

PERCEPTUALLY INSPIRED FEATURES FOR SPEAKER LIKABILITY CLASSIFICATION

Sira Gonzalez^{1,2} and Xavier Anguera¹*

¹ Telefonica Research, 08019 Barcelona, Spain

² Imperial College London, London SW7 2BT, UK

ABSTRACT

In this paper, we present a novel approach to the classification of speaker likability, that is, a measure of how pleasant a given speaker is to listen to. Instead of blindly extracting a large number of features, we identify a small set of features which represent perceptual speech characteristics. This set of features is sent to a linear support vector machine to perform speaker likability classification. We train and evaluate the performance of our algorithm on the Interspeech 2012 speaker trait challenge database and we show that our likability classifier achieves an absolute improvement of 3.2% over the baseline classifier developed for the challenge while considerably reducing the number of features needed.

Index Terms— Speaker traits, likability, classification

1. INTRODUCTION

Speech is the main means of communication between humans. The way we are perceived by others when we speak is an important indicator of who we are, and becomes even more relevant when speech is the only signal we receive from our counterpart (e.g. in a telephone conversation). Speaker likability classification can be used to automatically assess whether a person's voice is well or badly perceived by others. Such systems become very useful, for example, as a non-biased indicator for recruitment of telephone assistants or for self evaluation purposes.

The likability of a speaker is a subjective measure and it is difficult to determine objectively what makes a voice agreeable. Rules for breathing, articulation, tone or timing to be a good orator are described in [1]. Moreover, in [2], it was found that attractive voices were associated with lack of tension, presence of confidence and favorable personality ratings. A study to determine the characteristics of dysphonic and normal voices which are important for listeners is explained in [3]. Its findings indicate that naïve listeners relied primarily on the fundamental frequency for both voice sets, but also attended to abnormality and breathiness for the pathological voices and to resonance information for normal ones. In [4], the authors tried to determine the suitability of a speaker to serve as a model for building a concatenative

text-to-speech synthesis. They spotted some male-female differences and found a positive correlation of features related to the unvoiced speech power and spectral tilt. The effect of the speaker voice in advertising was studied in [5], where a preference for faster than normal syllable speed and low pitch was demonstrated. In [6], a number of acoustic features was studied and the authors concluded that temporal features were not related to the overall rating of a speaker, but features related to the fundamental frequency dynamics and fluency gave significant correlations.

In recent years, a number of speaker likability classifiers have been implemented. In [7], the authors first focused on the consistency of subjective evaluation of speech likability and found that all sentences from the same speaker were rated similarly but the agreement between different listeners for the same speaker was low. They extracted several parameters and found that the most significant results were mostly dependent on the gender of the speaker. For the classification, around a thousand features used for emotion and affect recognition were extracted, achieving a 69.66% weighted average recall using classification trees on a small database containing 90 speakers. In [8], also by extracting a large number of speech features, specially focused on auditory characteristics, a 67.6% unweighted accuracy was reported also with classification trees on a database later reorganized for the Interspeech 2012 speakers trait challenge [9]. A voice pleasantness classifier using a database containing 77 professional female speakers with radio, theater or other vocal experience is used in [10]. Up to 179 features are extracted from areas such as clinical, quality, intelligibility, naturalness and emotion. The final architecture is a support vector machine (SVM) classifier combined with a Gaussian mixture model (GMM)/Bayes methodology in a late fusion scheme. The authors found out that the best performance was achieved with only 6 features, reporting a classification accuracy of 90.9%. However, the small amount of speakers in the database did not allow a clear training/development/test set separation and the performance was measured using 10-fold cross-validation.

Recently, the speaker trait challenge proposed at Interspeech 2012 [9] gathered a lot of attention on speaker likability classification by providing a common database and baseline algorithm to work on. In the next section, we will briefly explain this challenge and the different ways in which it was

*S. Gonzalez performed the work while at Telefonica Research

approached. In this paper, we propose a new speaker likability classifier focused on understanding the features which make speech likable. We develop our algorithm on the database provided for the challenge. Instead of extracting a large number of features, we extract a small set of perceptually inspired meaningful features and train a linear SVM with them. We evaluate our algorithm on the test set of the database and we show an absolute improvement of 3.2% with respect to the baseline results from [9] while drastically reducing the number of features from 6125 to 7.

