

Reconciling categorical and gradient models of phonotactics

SCiL 2025 @ The University of Oregon

Connor Mayer

July 20th, 2025

Department of Language Science
University of California, Irvine

What is this talk about?

Question: Are phonotactic grammars categorical or gradient?

What is this talk about?

Question: Are phonotactic grammars categorical or gradient?

Answer: It depends on which monoid you use!

What is this talk about?

Two points I want to make:

What is this talk about?

Two points I want to make:

1. Gradient phonotactic models account for new data from a Turkish acceptability judgment task better than categorical models.

What is this talk about?

Two points I want to make:

1. Gradient phonotactic models account for new data from a Turkish acceptability judgment task better than categorical models.
2. This distinction turns out to be somewhat superficial if we think of models from a monoid-general perspective.

What is phonotactics?

What is phonotactics?

The legal ways in which sounds can be sequenced into words.

What is phonotactics?

The legal ways in which sounds can be sequenced into words.

This is (mostly) learned and language-specific:

What is phonotactics?

The legal ways in which sounds can be sequenced into words.

This is (mostly) learned and language-specific:

- /stik/ would be an ok English word; not a good Spanish word

What is phonotactics?

The legal ways in which sounds can be sequenced into words.

This is (mostly) learned and language-specific:

- /stik/ would be an ok English word; not a good Spanish word
- /bgera/ 'sound' is a fine Georgian word; not a good English word

Phonotactic acceptability judgments are gradient

A typical source of data is to ask speakers for *acceptability judgments*:

Phonotactic acceptability judgments are gradient

A typical source of data is to ask speakers for *acceptability judgments*:

- e.g. “on a scale of 1-7, how likely is ‘steek’ to become an English word?”

Phonotactic acceptability judgments are gradient

A typical source of data is to ask speakers for *acceptability judgments*:

- e.g. “on a scale of 1-7, how likely is ‘steek’ to become an English word?”

These judgments consistently display *gradience* [e.g. Chomsky and Halle, 1965, Coleman and Pierrehumbert, 1997, Scholes, 1966, Bailey and Hahn, 2001, Hayes and Wilson, 2008, Daland et al., 2011, a.o.].

What do we mean by gradient?

What do we mean by gradient?

poik

lvag

kip

What do we mean by gradience?

lvag $\ll$ poik $\ll$ kip

Where does this gradience come from?

Gradience in acceptability judgments can arise from performance factors such as misperception [e.g. Kahng and Durvasula, 2023]. However...

Where does this gradience come from?

Gradience in acceptability judgments can arise from performance factors such as misperception [e.g. Kahng and Durvasula, 2023]. However...

...patterns of gradient well-formedness often seem to be driven by the very same principles that govern absolute well-formedness... I conclude that the proposed attribution of gradient well-formedness judgments to performance mechanisms would be un insightful. Whatever “performance” mechanisms we adopted would look startlingly like the grammatical mechanisms that account for non-gradient judgments. [Hayes, 2000, p. 90]

Where does this gradience come from?

Gradience in acceptability judgments can arise from performance factors such as misperception [e.g. Kahng and Durvasula, 2023]. However...

...patterns of gradient well-formedness often seem to be driven by the very same principles that govern absolute well-formedness... I conclude that the proposed attribution of gradient well-formedness judgments to performance mechanisms would be un insightful. Whatever “performance” mechanisms we adopted would look startlingly like the grammatical mechanisms that account for non-gradient judgments. [Hayes, 2000, p. 90]

Typical modeling approach is to use a grammar that produces a gradient output.

- Often based on statistical frequencies in the lexicon [e.g. Hayes and Wilson, 2008].

Implementing a gradient phonotactic grammar

Our phonotactic grammars consist of a score function that assigns values to words.

Implementing a gradient phonotactic grammar

Our phonotactic grammars consist of a score function that assigns values to words.

$$\text{score} : \Sigma^* \rightarrow [0, 1]$$

Implementing a gradient phonotactic grammar

Our phonotactic grammars consist of a score function that assigns values to words.

$$\text{score} : \Sigma^* \rightarrow [0, 1]$$

Such a model can represent gradient acceptability judgments:

Implementing a gradient phonotactic grammar

Our phonotactic grammars consist of a score function that assigns values to words.

$$\text{score} : \Sigma^* \rightarrow [0, 1]$$

Such a model can represent gradient acceptability judgments:

$$\text{score}(\text{lvag}) = 0.01$$

Implementing a gradient phonotactic grammar

Our phonotactic grammars consist of a score function that assigns values to words.

$$\text{score} : \Sigma^* \rightarrow [0, 1]$$

Such a model can represent gradient acceptability judgments:

$$\text{score}(\text{lvag}) = 0.01 < \text{score}(\text{poik}) = 0.2$$

Implementing a gradient phonotactic grammar

Our phonotactic grammars consist of a score function that assigns values to words.

$$\text{score} : \Sigma^* \rightarrow [0, 1]$$

Such a model can represent gradient acceptability judgments:

$$\text{score}(\text{lvag}) = 0.01 < \text{score}(\text{poik}) = 0.2 < \text{score}(\text{kip}) = 0.4$$

Is phonotactics categorical?

Is phonotactics categorical?

Gorman [2013] argues that we have been premature in assuming the phonotactic grammar computes gradient outputs.

Is phonotactics categorical?

Gorman [2013] argues that we have been premature in assuming the phonotactic grammar computes gradient outputs.

- **Proposal:** grammar is categorical and gradient comes from other sources.

