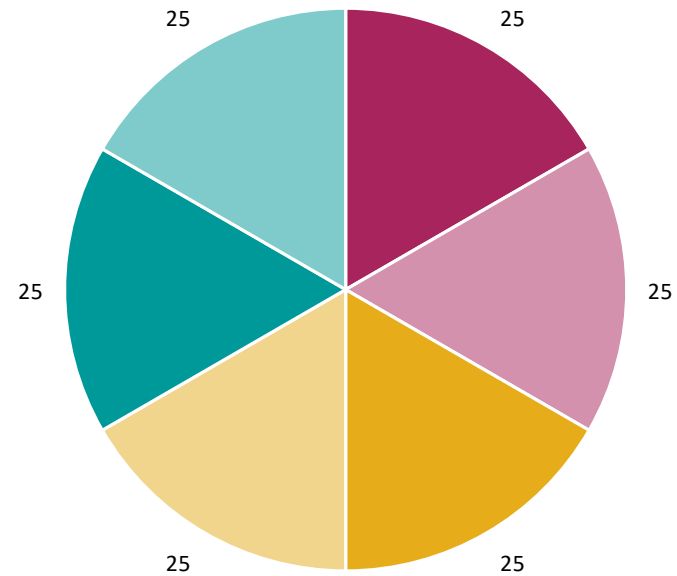
Sample of 150 vowels:

- checked monophthongs
- closed syllable with fortis coda (i.e. fortis fricative, plosive or affricate)
- monosyllabic words ($\leftrightarrow$ polysyllabic shortening)
- 29 speakers (mean utterance per speaker = 5.172414):
 - 13 female: mean utterance per speaker = 5
 - 16 male: mean utterance per speaker = 5.3125

MATERIALS AND METHOD

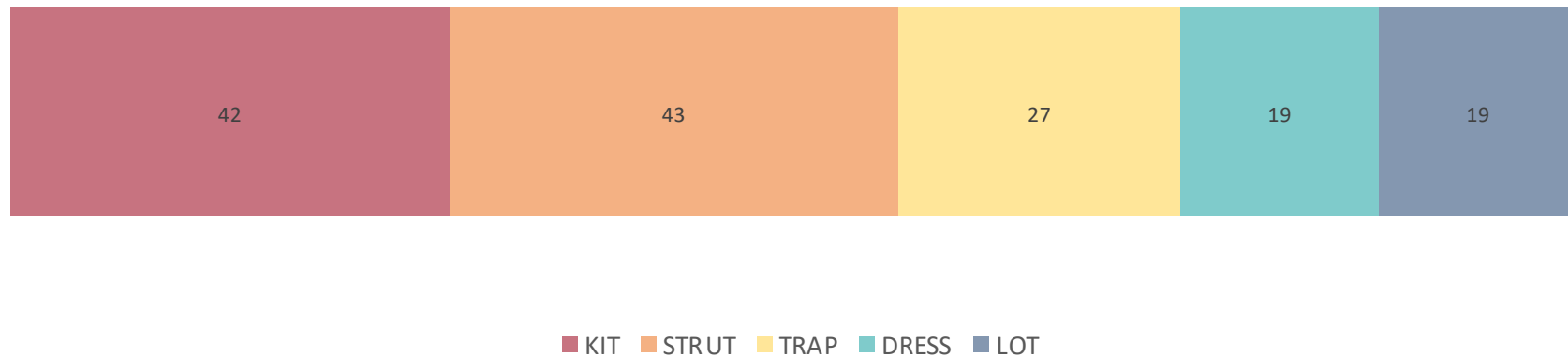
Sample of 150 vowels:

- equal distribution of **primary** accentuation (red), **secondary** accentuation (yellow) and **non-accentuation** (green)
- equal distribution (per accentuation status) of syllables preceding IU boundaries (light) and non-IU boundaries (dark)