2. INTERSPEECH SPEAKER TRAIT CHALLENGE

The speaker trait challenge proposed at Interspeech 2012 [9] provided a common ground (consisting on a database and a baseline algorithm) for ‘perceived’ speaker traits: personality, likability and intelligibility of pathologic speakers. The speaker likability database consisted on 800 speakers from a database originally recorded to study automatic age and gender recognition from telephone speech. The database was divided into a training set, a development set and a test set, containing 394, 178 and 228 speakers respectively. The features used for the classification, which included several functionals applied to them, were extracted using the openSMILE toolkit [11] and can be classified into three sets: energy, spectral and voicing features. The baseline classification accuracy for speaker likability using a total of 6125 features was 55.9% using linear SVM and 59.0% using random forests [9].

Many approaches to the challenge focused on the method used for classification and not so much on the feature set, using the standard features provided by the organizers. Stricted Boltzman machine and deep belief networks [12], anchor models [13] or a new machine learning algorithm based on Gaussian processes [14] were used for the classification. The highest classification accuracy, 64.0%, was achieved by combining the stricted Boltzman machine and deep belief networks.

Among the participants focusing on feature selection, [15] used the Fisher information metric and later on a genetic algorithm. Two methods for feature selection were implemented by [16]: one based on classification and one based on statistical dependence. In addition to all the features from the baseline approach, [17] used other four kinds of features: basic prosodic features, prosodic polynomial coefficients, spectral features and shifted delta cepstrum features. In [18], the authors achieved the best performance in the test set, with a classification accuracy of 65.8%, by combining pitch and intonation features with spectral features, using more than 10,000 features.

In [19], the authors tried to understand what makes some voices more likable than others. They subjectively evaluated six characteristics of the speech, finding that likable speakers exhibit almost no perceivable accent, command style or disfluencies.

Although the subjectivity of speech likability makes its classification challenging, the low overall performance of the likability classifiers at the challenge may also be related to the short duration of the utterances, which makes the extraction of reliable features difficult.

3. SPEECH FEATURES

As we have seen in the introduction, many approaches extract a large number of features to perform speaker likability classification and, sometimes, the initial set is reduced by feature selection. Although the number of features may drop drastically, the final selection may not be directly correlated to perceptual characteristics of the speech. In this paper, we focus on the extraction of a small but meaningful set of features which we consider related to the perceptual speaker likability. In some cases, basic functionals are applied to them to extract the underlying information.

The features are evaluated on the training set of the database to observe their correlation with the speaker likability. If positive results are found, the distribution on the development set is also used to verify this observation.

3.1. Active speech level

Active speech level provides a measure of the loudness of the speech. Under the assumption that all recordings have been made in the same conditions, differences in active speech level is only due to specific speakers.

By evaluating the active speech level of likable and non-likable speakers on the training set of the database using the ITU-T P.56 recommendation [20, 21], we observe a preference for moderate loudness, as seen in Fig. 1(a).

3.2. Variation of the speech power

How the speech power changes over time provides information about the speaker’s accentuation and emphasis. We investigate the total power change over time and the power change in five mel-spaced frequency bands, centered at 261, 621, 1114, 17913 and 2722 Hz. Two approaches are taken into account: power difference in consecutive frames and power derivative using neighboring frames. Frames of 90 ms duration with an inter-frame increment of 10 ms were used.

The most significant results on the training set are obtained for the standard deviation of the total power difference between consecutive frames, Fig. 1(b), and the standard deviation of the power derivative over 100 ms of the frequency band centered at 2.72 kHz, Fig. 1(c). While in the first case a higher standard deviation is preferred, i.e. higher range of variation, in the second case lower standard deviation is more likable. At higher frequencies most power changes are due to fricative sounds, and the lower standard deviation preference