Is phonotactics categorical?

Gorman [2013] argues that we have been premature in assuming the phonotactic grammar computes gradient outputs.

- **Proposal:** grammar is categorical and gradient comes from other sources.
- A categorical grammar labels words as either grammatical or ungrammatical

Is phonotactics categorical?

Gorman [2013] argues that we have been premature in assuming the phonotactic grammar computes gradient outputs.

- **Proposal:** grammar is categorical and gradience comes from other sources.
- A categorical grammar labels words as either grammatical or ungrammatical

$$\text{score} : \Sigma^* \rightarrow \{0, 1\}$$

Is phonotactics categorical?

Gorman [2013] argues that we have been premature in assuming the phonotactic grammar computes gradient outputs.

- **Proposal:** grammar is categorical and gradience comes from other sources.
- A categorical grammar labels words as either grammatical or ungrammatical

$$\text{score} : \Sigma^* \rightarrow \{0, 1\}$$

These models cannot capture a situation where $\text{lvag} \ll \text{poik} \ll \text{kip}$.

Past work on categorical grammars

It is claimed that **categorical models do as well as or better than gradient models** in predicting empirical phonotactic phenomena.

Past work on categorical grammars

It is claimed that **categorical models do as well as or better than gradient models** in predicting empirical phonotactic phenomena.

- English onset acceptability [Gorman, 2013, Durvasula, 2020, Dai, 2025]

It is claimed that **categorical models do as well as or better than gradient models** in predicting empirical phonotactic phenomena.

- English onset acceptability [Gorman, 2013, Durvasula, 2020, Dai, 2025]
- Polish onset acceptability [Kostyszyn and Heinz, 2022, Dai, 2025]

It is claimed that **categorical models do as well as or better than gradient models** in predicting empirical phonotactic phenomena.

- English onset acceptability [Gorman, 2013, Durvasula, 2020, Dai, 2025]
- Polish onset acceptability [Kostyszyn and Heinz, 2022, Dai, 2025]
- Turkish vowel distributions [Gorman, 2013, Dai, 2025]

It is claimed that **categorical models do as well as or better than gradient models** in predicting empirical phonotactic phenomena.

- English onset acceptability [Gorman, 2013, Durvasula, 2020, Dai, 2025]
- Polish onset acceptability [Kostyszyn and Heinz, 2022, Dai, 2025]
- Turkish vowel distributions [Gorman, 2013, Dai, 2025]
- English medial cluster distributions [Gorman, 2013]

Limitations of previous work

Limitations of previous work

1. Use a very small number of data sets, almost all about consonant clusters

Limitations of previous work

1. Use a very small number of data sets, almost all about consonant clusters
2. The gradient model used in (almost) all cases is the *UCLA Phonotactic Learner*

Limitations of previous work

1. Use a very small number of data sets, almost all about consonant clusters
2. The gradient model used in (almost) all cases is the *UCLA Phonotactic Learner*
3. (Authors have different definitions of “categorical”)

Limitation 2: The UCLA Phonotactic Learner?

Limitation 2: The UCLA Phonotactic Learner?

The UCLA Phonotactic Learner has become the poster child for gradient phonotactics [Hayes and Wilson, 2008].

Limitation 2: The UCLA Phonotactic Learner?

The UCLA Phonotactic Learner has become the poster child for gradient phonotactics [Hayes and Wilson, 2008].

- But it also has to learn constraints from the data!

Limitation 2: The UCLA Phonotactic Learner?

The UCLA Phonotactic Learner has become the poster child for gradient phonotactics [Hayes and Wilson, 2008].

- But it also has to learn constraints from the data!
- Makes poor predictions for attested structures [e.g. Daland et al., 2011, Wilson and Gallagher, 2018]

Limitation 2: The UCLA Phonotactic Learner?

The UCLA Phonotactic Learner has become the poster child for gradient phonotactics [Hayes and Wilson, 2008].

- But it also has to learn constraints from the data!
- Makes poor predictions for attested structures [e.g. Daland et al., 2011, Wilson and Gallagher, 2018]
- Its performance is sensitive to how it is parameterized.

Limitation 2: The UCLA Phonotactic Learner?

The UCLA Phonotactic Learner has become the poster child for gradient phonotactics [Hayes and Wilson, 2008].

- But it also has to learn constraints from the data!
- Makes poor predictions for attested structures [e.g. Daland et al., 2011, Wilson and Gallagher, 2018]
- Its performance is sensitive to how it is parameterized.
- Do categorical models outperform it because it is gradient? Because of its constraint selection process? Because it has been run with sub-optimal hyperparameters?

A simpler comparison

We'll compare the performance of three categorical boolean models of Turkish vowel phonotactics against a simple probabilistic bigram model with a similar structure.

A simpler comparison

We'll compare the performance of three categorical boolean models of Turkish vowel phonotactics against a simple probabilistic bigram model with a similar structure.

We'll evaluate how these models predict new experimental data from a Turkish nonce word acceptability judgment task.

A new dataset of Turkish acceptability judgments

Constraints on Turkish vowels

Constraints on Turkish vowels

BACKNESS HARMONY: $*[\alpha\text{back}] \dots [-\alpha\text{back}]$

- A vowel must agree in backness with the preceding vowel.

Constraints on Turkish vowels

BACKNESS HARMONY: $*[\alpha\text{back}] \dots [-\alpha\text{back}]$

- A vowel must agree in backness with the preceding vowel.

ROUNDING HARMONY: $*[\alpha\text{round}] \dots [-\alpha\text{round}, +\text{high}]$.

