

PP: *Phonesthemic Palimpsest*

Lawrence McGuire

February 2025

Gibberish'ing: The Computed Simultaneous Poetry

The idea to work with nonsense as the material for a poem came from listening to some recordings by the Toronto group 'The Four Horsemen' performing and reciting what is called a simultaneous poem; multiple voices reciting a nonsense poem based on phonetic elements that can be classified as originating from their respective native languages. Usually, these recitations were accompanied by a cacophony of electronically synthesized sounds. What is not mentioned in these Dada readings, is the element of space. It became obvious to me that a simultaneous poem would always involve the performers to be spread around a room, uttering their poems quite segregated spatially. In "PP," I wanted to compress this space into a singular point. One where all voices physically meet, are electronically summed and then almost point-wise diffused into space.

To generate this poem, I collected the 15 highest phoneme frequencies in 33 languages and arranged these into lists. This is followed by a normalization of these frequencies and are then used as weighted tables for a simple nested Markov chain. Firstly, by some chance procedures, the algorithm picked the language, it then subsequently chose a phoneme from the set of that language. For each language, a string was generated that contained an exponentially increasing number of phonemes depending on which language number we were at. I eventually called these strings the *poop-sequences*. Ultimately, each sequence, started with an equal chance of choosing between the 33 language sets, but as we go further into the phoneme count, the weights become lopsided towards one certain language. So, each poop-sequence eventually only contained phonemes originating from one phoneme set. The goal for generating this poem was for me to perform, record, and in the end resynthesize it by various procedures. So, in order for me to physically be able to perform these *poop-sequences*, I had to actually cluster phonemes into nonsense words, where I imagined each poop sequence to have a certain message and impose a certain prosody and speaking style on the collection of phonemes which was also influenced by the language

it was converging to. One sequence I imagined a nonsense-speaker of which the speech was transcribed to be out of breath and having to stop every two or three words, all the while holding a Spanish like-rhythm. Imposing an imaginary syntax also became a part of these operations. The chosen linguistic topologies also differed according to which language the poop-sequences were set to be in; verb-subject-object, subject-object-verb, or subject-verb-object. These, for example, had an impact on the size of phoneme clusters. If the language uses a zero-article system, this also impacted the grouping and imagined flow of the language. Each sequence converges to the most common sounds in a certain language, but remains arguably gibberish. By imposing a partly improvised, partly linguistically informed stratification on each string, I attempted to make the nonsense poem less nonsensical. A vast richness in possible meanings can occur even when working on the sub-morphemic level; these associations with certain phonemes are called phonesthemes, and opposes the idea that the relationship between sound and meaning is completely arbitrary: a sound symbolism in language.

Now, I ended up not performing all of these sequences, and instead, I decided to use artificial performers to recite these poems. For the 33 "performances," I did some additional editing to allow for longer pauses in the dialogues of the synthetic characters with their counterparts (these are not noted down, but imagined, personified, and vocalized in the editing process by myself). The clusters of IPA symbols (e.g., nonsense words) were then recited by a randomly chosen synthetic speaker (having a certain pronunciation due to their language's phonetic inventory). In the first section all of these synthetic, nonsense sentences are played back at the same time through two speakers. Important to note is that all sentences are summed to one channel, where the monophonic signal is sent to the DAC and respective speakers. In the third section, the poop-sequences are serialized and recited by a single speaker. This speaker is me, but using a speech-to-speech algorithm which takes the synthetic nonsense sentences as inputs and spits out a timbre-transferred version in my voice timbre. The serialization of the poop-sequences and the singularization of the vocal persona and language specificity essentially expresses an alternative hyper-compressed form of the usually unintelligible simultaneous nonsense poem.

