

Title	Two-stage dimensional emotion recognition by fusing predictions of acoustic and text networks using SVM
Author(s)	Atmaja, Bagus Tris; Akagi, Masato
Citation	Speech Communication, 126: 9-21
Issue Date	2020-11-19
Type	Journal Article
Text version	author
URL	http://hdl.handle.net/10119/18116
Rights	Copyright (C)2020, Elsevier. Licensed under the Creative Commons Attribution-NonCommercial-NoDerivatives 4.0 International license (CC BY-NC-ND 4.0). [http://creativecommons.org/licenses/by-nc-nd/4.0/] NOTICE: This is the author's version of a work accepted for publication by Elsevier. Changes resulting from the publishing process, including peer review, editing, corrections, structural formatting and other quality control mechanisms, may not be reflected in this document. Changes may have been made to this work since it was submitted for publication. A definitive version was subsequently published in Bagus Tris Atmaja, Masato Akagi, and Peter Birkholz, Speech Communication, 126, 2020, 9-21, https://doi.org/10.1016/j.specom.2020.11.003
Description	

Two-stage dimensional emotion recognition by fusing predictions of acoustic and text networks using SVM

Bagus Tris Atmaja^{a,b,*}, Masato Akagi^b

^a*Japan Advanced Institute of Science and Technology, 1-1 Asahidai, Nomi, Ishikawa
923-1292, Japan*

^b*Sepuluh Nopember Insitute of Technology, Sukolilo, Surabaya 60111, Indonesia*

Abstract

Automatic speech emotion recognition (SER) by a computer is a critical component for more natural human-machine interaction. As in human-human interaction, the capability to perceive emotion correctly is essential to taking further steps in a particular situation. One issue in SER is whether it is necessary to combine acoustic features with other data such as facial expressions, text, and motion capture. This research proposes to combine acoustic and text information by applying a late-fusion approach consisting of two steps. First, acoustic and text features are trained separately in deep learning systems. Second, the prediction results from the deep learning systems are fed into a support vector machine (SVM) to predict the final regression score. Furthermore, the task in this research is dimensional emotion modeling, because it can enable deeper analysis of affective states. Experimental results show that this two-stage, late-fusion approach, obtains higher performance than that of any one-stage processing, with a linear correlation from one-stage to two-stage processing. This late-fusion approach improves previous early fusion result measured in concordance correlation coefficients score.

Keywords: automatic speech emotion recognition, affective computing, late fusion, multimodal fusion, dimensional emotion

*Corresponding author, *E-mail address:* bagus@jaist.ac.jp

1. Introduction

Understanding human emotion is important for responding properly in a particular situation for both human-human communication and future machine-human communication. Emotion can be recognized from many modalities: facial expressions, speech, and motion of body parts. In the absence of visual features, speech is the only way to recognize emotion, as in the case of a telephone call or a call-center application (Petrushin, 1999). By identifying caller emotion automatically from a system, appropriate feedback can be applied quickly and precisely.

Speech is a modality in which both acoustic and verbal information can be extracted to recognize human emotion. Unfortunately, most speech emotion recognition (SER) systems use only acoustic features for predicting categorical emotion. In contrast, this research proposes to use both acoustic and text features to improve dimensional SER performance. Text can be extracted from speech, and it may contribute to emotion recognition. For example, an interlocutor can perceive emotion not only from prosodic information but also from semantics. Grice (2002) stated in his implicature theory that what is implied derives from what is said. For example, if someone says that he is angry but looks happy, then the implication is that he is indeed angry. Hence, it is necessary to use linguistic information to determine expressed emotion from speech. A fusion of acoustic and linguistic information from speech is viable since (spoken) text can be obtained from speech-to-text technology. This bimodal features fusion strategy may improve the performance of SER over acoustic-only SER.

Besides the categorical approach, emotion can also be analyzed via a dimensional approach. In dimensional emotion, affective states are lines in a continuous space. Some researchers have used a two-dimensional (2D) space comprising valance (positive or negative) and arousal (excited or apathetic). Other researchers have proposed a 3D emotional space by adding either dominance (degree of power over emotion) or liking/disliking. Although it is rare, a 4D emotional space has also been studied by adding expectancy or naturalness.

While some researchers, e.g., Russell (1980), argue that a 2D emotion model is enough to characterize all categorical emotions, in this research, we choose a 3D emotion model with valence, arousal, and dominance as the emotion dimensions/attributes.

