

LEARNING FROM SMALL DATA: FEATURES, LOCALITY AND DIRECTIONALITY

Magdalena Markowska



Stony Brook University



NAPhC 2025 @ Concordia

THE SUBREGULAR HYPOTHESIS

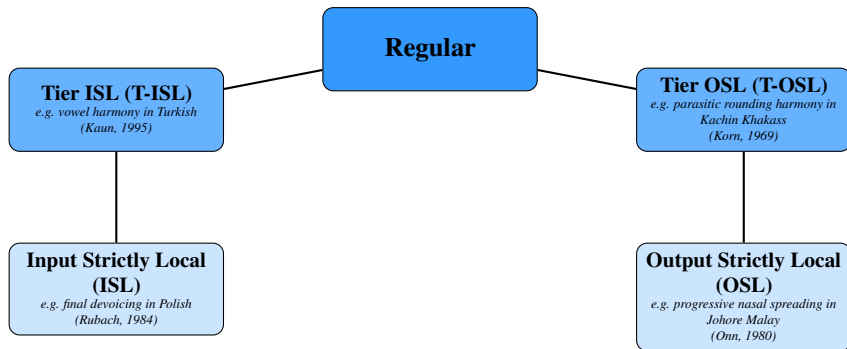
Subregular Phonology seeks to resolve the tension between theories of phonology that are *sufficiently expressive* to account for the cross-linguistic variation in phonological patterning and ones that are *sufficiently constrained* so that grammars can be learned from examples.

For processes these include the input and output strictly local maps and tier-based versions thereof.¹

¹Chandlee and Heinz (2018); Burness et al. (2021)

THE SUBREGULAR HYPOTHESIS

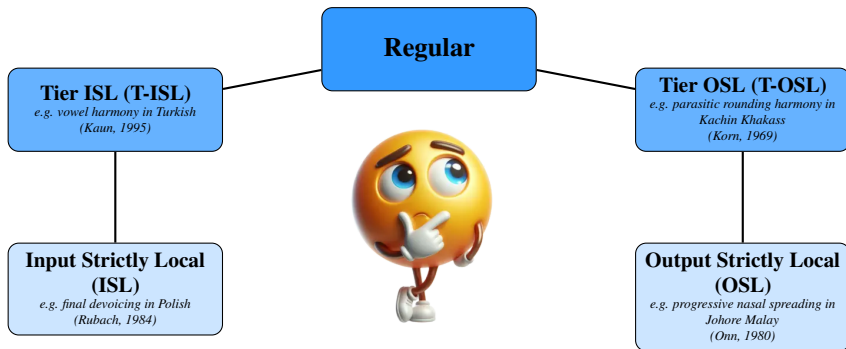
Subregular Phonology argues that most phonological patterns fall within a subregion of regular languages. Multiple classes within this region have been identified¹, for example:



¹Heinz (2018)

THE SUBREGULAR HYPOTHESIS

Subregular Phonology argues that most phonological patterns fall within a subregion of regular languages. Multiple classes within this region have been identified¹, for example:



Why does it matter?

¹Heinz (2018)

WHY THE SUBREGULAR HYPOTHESIS MATTERS

① Expressivity:

Captures both local and non-local phonological patterns²

② Provable learnability results:

Subregular classes are mathematically well-defined, which allows for efficient algorithms that effectively generalize the observed patterns³

③ Computationally tractable:

Many patterns can be represented with *subsequential functions*⁴ instantiated by *deterministic finite state transducers* (DFTs), making them suitable for various language technology applications

²Heinz (2010); Jardine (2016); Heinz (2018); McMullin and Hansson (2019)

³Heinz and Rogers (2013); Chandlee (2014); Chandlee et al. (2015a); Jardine and Heinz (2016)

⁴Sakarovitch (2009); Heinz and Lai (2013); Chandlee et al. (2014), Chandlee (2017); Chandlee and Heinz (2018)

THIS WORK

Over the past two decades, theoretical work in subregular phonology has led to the development of various learning **algorithms** that can successfully identify morphophonological patterns in the limit, given a characteristic sample.⁵ Here I focus on those that fall within the **ISL** class.

- ISLFLA: $O(n^2)$
- SOSFIA: $O(n)$

⁵Oncina et al. (1993); Oncina and Varó (1996); Chandlee et al. (2015b), Burness and McMullin (2019); Chandlee et al. (2014); Jardine et al. (2014)

THIS WORK

Over the past two decades, theoretical work in subregular phonology has led to the development of various learning **algorithms** that can successfully identify morphophonological patterns in the limit, given a characteristic sample.⁵ Here I focus on those that fall within the **ISL** class.

- ISLFLA: $O(n^2)$
- SOSFIA: $O(n)$



Practical Limitation: large amounts of training data.

⁵Oncina et al. (1993); Oncina and Varó (1996); Chandlee et al. (2015b), Burness and McMullin (2019); Chandlee et al. (2014); Jardine et al. (2014)

THIS WORK

Over the past two decades, theoretical work in subregular phonology has led to the development of various learning **algorithms** that can successfully identify morphophonological patterns in the limit, given a characteristic sample.⁵ Here I focus on those that fall within the **ISL** class.