SOFTWARE USED:

- *Supercollider*: generating poop-sequences, LPC analysis-resynthesis
- *Python*: various scripts for phoneme frequency analysis using data from the Cross-linguistic Phonological Frequencies Corpus (XPF) PHOIBLE (<https://phoible.org/parameters>)
- *Webtools*: <https://ipa-reader.com/>, <http://phonetictools.altervista.org/phonverter/>, <https://github.com/hogobogobogo/media/blob/main/DC%20collection%20Son%20students%20-%20SINGLE.xlsx>
- *Praat*, *Chipspeech*, and *NKOAPP* for sound generation and analyses

Appendix

Algorithm 1: Generate Poop-Sequences

Input: Number of sequences $N = 33$, initial length $L_0 = 1$, growth factor $\alpha = 2$
Output: List of sequences S
 $S \leftarrow \emptyset$;
for $i \leftarrow 1$ **to** N **do**
 targetLanguage $\leftarrow i$;
 ;
 currentLength $\leftarrow L_0 \times \alpha^{(i-1)}$;
 ;
 newSequence $\leftarrow$ empty string;
 ;
 for $j \leftarrow 1$ **to** *currentLength* **do**
 $p \leftarrow \frac{j}{\text{currentLength}}$;
 ;
 languageIndex $\leftarrow$ **weightedLanguageChoice**(p , targetLanguage);
 ;
 phonemeTable $\leftarrow$ phonemeTables[languageIndex];
 ;
 chosenPhoneme $\leftarrow$ **weightedRandomChoice**(phonemeTable);
 ;
 newSequence $\leftarrow$ newSequence + chosenPhoneme;
 ;
 end
 Append newSequence to S ;
end
return S ;

Algorithm 2: weightedLanguageChoice

Input: Linear progression parameter $p \in [0, 1]$, target language T
Output: Selected language index $L \in \{1, \dots, 33\}$
 $totalWeight \leftarrow 0$;
 $weights \leftarrow \emptyset$;
for $L \leftarrow 1$ **to** 33 **do**
 if $L = T$ **then**
 $w \leftarrow \frac{1-p}{33} + p$;
 else
 $w \leftarrow \frac{1-p}{33}$;
 end
 Append w to $weights$;
 $totalWeight \leftarrow totalWeight + w$;
end
 $randomValue \leftarrow \text{RANDOM_NUMBER}(0, totalWeight)$;
 $cumulativeWeight \leftarrow 0$;
for $L \leftarrow 1$ **to** 33 **do**
 $cumulativeWeight \leftarrow cumulativeWeight + weights[L]$;
 if $randomValue \leq cumulativeWeight$ **then**
 return L ;
 end
end

Algorithm 3: weightedRandomChoice

Input: Weighted table T , each element is a tuple (phoneme, weight)
Output: Selected phoneme
 $totalWeight \leftarrow 0$;
for each $(item, weight)$ **in** T **do**
 $totalWeight \leftarrow totalWeight + weight$;
end
 $randomValue \leftarrow \text{RANDOM_NUMBER}(0, totalWeight)$;
 $cumulativeWeight \leftarrow 0$;
for each $(item, weight)$ **in** T **do**
 $cumulativeWeight \leftarrow cumulativeWeight + weight$;
 if $randomValue \leq cumulativeWeight$ **then**
 return $item$;
 end
end

Below you can find 33 sets of weighted tables of phoneme frequencies in different languages.

Turkish			normalized
/a/	11.920%	0.11920	0.148515468
/e/	8.912%	0.08912	0.111037739
/i/	8.600%	0.08600	0.107150422
/r/	6.722%	0.06722	0.08375176
/l/	5.922%	0.05922	0.073784279
/ʊ/	5.114%	0.05114	0.063717123
/d/	4.706%	0.04706	0.058633708
/k/	4.683%	0.04683	0.058347142
/n/	4.487%	0.04487	0.05590511
/m/	3.752%	0.03752	0.046747486
/j/	3.336%	0.03336	0.041564396
/u/	3.235%	0.03235	0.040306002
/t/	3.014%	0.03014	0.037552485
/s/	3.014%	0.03014	0.037552485
/b/	2.844%	0.02844	0.035434395

