

Entropy of Sounds: Sonnets to Battle Rap

Jordan Ackerman (jackerman2@ucmerced.edu)

Cognitive and Information Sciences, 5200 Lake Rd
Merced, CA 95343

Abstract

Poetry and lyrics across cultures, from Sonnets to Rap, demonstrate an obvious human cognitive capacity for the perception and production of various multi-syllable sound patterns. Here we use entropy to measure discrete serialized representations of phones and to explore the complexity of these sound structures across genres of creative language arts. The present exploratory analysis has two main objectives. First, our aim is to broaden the scope of cognitive processes and data that are considered in statistical learning approaches to phonological learning and language acquisition. Second, we hope to provide a basis for more targeted computational and phonological investigations of these patterns. We compare the conditional entropy of sequences of phonological patterns in lyrics and find that, in general, Battle Rap and Sonnets maintain noticeably lower entropy than other genres across sequence sizes, while lyrics from Electronic music and Hip-Hop display relatively high entropy.

Keywords: Conditional Entropy; Phonology; Learning; Poetry; Music; Genres

Introduction

Background

Sound patterns that use stress, rhyme, assonance, and consonance are common in language art practices across cultures. As genres like Hip-hop, Rap, and improvisational rhyming trend toward fluent use of larger rhyming patterns than their literary cousins, many questions arise about the perception, production, and complexity of these structures.

C.E. Shannon estimated the source entropy of English characters using human guessing to be between 0.6 and 1.3 bits per [orthographic] character (Shannon, 1951). In 1965 Kolmogorov noted that while English characters (at the time) had an estimated source entropy of 1.9 bits per character, it is likely that works from artistic disciplines, such as Sonnets, would have more constraints (predictability) and therefore, should have a lower source entropy, between 1.0 to 1.2 bits per character. (Kolmogorov, 1965).

Since then many better estimates of the entropy of English have been calculated (MacKay, 2005; Cover & King, 1978), along with numerous linguistically driven information theoretic studies (Montemurro & Zanette, 2011). Work focusing on sequences of vowels and consonants has also been conducted demonstrating the interdependence of = constituent parts like vowels and consonants (Markov, 2006; Goldsmith & Riggle, 2012). Furthermore, the cognitive science of learnability has flourished, reinforcing the desire to explore realms

of human patterning in terms of perception, production, and statistical learning (Saffran, Newport, & Aslin, 1996). Finally, “it should be noted that the broader problem of measuring the information connected with creative human endeavor is of the utmost significance.” –Kolmogorov (1965)

Here, we measure the information associated with sound item sequences in lyrics and poetry as shown in Table 1. Orthographic representations of texts are collected, but instead of analyzing the serialized orthographic characters of a phrase like “The Atomic Bomb Designer”, words are serially encoded into ARPABET form (or some constituent parts: vowel, stress, consonants) to represent the phonological information of the text.

Encoding	Example
Words	THE ATOMIC BOMB DESIGNER
ARPABET	DH AH0 AH0 T AA1 M IH0...
Vowel	AH AH AA IH AA IH AY ER
Stress	0 0 1 0 1 0 1 0
Cons	DH T M K B M D Z N

Table 1: Categories of phonological items derived from orthography. ARPABET encoding, also referred to as ALL in this text, represents the full and faithful transcription from orthography to ARPABET

Purpose

Many are familiar with the rhyme and long range metrical constraints on language in the domains of poetry or iambic pentameter (Freeman, 2018). But as various multi-syllabic constraints have become common in arenas like Hip-Hop and Battle Rap, an analysis of the relative complexity of sound sequences across genre is increasingly relevant. Below is an excerpt from one participant in a rap battle, where two rappers take turns (1 to 3 minutes each), trying to ‘out rap’ each other.

*The atomic bomb designer
Dr. Robert Oppenheimer
And his squad of top advisors...*
-Bender (2012)

Notice the multi-syllabic patterning across and within lines. This, and other non-obvious or even unintentional sound patterns throughout language arts, often remain unin-

vestigated as the identification of marked patterns can be difficult, time consuming, and up to interpretation. Here we propose a targeted information theoretic approach that isolates various streams of sound symbols extracted from 14 genres of verbal art, ranging from sonnets, to musical lyrics, to a capella Battle Rap. We suspect that evidence of underlying sound sequences, as in the example above, and other repeated phonological patterns, will be detectable in their entropy.