- A high vowel must agree in roundness with the preceding vowel.

Constraints on Turkish vowels

BACKNESS HARMONY: $*[\alpha\text{back}] \dots [-\alpha\text{back}]$

- A vowel must agree in backness with the preceding vowel.

ROUNDING HARMONY: $*[\alpha\text{round}] \dots [-\alpha\text{round}, +\text{high}]$.

- A high vowel must agree in roundness with the preceding vowel.

These constraints govern suffix allomorphy, but their effect is also detectable in the lexicon and in acceptability judgment tasks [Zimmer, 1969].

Our data

The data we'll look at are acceptability judgments from a large, online study.

The data we'll look at are acceptability judgments from a large, online study.

- **Participants:** 85 native Turkish speakers (38F; mostly age 25-35) recruited on Prolific

The data we'll look at are acceptability judgments from a large, online study.

- **Participants:** 85 native Turkish speakers (38F; mostly age 25-35) recruited on Prolific
- **Task:** Wug word acceptability judgments

Stimuli: 576 wug words with CVCVC shape

Stimuli: 576 wug words with CVCVC shape

- Nine words for each unique pair of vowels (8×8 total pairs)

Stimuli: 576 wug words with CVCVC shape

- Nine words for each unique pair of vowels (8×8 total pairs)
- Probability of consonants controlled for within vowel groups

Stimuli: 576 wug words with CVCVC shape

- Nine words for each unique pair of vowels (8×8 total pairs)
- Probability of consonants controlled for within vowel groups
- Synthesized to speech using Google Cloud

Stimuli: 576 wug words with CVCVC shape

- Nine words for each unique pair of vowels (8×8 total pairs)
- Probability of consonants controlled for within vowel groups
- Synthesized to speech using Google Cloud
- Words and recordings vetted by two native Turkish speakers

Experiment task

Deney

Kalan süre

 Tekrar oynat

beyop

lvag

caçör

matan

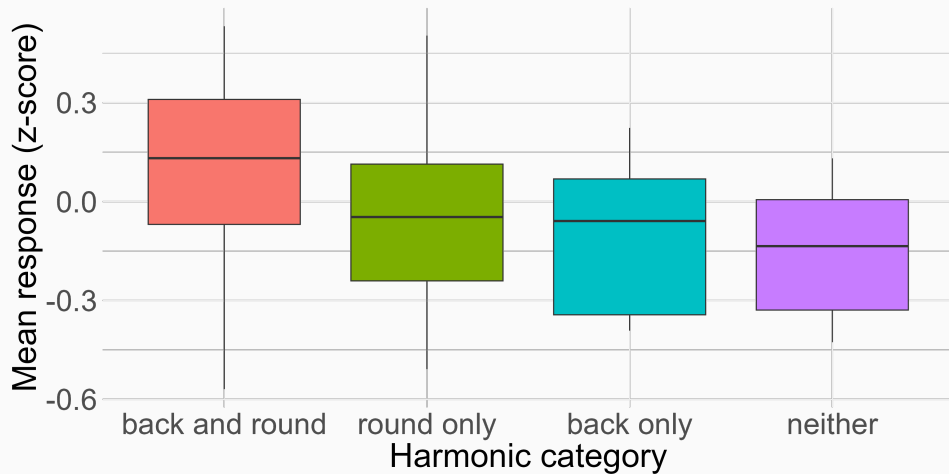
Each participant rated 192 tokens after training and attention checks: 16,320 tokens.

Each participant rated 192 tokens after training and attention checks: 16,320 tokens.

Responses are normalized to z-scores within participant

- Controls for differences in mean and spread between participants

Results



Defining our models

We'll test four simple models that have similar structures:

Value type	Constraint values
Log probability	Conditional probabilities
Boolean	Harmony [Gorman, 2013]
Boolean	Exception filtering [Dai, 2025]
Boolean	Threshold

All the models are TSL-2 grammars that operate on the vowel tier

All the models are TSL-2 grammars that operate on the vowel tier

- Informally, we **ignore consonants** and assign scores based on **vowel bigrams**

All the models are TSL-2 grammars that operate on the vowel tier

- Informally, we **ignore consonants** and assign scores based on **vowel bigrams**
- Constraints can reference start and end symbols $\bowtie$ and $\bowtie$

Scoring bigrams

Each model type has a Δ function that assigns a value to a bigram.

Scoring bigrams

Each model type has a Δ function that assigns a value to a bigram.

Boolean models

$$\Delta_b : \Sigma^2 \rightarrow \{0, 1\}$$

Scoring bigrams

Each model type has a Δ function that assigns a value to a bigram.

Boolean models

$$\Delta_b : \Sigma^2 \rightarrow \{0, 1\}$$

Log probability model

$$\Delta_p : \Sigma^2 \rightarrow (-\infty, 0]$$

Boolean model: score is 1 if a word contains only legal bigrams, 0 otherwise

$$\text{bool_score}(x_1, \dots, x_n) = \bigwedge_{i=1}^{n-1} \Delta_b(x_i, x_{i+1})$$

Boolean model: score is 1 if a word contains only legal bigrams, 0 otherwise

$$\text{bool_score}(x_1, \dots, x_n) = \bigwedge_{i=1}^{n-1} \Delta_b(x_i, x_{i+1})$$

Log probability model: score is the sum of the log probability of each bigram

$$\text{log_prob_score}(x_1, \dots, x_n) = \sum_{i=1}^{n-1} \Delta_p(x_i, x_{i+1})$$