35 Darwin argued that the biological category of a species, like emotion categories, does not have an essence due to the high variability of individuals (Charles et al., 1872). Mehrabian and Russell (1974) developed a pleasure, arousal, and dominance (PAD) model to assess environmental perception, experience, and psychological responses, as an alternative to categorical emotion.
40 The latter, also called as dimensional emotion, may represent human emotion better than categorical emotion. This dimensional emotion view is also known as the circumplex model of affect, and the pleasure dimension is often replaced by valence for the same meaning (the VAD model). Although most research used the 2D model (valence and arousal), recent research shows four dimensions
45 needed to represent the meaning of emotion words (Fontaine et al., 2017). However, current datasets lack from the availability of the fourth dimension label (i.e., expectancy). We evaluate the VAD emotion model since the datasets also present the labels in 3D space.

Deep neural networks (DNN) have recently gained more interest in modeling
50 human cognitive processing for several tasks. Fayek et al. (2017) evaluated some DNN architectures for categorical SER. They found fully-connected (FC) networks and recurrent neural network (RNN) worked well for SER task using acoustic features only. In neuropsychological science, the neural mechanism that integrates acoustic (verbal) and linguistic (non-verbal) information remains
55 unclear (Berckmoes and Vingerhoets, 2004). The paper also stated that "the various parameters of prosody [acoustics] are processed separately in specific brain areas" while no information is given for linguistic processing. In this understanding, separation of acoustic and linguistic/text processing is better modeled by a late fusion than an early fusion. This research makes use of
60 support vector machine (SVM) for a late-fusion prediction from DNN based acoustic and linguistic emotion recognitions. The small remaining test data

after used by DNNs is a reason to use SVM over DNN.

This study aims to evaluate the combination of acoustic and text features to improve the performance of dimensional automatic SER by using two-stage processing. Current research on pattern recognition has also shown that the use of multimodal features from audio, visual, and motion-capture data increases performance as compared to using a single modality (Hu and Flaxman, 2018; Yoon et al., 2018; Tripathi and Beigi, 2018). Meanwhile, research on big data has revealed that the use of more data will improve performance for results from the same algorithm (Halevy et al., 2009). By using both acoustic and text features, SER should obtain improved performance over acoustic-only and text-only recognition. This assumption is also motivated by the fact that human emotion perception uses multimodal sensing, peculiarly verbal and non-verbal information. Many technologies, such as human-robot interaction, can potentially benefit from such improvement in emotion recognition.

The main contributions of this study then are: (1) a proposal of two-stage processing for dimensional emotion recognition from acoustic and text features using LSTM and SVM, and a comparison of the results with unimodal results and another fusion method on the same metric and dataset scenario; (2) an evaluation of different acoustic and text features to find the best pair of acoustic-text pair based on evaluated features, including a frame-based acoustic feature and utterance-based statistical functions with and without silent pause features; (3) evaluation of speaker-dependent vs. speaker-independent scenarios in dimensional speech emotion recognition from text features; and (4) evaluation of using text features on a dataset that originally contains target sentences but removed to avoid the effect of these target sentences.

The rest of this paper is organized as follows. “Related work” reviews closely related work to this research including the difference of this study from previous research, “Datasets and features” outlines the datasets and feature sets used in this research, “Two-stage bimodal emotion recognition” explains the method to achieve the results, “Results and discussion” shows the results and its discussion, and finally “Conclusions” concludes this study and proposes future work.

2. Related work

Speech emotion recognition (SER) began to be seriously researched as part
95 of human-computer interaction with the work by e.g., Kleine-cosack (2006). The
amount of research has grown as datasets have become publicly available, includ-
ing the Berlin EMO-DB, IEMOCAP, MSP-IMPROV, and RAVDESS datasets.
To enable the analysis and comparison with previous research, we include the
following literature reviews of related work. We focus on comparing previous
100 work that used the same or similar datasets as this work does (specifically,
IEMOCAP, MSP-IMPROV, or both), and especially on research that focused
on dimensional rather than categorical emotion. While the focus here is on
bimodal emotion recognition using both acoustic and text data, some work on
speech-only or text-only emotion recognition is briefly described.

105 2.1. Acoustic emotion recognition

Recognition of emotion within speech signals has been actively developed
since the success of recognizing emotion via facial expressions. From categor-
ical emotion detection, the paradigm of SER has shifted to predicting degrees
of emotion attributes, or dimensional emotion. One of the earliest papers on
110 (categorical) SER (Petrushin, 1999) explored how well humans and computers
recognize emotion in speech. Since then, research on categorical emotion recog-
nition has grown following the development of affective research in psychology.