Czech		normalized
/o/	0.08124	0.110038061
/a/	0.07628	0.103319834
/ɪ/	0.07490	0.101450649
/ɛ/	0.06805	0.092172453
/v/	0.04843	0.065597529
/i:/	0.04494	0.06087039
/t/	0.04702	0.063687711
/r/	0.04487	0.060775576
/n/	0.04373	0.059231467
/l/	0.04058	0.054964851
/k/	0.03803	0.051510924
/t/	0.03780	0.051199393
/s/	0.03591	0.048639424
/p/	0.03227	0.043709112
/c/	0.02424	0.032832627

French		normalized
/a/	0.081	0.104247104
/ɪ/	0.069	0.088803089
/ʌ/	0.068	0.087516088
/ɛ/	0.068	0.087516088
/e/	0.065	0.083655084
/s/	0.058	0.074646075
/i/	0.056	0.072072072
/ə/	0.049	0.063063063
/t/	0.045	0.057915058
/k/	0.045	0.057915058
/p/	0.043	0.055341055
/d/	0.035	0.045045045
/m/	0.034	0.043758044
/ñ/	0.033	0.042471042
/n/	0.028	0.036036036

New Zealand English		normalized
/ə/	2	0.166944908
/ɪ/	1.3	0.10851419
/æ/	0.8	0.066777963
/t/	0.8	0.066777963
/n/	0.75	0.062604341
/ɪ/	0.74	0.061769616
/d/	0.7	0.058430718
/s/	0.69	0.057595993
/u:/	0.6	0.050083472
/ɔ:/	0.6	0.050083472
/k/	0.6	0.050083472
/m/	0.6	0.050083472
/ɛ/	0.6	0.050083472
/v/	0.6	0.050083472
/f/	0.6	0.050083472

Am. english		normalized
/ə/	0.1149	0.1596720400
/n/	0.0711	0.0988048916
/ɹ/	0.0694	0.0964424680
/t/	0.0691	0.0960255698
/ɪ/	0.0632	0.0878265703
/s/	0.0475	0.0660088938
/d/	0.0421	0.0585047248
/l/	0.0396	0.0550305725
/i/	0.0361	0.0501667593
/k/	0.0318	0.0441912173
/ð/	0.0295	0.0409949972
/ɛ/	0.0286	0.0397443024
/m/	0.0276	0.0383546415
/z/	0.0276	0.0383546415
/p/	0.0215	0.0298777098

Br. English		normalized
/ə/	0.1074	0.153890242
/n/	0.0711	0.101877060
/t/	0.0691	0.099011320
/ɹ/	0.0351	0.050293738
/ɪ/	0.0833	0.119358074
/s/	0.0475	0.068061327
/d/	0.0421	0.060323829
/l/	0.0396	0.056741654
/i:/	0.0361	0.051726608
/k/	0.0318	0.045565267
/ð/	0.0295	0.042269666
/ɛ/	0.0286	0.040980083
/m/	0.0276	0.039547213
/z/	0.0276	0.039547213
/p/	0.0215	0.030806706

Namibian		normalized
/t/	1	0.066666667
/d/	1	0.066666667
/k/	1	0.066666667
/g/	1	0.066666667
/f/	1	0.066666667
/v/	1	0.066666667
/s/	1	0.066666667
/z/	1	0.066666667
/ʃ/	1	0.066666667
/ʒ/	1	0.066666667
/h/	1	0.066666667
/n/	1	0.066666667
/m/	1	0.066666667
/ŋ/	1	0.066666667
/l/	1	0.066666667