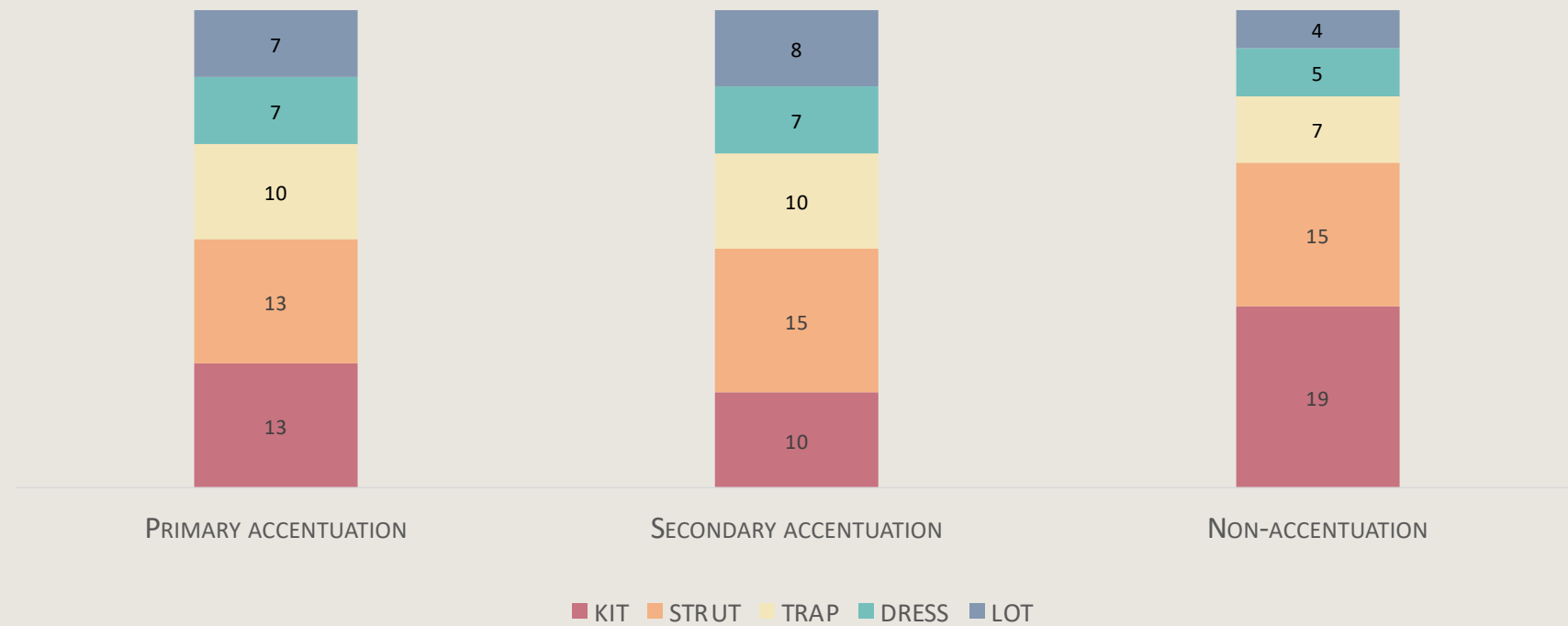


MATERIALS AND METHOD

Sample of 150 vowels: distribution of vowel types



VOWEL TYPES PER ACCENTUATION STATUS



$\chi^2(8) = 5.6422, p = 0.69$

VOWEL TYPES PER BOUNDARY STRENGTH



Non-IU BOUNDARY

IU BOUNDARY

KIT STRUT TRAP DRESS LOT

$\chi^2(4) = 6.4852, p = 0.17$

MATERIALS AND METHOD

Multivariate regression analysis in R

dependent variable

vowel length (measured instrumentally in Praat)

independent variables *linguistic*

proximity to IU boundary (in n° of syllables)

boundary strength: IU vs. non-IU boundary

accentuation status: – primary (tonic) accent
– secondary (non-tonic) accent
– unaccented

vowel type: KIT, STRUT, TRAP, DRESS, LOT

non-linguistic speaker
sex

MATERIALS AND METHOD

Inter-rater reliability

Pearson's correlation test

correlation coefficient = 0.7699107

$t = 5.9105$

$p\text{-value} < 0.001$

= individual measurements significantly correlated → sufficient inter-rater agreement



RESULTS

RESULTS

Multivariate linear regression

model_1 = lm(Length_vowel ~ Accentuation + Boundary + Proximity + VowelType + Participant + Sex)

	df	Sum Sq	Mean Sq	F Value	Pr (>F)	
Accentuation	2	0.150821	0.075411	47.9901	8.19E-16 ***	
Boundary	1	0.04691	0.04691	29.8526	2.81E-07 ***	
Proximity	1	0.004639	0.004639	2.9524	0.088489 .	→ insignificant
VowelType	4	0.022175	0.005544	3.528	0.009421 **	
Speaker	27	0.105891	0.003922	2.4958	0.000431 ***	
Sex	1	0.000734	0.000734	0.4674	0.495593	→ insignificant
Residuals	113	0.177566	0.001571			

+ high VIF scores for Speaker = suggests collinearity with other factor(s)