Jin and Wang (2005) reported a first trial on SER in categorical and two-
dimensional (2D) spaces. They found that acoustic features are helpful in de-
115 scribing and distinguishing emotion through the concept of emotion modeling
(2D space). In 2009, Giannakopoulos et al. (2009) re-investigated the associa-
tion of speech signals with an emotion wheel (continuous space). They proposed
a method to estimate the degrees of valence and arousal. Their method, includ-
ing a proposed feature set, could estimate both valence and arousal, with error
120 close to that of average human annotation. Grimm et al. (2007) used a fuzzy-
logic estimator and a rule base derived from acoustic features in speech, such as

pitch, energy, speaking rate, and spectral characteristics, to describe emotion primitives (valence, arousal, and dominance). They obtained moderate to high correlation ($0.42 < r < 0.85$) between their method and human annotation.

125 Using the IEMOCAP dataset, Parthasarathy and Busso tried to train a neural network system to predict valence, arousal, and dominance simultaneously (Parthasarathy and Busso, 2017). They proposed a multitask learning (MTL) system based on the mean squared error (MSE) to balance prediction of the three emotion dimensions. They found that, by combining a shared layer and
130 an independent layer, the MTL system’s best performance exceeded that of the traditional single-task learning (STL) method.

AbdelWahab and Busso (2018) proposed using a domain-adversarial neural network (DANN) to solve the problem of mismatch between training and test data in dimensional SER. Using the DANN, they obtained performance that
135 significantly improved that of a source-trained DNN. Thus, they addressed the importance of minimizing the mismatch between the source (training) and target (test) data. Furthermore, using the DANN showed that creating a flexible, discriminant feature representation can reduce the gap in the feature space between the source and target domains.

140 Some of the above results on dimensional SER showed that recognizing valence is more difficult than recognizing arousal. To overcome this issue, Sridhar et al. (2018) used higher regularization (dropout) for valence than for the other dimensions when training SER through a DNN. Their system analysis showed that higher dropout is needed for predicting valence. By using higher regular-
145 ization, models could identify more general acoustic patterns that were observed across speakers. Elbarougy and Akagi (2014) used a three-layer model based on human perception for the same purpose. Li and Akagi (2019) improved on that work by combining acoustic features for multilingual emotion recognition.

Although some improvements have been achieved, El Ayadi et al. (2011)
150 addressed the SER issue of whether it suffices to use acoustic features for modeling emotions or it is necessary to combine them with other types of features, such as linguistic discourse information or facial features. Text features can

be obtained through automatic speech recognition (ASR) and may be helpful in significantly improving SER performance. In particular, text features are
155 expected to improve the performance of valence recognition, for which acoustic features have typically failed to achieve high performance. Moreover, text features are commonly used for sentiment analysis, which is similar to valence prediction.

2.2. Text emotion recognition

160 As mentioned above, one area of research on text processing focuses on sentiment analysis. This area is closely related to recognizing valence, i.e., the polarity or semantic orientation of an event, object, or situation (Jurafsky and Martin, 2017). Although the early research sought to recognize sentiment in text, extension to recognize categorical and dimensional emotion has been
165 attempted in recent years. As in other areas of research on pattern recognition, some researchers in text processing have used unsupervised learning to detect emotion in text (Mäntylä et al., 2016; Mohammad, 2016), while others have used supervised learning based on machine learning (Alm et al., 2005; Yang et al., 2016).

170 Alm et al. (2005) used a bag-of-words (BoW) model and other text features from text datasets to predict emotion within those datasets. Using a multiclass linear classifier, they obtained encouraging results that suggest a potential direction for future research. Their proposed method could predict basic emotion from text with accuracy close to 70%.

175 Kim et al. (2010) used unsupervised learning to predict categorical emotion from three different datasets: SemEval, ISEAR, and Fairy Tales. Using three different techniques, they found that the best performance was achieved with categorical classification based on non-negative matrix factorization (NMF). Atmaja and Akagi (2019) used a deep-learning-based classification model and
180 improved the precision, recall, and F-score results for the ISEAR dataset from 0.528, 0.417, and 0.372 to 0.56, 0.54, and 0.54, respectively. They showed the effectiveness of the deep-learning-based method for categorical emotion recogni-

tion on a larger dataset, while on a smaller dataset, the unsupervised approach achieved better results. Apart from categorical emotion recognition, Atmaja and Akagi (2019) also performed dimensional emotion recognition on the same dataset used for the categorical task. Similar to the categorical task, the results showed fewer errors when the size of the training set was increased.