Entropy provides a tool to measure the amount of uncertainty or surprise associated with some message. Information Theory tells us that sequences of items with lower conditional entropy (conditioned on some context i.e. n-gram) are indicative of higher predictability of the elements involved. So intuitively, analysis of vowel, stress, or consonant items here can be approximated to describe the predictability of varying sizes of sounds sequences. Reductions in entropy can be understood as 'information gain'.

Approach

A number of linguistic constraints (Semantics, Syntax, Morphology, Articulation, etc. . .) guide word choice in the normal output of natural language. But in verbal art, phonological patterning can become paramount, giving rise to a variety of perceptually interesting patterns (rhyme, assonance, repetition)

In language arts like lyrics and poetry, multi-term sound patterns do not constrain the entirety of the signal, and authors often maintain commitments to an array of other linguistic constraints. Here our interest in artistic sound patterns naturally focuses our investigation towards the predictability of multi-term sound structures within lyrics, represented as shown in Table 1. We predict that genres suspected to have the most formal constraints would contain more phonological regularities, and therefore, should have lower conditional entropy in these domains. So when measuring the conditional entropy of discrete sound items (ARPABET ALL, stress, vowels, consonants) across genres, we hypothesize that the statistical regularities of phonological patterns in a language should be captured together with whatever additional phonological predictability is specific to a given genre, artist, or work. Lower relative entropy could point to the presence of more formal constraints that exist in different streams of phonological information.

$$H(X) = \sum_{x \in A_x} p(x) \log_2 \frac{1}{p(x)}$$

Figure 1: Shannon Entropy

Assumptions

Because we are comparing genres, we consider the genre that each artist produced their work within as part of the process

that generates sounds with some particular transition probabilities. An author, or even a language itself, is often considered an approximately ergodic source, satisfying an important assumption of information theoretic analysis. Here we treat each genre of expression as an approximately ergodic source in order to explore the creative sound structures that vary between them.

Entropy measures can be useful for the comparison of the broad structural complexities of sounds in language. And conveniently, we do not need to pre-identify marked patterns by human coders or use unsupervised pattern discovery (Addanki & Wu, 2013; Reddy & Knight, 2011), although this second effort could provide a foundation for scaling further analysis.

Methods

For the scope of this study, we focus on the complexity of the basic elements of sound patterns across genres. We do this in order to identify sequence sizes (phonological structures) that may be interestingly different and merit targeted investigation. Simple information content or Shannon entropy measures (Figure 1) can be appropriate for exploring the complexity associated with individual items, or averages over individual items. This gives a framework for describing complexity based on the probability distribution of a variable X , comprised of a list of items x from an alphabet A . Some such studies were recently conducted focusing on the Shannon entropy and vocabulary of phenomena like improvised jazz (Simon, 2005, 2007) and humpback whale songs (Suzuki, Buck, & Tyack, 2006).

But here we are largely interested in signals associated with multi-term sound patterns, so we utilize conditional entropy measures (Figure 2) and multiple block sizes to explore larger and larger sequences of sound items upon which to condition the prediction of the next random variable (sound item). This is a proxy for asking not just about the predictability of individual items, but about predictability of individual items in context, i.e. sequences of items. Following from this, we are interested in comparing entropy scores by group (genre) and across sound item types (stress, vowels, consonants, ALL). We investigate a wide range of literary genres to get a sense of the differences in entropy across human creative language endeavors.