RESULTS

Multivariate linear regression

model_2 = lm(Length_vowel ~ Accentuation + Boundary + Proximity + VowelType)

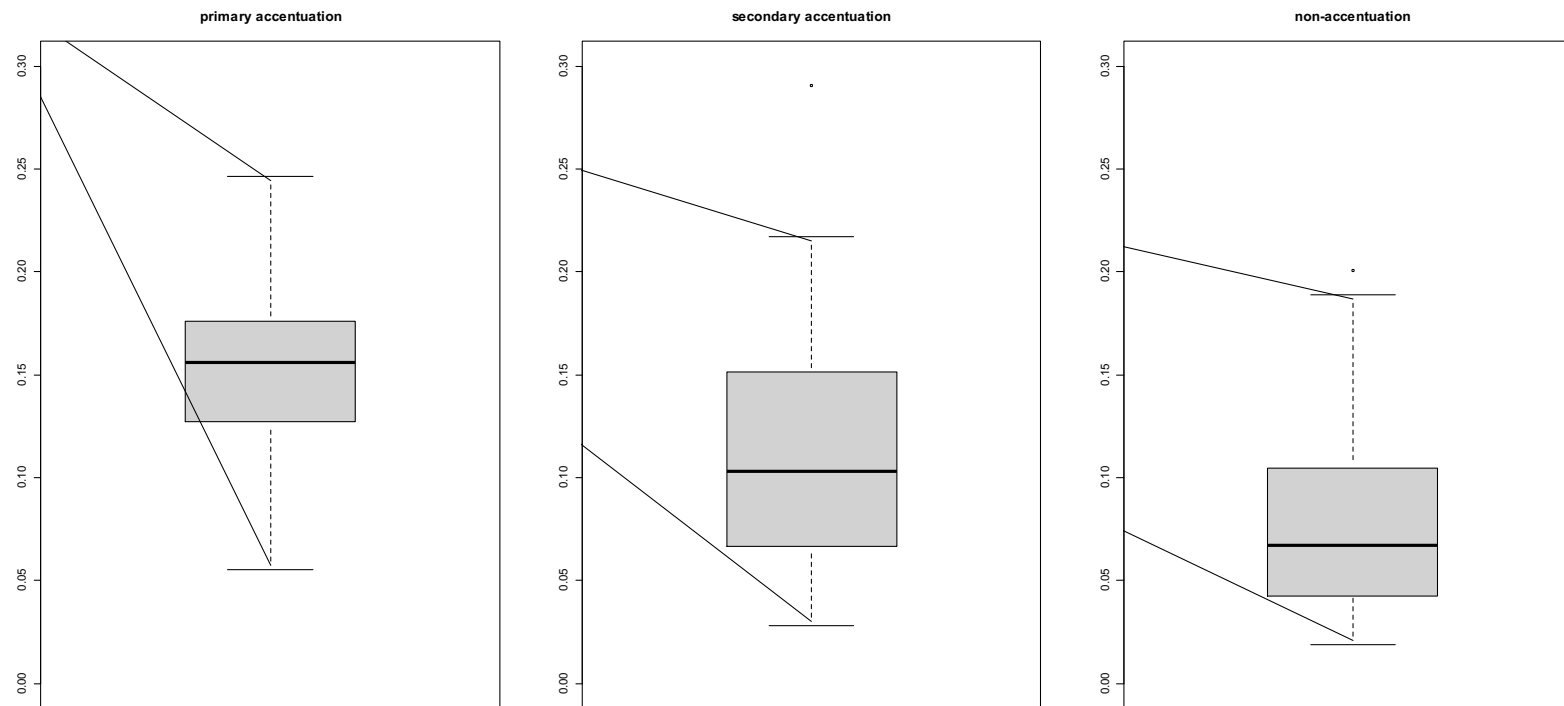
	df	Sum Sq	Mean Sq	F Value	Pr (>F)
Accentuation	2	0.150821	0.075411	37.4146	9.23E-14***
Boundary	1	0.04691	0.04691	23.274	3.60E-06***
Proximity	1	0.004639	0.004639	2.3018	0.13146
VowelType	4	0.022175	0.005544	2.7505	0.03056*
Residuals	141	0.284192	0.002016		