Italian		normalized
/i/	0.1128	0.129342965
/e/	0.1036	0.118793716
/a/	0.1020	0.116959064
/ɔ/	0.0983	0.112716432
/u/	0.0833	0.095516569
/n/	0.0688	0.078890036
/t/	0.0665	0.076252723
/l/	0.0651	0.074647403
/s/	0.0480	0.05503956
/m/	0.0302	0.034629056
/d/	0.0288	0.033023736
/k/	0.0280	0.03210641
/g/	0.0164	0.018805183
/f/	0.0104	0.011925238
/v/	0.0099	0.011351909

Polish		normalized
/a/	8.12%	0.130714746
/o/	7.76%	0.124919511
/e/	7.74%	0.124597553
/r/	3.96%	0.063747585
/t/	3.86%	0.062137798
/n/	3.67%	0.059079202
/i/	3.61%	0.058113329
/v/	3.24%	0.052157115
/ɹ/	3.24%	0.052157115
/j/	3.06%	0.049259498
/p/	2.95%	0.047488731
/s/	2.91%	0.046844816
/u/	2.83%	0.045556986
/d/	2.60%	0.041854475
/k/	2.57%	0.041371539

Hebrew		normalized
/h/	10.87%	0.131789525
/j/	11.06%	0.134093113
/f/	10.38%	0.125848691
/t/	6.90%	0.083656644
/n/	6.50%	0.078806984
/ʁ/	5.61%	0.068016489
/k/	5.60%	0.067895247
/m/	4.80%	0.058195926
/s/	3.90%	0.047284190
/l/	4.10%	0.049709020
/b/	3.70%	0.044859360
/g/	3.20%	0.038797284
/ħ/	2.48%	0.030067895
/χ/	2.14%	0.025945684
/ts/	1.24%	0.015033948

Swiss German (most common uncommon phonemes)	normalized
/R/	0.85
/ts/	0.8
/χ/	0.75
/u:/	0.75
/y:/	0.7
/œ/	0.65
/v/	0.65
/uə/	0.6
/iə/	0.55
/ŋ/	0.55
/ø/	0.5
/æ:/	0.5
/æi/	0.45
/tʃ/	0.4
/ʌ/	0.3
/t'/	0.25
/ʃ/	0.2

Taiwanese Hokkien	normalize
a	0.12
	0.11
i	0.1
u	0.095
t	0.085
k	0.08
p	0.075
ts	0.07
n	0.065
m	0.06
o	0.058
e	0.055
	0.05
ai	0.045
au	0.045

slovenian		normalized
a	0.115	0.105215005
i	0.105	0.096065874
o	0.095	0.086916743
e	0.09	0.082342177
u	0.085	0.077767612
n	0.08	0.073193047
r	0.075	0.068618481
t	0.07	0.064043916
s	0.065	0.05946935
l	0.06	0.054894785
k	0.058	0.053064959
m	0.055	0.05032022
d	0.05	0.045745654
v	0.045	0.041171089
p	0.045	0.041171089

russian		normalized
a	8.5	0.105721393
i	7.2	0.089552239
n	6.8	0.084577114
t	6.5	0.080845771
s	5.9	0.073383085
r	5.7	0.070895522
v	5.3	0.065920398
l	5.1	0.063432836
k	4.8	0.059701493
d	4.6	0.05721393
m	4.4	0.054726368
p	4.2	0.052238806
u	4	0.049751244
j	3.8	0.047263682
b	3.6	0.044776119

s korean		normalized
a	10.2	0.103658537
i	9.5	0.096544715
k	8.7	0.088414634
n	7.9	0.080284553
s	7.3	0.074186992
u	6.8	0.069105691
t	6.5	0.066056911
m	6.2	0.06300813
p	5.9	0.05995935
h	5.6	0.056910569
o	5.3	0.053861789
	5	0.050813008
e	4.8	0.048780488
t	4.5	0.045731707
	4.2	0.042682927