Data

We collect text from three sources, one poetry data set mined from poetryfoundation.org, song lyrics data from lyricsfreak.com, and 100 rap battles from battlerap.com. The 100 battle rap texts in question were transcribed to orthography

$$H(X|Y) = \sum_{x \in A_x} \sum_{y \in A_y} p(x, y) \log_2 \frac{1}{p(x|y)}$$

Figure 2: Conditional Entropy

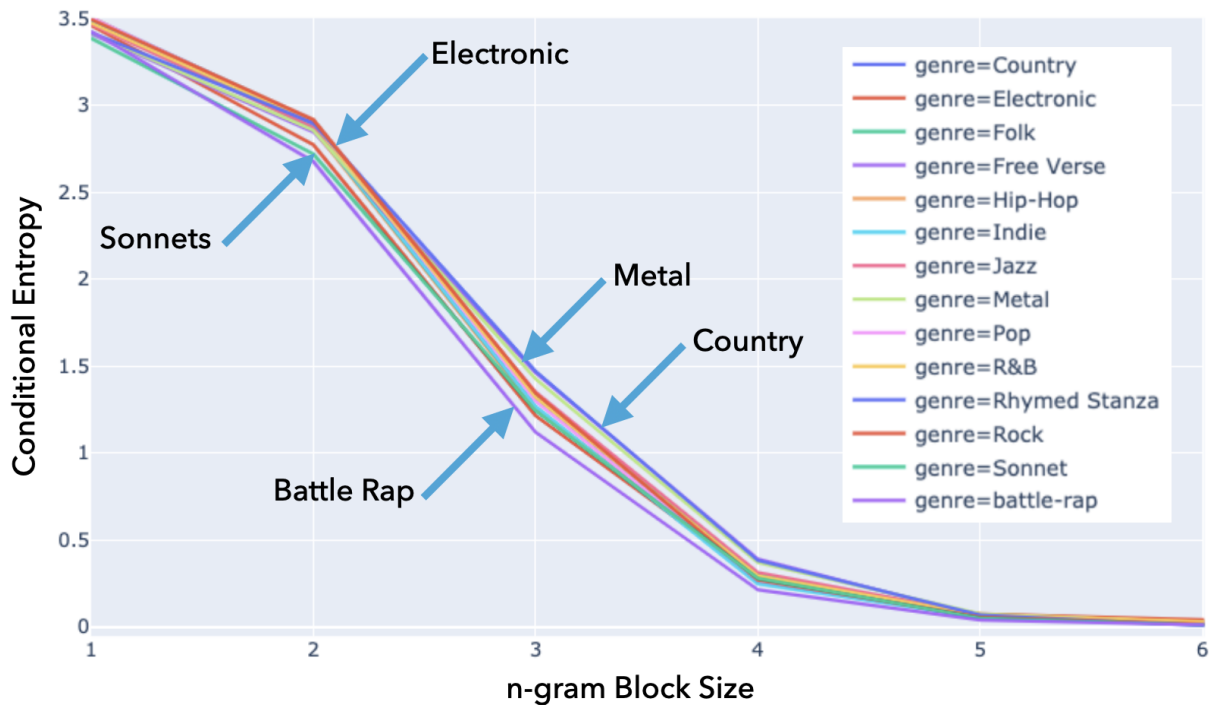


Figure 3: Conditional entropy of vowel sound items by genre: Single 3600 item sample per genre (concatenation of 36 samples of 100 sound items per genre). Block Sizes 1-6

either by battlerap.com or the performers themselves.

Processing

Sonnets and Battle Rap are our subjects of interest largely because humans can observe multi-term repeated sequences within them. However, these genres are also limiting factors in sampling for two reasons. On the one hand, very few rap battle performances have a corresponding transcript, so the number of transcribed works in this category is quite low (100 to 200 works). On the other hand, although we obtain thousands of sonnets, they have an average of only 165 syllables. To put this in perspective, most genres average 220-400 syllables per poem or song, with Hip-Hop coming in at 499 and battle rap at 4311. For the sake of reasonable comparison across genres, and with the understanding that vastly different sample lengths and alphabet sizes impact entropy scores, we report results below on the basis of data prepared as follows. We randomly select 36 works (song/poem/rap) from each of our 14 genres. From each work we extract the first 100 consecutive phone items, and repeat this for each sound item type (Vowels, Cons, Stress, ALL). We use CMU Phonetic Dictionary (Weide R. L, 1998) to transcribe orthographies to ARPABET representations. This allows for a simple comparison of their information across genres within a set sample size. Limiting ourselves at 100 phone items may not allow us to capture certain long range patterns relevant to the structure of some of these genres (Ebeling & Poeschel, 1994). But this trade-off seems acceptable, as our focus here is on the predictability of sequences of short and medium length rele-

vant for perceptually interesting or phonologically patterned language.