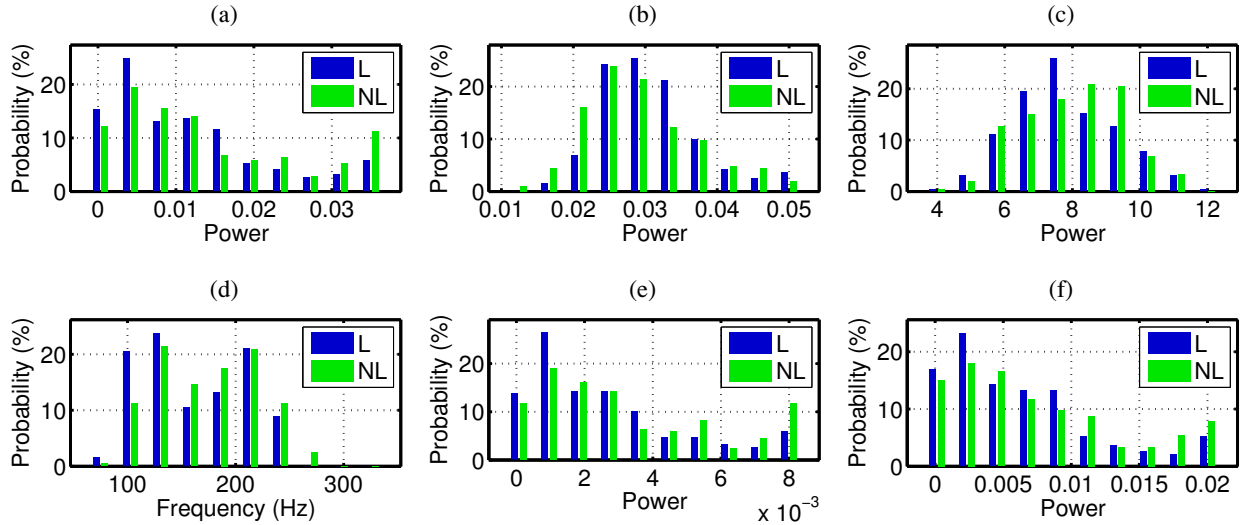


Fig. 1. Histogram for likable (blue) and non-likable (green) voices on the training set [9] for: (a) speech active level, (b) standard deviation of the total power derivative, (c) standard deviation of the power derivative over 100 ms of the frequency band centered at 2.72 kHz, (d) median pitch, (e) mean speech variance, and (f) speech variance standard deviation.

may indicate a predilection for smooth fricative transitions or for lower fricative power.

3.3. Fundamental frequency

The fundamental frequency and its functionals have been shown to be an important characteristic in different studies [3, 5, 6]. Therefore, we decide to evaluate its impact on speech likability classification by analyzing the median (for robustness to the fundamental frequency estimator errors), the standard deviation, the mean of the first derivative, and the standard deviation of the first derivative. The PEFAC algorithm is used for pitch extraction [22, 21]. We observe that only the median pitch, Fig. 1(d), provides a different distribution between likable and non-likable speakers, illustrating a preference for low fundamental frequencies, related to male speakers. This matches the results showed in [7], which found that higher likability rates are given to male speakers with low fundamental frequency.

3.4. Speech variance

The mean speech variance, a measure of how spread out the speech signal is over a specific period of time, has been previously used for speech quality [23] and intelligibility [24] estimation. It is calculated by segmenting the speech into overlapping frames and measuring the variance in each frame. The frame length was set to 20 ms with 5 ms overlap. We calculate the mean and the standard deviation of this variance over the whole utterance.

The results on the training set show that this feature can provide useful information for the classification, a lower mean, Fig 1(e), and standard deviation, Fig. 1(f), being preferable. This means that it is more likable for the speech amplitude to be uniformly distributed across each time frame.

3.5. Silent segments per second

By measuring the average silent time per second of the speaker we can infer information regarding the amount of pauses made by a speaker when speaking. A speech frame is considered to be silent if its power falls below a specific threshold. Initially, we normalized the speech using the ITU-T P.56 recommendation [20, 21] and we divided the speech into frames with a duration of 90 ms and with an inter-frame increment of 10 ms. The power threshold was empirically set to $7 \cdot 10^{-4}$.

We find that a lower silent time per second is more agreeable, Fig. 2(a). By looking at the spectrogram of likable and non-likable utterances, we observe that the main difference is that while likable speakers tend to keep some power in word transitions, least likable speakers do not usually have this cadence.

3.6. Correlation between neighboring frames

The smoothness of a speech utterance can be seen as the cross-correlation between neighboring frequencies. We decided to measure the peak of the cross-correlation for frames with a duration of 20 ms separated from 5 ms to 30 ms.