austrian		normalized
	0.12	0.107526882
a	0.11	0.098566308
n	0.1	0.089605735
	0.095	0.085125448
t	0.085	0.076164875
	0.08	0.071684588
s	0.075	0.067204301
l	0.07	0.062724014
m	0.065	0.058243728
k	0.06	0.053763441
d	0.058	0.051971326
	0.055	0.049283154
p	0.05	0.044802867
o	0.048	0.043010753
e	0.045	0.040322581

german (already normalized)		normalized
n	0.1449	0.1449
t	0.0976	0.0976
	0.0965	0.0965
	0.081	0.081
d	0.0778	0.0778
	0.0703	0.0703
a	0.057	0.057
	0.0537	0.0537
s	0.0534	0.0534
l	0.0527	0.0527
e:	0.049	0.049
i:	0.0482	0.0482
f	0.0424	0.0424
m	0.0389	0.0389
u	0.0365	0.0365

catalan		normalized
	1037062	0.227255834
i	413521	0.090616626
s	352822	0.077315395
n	336094	0.073649716
l	311592	0.068280488
t	283283	0.062077016
u	275223	0.060310794
a	252160	0.0552569
k	244778	0.053639251
r	202063	0.04427893
m	196537	0.043067994
z	172096	0.037712133
p	166204	0.036420994
e	162151	0.035532843
	157826	0.034585087

farsi		normalized
	0.12	0.133333333
æ	0.10	0.111111111
	0.09	0.1
n	0.08	0.088888889
	0.07	0.077777778
d	0.06	0.066666667
m	0.06	0.066666667
s	0.05	0.055555556
t	0.05	0.055555556
i	0.05	0.05
e	0.04	0.044444444
k	0.04	0.044444444
	0.04	0.038888889
l	0.03	0.033333333
x	0.03	0.033333333

norwegian		normalized
e	0.121	0.121
a	0.108	0.108
i	0.094	0.094
o	0.081	0.081
u	0.074	0.074
æ	0.067	0.067
t	0.064	0.064
n	0.061	0.061
r	0.057	0.057
s	0.054	0.054
d	0.051	0.051
l	0.047	0.047
k	0.044	0.044
m	0.04	0.04
p	0.037	0.037

japanese	count	normalized
a	153,135	0.135665971
i	119,926	0.106245321
u	108,339	0.09598012
e	87,689	0.077685789
o	83,762	0.074206766
n	75,493	0.06688106
t	69,302	0.061396305
s	65,217	0.057777305
r	60,914	0.053965174
k	58,201	0.051561663
m	55,843	0.049472654
p	50,482	0.044723215
l	48,917	0.043336744
d	46,730	0.041399228
	44,815	0.039702684

latvian		normalized
a	11.78	0.149853708
i	9.33	0.11868719
s	8.06	0.102531485
e	6.15	0.078234321
t	6.04	0.076835008
r	5.58	0.070983335
u	5.16	0.065640504
n	4.3	0.05470042
k	3.92	0.049866429
m	3.54	0.045032439
l	3.44	0.043760336
d	3.03	0.038544714
p	2.89	0.036763771
v	2.84	0.036127719
j	2.55	0.032438621

maltese	normalized (no data)
	1 0.142857143
	1 0.142857143
t	1 0.142857143
n	1 0.142857143
r	1 0.142857143
	1 0.142857143
s	1 0.142857143

macanese		normalized
a	12	0.13333333
i	10	0.11111111
u	9	0.1
	8	0.08888889
n	7	0.07777778
s	6	0.06666667
t	6	0.06666667
d	5	0.05555556
k	5	0.05555556
m	4.5	0.05
l	4	0.04444444
p	4	0.04444444
	3.5	0.03888889
e	3	0.03333333
o	3	0.03333333

estonian		normalized
a	12	0.125
i	10	0.10416667
s	9	0.09375
t	8	0.08333333
n	7	0.07291667
e	7	0.07291667
l	6	0.0625
k	6	0.0625
d	5	0.05208333
m	5	0.05208333
u	5	0.05208333
o	4.5	0.046875
	4	0.04166667
v	4	0.04166667
p	3.5	0.03645833