Analysis

We use markov models of lyrics encoded as ARPABET sound items (stress, vowels, etc...) at many block sizes to extract transitional probabilities and then calculate their conditional entropy from the equation in Figure 2. This allows us to model the complexity of sounds as we increase the sound 'context' or sequence size upon which we condition. We compare the entropy of each genre's 36 samples of 100 phone items in two ways. First, as in Figure 3, all 36 samples of 100 items for each genre are concatenated and entropy measures taken from the resulting 3600 item sample in each of the 14 genres. Alternatively, we individually take the conditional entropy of each of the 36 samples in each genre and average them to arrive at a mean conditional entropy per genre. Finally, we compare genre entropy scores in pairwise fashion using Tukey HSD pairwise tests and Jensen-Shannon distance metrics and report representative results.

Results

Due to the fact that we are comparing entropy across three relatively large dimensions, phone item type, block size, and genre, we are not able to report the complete results, but relay representative trends and summaries. Analysis was run up to block sizes of 20, but only between 3 and 10 are reported here.

Stress Sequence Entropy

In the case of stress sequences, at low n-gram size (1-3) Sonnets and Free Verse maintain the highest entropy, and at larger sequence sizes (n), they display the lowest entropy. This indicates a larger entropy reduction and therefore larger information gain in these genres than others. This is consistent with the notion that when there are larger patterns in a text than a given n-gram size can account for, entropy tends to be overestimated (Pierce, John R., 1980). This would explain why Sonnet stress entropy begins higher, but as ngram sizes increase towards the size of patterns like iambic pentameter (blocks of size 10), Sonnet stress patterns are relatively more predictable, reflected by their lower entropy. It should also be noted that while the error rates in phone encoding from CMU Phonetic Dictionary are generally not subject to increased error through concatenation of words, stress encodings are. Our phonetic encoding scheme does not take into account the change in stress (rhythm) patterns that occurs when joining words together.

Phone Sequence Entropy

Figure 3 Shows the decreasing conditional entropy by genre as block (sequence) size increases for the vowel sound items. We can see that all genres begin with high entropy at block size 1, but as blocks increase in size, their entropy is reduced, i.e. information is gained. The vertical spread between genres indicates that across some sequence sizes, certain genres have relatively more predictability and/or information gain relative to other genres. It is also interesting to note the sigmoid-like (decreasing function) pattern that all 14 genres follow as a group. The information gain (entropy reduction or $f'(x)$) from sequences of size 1 to 2 is not as great as the reduction from sequences of sizes 2 to 3 or 3 to 4. However beyond vowel sequences of 4 items, information gain slows down dramatically. This last point is unsurprising as we tend not to find strong dependencies between phonological items at large distances. But traditionally, source entropy plots over block size tend to follow a strictly decreasing function where $f'(x) < 0$ for all values of x . In other words, strictly decreasing functions have negative slopes that transition from steep to less steep, rather than from less steep, to more steep, to less steep again.

Corresponding plots of stress, vowel+stress, as well as consonant items (in a few genres) also display a similar sigmoidally decreasing function, while all other item categories (vowel-consonant, consonant-vowel-consonant, vowel-stress-cons, and ALL individual phone items) show the expected strictly decreasing source entropy functions.

Pairwise ANOVAs & Tukey HSD

We might continue counting and displaying raw entropy scores for each genre, and there is much more to consider in this arena, but here we aim to compare entropic measures in order to identify significant differences between genres.

For instance, pairwise ANOVA results, as shown in Figure 4, demonstrate a simple group-wise comparison of the aver-

age conditional entropy of 36 samples (of 100 items) from each genre. Of these 182 (14x13) pairwise tests conducted on the ALL phones items category at block size 3, only the 5 shown in Figure 4 were significant.

	sum_sq	df	F	PR(>F)
genre	0.106914	13.0	3.51787	0.000028
Residual	1.145531	490.0	NaN	NaN
	coef	t	pvalue-hs	
battle_rap-Electronic	-0.042102	-3.694279	0.021362	
Hip_Hop-Free Verse	0.045361	3.980280	0.007029	
R&B-Free Verse	0.041507	3.642093	0.025702	
battle_rap-Hip_Hop	-0.052947	-4.645925	0.000396	
battle_rap-R&B	-0.049093	-4.307738	0.001793	

Figure 4: Significant pairs from genre pairwise ANOVA - ALL Sound Items - Block Size 3

To avoid the build up of error by repeatedly performing ANOVA tests of this kind across genre and block sizes, we transition to the Tukey HSD test which allows us retain statistical soundness while conducting many pairwise significance tests.