Mäntylä et al. (2016) used emotion words from an affective lexicon to mine valence, arousal, and dominance in text communication. Specifically, they used text communication data from a software development situation, including issues and comments captured through issue repository technology. They used the measure of valence, arousal, and dominance (VAD) to detect the productivity and burnout of the software developers. The results showed that increased emotions in terms of VAD correlated with increased productivity. Their results also complemented previous results showing that VAD can be measured from text, though at first only the sentiment (valence) was used in text processing.

Research on text emotion recognition has usually used written language (from chats, Twitter, forum threads, etc.), which differs from spoken language. Also, most such work on text processing has detected only the valence, i.e., only one emotion dimension. Because speech transcription converts spoken language to a written form, it should contain more emotion information than written plain text does. Evaluation on the other emotion dimensions (arousal and dominance) is also necessary to determine the impact on those dimensions, along with evaluation on the combination with acoustic features for that purpose.

2.3. Bimodal emotion recognition

Using bimodal or multimodal features for emotion recognition is not new. Among many modalities, audio and visual features are the most used for extracting emotion information. When only speech is conveyed, however, two types of information can be extracted: acoustic and text features. Among many research papers, the reports by Eyben et al. (2010), Karadogan and Larsen (2012), Ye and Fan (2014), Jin et al. (2015), Aldeneh et al. (2017), Yoon et al. (2018), Atmaja et al. (2019), and Zhang et al. (2019) are the most related to this paper.

Eyben et al. (2010) proposed an online method to detect not only valence and arousal but also the time when those emotion attributes are detected. They used
 215 a recurrent neural network (RNN) based on long short-term memory (LSTM)
 to recognize a framewise valence-arousal continuum with time. By adding a
 keyword spotter, they were able to improve the performance by using regression
 analysis. The results were measured in Pearson correlation coefficient (PCC).
 They also found that keywords like “again,” “angry,” “assertive,” and “very”
 220 were related to activation, while typical keywords correlated to valence were
 “good,” “great,” “lovely,” and “totally.” Similar to that idea, Karadogan and
 Larsen (2012) used affective words from Affective Norms for English Words
 (ANEW) to determine a valence-arousal value and combine it with a result
 from acoustic features. The latter paper also obtained similar improvement
 225 over using a single modality.

Ye and Fan (2014) used bimodal features from acoustic and text information
 to recognize emotion within speech. The acoustic features were trained in two
 parallel classifiers: an SVM and a backpropagation network. The text features
 were trained in two serial classifiers, which were both Naive Bayes classifiers.
 230 The second classifier acted as a filter for unreliable parts from the first classi-
 fier. Decision-level fusion (late fusion) was then implemented by combining the
 acoustic and text features with tree-weighting factors for the SVM, backpropa-
 gation network, and text classifiers. The resulting fusion method obtained 93%
 accuracy, as compared to 83% from the acoustic features only and 89% from the
 235 text features only. The task was categorical emotion detection from a Chinese
 database. Similar to that approach for a categorical task, Jin et al. (2015) used
 the IEMOCAP dataset to test combinations of acoustic and text features for
 SER. The novelty of their method was the use of an emotion vector for lexical
 features, which improved the accuracy in four-class emotion recognition from
 240 53.5% (acoustic) and 57.4% (text) to 69.2% (acoustic + text).

Aldeneh et al. (2017) used acoustic and lexical features to detect the degree
 of valence from speech. They used 40 mel-filterbanks (MFBs) as acoustic fea-
 tures and word vectors as text features. Continuous valence values were then

converted to three categorical classes: negative, neutral, and positive. Using
245 that approach, they improved the weighted accuracy from 64.5% (text) and
58.9% (acoustic) to 69.2% (acoustic + text).

Yoon et al. (2018) used audio and text networks to predict emotion classes
from the IEMOCAP dataset. Both networks used RNNs with inputs of mel-
frequency cepstral coefficients (MFCCs) for audio and word vectors for text.
250 The proposed multimodal dual recurrent encoder (MDRE) improved on the
single-modality RNNs from 54.6% (audio) and 63.5% (text) to 71.8% (audio +
text). Atmaja et al. (2019) obtained a better result by using 34 acoustic features
after silence removal and combining them with word embeddings. With LSTM
used for the text and dense networks for speech, the latter paper obtained an
255 accuracy of 75.49% on the same dataset and task.