- ISLFLA: $O(n^2)$
- SOSFIA: $O(n)$



Practical Limitation: large amounts of training data.

For example:

- a characteristic sample that guarantees success in learning vowel nasalization in English consists of approximately 18,278 UR-SR pairs.

⁵Oncina et al. (1993); Oncina and Varó (1996); Chandlee et al. (2015b), Burness and McMullin (2019); Chandlee et al. (2014); Jardine et al. (2014)

THIS WORK

Over the past two decades, theoretical work in subregular phonology has led to the development of various learning **algorithms** that can successfully identify morphophonological patterns in the limit, given a characteristic sample.⁵ Here I focus on those that fall within the **ISL** class.

- ISLFLA: $O(n^2)$
- SOSFIA: $O(n)$



Practical Limitation: large amounts of training data.

For example:

- a characteristic sample that guarantees success in learning vowel nasalization in English consists of approximately 18,278 UR-SR pairs. (considering 26 out of 44 phonemes)!

⁵Oncina et al. (1993); Oncina and Varó (1996); Chandlee et al. (2015b), Burness and McMullin (2019); Chandlee et al. (2014); Jardine et al. (2014)

THIS WORK

Over the past two decades, theoretical work in subregular phonology has led to the development of various learning **algorithms** that can successfully identify morphophonological patterns in the limit, given a characteristic sample.⁶ Here I focus on those that fall within the **ISL** class.

- ISLFLA: $O(n^2)$
- SOSFIA: $O(n)$

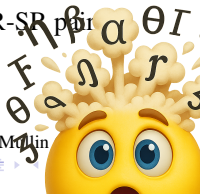


Practical Limitation: large amounts of training data.

For example:

- a characteristic sample that guarantees success in learning vowel nasalization in English consists of approximately 18,278 UR-SR pairs (considering 26 out of 44 phonemes)!

⁶Oncina et al. (1993); Oncina and Varó (1996); Chandlee et al. (2015b), Burness and McMullin (2015); Chandlee et al. (2014); Jardine et al. (2014)



A SOLUTION PROPOSAL

In this talk I show:

- how to switch from **segment-based** to **feature-based** representations to significantly reduce the amount of data required for learning.
- why the **directionality** of string processing (left-to-right vs. right-to-left) also plays a crucial role in data efficiency.
- how learning k -values for individual features can help reduce data size.

ROADMAP

- ① k -ISL functions (Chandlee, 2014):
 - ▶ definition and representation
 - ▶ training data requirement that guarantees success → why so large?
- ② The role of features in learning phonological functions
- ③ Right-to-left vs. left-to-right processing
- ④ Learning k value for individual features
- ⑤ Case studies:
 - ▶ vowel nasalization in English (6,157 → 44)
 - ▶ vowel shortening in Yawelmani (16,583 → 164)
 - ▶ ə-epenthesis in Chukchi (17,181 → 2,331)
 - ▶ opaque interaction of final devoicing and o -raising in Polish (16,583 → 8,211)

k -ISL FUNCTIONS AND k -ISL TRANSDUCERS

k -ISL functions can be instantiated with a *deterministic finite-state transducer (DFT)*, whose states and transitions are organized in such a way that the current state of the machine is always determined by the previous $k - 1$ symbols on the input.

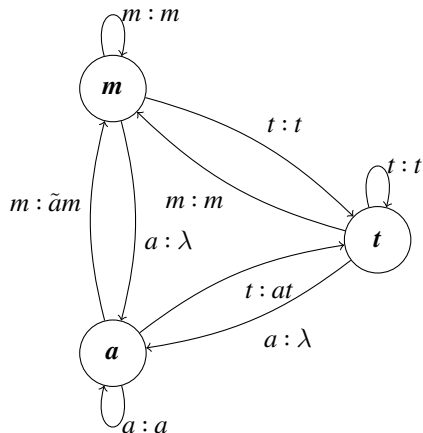
Crucially, every k -ISL DFTs has **the same structure**.

- vowel nasalization in English ($VN \rightarrow \tilde{V}N$): $k = 2$
- final devoicing in Polish ($D\bowtie \rightarrow T\bowtie$): $k = 2$

VOWEL NASALIZATION AS A 2-ISL PROCESS

$$\Sigma = \{a, m, t\}$$

$$\Delta = \Sigma \cup \{\tilde{a}\}$$

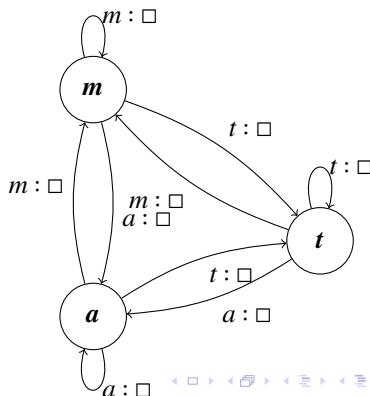


HOW DO WE LEARN SUCH FUNCTIONS?