In order to see the quality of these significant relations we can also represent this pairwise significance information in network form as in Figure 5. Nodes represent genres that passed some significance test. A pairwise similarity matrix lets us visualize the density and quality of significance relations, where edges indicate specific significant pairs and their directed edges denote the low-high entropy relation. For reference, values within nodes convey averages of conditional entropy across the 36 100-item samples in each genre. Figure 5 illustrates the significant pairwise differences from these tests, but only on vowel items of sequence size 4, instead of ALL individual item sequences of size 3, as shown in Figure 4. This approach can be used to visualize important relations across sound items types and block sizes. A fully connected graph with 14 nodes would indicate that the entropy in each genre is significantly different from every other genre. For vowel items at block size 4, Electronic and Hip-Hop have the relatively higher entropy, forming various pairwise differences with lower entropy genres, Sonnets, Battle Rap, Free Verse, and Rhymed Stanza.

Stepping back to visualize a broader picture of the range of differences across vowel sequence length, we run pairwise Tukey HSD across each genre and block size and report the counts (by genre) of pairwise significance tests passed in Figure 6. So then, column 4 of Figure 6 represents the counts of Tukey HSD pairwise tests passed and displayed in Figure 5. Each cell in Figure 6 can have a value of up to 13. A score of 13 would mean that a given genre was statistically different

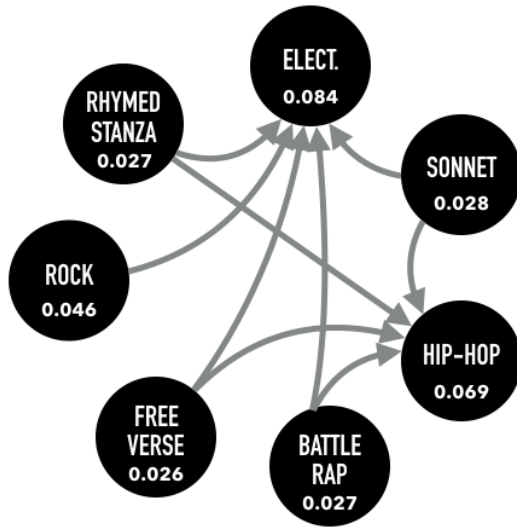


Figure 5: Network of passed mean conditional entropy significance tests on vowel sound items, Block Size 4. Edge origin indicates lower entropy of the significant pair, arrow's head indicates the higher entropy. Values in each node show means of conditional entropy across the 36 samples (each of 100 items) per genre. Connection counts same as column 4 of Figure 6.

from all other 13 genres at that block size (column).

Using this approach, we lose dimensionality about the valence of specific pairwise relations between genres, but we are better able to see patterns in the counts of significance tests passed by each genre as block size increases. This can provide information about which genres and which sequence lengths may stand out as interestingly different.

Figure 6 shows the two main groups of genres that emerge from counts of their pairwise significance tests. In general, Electronic, Free Verse, Rhymed Stanza, Sonnet, Hip-Hop, Battle Rap participate in many significant pairs. Electronic, and Hip-Hop fall on the higher entropy side while Free Verse, Rhymed Stanza, Sonnets, and Battle Rap display lower entropy. Figure 6 also makes clear that there are block sizes and genres that do not accommodate many significant cross-genre relations, notably block sizes of 3 and genres with mostly low counts. Conditional entropy differences across genres seem to be described by these two clusters. One where genres find very few significant differences to any others (Country, Folk, Indie, Jazz, Metal, Pop, RB, Rock), and one where multiple pairwise differences occur, because of either high or low entropy.