Instead of using lexical features, Zhang et al. (2019) used phonemes and
combined them with acoustic features to recognize valence in speech. They
used 39 unique phonemes from the IEMOCAP and MSP-IMPROV datasets
and a 40-dimensional log-scale MFB energy for the acoustic features. Using a
260 scaled version of valence, converted from a 5-point scale to three categorical
classes, they showed that their multistage fusion model outperformed all other
models on both IEMOCAP and MSP-IMPROV.

2.4. Multimodal emotion recognition

The term multimodal used in this subsection refers to the use of three or more
265 modalities for emotion recognition. These modalities are usually visual, audio,
text, gesture, and eye gaze. In a recent survey (Yang and Hirschberg, 2018),
it is confirmed that the use of multimodal classifiers can outperform unimodal
classifiers. Motivated by human multimodal experiences, multimodal processing
not only applies for emotion recognition, but also for other applications, e.g.,
270 audio-visual speech recognition, event detection, and media summarization.

Zadeh et al. (2017) used audio, visual, and text for sentiment analysis.
They proposed a model, termed as tensor fusion network (TFN), by posing
intra-modality and inter-modality dynamics of multimodalities. The proposed

method represents interactions between unimodal, bimodal, and trimodal features. The experiment on publicly available CMU-MOSI dataset produced state-of-the-art results for sentiment analysis.

Tripathi and Beigi (2018) used audio, text, and motion capture (mocap) to identify the categorical emotion of the IEMOCAP dataset. Using unimodal features, they obtained an accuracy of 55.65%, 64.78%, and 51.11% for audio, text, and mocap, respectively. Using multimodal features, they improved the accuracy of 71.04%. They concatenated different classifiers (LSTM for text and audio, CNN for mocap) to obtain the final categorical emotion prediction.

The lack of multimodal emotion recognition itself is the necessity of several modalities, mainly audio and video. In some cases, only audio data are available, e.g., telephone calls, voice assistants, and audio messages. Due to privacy and other reasons, the use of videos and other modalities may limit obtained data for predicting emotional states. In the case of speech, acoustic and linguistic features can be extracted back to obtain bimodal features.

Apart from the advantages and disadvantages of bimodal/multimodal recognition over using a single modality, there is a need to develop a new method for SER. The reasons are (1) some prior research did not predict all emotion attributes (Eyben et al., 2010; Karadogan and Larsen, 2012; Zhang et al., 2019), while other studies predicted emotion categories instead of attributes; (2) instead of predicting continuous emotion attribute scores, some studies switched to a categorical task for simplicity (Zhang et al., 2019; Aldeneh et al., 2017); and (3) some of the reported results are not up to date and showed low improvement (Eyben et al., 2010; Karadogan and Larsen, 2012).

Although we have limited our work to using both acoustic and text features, other researchers have already proposed another solution to solve the issues above, namely, using audiovisual emotion recognition. Nevertheless, there is still a need to propose and evaluate methods using both acoustic and text features, because some target applications only involve speech data. In these voice-based applications, no visual or written-text information is acquired. To maximize the resources for extracting emotion within speech, this paper exploits both

305 acoustic and text features (obtained via speech transcription) and combines
 them for dimensional emotion regression. By using both kinds of information in
 a two-stage process, we expect the proposed method’s performance to be close
 to or exceed the performance obtained by using visual information. Note here
 that visual information, particularly facial expressions, has been reported to
 310 have more influence on dimensional emotion than other modalities do (Fabien
 Ringeval et al., 2018).

3. Data and feature sets

3.1. Datasets

Datasets for investigating our proposal to use two-stage processing for di-
 315 mensional SER must meet certain requirements. The requirements are that (1)
 the dataset has both speech data and text transcription (to speed up text data
 acquisition), (2) the dataset is already annotated with dimensional labels, and
 (3) the dataset is publicly available. The following two datasets satisfy these
 requirements.

320 1. IEMOCAP

IEMOCAP, which stands for interactive emotional dyadic motion capture database,
 contains recordings of dyadic conversations with markers on the face, head, and
 hands. The recordings thus provide detailed information about the actors’s
 facial expressions and hand movements during both scripted and spontaneous
 325 spoken communication scenarios (Busso et al., 2008). This research only uses
 the acoustic and text features, because the goal is bimodal speech emotion recog-
 nition. The IEMOCAP dataset is freely available upon request, including its
 labels for categorical and dimensional emotion. We use the dimensional emo-
 tion labels, which are average scores for two evaluators, because they enable
 330 deeper analysis of emotional states. The dimensional emotion scores, for va-
 lence, arousal, and dominance, are meant to range from 1 to 5 as a result of
 Self-Assessment Manikin (SAM) evaluation. We have found some labels with
 scores lower than 1 or higher than 5; however, we remove those data (seven sam-

ples). All labels are then converted from the 5-point scale to a floating-point
335 values in range $[-1, 1]$ when they are fed to a DNN system.