Residual standard error: 0.04489 on 141 degrees of freedom

Multiple R-squared: 0.4414, Adjusted R-squared: 0.4097

F-statistic: 13.93 on 8 and 141 DF, p-value: 8.556e-15

RESULTS: ACCENTUATION STATUS



RESULTS: ACCENTUATION STATUS

Tukey post-hoc test, based on ANOVA for accentuation status

```
Tukey multiple comparisons of means  
95% family-wise confidence level
```

```
Fit: aov(formula = Length_vowel ~ ACC, data = data)
```

```
$ACC
```

		diff	lwr	upr	p adj
SEC-PRIM	-0.04309699	-0.06646315	-0.01973083	0.0000699	
UNACC-PRIM	-0.07750943	-0.10087559	-0.05414328	0.0000000	
UNACC-SEC	-0.03441244	-0.05777860	-0.01104629	0.0018546	

RESULTS: ACCENTUATION STATUS

Moreover, in subsets of syllables before non-IU vs. IU boundary, effect of accentuation status remains:

Before **non-IU** boundary

	df	Sum Sq	Mean Sq	F Value	Pr (>F)
Accentuation	2	0.0558	0.027899	1.12E+01	5.88E-05 ***
Residuals	72	0.1796	0.002494		

Before **IU** boundary

	df	Sum Sq	Mean Sq	F Value	Pr (>F)
Accentuation	2	0.0558	0.027899	1.12E+01	5.88E-05 ***
Residuals	72	0.1796	0.002494		

RESULTS: ACCENTUATION STATUS

But: Tukey post-hoc tests

Before **non-IU** boundary

	diff	lwr	upr	p adj
SEC-PRIM	-0.05105161	-0.07892584	-0.023177392	0.0001149
UNACC-PRIM	-0.08823713	-0.11611135	-0.060362905	0.0000000
UNACC-SEC	-0.03718551	-0.06505974	-0.009311291	0.0058612

Before **IU** boundary

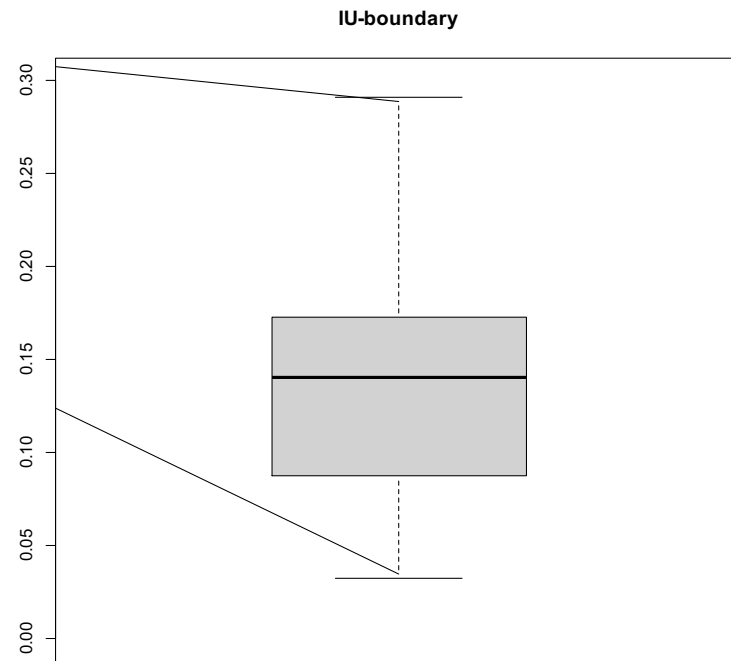
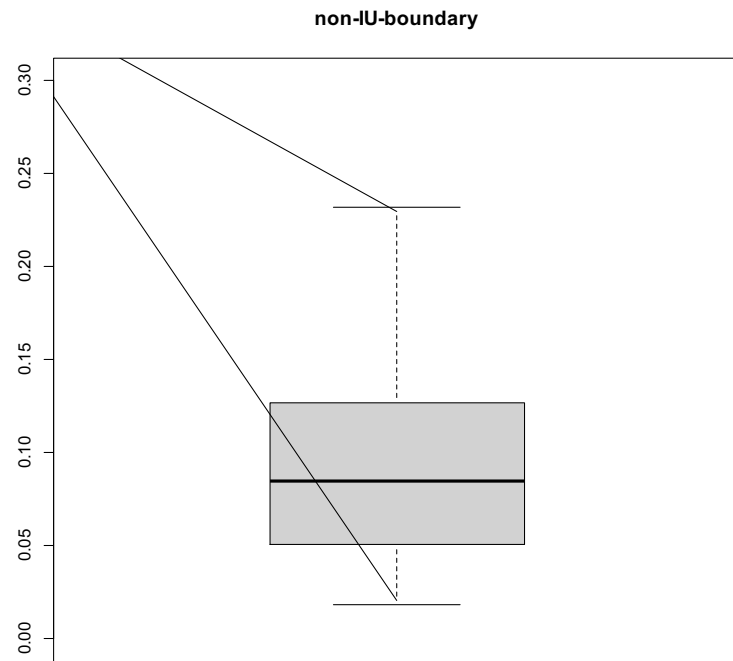
	diff	lwr	upr	p adj
SEC-PRIM	-0.03514237	-0.06894643	-0.001338301	0.0397619
UNACC-PRIM	-0.06678174	-0.10058580	-0.032977672	0.0000324
UNACC-SEC	-0.03163937	-0.06544344	0.002164695	0.0713230