Jensen-Shannon Distance

In order to arrive at an entropy based measure of similarity between genres, we use Jensen-Shannon Divergence (JSD), Figure 7b. JSD is a symmetric measure based on the asymmetric Kullback-Leibler Divergence (KLD) shown in 7a, which allows comparison of two probability distributions P and Q.

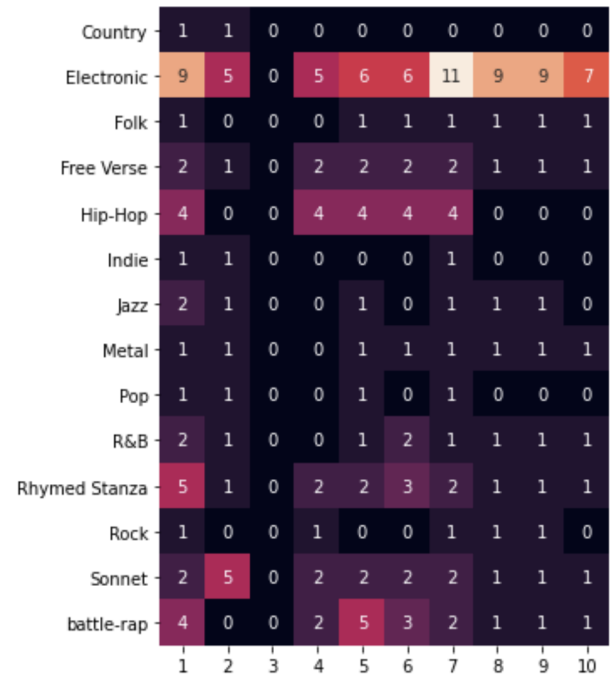


Figure 6: Counts of significant pairwise Tukey HSD tests - Vowel sound items - Counts represent number of pairwise tests passed, columns are block sizes

Symmetry is important here because we want each given entropy metric to be the same when calculating P vs Q and Q vs P (e.g. Folk vs Country and Country vs Folk). Jensen-Shannon Divergence smooths and makes symmetric the KLD where M is $(P + Q)/2$. Lastly, in 7c the square root of JSD is taken to arrive at the Jensen-Shannon distance metric.

Applying this measure to all genres in a pairwise fashion, Figure 8 shows the means of Jensen-Shannon Distance from each genre to all other 13 genres, for a given item type and block size. For example, if we individually calculate the Jensen-Shannon Distance of unigram stress items between Country and each of the other 13 genres, we get 13 distance measures, averaging them results in a score of 0.038, as shown in the top left of the figure. This process is repeated for each genre and for 1-3 grams sequences across both stress and vowels items. It provides us with an entropy based measure that allows us to notice which genres are, on average, more different from other genres.

The yellow highlighted regions of Figure 8 indicate genres with the highest average Jensen-Shannon distance from other genres. In the realm of 1-3 grams stress sequences, Free Verse, Rhymed Stanza, and Sonnets are most different from the other genres, while with respect to 1-3 gram vowel sequences, Electronic, Free Verse, Sonnets, Hip-Hop, and Battle Rap are most differentiated. However, much like in our pairwise Tukey HSD comparisons, these results demonstrate that there is a difference, and not the valence of the difference. For instance, both Electronic and Battle Rap vowel sequences

a) Kullback-Leibler Divergence

$$D_{KL}(P \parallel Q) = \sum_{x \in \mathcal{X}} P(x) \log \left(\frac{P(x)}{Q(x)} \right)$$

b) Jensen-Shannon Divergence

$$\frac{1}{2} D(P \parallel M) + \frac{1}{2} D(Q \parallel M)$$

c) Jensen-Shannon Distance

$$\sqrt{\frac{1}{2} D(P \parallel M) + \frac{1}{2} D(Q \parallel M)}$$

Figure 7: Expressions for Kullback-Leibler Divergence, Jensen-Shannon Divergence, & Jensen-Shannon Distance

have a relatively high average Jensen-Shannon distance from other genres, but for different reasons. As we saw before, while vowels in Electronic lyrics systematically display relatively high entropy, vowels from Battle Rap reflect relatively low entropy.