The total length of the IEMOCAP dataset is about 12 hours, or 10039
turns/utterances, from ten actors in five dyadic sessions (two actors each). The
speech modality used to extract acoustic features is a set of files in the dataset
with a single channel per sentence. The sampling rate of the speech data was
340 16 kHz. For text data, we use the manual transcription in the dataset without
additional preprocessing.

2. MSP-IMPROV

MSP-IMPROV, developed by the Multimodal Signal Processing (MSP) Lab at
the University of Texas, Dallas, is a multimodal emotional database obtained
345 by applying lexical and emotion control in the recording process while also pro-
moting naturalness. The dataset provides audio and visual recordings, while
text transcriptions are obtained via automatic speech recognition (ASR) pro-
vided by the authors. As with IEMOCAP, we use the speech and text data
with dimensional emotion labels. The annotation method for the recordings
350 was the same as for IEMOCAP, i.e., SAM evaluation, with rating by at least
five evaluators. We treat missing evaluations as neutral speech (i.e., a score of
3 for valence, arousal, and dominance). Also as with IEMOCAP, all labels are
converted to floating-point values in range $[-1, 1]$ from the original 5-point scale.

The MSP-IMPROV dataset was designed within a dialogue framework to
355 elicit target sentences which had the same semantic content, but were produced
with different emotional expressions. In one recording, the target sentences were
produced ad lib; for another recording, the target sentences were read. These
two recordings are referred to as “Target-improvised” and “Target-read”, re-
spectively. For our purposes, since our goal is to examine the effect of both
360 linguistic and acoustic information on emotional ratings, these recordings were
not appropriate for our study. However, we were able to use two sets of record-
ings which did not have the same semantic content, which is called “Other-
improvised” and “Natural-interaction”. The former included conversations of
the actors during improvisation sessions; the latter included the exchanges dur-

365 ing the breaks, while the actors were not acting, which also were being recorded.
 A similar protocol was used by Zhang et al. (2019), and we follow their lead in
 referring to this subset of the MSP-IMPROV dataset as MSP-I+N (MSP im-
 proved and natural interaction), or MSPIN. In our work, we included the text
 transcriptions used by Zhang et al. (transcriptions are provided by the authors
 370 of the dataset); for the additional utterances not included in the Zang study,
 transcriptions were obtained using Mozilla’s DeepSpeech (Mozilla, 2019). We
 thus use 7166 utterances from a total of 8438. The speech data in the dataset
 was sampled in mono at 44.1 kHz, with one file per utterance/sentence.

We split each dataset into two partitions to observe any differences between
 375 a speaker-dependent (SD) partition and a speaker-independent partition made
 by leaving one session out (LOSO) for each dataset. For example, for the IEMO-
 CAP dataset, the last session (i.e., session 5) which is recorded from two different
 actors (out of 10) is only used for testing. Similarly, for MSP-I+N, all utterances
 from session 6 (two speakers out of 12) are used for the test set. Our rule for data
 380 splitting is to divide between the training + development and test sets in a ratio
 close to 80:20. This rule is applied for both the SD and LOSO partitions. Then,
 of the training + development data, 80% is used for training and the remaining
 20% is used for development, as shown in Figure 1. Both methods are evaluated
 with the same unseen test sets to compare the performance and measure the
 385 improvement. Note that we did not use cross validation (but instead divided
 into training and test data) for evaluation since the number of samples for both
 datasets is adequate (10039 and 7166 samples). This strategy is also utilized to
 keep the same test set for LSTM (one-stage processing) and SVM (two-stage
 processing) which is difficult if the samples are shuffled/cross-validated.