SOSFIA

Given a characteristic sample of input-output pairs and an output empty DFT for the relevant class, SOSFIA is **guaranteed** to calculate the right outputs, for every transition (i.e. learn the function).

...	→	...
/kæm/	→	[k ^{green} æm]
/kkm/	→	[k ^{blue} km]
/mæm/	→	[m ^{green} æm]
/mæt/	→	[m ^{blue} æt]
...		...



DATA REQUIREMENTS FOR SOSFIA

The size of *k*-ISL DFT depends directly on the window size *k* and the size of the input alphabet Σ .

The state space for such machines can be expressed as $|Q| = |\Sigma|^{k-1}$.

Because SOSFIA needs to see strings that correspond to paths for every transition, the **characteristic sample** *S* that guarantees success will be, roughly speaking, proportional to $|Q|$. Let's consider some numbers:

	$ \Sigma $	$ Q $	$ S $
<i>k</i> = 2	4	4	84
	10	10	1,110
	50	50	127,550
<i>k</i> = 3	4	20	1,365
	10	110	111,100
	50	2550	318,877,550

EXAMPLE

Transition: $q = (\text{current_state}, \text{input}, \text{output}, \text{next_state})$

$f_1: k=2, V \rightarrow \tilde{V} \text{ --- } M$

$f_2: k=3, V \rightarrow \tilde{V} \text{ --- } MM$

Transition:	Required data:
f_1	
$(\times, a, \lambda, a)$	$\{aa-aa, am-\tilde{a}m, at-at\}$
(a, a, a, a)	$\{aaa-aaa, aat-aat, aam-a\tilde{a}m\}$
$(a, m, \tilde{a}m, m)$	$\{ama-\tilde{a}ma, amt-\tilde{a}mt, amm-\tilde{a}mm\}$
$(a, \times, a, \times)$	$\{a-a\}$
...	...
f_2	
(a, a, a, aa)	$\{aaa-aaa, aam-a\tilde{a}m, aaaa-aaaa, aaam-aa\tilde{a}m, \dots\}$
$(aa, m, \tilde{a}m, am)$	$\{aama-a\tilde{a}ma, aammm-a\tilde{a}mmm, aammt-a\tilde{a}mmt, \dots\}$
(am, t, t, mt)	$\{amta-\tilde{a}mta, amtaa-\tilde{a}mtaa, amtat-\tilde{a}mtat \dots\}$
$(mt, \times, \lambda, \times)$	$\{mt-mt\}$
...	...

SOME ASSUMPTIONS AND OBSERVATIONS

- If no data available for a transition, the output is assumed to be *faithful* to the input.
- For *non-identical* transitions, the learner has to be exposed to very particular strings.
 - ▶ **The more non-identical transitions, the more data required.**
- The structure of a k -ISL DFT also directly affects the data size:
 - ▶ The bigger the k the bigger the k -ISL DFT.
 - ▶ The bigger the alphabet the bigger the k -ISL DFT.
 - ▶ **The bigger the k -ISL DFT, the more data the learner needs.**

SOME ASSUMPTIONS AND OBSERVATIONS

- If no data available for a transition, the output is assumed to be *faithful* to the input.
- For *non-identical* transitions, the learner has to be exposed to very particular strings.
 - ▶ **The more non-identical transitions, the more data required.**
- The structure of a k -ISL DFT also directly affects the data size:
 - ▶ The bigger the k the bigger the k -ISL DFT.
 - ▶ The bigger the alphabet the bigger the k -ISL DFT.
 - ▶ **The bigger the k -ISL DFT, the more data the learner needs.**

The goal is then to:

- I. Reduce the machine size
- II. Reduce the amount of non-identical transitions

SOME ASSUMPTIONS AND OBSERVATIONS

- If no data available for a transition, the output is assumed to be *faithful* to the input.
- For *non-identical* transitions, the learner has to be exposed to very particular strings.
 - ▶ **The more non-identical transitions, the more data required.**
- The structure of a k -ISL DFT also directly affects the data size:
 - ▶ The bigger the k the bigger the k -ISL DFT.
 - ▶ The bigger the alphabet the bigger the k -ISL DFT.
 - ▶ **The bigger the k -ISL DFT, the more data the learner needs.**

The goal is then to:

- I. Reduce the machine size → FEATURES + k -value
- II. Reduce the non-identical transitions → DIRECTIONALITY

THE MAIN IDEA

Decompose the learning problem by generalizing over **features**.

- A learner predicts next feature values, as opposed to next symbol.

Let F be a finite set of features with binary or ternary values. If $f \in F$, ϕ_f is a homomorphism from Σ to the set $\{+, -, 0\}$.

- $\phi_{[\text{cor}]}(\text{mat}) = [- \ 0 \ +]$
- $\phi_{[\text{cons}]}(\text{mat}) = [+ \ - \ +]$

If Φ is an ordered set of features, then Φ is a pointwise product of its homomorphisms.

- $\Phi_{\begin{bmatrix} \text{cor} \\ \text{cons} \end{bmatrix}}(\text{mat}) = \begin{bmatrix} - & 0 & + \\ + & - & + \end{bmatrix}$

Finally, if S is a sample of input/output pairs composed of segments, $S_{(F,f)}$ is sample for feature f .