Discussion

The conventions of language place limits on its structure, and therefore, constrain the space of likely possible messages, resulting in lower entropy. In some genres, explicit constraints are clearly defined. This suggests that there may be genre specific sound patterning constraints that are demonstrated in the predictability of their sound transitions. Although it might have been expected that sonnets and battle rap have low entropy, it is a surprise that free verse, which traditionally does not rhyme or have a regular meter, would have similarly low entropy. These trends are also broadly mirrored when considering consonant and ALL item categories.

We cannot reasonably estimate the source entropy rates in the limit with samples this small as estimates become deterministic at low values of n . However, the simple perceptually interesting multi-syllable sound patterns we are interested in comparing across genres may be reasonably represented at these low block sizes 1-10. Even relatively large structures like iambic pentameter (10 syllables per line) seem to be, at least partially, captured in the relatively lower entropy scores displayed by sonnet stress.

Using this sampling approach it is clear that some genres do exhibit lower entropy than others, depending on the circumstance. This is all the more interesting because they employ drastically different sound patterning conventions. Sonnets often have some iambic meter and end rhyme such as 'ABAB' or 'AABB' constraints. While in Battle Rap, patterns may be large and imperfect, they do not follow a standardized metrical or rhyming structure as sonnet do. How-

genre	1-3 stress ngrams			1-3 vowel ngrams		
	1	2	3	1	2	3
Country	0.038	0.064	0.086	0.061	0.179	0.489
Electronic	0.043	0.073	0.102	0.063	0.200	0.509
Folk	0.033	0.056	0.079	0.058	0.171	0.484
Free Verse	0.075	0.118	0.151	0.091	0.207	0.495
Hip-Hop	0.034	0.060	0.083	0.075	0.197	0.500
Indie	0.038	0.065	0.089	0.067	0.183	0.491
Jazz	0.038	0.067	0.092	0.060	0.179	0.489
Metal	0.049	0.078	0.102	0.074	0.184	0.481
Pop	0.047	0.076	0.102	0.067	0.182	0.489
R&B	0.035	0.061	0.086	0.068	0.185	0.487
Rhymed Stanza	0.061	0.096	0.123	0.086	0.192	0.485
Rock	0.036	0.061	0.085	0.057	0.169	0.479
Sonnet	0.064	0.101	0.130	0.093	0.218	0.514
battle-rap	0.033	0.067	0.092	0.069	0.203	0.528

Figure 8: Mean Jensen-Shannon Distances from each genre to all other genres. Columns represent n -gram block sizes.

ever, Hip-Hop rap lyrics, which one might expect are similar to Battle Rap, consistently exhibit relatively high entropy in both vowel and consonant item categories. This could be due to Hip-Hop's use of relatively less internal rhyme (within a line) than Battle Rap, in favor of end-of-line rhyme. Finally, Electronic lyrics stand out as the genre with highest entropy in the case of vowels.

It is also noted that most of the genres that participate in fewer significantly different pairs have similar historical roots (Blues). This admittedly anecdotal observation may open a door to exploring sound entropy in terms of genealogy and the development of lyrical sound patterns diachronically.

Learnability

On the one hand, simple sound patterns (rhyme, repetition) seem to aid in memory and learning. On the other hand, when patterns get large or complex, learning to produce or perceive a pattern may require the acquisition of a specific vocabulary or grammar. This is especially true in arenas that employ multi-term conditional patterning and where the size and reliability of sequences may vary more dramatically than those set by iambic meter specifications, for example.

Future Directions

Follow-ups to this study should include use of larger datasets to compare the complexity of literary genres and individual artists. As we were limited by the volume of previously transcribed content, additional rap battles and improvisational performances should also be transcribed to enable further analysis. In the end, some of our qualitative descriptions of

these genre-based differences must be more rigorously established with larger samples and more exhaustive modeling.

It should be mentioned that the 100 rap battles considered, just like the other samples, represent written, and not improvised content. However, it would be of particular interest in terms of learnability to directly compare findings in the realm of pre-written lyrics to those in similar spontaneous or improvised verbal expression. This could help to tease apart the complexity of improvised vocabularies of creative language from those involving large amounts of human engineering (i.e. explicitly contriving and following some pattern without time constraints).