390 3.2. Feature sets

Most research on SER, or generally on pattern classification, focuses on
 two main topics: feature extraction, as in Batliner et al. (2011), and classi-
 fication/regression methods, as in Albornoz et al. (2011). While this research
 focuses on the second topic, we also evaluate state-of-the-art feature sets used for

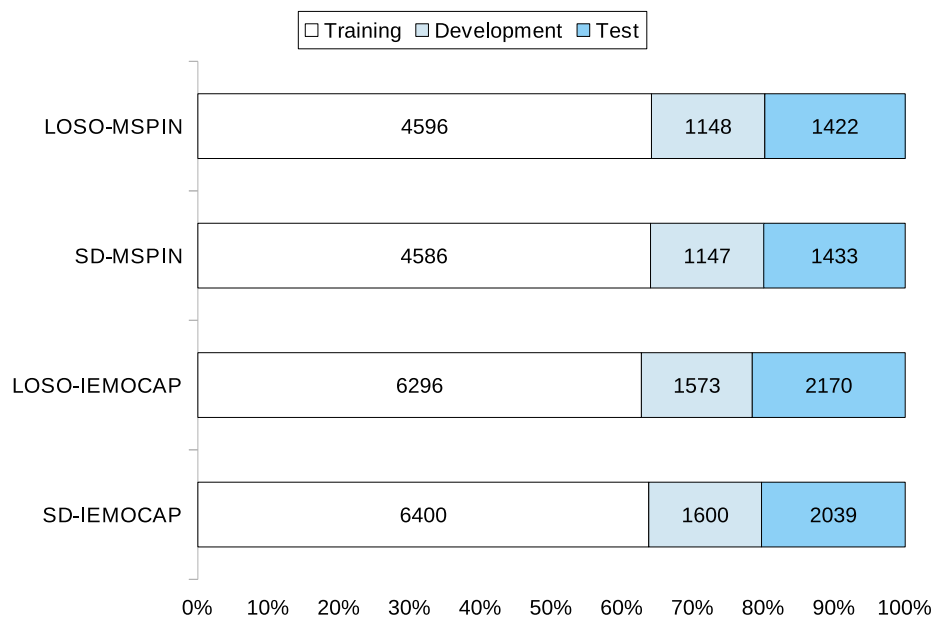


Figure 1: Proportions of data splitting for each partition of each dataset. In one-stage LSTM processing, the outputs of the model are both development and test data. In the second stage SVM processing, the input data is the prediction from the development set of the previous stage and the output is the prediction of test data.

395 SER. For acoustic features, we evaluate three feature sets: the Geneva Minimal-
 istic Acoustic Parameter Set (GeMAPS), statistical functions from GeMAPS,
 and the same functions from GeMAPS with a silence feature. For text features,
 aside from the original word vectors extracted from the text transcription, we
 also evaluate two word embeddings that are pretrained on a larger corpus: the
 400 Word2Vec embedding (Mikolov et al., 2013) and GloVe embedding (Pennington
 et al., 2014). These feature sets are explained below.

Acoustic features. The type of acoustic features extracted from a speech signal
 is the most important part of an SER system. GeMAPS is an effort to stan-
 dardize the acoustic features used for voice research and affective computing
 405 (Eyben et al., 2016a). The feature set consists of 23 acoustic low-level descrip-
 tors (LLDs) such as fundamental frequency (f_0), jitter, shimmer, and formants,
 as listed in Table 1. As an extension of GeMAPS, eGeMAPS includes statistical
 functions derived from the LLDs, such as the minimum, maximum, mean, and
 other values. Since these features are extracted on frame-based processing, the
 410 size of feature becomes large for one utterance (e.g., 3409×23 for IEMOCAP),
 which is suitable for deep learning methods like LSTM. Including the LLDs,
 the total number of features in eGeMAPS is 88. These statistical values are
 often called high-level statistical functions (HSF). Schmitt and Schuller (2018)
 found, however, that using only the mean and standard deviation (std) from the
 415 LLDs achieved a better result than using eGeMAPS and audiovisual features.
 These global features may represent more emotion information within speech
 than frame-based features. We thus coded these two statistical functions (47
 values) from the LLDs as the HSF1 feature set. We also investigate the effect of
 including a silence feature in this SER research, as explained below. We define
 420 the combination of HSF1 with the silence feature as HSF2.

Silence, in this paper, is defined as the proportion of silent frames among all
 frames in an utterance. In human communication, the proportion of silence in
 speaking depends on the speaker’s emotion. For example, a happy speaker may
 have fewer silences (or pauses) than a sad speaker. The proportion of silence in

an utterance can be calculated as

$$S = \frac{N_s}{N_t}, \quad (1)$$

where N_s is the number of frames categorized as silence (silent frames), and N_t is the total number of frames. A frame is categorized as silent if it does not exceed a threshold value defined by multiplying a factor by a root mean square (RMS) energy, X_{rms} . Mathematically, this is formulated as

$$th = \alpha \times \overline{X_{rms}}, \quad (2)$$

where X_{rms} is defined as

$$X_{rms} = \sqrt{\frac{1}{n} \sum_{i=1}^n x[i]^2}. \quad (3)$$

This silence feature is similar to the disfluency feature proposed in Moore et al. (2014). In that paper, the author divided the total duration of disfluency by the total utterance length for n words. Figure 2 illustrates the calculation of our silence feature. If X_{rms} from a frame is below th , then it is categorized as silent, and the calculation of equation 1 is applied.