- $S_{(F,f)} = \{(\Phi_F(x), \Phi_f(y)) \mid (x, y) \in S\}$

FACTORING BY FEATURE (FxF)

Given a characteristic sample S , for each $f \in F$:

STEP I. Project $S_{([f],f)}$.

STEP II. Check if $S_{([f],f)}$ is *functional*.

- A. If yes, given a $\tau_{\square}$, extract a k -ISL DFT_f and calculate the outputs with a learner (e.g. SOSFIA)
- B. If not, pick another $f' \in F$ and project $S_{([f,f'],f)}$.

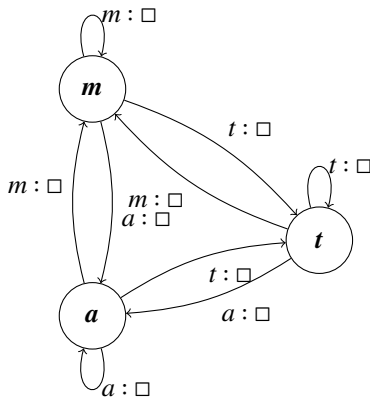
Repeat Step II until a functional sample is obtained.

AN EXAMPLE FROM ENGLISH

$k = 2$

$F = \{\text{cons, nasal, cor, lat, ...}\}$

...	→	...
/kæm/	→	[k $\tilde{a}$ m]
/kkm/	→	[k $\tilde{k}$ m]
/mæn/	→	[m $\tilde{a}$ m]
/mæt/	→	[m $\tilde{a}$ t]
...		...



FEATURE CONSONANTAL

$$k = 2$$

$$f = [\text{consonantal}]$$

...		...
/+--+/	→	[+--+]
/+++ /	→	[+++]
/+--- /	→	[+---]
/--- /	→	[---]
...		...

FEATURE CONSONANTAL

$$k = 2$$
$$f = [\mathbf{consonantal}]$$

...		...
/+-+/	→	[+-+]
/+++/	→	[+++]
/+--/	→	[+--]
/---/	→	[---]
...		...

[check_if_functional](#)

FEATURE CONSONANTAL

$$k = 2$$
$$f = [\text{consonantal}]$$

...		...
/+++/	→	[++]
/+++/	→	[+++]
/+--/	→	[+--]
/---/	→	[---]
...		...

check_if_functional

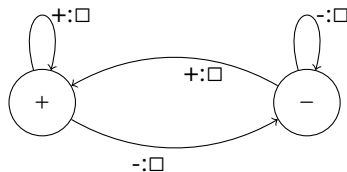
True

FEATURE CONSONANTAL

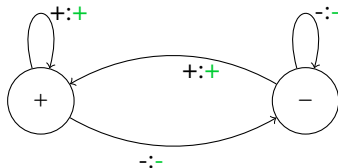
$$k = 2$$

$$f = [\text{consonantal}]$$

...		...
/+-+ /	→	[+-+]
/+++ /	→	[+++]
/+-- /	→	[+--]
/--- /	→	[---]
...		...



SOSFIA FxF'S OUTPUT FOR [CONSONANTAL]



FEATURE NASAL

$$k = 2$$

$$f = [\mathbf{nasal}]$$

$$\begin{array}{ccc}
 \dots & & \dots \\
 /--+/ & \rightarrow & [-++] \\
 /--+/ & \rightarrow & [--+] \\
 /+-+/ & \rightarrow & [+++] \\
 /+-+/ & \rightarrow & [+--+] \\
 \dots & & \dots
 \end{array}$$

FEATURE NASAL

$$k = 2$$

$$f = [\mathbf{nasal}]$$

$$\begin{array}{ccc}
 \dots & & \dots \\
 /--+/ & \rightarrow & [-++] \\
 /--+/ & \rightarrow & [--+] \\
 /+-+/ & \rightarrow & [+++] \\
 /+-+/ & \rightarrow & [+-+] \\
 \dots & & \dots
 \end{array}$$

`check_if_functional`

FEATURE NASAL

$$k = 2$$

$$f = [\mathbf{nasal}]$$

...		...
/ --+ /	→	[-++]
/ --+ /	→	[--+]
/ +-+ /	→	[+++]
/ +-+ /	→	[+-+]
...		...

check_if_functional

False

FEATURES NASAL AND CONSONANTAL

$$k = 2$$

$$f, f' = [\mathbf{nasal}, \text{consonantal}]$$

...

...

$$/ \begin{bmatrix} - \\ + \end{bmatrix} \begin{bmatrix} - \\ - \end{bmatrix} \begin{bmatrix} + \\ + \end{bmatrix} / \rightarrow \begin{bmatrix} - & + & + \end{bmatrix}$$

$$/ \begin{bmatrix} - \\ + \end{bmatrix} \begin{bmatrix} - \\ + \end{bmatrix} \begin{bmatrix} + \\ + \end{bmatrix} / \rightarrow \begin{bmatrix} - & - & + \end{bmatrix}$$

$$/ \begin{bmatrix} + \\ + \end{bmatrix} \begin{bmatrix} - \\ - \end{bmatrix} \begin{bmatrix} + \\ + \end{bmatrix} / \rightarrow \begin{bmatrix} + & + & + \end{bmatrix}$$

$$/ \begin{bmatrix} + \\ + \end{bmatrix} \begin{bmatrix} - \\ - \end{bmatrix} \begin{bmatrix} - \\ + \end{bmatrix} / \rightarrow \begin{bmatrix} - & - & + \end{bmatrix}$$

...