Many salient questions remain outstanding. For example, what are these specific phonological constraints, what is their vocabulary, complexity, and how are they perceived, learned, and produced? This work should be done in conjunction with various phonological, behavioral, and computational investigations. It is also important to understand how the constraints in creative sound sequences interact with other components of language, like vocabulary, morphology, or phonetic inventory, to produce complex dynamics.

As noted in previous work (Markov, 2006; Goldsmith & Riggle, 2012), there are interactions between chains of vowels, consonants, and presumably stress as well. This leads to the natural question, how do the patterns presented here hold when transitioning from simple phonological items (stress, vowels, consonants) to more complex phonological items like syllables (consonant-vowel-consonant - CVC), rhymes (vowel-consonant - VC), or specific variations such as masculine and feminine rhyme.

Conclusion

The present work has described a methodology for systematically investigating entropy of sound along three dimensions, source (genre), sequence size (n-gram blocks), and phonological items, vowels, consonants, and stress patterns. We demonstrated use of this methodology, finding that some decompositions of phonological sequences systematically display strictly decreasing functions, while some do not. In sum, we have shown that some genres have relatively more predictable sound sequences (lower entropy) than other genres at certain n-gram sizes, and other genres, notably Electronic and Hip-Hop, exhibit markedly less predictability.

References

Addanki, K., & Wu, D. (2013). Unsupervised rhyme scheme identification in hip hop lyrics using hidden Markov models. In *International Conference on Statistical Language and Speech Processing* (pp. 39–50). Springer.

Cover, T., & King, R. (1978, July). A convergent gambling estimate of the entropy of English. *IEEE Transactions on Information Theory*, 24(4), 413–421. doi: 10.1109/TIT.1978.1055912

Ebeling, W., & Poeschel, T. (1994, May). Entropy and Long range correlations in literary English. *Europhysics Letters*

(EPL), 26(4), 241–246. (arXiv: cond-mat/0204108) doi: 10.1209/0295-5075/26/4/001

Freeman, J. (2018). Vowel transitions in the sonnets of Shakespeare: an information theoretic analysis. *University of Tennessee at Chattanooga*, 26.

Goldsmith, J., & Riggle, J. (2012). Information theoretic approaches to phonological structure: the case of Finnish vowel harmony. *Natural Language Linguistic Theory*, 30(3), 859–896.

Kolmogorov, A. N. (1965). Three approaches to the quantitative definition of information. *Problemy Peredachi Informatsii*, 1(1), 3–11.

MacKay, D. J. C. (2005). *Information Theory, Inference, and Learning Algorithms*. Cambridge University Press.

Markov, A. A. (2006). An example of statistical investigation of the text Eugene Onegin concerning the connection of samples in chains. *Science in Context*, 19(4), 591–600.

Montemurro, M., & Zanette, D. (2011). Universal entropy of word ordering across linguistic families. *PLOS One*, 6(5). doi: 10.1371/journal.pone.0019875

Pierce, John R. (1980). *An introduction to information theory: Symbols, signals and noise*.

Reddy, S., & Knight, K. (2011). Unsupervised discovery of rhyme schemes. In *Proceedings of the 49th Annual Meeting of the Association for Computational Linguistics: Human Language Technologies: short papers-Volume 2* (pp. 77–82). Association for Computational Linguistics. Retrieved from <http://dl.acm.org/citation.cfm?id=2002754>

Saffran, J. R., Newport, E. L., & Aslin, R. N. (1996, August). Word Segmentation: The Role of Distributional Cues. *Journal of Memory and Language*, 35(4), 606–621. doi: 10.1006/jmla.1996.0032

Shannon, C. E. (1951). Prediction and entropy of printed English. *Bell System Technical Journal*, 30, 50–64.

Simon, S. J. (2005). A multi-dimensional entropy model of jazz improvisation for music information retrieval. *University of North Texas*, 196.

Simon, S. J. (2007). Measuring Information In Jazz Improvisation..

Suzuki, R., Buck, J. R., & Tyack, P. L. (2006, February). Information entropy of humpback whale songs. *The Journal of the Acoustical Society of America*, 119(3), 1849–1866. doi: 10.1121/1.2161827

Weide R. L. (1998). Carnegie Mellon Pronouncing Dictionary. <http://www.speech.cs.cmu.edu/cgi-bin/cmudict>, release 0.7b.