An example

An example

Boolean model

$$\text{bool_score}([oi]) = \Delta_b(\times o) \wedge \Delta_b(oi) \wedge \Delta_b(i \times)$$

An example

Boolean model

$$\begin{aligned}\text{bool_score}([oi]) &= \Delta_b(\times o) \wedge \Delta_b(oi) \wedge \Delta_b(i \times) \\ &= 1 \wedge 0 \wedge 1\end{aligned}$$

An example

Boolean model

$$\begin{aligned}\text{bool_score}([oi]) &= \Delta_b(\times o) \wedge \Delta_b(oi) \wedge \Delta_b(i \times) \\ &= 1 \wedge 0 \wedge 1 \\ &= 0\end{aligned}$$

An example

Boolean model

$$\begin{aligned}\text{bool_score}([oi]) &= \Delta_b(\times o) \wedge \Delta_b(oi) \wedge \Delta_b(i \times) \\ &= 1 \wedge 0 \wedge 1 \\ &= 0\end{aligned}$$

Log probability model

An example

Boolean model

$$\begin{aligned}\text{bool_score}([oi]) &= \Delta_b(\times o) \wedge \Delta_b(oi) \wedge \Delta_b(i \times) \\ &= 1 \wedge 0 \wedge 1 \\ &= 0\end{aligned}$$

Log probability model

$$\text{log_prob_score}([oi]) = \Delta_p(\times o) + \Delta_p(oi) + \Delta_p(i \times)$$

An example

Boolean model

$$\begin{aligned}\text{bool_score}([oi]) &= \Delta_b(\times o) \wedge \Delta_b(oi) \wedge \Delta_b(i \times) \\ &= 1 \wedge 0 \wedge 1 \\ &= 0\end{aligned}$$

Log probability model

$$\begin{aligned}\text{log_prob_score}([oi]) &= \Delta_p(\times o) + \Delta_p(oi) + \Delta_p(i \times) \\ &= -2.52 + -2.24 + -0.78\end{aligned}$$

An example

Boolean model

$$\begin{aligned}\text{bool_score}([oi]) &= \Delta_b(\times o) \wedge \Delta_b(oi) \wedge \Delta_b(i \times) \\ &= 1 \wedge 0 \wedge 1 \\ &= 0\end{aligned}$$

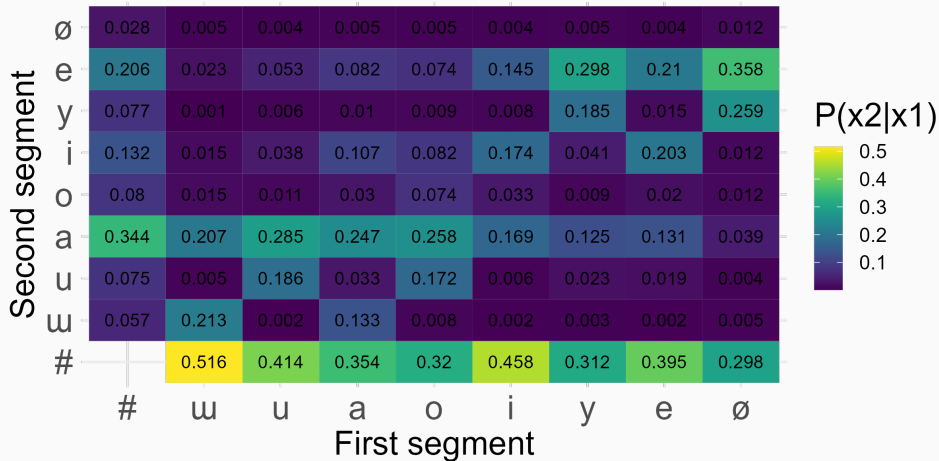
Log probability model

$$\begin{aligned}\text{log_prob_score}([oi]) &= \Delta_p(\times o) + \Delta_p(oi) + \Delta_p(i \times) \\ &= -2.52 + -2.24 + -0.78 \\ &= -5.54\end{aligned}$$

How do we define Δ for each model?

Conditional probability model

The probability model uses add-one smoothed conditional log probabilities derived from 18,472 citation forms in the TELL database [Inkelas et al., 2000].



Boolean harmony model [Gorman, 2013]

Words are grammatical if they satisfy both rounding and backness harmony.

Second segment

∅	TRUE	FALSE	FALSE	FALSE	FALSE	TRUE	TRUE	TRUE	TRUE
e	TRUE	FALSE	FALSE	FALSE	FALSE	TRUE	TRUE	TRUE	TRUE
y	TRUE	FALSE	FALSE	FALSE	FALSE	FALSE	TRUE	FALSE	TRUE
i	TRUE	FALSE	FALSE	FALSE	FALSE	TRUE	FALSE	TRUE	FALSE
o	TRUE	TRUE	TRUE	TRUE	TRUE	FALSE	FALSE	FALSE	FALSE
a	TRUE	TRUE	TRUE	TRUE	TRUE	FALSE	FALSE	FALSE	FALSE
u	TRUE	FALSE	TRUE	FALSE	TRUE	FALSE	FALSE	FALSE	FALSE
ʊ	TRUE	TRUE	FALSE	TRUE	FALSE	FALSE	FALSE	FALSE	FALSE
#		TRUE	TRUE	TRUE	TRUE	TRUE	TRUE	TRUE	TRUE
	#	ʊ	u	a	o	i	y	e	∅