...

FEATURES NASAL AND CONSONANTAL

$$k = 2$$

$$f, f' = [\mathbf{nasal}, \text{consonantal}]$$

...

...

$$/ \begin{bmatrix} - \\ + \end{bmatrix} \begin{bmatrix} - \\ - \end{bmatrix} \begin{bmatrix} + \\ + \end{bmatrix} / \rightarrow \begin{bmatrix} - & + & + \end{bmatrix}$$

$$/ \begin{bmatrix} - \\ + \end{bmatrix} \begin{bmatrix} - \\ + \end{bmatrix} \begin{bmatrix} + \\ + \end{bmatrix} / \rightarrow \begin{bmatrix} - & - & + \end{bmatrix}$$

$$/ \begin{bmatrix} + \\ + \end{bmatrix} \begin{bmatrix} - \\ - \end{bmatrix} \begin{bmatrix} + \\ + \end{bmatrix} / \rightarrow \begin{bmatrix} + & + & + \end{bmatrix}$$

$$/ \begin{bmatrix} + \\ + \end{bmatrix} \begin{bmatrix} - \\ - \end{bmatrix} \begin{bmatrix} - \\ + \end{bmatrix} / \rightarrow \begin{bmatrix} - & - & + \end{bmatrix}$$

...

...

check_if_functional

FEATURES NASAL AND CONSONANTAL

$$k = 2$$

$$f, f' = [\mathbf{nasal}, \text{consonantal}]$$

...

...

$$/ \begin{bmatrix} - \\ + \end{bmatrix} \begin{bmatrix} - \\ + \end{bmatrix} \begin{bmatrix} + \\ + \end{bmatrix} / \rightarrow \begin{bmatrix} - & + & + \end{bmatrix}$$

$$/ \begin{bmatrix} - \\ + \end{bmatrix} \begin{bmatrix} - \\ + \end{bmatrix} \begin{bmatrix} + \\ + \end{bmatrix} / \rightarrow \begin{bmatrix} - & - & + \end{bmatrix}$$

$$/ \begin{bmatrix} + \\ + \end{bmatrix} \begin{bmatrix} - \\ + \end{bmatrix} \begin{bmatrix} + \\ + \end{bmatrix} / \rightarrow \begin{bmatrix} + & + & + \end{bmatrix}$$

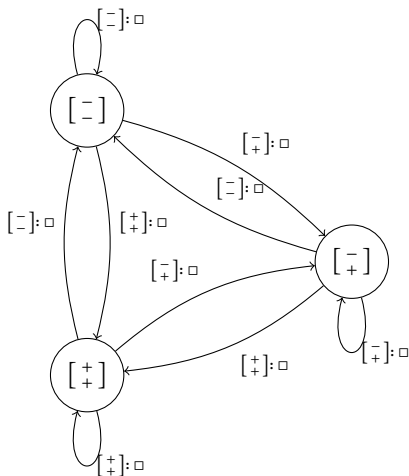
$$/ \begin{bmatrix} + \\ + \end{bmatrix} \begin{bmatrix} - \\ + \end{bmatrix} \begin{bmatrix} - \\ + \end{bmatrix} / \rightarrow \begin{bmatrix} + & - & - \end{bmatrix}$$

...

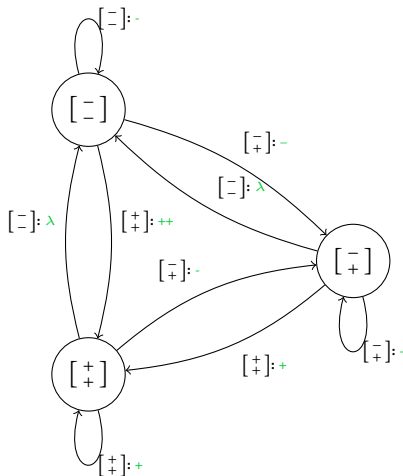
...

check_if_functional

True

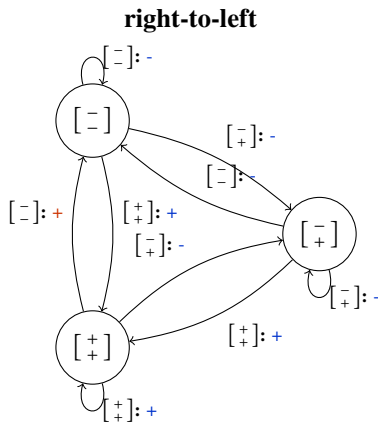
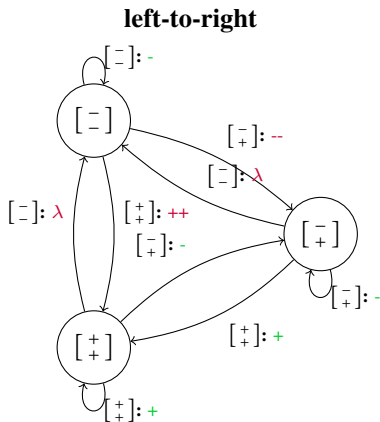
OUTPUT EMPTY DFT<sub>*nasal*
cons</sub>