RESULTS: ACCENTUATION STATUS

But: Tukey post-hoc tests

- before non-IU boundaries, significant differences between all accentuation statuses
- before IU boundaries:
 - clear difference between primary accented and unaccented syllables ($p < 0.001$)
 - smaller but significant difference between primary and secondary accented syllables ($p = 0.03976$)
 - no difference between secondary accented and unaccented syllables ($p = 0.07$)

RESULTS: BOUNDARY

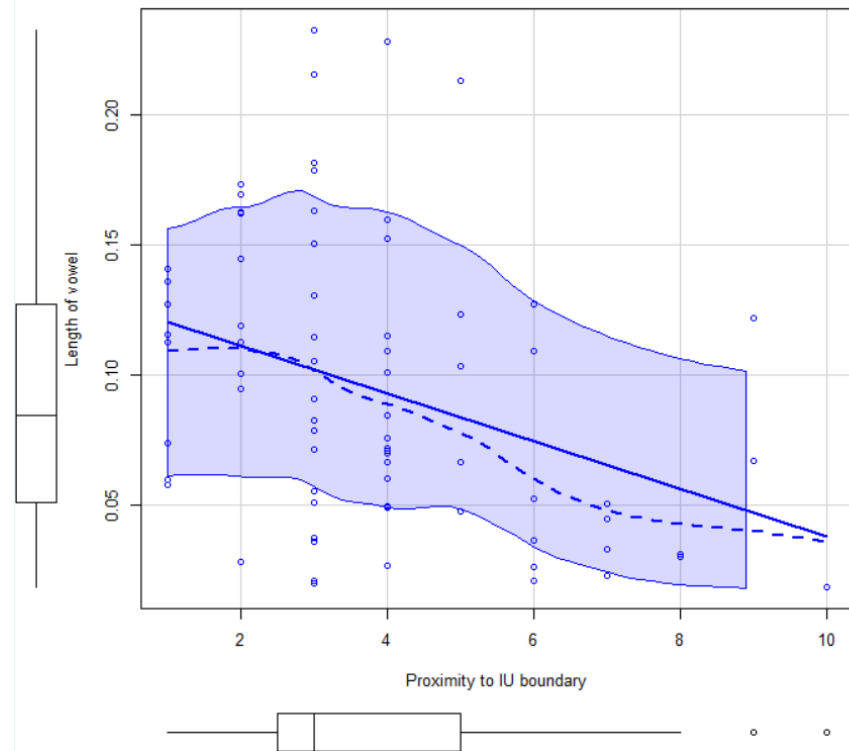


RESULTS: PROXIMITY

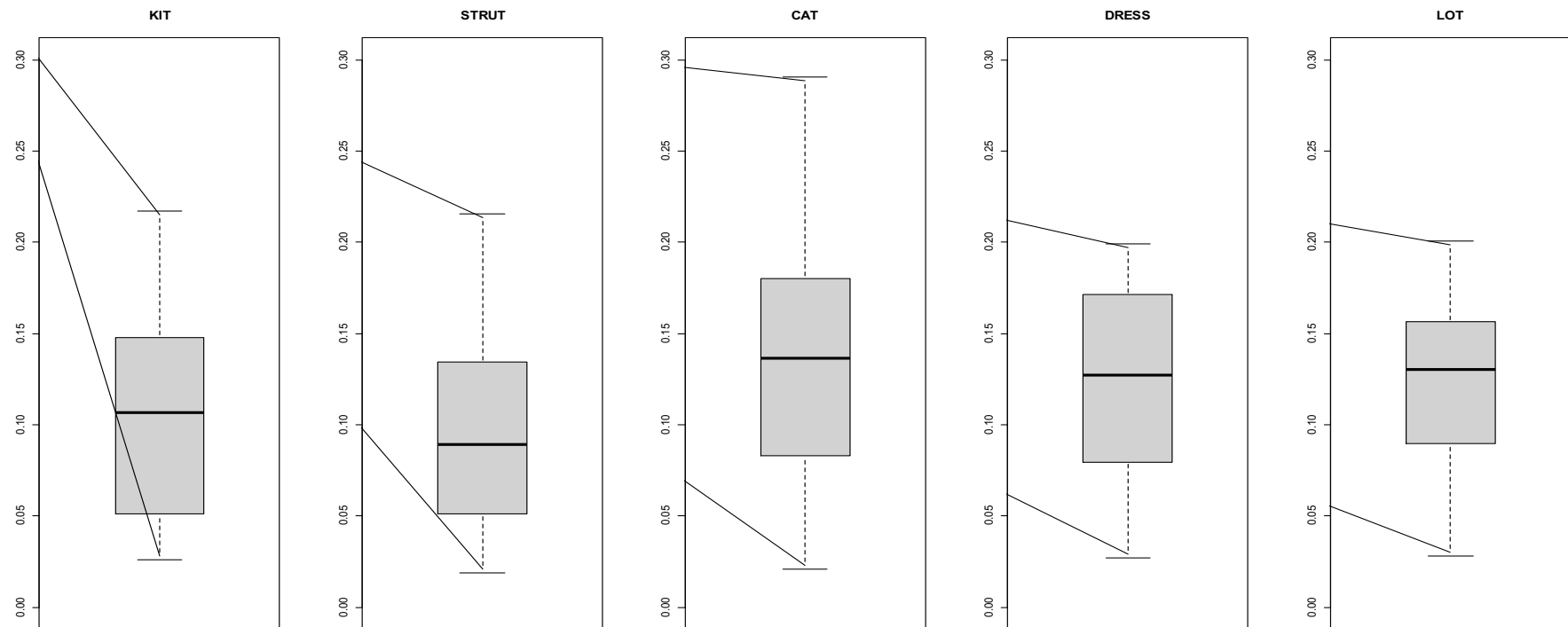
No significant result, but trend for vowel length to decrease further away from (next) IU boundary

But:

- few datapoints for vowels removed further than 4 syllables from IU boundary
- results for syllables at 1-3 syllables from IU boundary relatively similar (cf. dashed line)



RESULTS: VOWEL TYPE



RESULTS: VOWEL TYPE

Tukey post-hoc test, based on ANOVA for vowel type

```
Tukey multiple comparisons of means
 95% family-wise confidence level
```

```
Fit: aov(formula = Length_vowel ~ Vowel, data = data)
```

```
$Vowel
```