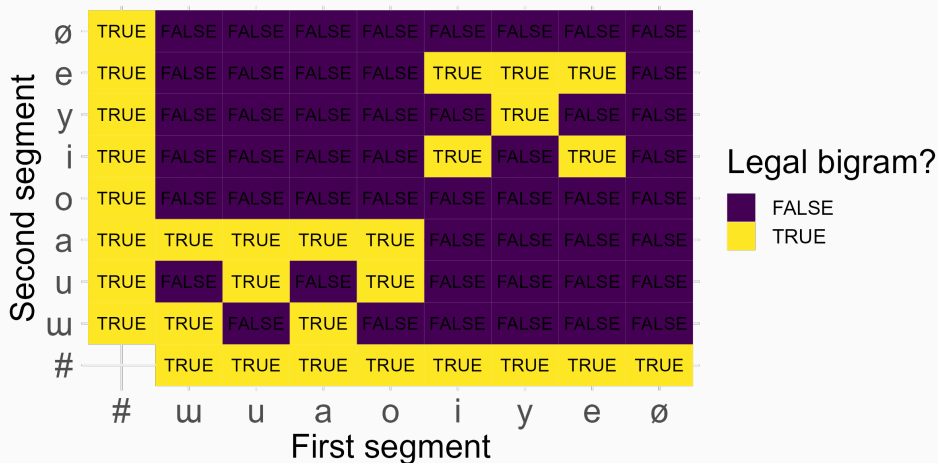
First segment

Legal bigram?

FALSE
TRUE

Boolean exception filtering model [Dai, 2025]

Categorical Turkish phonotactic grammar from Dai [2025] learned via an exception filtering process based on lexical frequency.



Threshold constraints

If $P(x_{i+1}|x_i)$ is above the 40th percentile, then $\Delta_b(x_i, x_{i+1}) = 1$

Second segment

	#	u	u	a	o	i	y	e	ø
ø	FALSE	FALSE	FALSE	FALSE	FALSE	FALSE	FALSE	FALSE	FALSE
e	TRUE	FALSE	TRUE	TRUE	TRUE	TRUE	TRUE	TRUE	TRUE
y	TRUE	FALSE	FALSE	FALSE	FALSE	FALSE	TRUE	FALSE	TRUE
i	TRUE	FALSE	FALSE	TRUE	TRUE	TRUE	TRUE	TRUE	FALSE
o	TRUE	FALSE	FALSE	FALSE	TRUE	FALSE	FALSE	FALSE	FALSE
a	TRUE	TRUE	TRUE	TRUE	TRUE	TRUE	TRUE	TRUE	FALSE
u	TRUE	FALSE	TRUE	FALSE	TRUE	FALSE	FALSE	FALSE	FALSE
u	TRUE	TRUE	FALSE	TRUE	FALSE	FALSE	FALSE	FALSE	FALSE
#	TRUE	TRUE	TRUE	TRUE	TRUE	TRUE	TRUE	TRUE	TRUE

First segment

Legal bigram?

FALSE
TRUE

Results

Let's look at correlations between model score for each vowel pair compared with its mean acceptability judgment across words and participants.

Results

Let's look at correlations between model score for each vowel pair compared with its mean acceptability judgment across words and participants.

Value type	Constraint set	r	τ	ρ
Log probability	Conditional probabilities	0.54	0.36	0.50
Boolean	Threshold (40th percentile)	0.46	0.37	0.45
Boolean	Harmony [Gorman, 2013]	0.38	0.30	0.37
Boolean	Exception filtering [Dai, 2025]	0.36	0.27	0.33

Results

Let's look at correlations between model score for each vowel pair compared with its mean acceptability judgment across words and participants.

Value type	Constraint set	r	τ	ρ
Log probability	Conditional probabilities	0.54	0.36	0.50
Boolean	Threshold (40th percentile)	0.46	0.37	0.45
Boolean	Harmony [Gorman, 2013]	0.38	0.30	0.37
Boolean	Exception filtering [Dai, 2025]	0.36	0.27	0.33

The simple probabilistic model outperforms the other models

- Closest competitor is the threshold model derived from the probabilistic model

Reconciling categorical and gradient models using monoids

The reconciliation begins



Commonalities between boolean and probabilistic models

Probabilistic and boolean TSL-2 models differ:

Commonalities between boolean and probabilistic models

Probabilistic and boolean TSL-2 models differ:

- **Boolean:** Assigns booleans to segmental bigrams, combines them using $\wedge$.

Commonalities between boolean and probabilistic models

Probabilistic and boolean TSL-2 models differ:

- **Boolean:** Assigns booleans to segmental bigrams, combines them using $\wedge$.
- **Probabilistic:** Assigns log probabilities to segmental bigrams, combines using $+$.

Commonalities between boolean and probabilistic models

Probabilistic and boolean TSL-2 models differ:

- **Boolean:** Assigns booleans to segmental bigrams, combines them using $\wedge$.
- **Probabilistic:** Assigns log probabilities to segmental bigrams, combines using $+$.

But the basic structure of each model is the same:

Commonalities between boolean and probabilistic models

Probabilistic and boolean TSL-2 models differ:

- **Boolean:** Assigns booleans to segmental bigrams, combines them using $\wedge$.
- **Probabilistic:** Assigns log probabilities to segmental bigrams, combines using $+$.

But the basic structure of each model is the same:

- We assign some **value** to each segmental bigram

Commonalities between boolean and probabilistic models

Probabilistic and boolean TSL-2 models differ:

- **Boolean:** Assigns booleans to segmental bigrams, combines them using $\wedge$.
- **Probabilistic:** Assigns log probabilities to segmental bigrams, combines using $+$.