cantonese	count	normalized
i	18413	0.099629898
	17947	0.097108444
	15581	0.084306384
a	14035	0.075941217
	12148	0.065730951
k	19370	0.104808077
ts	13566	0.07340353
l	13354	0.072256431
j	13328	0.072115749
h	12893	0.069762031
t	11234	0.060785438
s	8759	0.047393596
m	6917	0.037426818
w	4478	0.024229766
	2791	0.01510167

nepalese		normalized
a	15	0.141509434
i	12	0.113207547
n	10	0.094339623
k	9	0.08490566
t	8	0.075471698
	7	0.066037736
m	7	0.066037736
s	6	0.056603774
d	6	0.056603774
p	5	0.047169811
u	5	0.047169811
l	4.5	0.04245283
b	4	0.037735849
e	4	0.037735849
o	3.5	0.033018868

mandarin	count	normalized
i	108084	0.157476277
a	87378	0.127308039
e	79312	0.115556035
u	42090	0.061324308
o	36786	0.053596483
j	62200	0.090624185
w	45460	0.066234332
t	45215	0.065877372
	32357	0.047143517
t	31845	0.046397543
t	28244	0.041150956
	23199	0.03380049
l	22719	0.033101139
k	21727	0.031655815
x	19735	0.02875351

portugese		normalized
a	14.00%	0.140986908
i	10.50%	0.105740181
	9.80%	0.098690836
o	9.20%	0.09264854
e	8.70%	0.087613293
s	6.50%	0.065458207
u	6.00%	0.060422961
r	5.80%	0.058408862
n	5.20%	0.052366566
t	4.80%	0.048338369
l	4.30%	0.043303122
d	4.10%	0.041289023
m	3.90%	0.039274924
k	3.50%	0.035246727
p	3.00%	0.03021148

egyptian arabic		normalized
a	14	0.133333333
i	12	0.114285714
l	10	0.095238095
m	9	0.085714286
n	8	0.076190476
	7	0.066666667
k	7	0.066666667
t	6	0.057142857
d	6	0.057142857
s	5	0.047619048
u	5	0.047619048
b	4.5	0.042857143
	4	0.038095238
h	4	0.038095238
	3.5	0.033333333

finnish		normalized
a	14	0.122807018
i	12	0.105263158
t	11	0.096491228
n	10	0.087719298
s	9	0.078947368
l	8	0.070175439
k	8	0.070175439
e	7	0.061403509
r	7	0.061403509
o	6	0.052631579
m	5	0.043859649
u	5	0.043859649
p	4.5	0.039473684
j	4	0.035087719
	3.5	0.030701754

slovak		normalized
a	14	0.11965812
o	12	0.102564103
i	11	0.094017094
e	10	0.085470085
t	9	0.076923077
n	9	0.076923077
s	8	0.068376068
l	7	0.05982906
k	7	0.05982906
r	6	0.051282051
m	6	0.051282051
d	5	0.042735043
u	5	0.042735043
p	4.5	0.038461538
	3.5	0.02991453

dutch		normalized
	20	0.209424084
a	12	0.12565445
i	10	0.104712042
o	8	0.083769634
u	7	0.073298429
i	6	0.062827225
	5	0.052356021
e	5	0.052356021
u	4	0.041884817
s	4	0.041884817
n	3.5	0.036649215
r	3	0.031413613
t	3	0.031413613
m	2.5	0.02617801
k	2.5	0.02617801

spanish		normalized
e	13	0.13
a	12.5	0.125
o	11.5	0.115
s	8	0.08
i	7.5	0.075
n	7	0.07
r	6.5	0.065
l	6	0.06
d	5.5	0.055
t	5	0.05
k	4.5	0.045
b	4	0.04
g	3.5	0.035
x	3	0.03
m	2.5	0.025