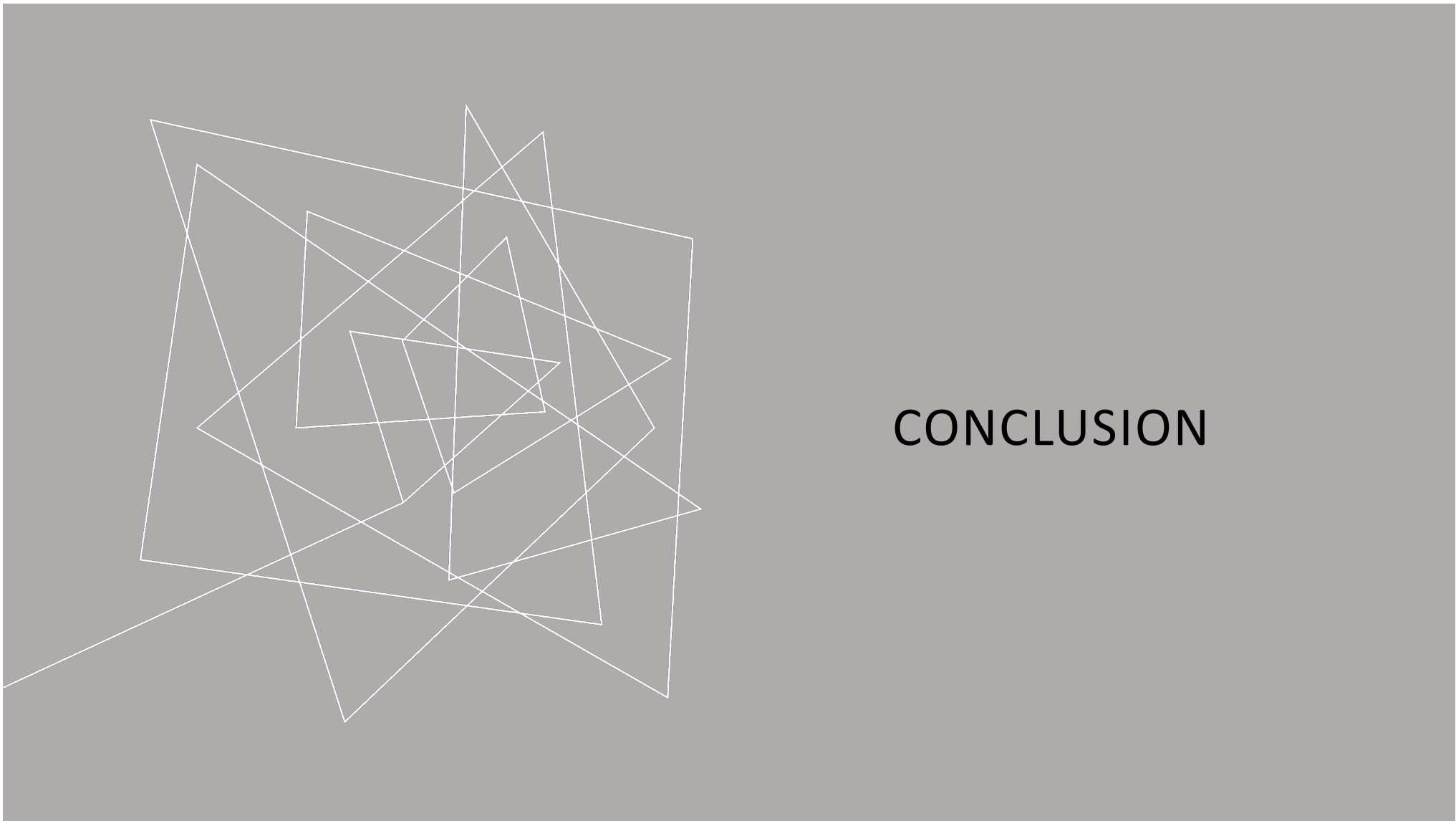
	diff	lwr	upr	p adj
DRESS-CAT	-0.013198248	-0.06043774	0.034041248	0.9383243
KIT-CAT	-0.035870052	-0.07478380	0.003043695	0.0862068
LOT-CAT	-0.014928990	-0.06216849	0.032310507	0.9063673
STRUT-CAT	-0.039811194	-0.07854748	-0.001074912	0.0407365
KIT-DRESS	-0.022671804	-0.06628807	0.020944464	0.6055892
LOT-DRESS	-0.001730741	-0.05291345	0.049451962	0.9999826
STRUT-DRESS	-0.026612946	-0.07007096	0.016845065	0.4423177
LOT-KIT	0.020941062	-0.02267521	0.064557330	0.6753262
STRUT-KIT	-0.003941142	-0.03816552	0.030283240	0.9977590
STRUT-LOT	-0.024882204	-0.06834021	0.018575807	0.5115548

BAD-LAD split: /æ/-lengthening ?

RESULTS: PARTICIPANT

Tukey post-hoc test for participant actually revealed no significant differences between individual speakers, despite overall significance

Lowest p -value between two individual speakers: 0.057057



CONCLUSION

CONCLUSION

Convincing findings for both accentuation status and boundary strength:

- **Accentuation:** three-way distinction between primary accentuation, secondary accentuation and non-accentuation is reflected phonetically in vowel production
 - clearest difference before non-IU boundaries
 - before IU boundary, smaller difference between primary and secondary accented syllables + difference between secondary accented and non-accented syllables disappears
- **Boundary:** vowels in IU-final syllable phonetically more salient than vowels in non-IU-final position

CONCLUSION

Inconclusive findings for more fine-grained levels of **proximity** to IU boundary:

- insignificant result in regression model, but possibly due to scarcity of data points for syllables at further distances from IU-final boundary
- alternatively, phonetic strengthening effect immediately before IU boundary rather than gradual trend