But the basic structure of each model is the same:

- We assign some **value** to each segmental bigram
- We **aggregate** those values to get a score for the word

Commonalities between categorical and structural models

We can abstract away from specific values/aggregators:

Commonalities between categorical and structural models

We can abstract away from specific values/aggregators:

$$\Delta: \Sigma^2 \longrightarrow \mathcal{R}$$

where $\mathcal{R}$ is some set of values

Commonalities between categorical and structural models

We can abstract away from specific values/aggregators:

$$\Delta: \Sigma^2 \longrightarrow \mathcal{R}$$

$$\text{score}(x_1 \dots x_n) = \bigotimes_{i=1}^{n-1} \Delta(x_i, x_{i+1})$$

where $\mathcal{R}$ is some set of values and $\bigotimes$ is some binary operator over $\mathcal{R}$.

Other values of $\mathcal{R}$ and $\bigwedge$

We can make these simple models compute even more interesting quantities!

Other values of $\mathcal{R}$ and $\bigcirc$

We can make these simple models compute even more interesting quantities!

What does it compute?	$\mathcal{R}$	$\bigcirc$
Boolean scores [Gorman, 2013, Kostyszyn and Heinz, 2022, Dai, 2025]	$\{0, 1\}$	$\wedge$
Log probabilities	$(-\infty, 0]$	$+$
Probabilities	$[0, 1]$	$\times$

Other values of $\mathcal{R}$ and $\bigcirc$

We can make these simple models compute even more interesting quantities!

What does it compute?	$\mathcal{R}$	$\bigcirc$
Boolean scores [Gorman, 2013, Kostyszyn and Heinz, 2022, Dai, 2025]	$\{0, 1\}$	$\wedge$
Log probabilities	$(-\infty, 0]$	$+$
Probabilities	$[0, 1]$	$\times$
Integer scores [Durvasula, 2020, Kostyszyn and Heinz, 2022]	$\mathbb{N}$	$+$

Other values of $\mathcal{R}$ and $\bigwedge$

We can make these simple models compute even more interesting quantities!

What does it compute?	$\mathcal{R}$	$\bigwedge$
Boolean scores [Gorman, 2013, Kostyszyn and Heinz, 2022, Dai, 2025]	$\{0, 1\}$	$\wedge$
Log probabilities	$(-\infty, 0]$	$+$
Probabilities	$[0, 1]$	$\times$
Integer scores [Durvasula, 2020, Kostyszyn and Heinz, 2022]	$\mathbb{N}$	$+$
Constraint violation profiles	$\mathbb{N}^k$	$+$

Other values of $\mathcal{R}$ and $\bigwedge$

We can make these simple models compute even more interesting quantities!

What does it compute?	$\mathcal{R}$	$\bigwedge$
Boolean scores [Gorman, 2013, Kostyszyn and Heinz, 2022, Dai, 2025]	$\{0, 1\}$	$\wedge$
Log probabilities	$(-\infty, 0]$	$+$
Probabilities	$[0, 1]$	$\times$
Integer scores [Durvasula, 2020, Kostyszyn and Heinz, 2022]	$\mathbb{N}$	$+$
Constraint violation profiles	$\mathbb{N}^k$	$+$
Left SL-2 string transduction	Σ^*	$+$

What's going on here?

What's going on here?

This definition of our TSL-2 models is in **monoid-general terms**

$$\Delta: \Sigma^2 \longrightarrow \mathcal{R}$$

$$\text{score}(x_1 \dots x_n) = \bigwedge_{i=1}^{n-1} \Delta(x_i, x_{i+1})$$

What's going on here?

This definition of our TSL-2 models is in **monoid-general terms**

$$\Delta: \Sigma^2 \longrightarrow \mathcal{R}$$

$$\text{score}(x_1 \dots x_n) = \bigwedge_{i=1}^{n-1} \Delta(x_i, x_{i+1})$$

We can parameterize our model with different monoids that provide implementations of $\mathcal{R}$ and $\bigwedge$.

What's a monoid?

A monoid is an algebraic structure.

What's a monoid?

A monoid is an algebraic structure.

Monoid: a set $\mathcal{R}$ closed under a binary relation $\otimes$ such that:

What's a monoid?

A monoid is an algebraic structure.

Monoid: a set $\mathcal{R}$ closed under a binary relation $\otimes$ such that:

- $\otimes$ is associative

What's a monoid?

A monoid is an algebraic structure.

Monoid: a set $\mathcal{R}$ closed under a binary relation $\otimes$ such that:

- $\otimes$ is associative
- There's an identity element $\top$ in $\mathcal{R}$ such that $a \otimes \top = \top \otimes a = a$

Why are monoids interesting?

Why are monoids interesting?

The models we work with in FLT (SL, SP, TSL, FSA, CFG, etc.) can be expressed in monoid-general terms (or, in some cases, semiring-general terms).

Why are monoids interesting?

The models we work with in FLT (SL, SP, TSL, FSA, CFG, etc.) can be expressed in monoid-general terms (or, in some cases, semiring-general terms).

- In terms of $\mathcal{R}$ and $\bigcirc$ rather than specific values and operators

Why are monoids interesting?

The models we work with in FLT (SL, SP, TSL, FSA, CFG, etc.) can be expressed in monoid-general terms (or, in some cases, semiring-general terms).