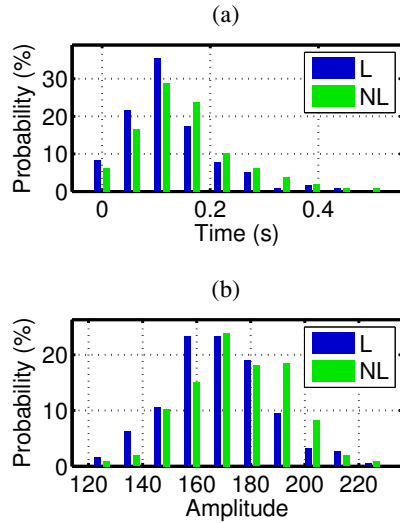


Fig. 2. Histogram for likable (blue) and non-likable (green) voices on the training set [9] for: (a) silent time per second, and (b) mean cross-correlation peak for consecutive frames

In Fig. 2(b) we can observe the histogram for likable and non-likable speakers. Opposite to what we were expecting, lower correlation is preferable between consecutive frames. An hypothesis is that while likable speech has smooth and longer transitions, therefore constantly changing its waveform, less likable speech remains longer in a phone (high inter-frame correlation) and then changes abruptly.

4. LIKABILITY CLASSIFIER

To select the best features for the classifier, we identify the features which provided a good separation between likable and non-likable speakers on the training set and evaluate that these distributions are maintained in the development set. Table 1 shows the seven selected features. As no difference

Table 1. Likability classifier features

1	Speech active level
2	Standard deviation of the power derivative
3	Standard deviation of the power derivative over 100 ms of the frequency band centered at 2.72 kHz
4	Mean speech variance
5	Speech variance standard deviation
6	Silent time per second
7	Mean cross-correlation peak amplitude for consecutive frames

Table 2. Likability classification accuracy results on the test set of the database

	UA(%)	# features	Classifier
Baseline [9]	59.0	6125	Random Forests
Montancie et al. [18]	65.8	> 10000	linear SVM
Proposed	62.2	7	linear SVM

is made between male and female speakers, the fundamental frequency and its functionals are not taken into account in our algorithm.

The classifier used was a linear SVM with sequential minimal optimization. The training was done on the training set of the database and the soft margin was chosen to obtain the best performance in the development set. The best unweighted average (UA) recall in the development set is 61.6% for a soft margin of 4.92. For the same margin, the accuracy on the training set is 63.2%.

For the results on the test set, we trained a linear SVM on the training and development sets using the soft margin optimized on the development set. The results obtained for our classifier on the test set of the database are shown in Table 2, together with the baseline performance and the best performance obtained at the Interspeech challenge. We can observe how our classifier is able to have an absolute increment of 3.2% over the baseline performance while reducing the number the number of features from 6125 to 7. The best result obtained in the challenge, 65.8%, was achieved by the fusion of three classifiers, each containing a different set of features, some of them in common. Therefore, it is difficult to determine how many features were actually used to achieve the final classification rate. The lowest limit shown in the table is taken from the individual classifier with the maximum number of features, 10342. Although the performance of this algorithm provides a better speaker likability classification, the high number of features used makes the identification of specific factors that make speech likable difficult. The seven selected features of our algorithm, however, provide information about specific aspects of speech which makes it more or less likable while improving the baseline classification.

5. CONCLUSIONS

In this paper, we have extracted and evaluated perceptually inspired features for speaker likability classification. We have shown that these features can provide an absolute improvement of 3.2% with respect to the baseline results of the Interspeech 2012 speaker traits challenge while decreasing the number of features from 6125 to 7. Moreover, the selection of this small but meaningful set of features can also provide an insight of speech characteristics that could be modified to make speech more likable.