SOSFIA'S OUTPUT FOR $\begin{bmatrix} \textit{nasal} \\ \textit{cons} \end{bmatrix}$



DIRECTIONALITY AND k -ISL FUNCTIONS

Whether a string is processed from left-to-right or right-to-left does not affect the **structure** of the machine (Chandlee, 2014).



DIRECTIONALITY AND k -ISL FUNCTIONS

The comparison from before showed us that changing the direction of the string affects only the outputs: only the transitions leaving nasal consonant state ($\begin{bmatrix} + \\ + \end{bmatrix}$) will be different from the inputs.

Essentially, **changing the direction can reduce** the number of non-identical transitions and thus reduce the **data** needed.

WHICH DIRECTION TO GO?

Changing direction can be achieved by simply reversing the data; the structure of the machine remains unchanged.

But how to know which direction to take?



WHICH DIRECTION TO GO?

Changing direction can be achieved by simply reversing the data; the structure of the machine remains unchanged.

But how to know which direction to take?

- flip a coin



WHICH DIRECTION TO GO?

Changing direction can be achieved by simply reversing the data; the structure of the machine remains unchanged.

But how to know which direction to take?

- flip a coin
- study the data



WHICH DIRECTION TO GO?

Changing direction can be achieved by simply reversing the data; the structure of the machine remains unchanged.

But how to know which direction to take?



- flip a coin
- study the data

Given the data, consider only the unfaithful input-output pairs and **calculate the longest common prefix (lcp)** for all substrings starting from the edge $\bowtie$.

- Change directionality if for any substring $\bowtie\sigma^+$, $\sigma \in \Sigma$ $\text{lcp} = \lambda$

HOW CAN k -VALUE BE LEARNED?

As many phonological processes are *feature-filling* or *feature-changing*, there often exists a subset of features that do not undergo any changes.

HOW CAN k -VALUE BE LEARNED?

As many phonological processes are *feature-filling* or *feature-changing*, there often exists a subset of features that do not undergo any changes.

This observation allows us for a possibility to adjust the memory window, k , for each feature individually.

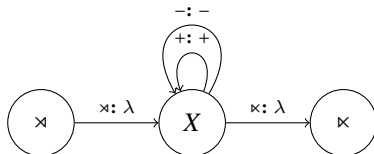
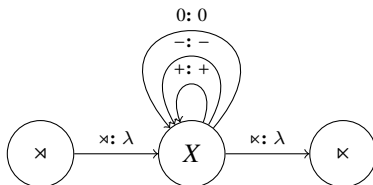
HOW CAN k -VALUE BE LEARNED?

As many phonological processes are *feature-filling* or *feature-changing*, there often exists a subset of features that do not undergo any changes.

This observation allows us for a possibility to adjust the memory window, k , for each feature individually.

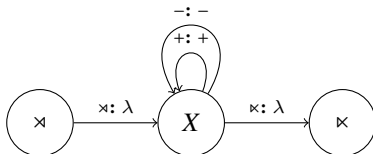
As a result, functions for the faithful features can be represented and learned from **one-state** machines.

AN EXAMPLE FROM ENGLISH

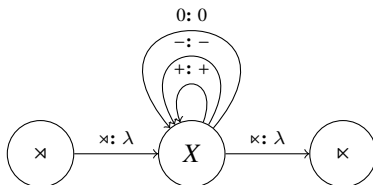
 $f = [\text{voice}]$

 $f = [\text{coronal}]$


AN EXAMPLE FROM ENGLISH

$f = [\text{voice}]$



$f = [\text{coronal}]$



Regardless of the number of feature values, the size of 1-ISL machine will remain the same.

PUTTING EVERYTHING TOGETHER

The notions of features, locality, and directionality create multiple possibilities.

The goal of a learner then is to start with the simplest assumption, i.e. a feature can be learned from itself with no memory stored ($k = 1$). If unsuccessful, the learner has three options:

- ➊ add another feature to the input
- ➋ increase k
- ➌ reverse the direction

PUTTING EVERYTHING TOGETHER

The notions of features, locality, and directionality create multiple possibilities.

The goal of a learner then is to start with the simplest assumption, i.e. a feature can be learned from itself with no memory stored ($k = 1$). If unsuccessful, the learner has three options:

- ① add another feature to the input
- ② increase k
- ③ reverse the direction

Soon coming 4th dimension: **tiers**

ALPHABET AND CHARACTERISTIC SAMPLE

- 2-ISL function:
 - ▶ $|\Sigma| = 26$
 - ▶ $|S| = 18,278$
- 3-ISL function:
 - ▶ $|\Sigma| = 7$
 - ▶ $|S| = 19,607$

EXPERIMENTS

Each function is learned:

- over **segments**, tested for learnability
- separately for each **feature** $f \in F$, also tested for success
- in both **left-to-right** and **right-to-left** processing directions
- $k = 1$ by default and increased when necessary

We also introduce a minimal sample search procedure: **AM β A**, which identifies a minimal subset of S sufficient for learning. We draw batches without replacement and test for success.