- In terms of $\mathcal{R}$ and $\bigotimes$ rather than specific values and operators

Different monoids allow the same underlying model to compute different quantities.

Why are monoids interesting?

The models we work with in FLT (SL, SP, TSL, FSA, CFG, etc.) can be expressed in monoid-general terms (or, in some cases, semiring-general terms).

- In terms of $\mathcal{R}$ and $\bigwedge$ rather than specific values and operators

Different monoids allow the same underlying model to compute different quantities.

- Unifies superficially different models [Goodman, 1999].

Why are monoids interesting?

The models we work with in FLT (SL, SP, TSL, FSA, CFG, etc.) can be expressed in monoid-general terms (or, in some cases, semiring-general terms).

- In terms of $\mathcal{R}$ and $\bigwedge$ rather than specific values and operators

Different monoids allow the same underlying model to compute different quantities.

- Unifies superficially different models [Goodman, 1999].

We can separate the structure of the model from the values it computes.

Why is this useful for us as phonologists?

Why is this useful for us as phonologists?

Monoids allow us to relate the grammar to different domains or contexts

Why is this useful for us as phonologists?

Monoids allow us to relate the grammar to different domains or contexts

- Giorgolo and Asudeh [2014] apply different semirings (cf. monoids) to the same underlying semantic model to capture differences in heuristic vs. mathematical reasoning.

Connecting the grammar to different domains

There's perhaps an analogy to be made to Turkish.

Connecting the grammar to different domains

There's perhaps an analogy to be made to Turkish.

- Harmony is essentially categorical when determining suffix allomorphy

Connecting the grammar to different domains

There's perhaps an analogy to be made to Turkish.

- Harmony is essentially categorical when determining suffix allomorphy

'cat-PL'

Connecting the grammar to different domains

There's perhaps an analogy to be made to Turkish.

- Harmony is essentially categorical when determining suffix allomorphy

'cat-PL' kedi-ler ✓

Connecting the grammar to different domains

There's perhaps an analogy to be made to Turkish.

- Harmony is essentially categorical when determining suffix allomorphy

'cat-PL' kedi-ler ✓ kedi-lar ✗

- Harmony is a gradient property of roots
 - 66% of citation forms in TELL satisfy backness harmony
 - 70% satisfy rounding harmony

Connecting the grammar to different domains

There's perhaps an analogy to be made to Turkish.

- Harmony is essentially categorical when determining suffix allomorphy

'cat-PL' kedi-ler ✓ kedi-lar ✗

- Harmony is a gradient property of roots
 - 66% of citation forms in TELL satisfy backness harmony
 - 70% satisfy rounding harmony
- **But both sensitive to the same configurations!**

Connecting the grammar to different domains

Regardless of monoid, both the categorical and probabilistic grammars we saw here

Connecting the grammar to different domains

Regardless of monoid, both the categorical and probabilistic grammars we saw here

- are sensitive only to bigram constraints

Connecting the grammar to different domains

Regardless of monoid, both the categorical and probabilistic grammars we saw here

- are sensitive only to bigram constraints
- use segmental representations

Connecting the grammar to different domains

Regardless of monoid, both the categorical and probabilistic grammars we saw here

- are sensitive only to bigram constraints
- use segmental representations
- operate on the vowel tier

Connecting the grammar to different domains

Regardless of monoid, both the categorical and probabilistic grammars we saw here

- are sensitive only to bigram constraints
- use segmental representations
- operate on the vowel tier

These are segmental TSL-2 grammars, regardless of the values they assign.

Connecting the grammar to different domains

Regardless of monoid, both the categorical and probabilistic grammars we saw here

- are sensitive only to bigram constraints
- use segmental representations
- operate on the vowel tier

These are segmental TSL-2 grammars, regardless of the values they assign.

The same applies to other representations or grammars.

Durvasula [2020] implores us to prioritize work on categorical models so we can

Durvasula [2020] implores us to prioritize work on categorical models so we can

- “focus on what’s a possible constraint or rule”; and

Durvasula [2020] implores us to prioritize work on categorical models so we can

- “focus on what’s a possible constraint or rule”; and
- “commit to a specific set of representations”

We can have our phonotactic cake and eat it too

This is a false dichotomy

We can have our phonotactic cake and eat it too

This is a false dichotomy

- Constraints and representations in the grammar can be studied independently of the values the grammar assigns.

We can have our phonotactic cake and eat it too

This is a false dichotomy

- Constraints and representations in the grammar can be studied independently of the values the grammar assigns.
- Insight into the structure of the grammar can come from both gradient and categorical analyses!

We can have our phonotactic cake and eat it too

This is a false dichotomy

- Constraints and representations in the grammar can be studied independently of the values the grammar assigns.
- Insight into the structure of the grammar can come from both gradient and categorical analyses!
- This flexibility allows our models to engage with a broader range of empirical phenomena.

Thank you!

Thanks to Huteng Dai, Karthik Durvasula, Richard Futrell, Jeff Heinz, Jon Rawski, Megha Sundara, Bernard Tranel, and three anonymous reviewers for their useful feedback and discussions.

Thanks also to my Turkish consultants Cem Babalik and Defne Bilhan.