6. REFERENCES

- [1] J. E. Frobisher, *Acting and oratory: designed for public speakers, teachers, actors, etc.* New York: College of Oratory and Acting, 1879.
- [2] M. Zuckerman, H. Hodgins, and K. Miyake, "The vocal attractiveness stereotype: Replication and elaboration," *Journal of Nonverbal Behavior*, vol. 14, pp. 97–112, 1990.
- [3] J. Kreiman, B. Gerratt, and K. Precoda, "Listener experience and perception of voice quality," *Journal of Speech, Language and Hearing Research*, vol. 33, no. 1, pp. 103–115, 1990.
- [4] A. Syrdal, A. Conkie, Y. Stylianou *et al.*, "Exploration of acoustic correlates in speaker selection for concatenative synthesis," in *Proceedings of ICSLP*, vol. 98, 1998, pp. 2743–2746.
- [5] A. Chattopadhyay, D. W. Dahl, R. J. B. Ritchie, and K. N. Shahin, "Hearing voices: The impact of announcer speech characteristics on consumer response to broadcast advertising," *Journal of Consumer Psychology*, vol. 13, no. 3, pp. 198–204, 2003.
- [6] E. Strangert and J. Gustafson, "What makes a good speaker? Subject ratings, acoustic measurements and perceptual evaluations," in *Proc. Interspeech Conf.*, Brisbane, Australia, 2008, pp. 1688–1692.
- [7] B. Weiss and F. Burkhardt, "Voice attributes affecting likability perception," in *Proc. Interspeech Conf.*, Makuhari, Japan, Sep. 2010, pp. 2014–2014.
- [8] F. Burkhardt, B. Schuller, B. Weiss, and F. Weninger, "'Would you buy a car from me?'—On the likability of telephone voices," in *Proc. Interspeech Conf.*, 2011, pp. 1557–1560.
- [9] B. Schuller, S. Steidl, A. Batliner, E. Nöth, A. Vinciarelli, F. Burkhardt, R. van Son, F. Weninger, F. Eyben, T. Bocklet *et al.*, "The interspeech 2012 speaker trait challenge," in *Proc. Interspeech Conf.*, Portland, USA, Sep. 2012.
- [10] L. Pinto-Coelho, D. Braga, M. Sales-Dias, and C. Garcia-Mateo, "On the development of an automatic voice pleasantness classification and intensity estimation system," *Computer Speech & Language*, 2012.
- [11] F. Eyben, M. Wöllmer, and B. Schuller, "Opensmile: the munich versatile and fast open-source audio feature extractor," in *Proceedings of the international conference on Multimedia*. ACM, 2010, pp. 1459–1462.
- [12] R. Bruecker and B. Schuller, "Likability classification - a not so deep neural network approach," in *Proc. Interspeech Conf.*, Portland, USA, Sep. 2012.
- [13] Y. Attabi and P. Dumouchel, "Anchor models and WCCN normalization for speaker trait classification," in *Proc. Interspeech Conf.*, Portland, USA, Sep. 2012.
- [14] D. Lu and F. Sha, "Predicting likability of speakers with Gaussian processes," in *Proc. Interspeech Conf.*, Portland, USA, Sep. 2012.
- [15] D. Wu, "Genetic algorithm based feature selection for speaker trait classification," in *Proc. Interspeech Conf.*, Portland, USA, Sep. 2012.
- [16] J. Pohjalainen, S. Kadiogu, and O. Rasanen, "Feature selection for speaker traits," in *Proc. Interspeech Conf.*, Portland, USA, Sep. 2012.
- [17] M. H. Sanchez, A. Lawson, D. Vergyri, and H. Bratt, "Multi-system fusion of extended context prosodic and cepstral features for paralinguistic speaker trait classification," in *Proc. Interspeech Conf.*, Portland, USA, Sep. 2012.
- [18] C. Montacie and M. Caraty, "Pitch and intonation contribution to speakers's trait classification," in *Proc. Interspeech Conf.*, Portland, USA, Sep. 2012.
- [19] B. Weiss and F. Burkhardt, "Is not bad good enough? Aspects of unknown voices' likability," in *Proc. Interspeech Conf.*, Portland, USA, Sep. 2012.
- [20] ITU-T, *Objective Measurement of Active Speech Level*, International Telecommunications Union (ITU-T) Recommendation P.56, Mar. 1993.
- [21] D. M. Brookes, "VOICEBOX: A speech processing toolbox for MATLAB," <http://www.ee.imperial.ac.uk/hp/staff/dmb/voicebox/voicebox.html>, 1997.
- [22] S. Gonzalez and M. Brookes, "A pitch estimation filter robust to high levels of noise (PEFAC)," in *Proc. European Signal Processing Conf. (EUSIPCO)*, Barcelona, Aug. 2011.
- [23] V. Grancharov, D. Y. Zhao, J. Lindblom, and W. B. Kleijn, "Low-complexity, nonintrusive speech quality assessment," *IEEE Trans. Audio, Speech, Lang. Process.*, vol. 14, no. 6, pp. 1948–1956, Nov. 2006.
- [24] D. Sharma, G. Hilkuysen, N. D. Gaubitch, P. A. Naylor, M. Brookes, and M. Huckvale, "Data driven method for non-intrusive speech intelligibility estimation," in *Proc. European Signal Processing Conf. (EUSIPCO)*, Denmark, Aug. 2010.