This yields two sample types:

- M : **segmental sample**, batch size = 50
- $M_{(F,f)}$: **featural sample** for feature f , batch size = 1
- We run **10 iterations** and report the average and standard deviation (SD)

Finally, we compute a **synthesis** of all individual featural samples to allow direct comparison with segmental performance.

ENGLISH: RESULTS

Feature f	Left-to-Right				Right-to-Left			
	Set $k = 2$		Learned k		Set $k = 2$		Learned k	
	Avg.	SD	Avg.	SD	Avg.	SD	Avg.	SD
[nas <i>cons</i>]	31	3.88	29	10.10	32	4.33	27	3.90
[round]	8	2.27	3	0	6	3.57	3	0
[cor]	23	4.76	4	0	16	11.73	4	0
...	...	...	...	...	...	...	...	...
Synthesis (Avg.)	189	10.06	46	3.71	137	10.50	44	1.89
Segments (Avg. $ M $)	6,157	650.21	–	–	–	–	–	–
$ S $	18,278							

YAWELMANI: RESULTS

Feature f	Left-to-Right				Right-to-Left			
	Set $k = 3$		Learned k		Set $k = 3$		Learned k	
	Avg.	SD	Avg.	SD	Avg.	SD	Avg.	SD
[high]	148	27.12	4	0	147	23.42	4	0
[long]	317	49.02	251	63.26	169	17.58	155	21.02
[cons]	36	5.47	3	0	30	11.52	3	0
...	...	...	...	...	...	...	...	...
Synthesis (Avg.)	952	49.90	258	62.95	851	56.51	164	21.24
Segments (Avg. $ M $)	16,583	702.06	—	—	—	—	—	—
$ S $	18,278							

CHUKCHI: RESULTS

Feature f	$ M_{(F,f)} $			
	Left-to-Right		Right-to-Left	
	Avg.	SD	Avg.	SD
[round]	165	30.34	192	51.10
[son]	225	30.28	227	35.57
[low]				
[cont]	260	12.77	244	19.54
[tense]				
[delrel]	924	57.82	892	41.75
[long]				
[distr]	801	14.98	917	55.02
[cor]				
[tense]	32	5.69	36	9.56
...	...	...	...	...
Synthesis (Avg.)	2,331	46.27	2,352	43.24
Segments (Avg. M)	17,181	28.87	—	—
 S 	19,607			

POLISH: RESULTS

Feature f	Left-to-Right				Right-to-Left			
	Set $k = 3$		Learned k		Set $k = 3$		Learned k	
	Avg.	SD	Avg.	SD	Avg.	SD	Avg.	SD
$\begin{bmatrix} \text{voice} \\ \text{son} \\ \text{high} \end{bmatrix}$	214	10.73	193	29	233	85.93	101	38.13
$\begin{bmatrix} \text{voice} \\ \text{son} \\ \text{low} \end{bmatrix}$	7,119	907.70	6,984	1,262.72	7,360	630.27	7,156	695.43
$\begin{bmatrix} \text{round} \end{bmatrix}$	137	51.76	4	0	173	45.04	4	0
$\begin{bmatrix} \text{dor} \end{bmatrix}$	32	4.33	3	0	31	11.15	3	0
...	...	...	...	...	...	...	...	...
Synthesis (Avg.)	8,235	208.39	8,211	139.89	8,114	145.28	8,068	202.83
Segments (Avg. M)	16,583	702.06	—	—	—	—	—	—
$ S $	19,607							

POLISH: BIGGER ALPHABET

Feature f	$ M_{(F,f)} $					
	Left-to-Right			Right-to-Left		
	$\Sigma = 8$	$\Sigma = 9$	$\Sigma = 10$	$\Sigma = 8$	$\Sigma = 9$	$\Sigma = 10$
$\begin{bmatrix} \text{voice} \\ \text{son} \\ \text{high} \end{bmatrix}$	215	239	179	76	126	–
$\begin{bmatrix} \text{voice} \\ \text{son} \\ \text{low} \\ \text{high} \end{bmatrix}$	7,506	–	–	8,186	–	–
$\begin{bmatrix} \text{voice} \\ \text{nasal} \\ \text{low} \end{bmatrix}$	–	8,827	8,007	–	8,982	–
$\begin{bmatrix} \text{round} \end{bmatrix}$	137	173	134	144	158	–
$\begin{bmatrix} \text{dor} \end{bmatrix}$	35	29	26	35	41	–
...	...	...	...	...	...	...
Synthesis (Avg.)	8,225	8,851	8,911	8,844	8,823	–
Segments (Avg. M)	30,344	52,161	79,540	–	–	–
 S 	37,448	66,429	111,110	37,448	66,429	111,110

CONCLUSION

Key findings

- Learning phonological patterns is sensitive to how we represent the data — switching from segments to **features** drastically reduces data requirements.
- Features allow for flexible k values and help to keep track of what is changing exactly and what remains the same.
- **Directionality matters**: in many cases, right-to-left processing leads to smaller and more efficient models.
- A sample search algorithm (**AMβA**.) successfully identifies a smaller sufficient sample. (We can definitely do better!)
- These results offer a roadmap for building interpretable, data-efficient learning models grounded in phonological structure.