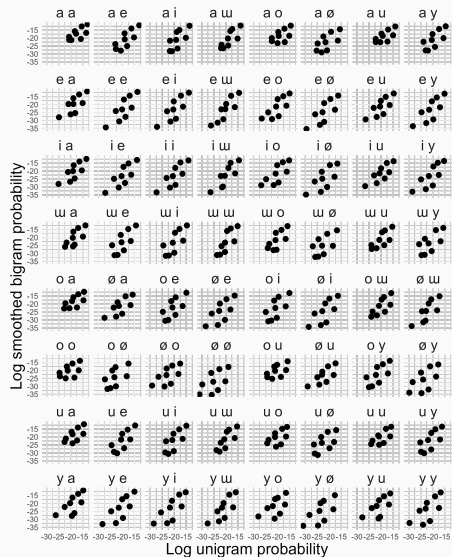
This work was supported by NSF Award #2214017.

References

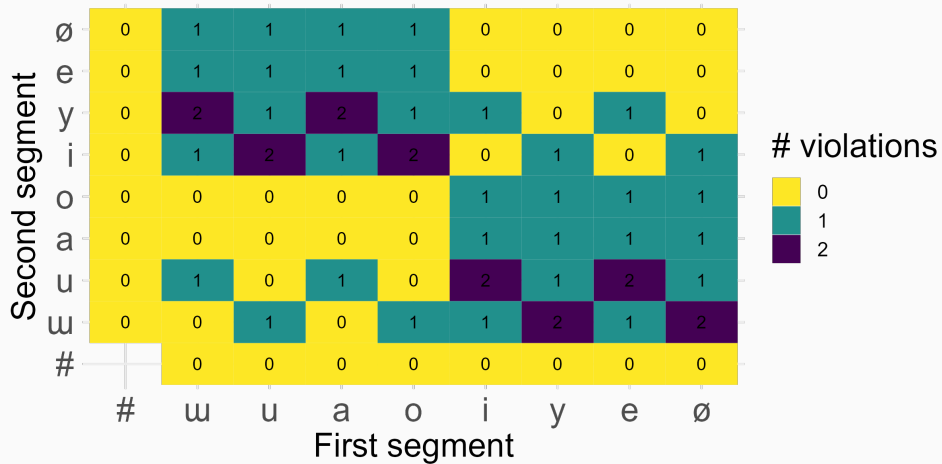
- Noam Chomsky and Morris Halle. Some controversial questions in phonological theory. *Journal of Linguistics*, 1(2):97–138, 1965.
- John Coleman and Janet Pierrehumbert. Stochastic phonological grammars and acceptability. In John Coleman, editor, *Proceedings of the 3rd Meeting of the ACL Special Interest Group in Computational Phonology*, pages 49–56. Association for Computational Linguistics, Somerset, NJ, 1997.
- Robert Scholes. *Phonotactic grammaticality*. Mouton, The Hague, 1966.
- Todd M. Bailey and Ulrike Hahn. Determinants of wordlikeness: Phonotactics or lexical neighborhoods. *Journal of Memory and Language*, 44: 568–591, 2001.
- Bruce Hayes and Colin Wilson. A maximum entropy model of phonotactics and phonotactic learning. *Linguistic inquiry*, 39(3):379–440, 2008.
- Robert Daland, Bruce Hayes, James White, Marc Garellek, Andreas Davis, and Ingrid Normann. Explaining sonority projection effects. *Phonology*, 28:197–234, 2011.
- Jimin Kahng and Karthik Durvasula. Can you judge what you don't hear? perception as a source of gradient wordlikeness judgments. *Glossa*, 8(1), 2023. doi: [url{https://doi.org/10.16995/glossa.9333}](https://doi.org/10.16995/glossa.9333).
- Bruce Hayes. *Gradient well-formedness in Optimality Theory*, pages 88–120. Oxford University Press, 2000.
- Kyle Gorman. *Generative phonotactics*. PhD thesis, University of Pennsylvania, 2013.
- Karthik Durvasula. O gradience, whence do you come?, 2020. Keynote presentation at the 2020 Annual Meeting on Phonology.

- Huteng Dai. An exception-filtering approach to phonotactic learning. *Phonology*, 42:e5, 2025.
- Kalina Kostyszyn and Jeffrey Heinz. Categorical account of gradient acceptability of word-initial Polish onsets. In *Proceedings of AMP 2021*. 2022.
- Colin Wilson and Gillian Gallagher. Accidental gaps and surface-based phonotactic learning: A case study of South Bolivian Quechua. *Linguistic Inquiry*, 49(3):610–623, 2018.
- K.E. Zimmer. Psychological correlates of some Turkish morpheme structure conditions. *Language*, pages 309–321, 1969.
- S. Inkelas, A. Küntay, C.O. Orgun, and R. Sprouse. Turkish electronic living lexicon (TELL): A lexical database. In *Proceedings of the Second International Conference on Language Resources and Evaluation (LREC'00)*. European Language Resources Association (ELRA), Athens, Greece, 2000.
- Josh Goodman. Semiring parsing. *Computational Linguistics*, 25(4):573–605, 1999.
- G. Giorgolo and A. Asudeh. One semiring to rule them all. In *Proceedings of the 36th Annual Meeting of the Cognitive Science Society*, pages 208–226. Cognitive Science Society, Québec City, 2014.

Stimuli structure



Cost semiring



Results

Value type	Constraint set	r	τ	ρ
Probability	UCLA Learner	0.56	0.37	0.54
Probability	Conditional probabilities	0.38	0.36	0.50
Log probability	Conditional probabilities	0.54	0.36	0.50
Integer	Harmony	0.38	0.30	0.38
Boolean	Threshold (40th percentile)	0.46	0.37	0.45
Boolean	Harmony [Gorman, 2013]	0.38	0.30	0.37
Boolean	Exception filtering [Dai, 2025]	0.36	0.27	0.33