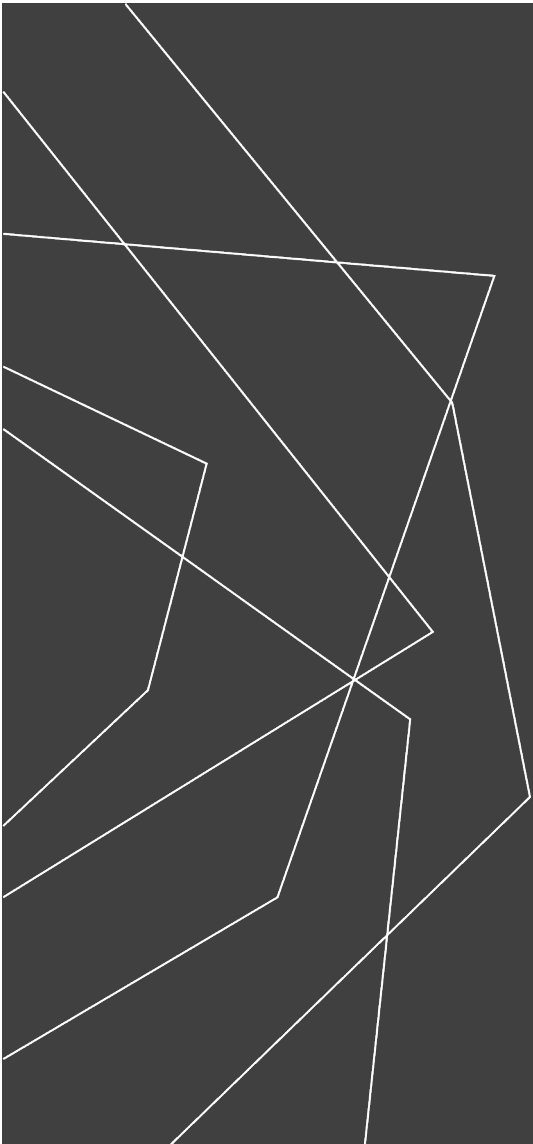
No sustained effect for different vowel types:

significant difference only between STRUT (shortest) and TRAP (longest) vowels: possibly due to /æ/-lengthening

CONCLUSION

Implications

Importance of distinguishing not just between pitch accented vs. unaccented syllables (ToBI), but also between primary and secondary pitch accented syllables (nuclear tone approach)



THANK YOU

Gil VERBEKE

gil.verbeke@ugent.be

Wout VAN PRAET

wout.vanpraet@uclouvain.be

REFERENCES

- Cho, T., and Keating, P.A. 2009. Effects of initial position versus prominence in English. *Journal of Phonetics* 37: 466-485.
- Edwards, J.E., Beckman, M.E., and Fletcher, J. 1991. The articulatory kinematics of nasal lengthening. *J. Acoust. Soc. Am.* 89: 369-382.
- Fougeron, C., and Keating, P.A. 1997. Articulatory strengthening at edges of prosodic domains. *J. Acoust. Soc. Am.* 101: 3728-3740.
- Wightman, C.W., Shattuck-Hufnagel, S., Ostendorf, M., and Price, P.J. 1992. Segmental durations in the vicinity of prosodic phrase boundaries. *J. Acoust. Soc. Am.* 91: 1707-1717.:

```

> summary(model.rand)
Linear mixed model fit by REML ['lmerMod']
Formula: Length_vowel ~ ACC + Vowel + BI + Prox + (1 | Participant)
Data: data

REML criterion at convergence: -453.7

Scaled residuals:
    Min       1Q   Median       3Q      Max
-1.9890 -0.6281 -0.1129  0.5325  3.0039

Random effects:
 Groups      Name      Variance Std.Dev.
Participant (Intercept) 0.0004914 0.02217
Residual          0.0016122 0.04015
Number of obs: 150, groups: Participant, 29

Fixed effects:
              Estimate Std. Error t value
(Intercept)  0.162109   0.013310  12.180
ACCSEC       -0.037768   0.008862  -4.262
ACCUNACC     -0.064930   0.009244  -7.024
VowelDRESS   -0.008490   0.012650  -0.671
VowelKIT     -0.026219   0.010569  -2.481
VowelLOT     -0.004178   0.012535  -0.333
VowelSTRUT   -0.015811   0.010670  -1.482
BI4           0.020618   0.011555   1.784
Prox         -0.004899   0.002446  -2.003

Correlation of Fixed Effects:
      (Intr) ACCSEC ACCUNA VDRESS Vw1KIT Vw1LOT VSTRUT BI4
ACCSEC      -0.195
ACCUNACC    -0.170  0.525
VowelDRESS  -0.398  0.013 -0.033
VowelKIT    -0.415  0.084 -0.089  0.491
VowelLOT    -0.381 -0.010  0.018  0.409  0.477
VowelSTRUT  -0.396  0.037 -0.030  0.509  0.603  0.509
BI4         -0.607 -0.196 -0.130  0.002 -0.084 -0.005 -0.065
Prox        -0.507 -0.203 -0.192 -0.018 -0.103 -0.070 -0.209  0.772
> |

```

```

> anova(model.rand)
Analysis of Variance Table

      npar    Sum Sq Mean Sq F value
ACC      2 0.107778  0.053889 33.4264
Vowel    4 0.017858  0.004464  2.7693
BI        1 0.044152  0.044152 27.3871
Prox      1 0.006465  0.006465  4.0102

```