FUTURE STEPS

- ① Evaluate the approach on more **naturalistic datasets**, including textbook-style data sets
- ② Extend the analysis to other classes: **TISL, OSL, TOSL**
- ③ Design a **searchable model space** with controlled variables — k , features, directionality, tiers — to identify the most data-efficient configurations
- ④ Incorporate **noisy and incomplete data** to evaluate robustness
- ⑤ Integrate **learning lexicon**

REFERENCES I

- Burness, P. and McMullin, K. (2019). Efficient learning of output tier-based strictly 2-local functions. In de Groote, P., Drewes, F., and Penn, G., editors, Proceedings of the 16th Meeting on the Mathematics of Language, pages 78–90, Toronto, Canada. Association for Computational Linguistics.
- Burness, P. A., McMullin, K. J., and Chandlee, J. (2021). Long-distance phonological processes as tier-based strictly local functions. Glossa: a journal of general linguistics, 6(1).
- Chandlee, J. (2014). Strictly Local Phonological Processes. PhD thesis, University of Delaware.
- Chandlee, J. (2017). Computational locality in morphological maps. Morphology, 27(4):599–641.
- Chandlee, J., Eyraud, R., and Heinz, J. (2014). Learning strictly local subsequential functions. Transactions of the Association for Computational Linguistics, 2:491–504.
- Chandlee, J., Eyraud, R., and Heinz, J. (2015a). Input strictly local functions. Transactions of the Association for Computational Linguistics, 3:117–128.
- Chandlee, J., Eyraud, R., and Heinz, J. (2015b). Output strictly local functions. In Proceedings of the 14th Meeting on the Mathematics of Language (MoL 2015), pages 112–125, Chicago, USA. Association for Computational Linguistics.
- Chandlee, J. and Heinz, J. (2018). Strict locality and phonological maps. Linguistic Inquiry, 49(1):23–60.
- Heinz, J. (2010). Learning long-distance phonotactics. Linguistic Inquiry, 41(4):623–661.
- Heinz, J. (2018). The computational nature of phonological generalizations. In Hyman, L. and Plank, F., editors, Phonological Typology, Phonetics and Phonology, chapter 5, pages 126–195. De Gruyter Mouton.
- Heinz, J. and Lai, R. (2013). Vowel harmony and subsequentiality. In Kornai, A. and Kuhlmann, M., editors, Proceedings of the 13th Meeting on Mathematics of Language, Sofia, Bulgaria.

REFERENCES II

- Heinz, J. and Rogers, J. (2013). Learning subregular classes of languages with grammatical inference. In Proceedings of the 14th Meeting on the Mathematics of Language (MoL), pages 1–11.
- Jardine, A. (2016). Locality and non-locality in tonal patterns. PhD thesis, University of Delaware.
- Jardine, A., Chandlee, J., Eyraud, R., and Heinz, J. (2014). Very efficient learning of structured classes of subsequential functions from positive data. In Clark, A., Kanazawa, M., and Yoshinaka, R., editors, Proceedings of the Twelfth International Conference on Grammatical Inference (ICGI 2014), volume 34, pages 94–108. JMLR: Workshop and Conference Proceedings.
- Jardine, A. and Heinz, J. (2016). Learning tier-based strictly local languages. Transactions of the Association for Computational Linguistics, 4:87–98.
- Kaun, A. C. (1995). The typology of rounding harmony: An optimality theoretic approach. Ph.d. dissertation, UCLA.
- Korn, E. (1969). Dinka vowel mutation and phonology. Studies in African Linguistics, 1(1):55–72.
- McMullin, K. and Hansson, G. (2019). Inductive learning of locality relations in segmental phonology. Phonology, 36(2):345–382.
- Oncina, J., García, P., and Vidal, E. (1993). Learning subsequential transducers for pattern recognition tasks. IEEE Transactions on Pattern Analysis and Machine Intelligence, 15:448–458.
- Oncina, J. and Varó, M. A. (1996). Using domain information during the learning of a subsequential transducer. Lecture Notes in Computer Science - Lecture Notes in Artificial Intelligence, pages 313–325.
- Onn, F. M. (1980). Aspects of Malay phonology and morphology: A generative approach. Ph.d. dissertation, University of Illinois at Urbana-Champaign.
- Rubach, J. (1984). Cyclic and lexical phonology. De Gruyter Mouton, Berlin, New York.
- Sakarovitch, J. (2009). Elements of Automata Theory. Cambridge University Press. Translated by Reuben Thomas from the 2003 edition published by Vuibert, Paris.