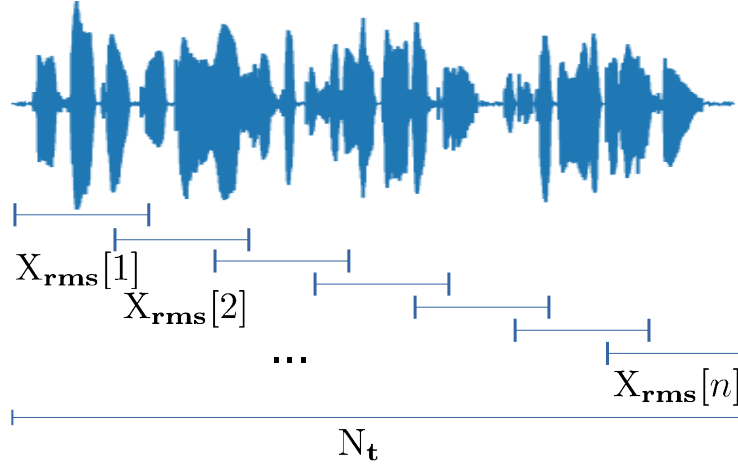


Figure 2: Moving frame to calculate the silence feature.

Table 1: Acoustic feature sets derived from the GeMAPS features by Eyben et al. (2016a) and the statistical functions used for dimensional SER in this research.

LLDs	HSF1	HSF2
intensity, alpha ratio, Hammarberg index, spectral slope 0-500 Hz, spectral slope 500-1500 Hz, spectral flux, 4 MFCCs, F0, jitter, shimmer, harmonics-to-noise ratio (HNR), harmonic difference H1-H2, harmonic difference H1-A3, F1, F1 bandwidth, F1 amplitude, F2, F2 amplitude, F3, and F3 amplitude.	mean (of LLDs), standard deviation (of LLDs)	mean (of LLDs), standard deviation (of LLDs), silence

Text features. To process a word sequence in a computational model, the text must be converted to numerical values. The resulting text feature is commonly known as a word embedding and is a vector representation of a word. Numerical values in the form of a vector are used to enable a computer to process text data, as it can only process numerical values. The values are points (numeric data) in a space whose number of dimensions is equal to the vocabulary size. The word representations embed those points in a feature space of lower dimension. In the original space, every word is represented by a one-hot vector, with a value of 1 for the corresponding word and 0 for other words. The element with a value of 1 is converted to a point in the range of the vocabulary size.

In addition to directly converting the text in the transcriptions of the datasets (IEMOCAP and MSP-I+N) to sequences, two pretrained word embeddings are used to weight the original word embeddings. As mentioned above, the two word-embedding models are Word2Vec (Mikolov et al., 2013) and GloVe (Pennington et al., 2014). Hence, we have three different text features for word embeddings: (1) a text sequence, or word embedding (WE), without any weighting, (2) a WE weighted by a pretrained Word2Vec model, and (3) a WE weighted by

a pretrained GloVe embedding model. Word2Vec is trained on large datasets to model its vector representation by using either continuous bag-of-words or continuous skip-gram model. GloVe is another model to represent vector in words by analyzing linear direction by using global log-bilinear regression method. All these features are fed into the same embedding layer in the text network.

4. Two-stage bimodal emotion recognition

4.1. Acoustic emotion recognition system

Most SER research uses only acoustic features. Our approach to acoustic SER is similar to that research. The contribution of our acoustic network is the evaluation of mean + std features at the utterance level and the use of a silence feature with the statistical functions to investigate any improvement. This evaluation is a continuation of (Tris Atmaja and Akagi, 2020) with extension on different feature sets and datasets. As explained in section 3.2, we evaluate three acoustic feature sets: LLDs, HSF1, and HSF2.

The LLD features are the 23 acoustic features listed in Table 1. For each frame (25 ms), these 23 acoustic features are extracted. With a hop size of 10 ms, the maximum number of sequences is 3409 for the IEMOCAP dataset and 3812 for the MSP-I+N dataset. Hence, the size of the input is 3409×23 for IEMOCAP and 3812×23 for MSP-I+N. The extraction process uses the openSMILE toolkit (Eyben et al., 2016b).

Figure 3 shows an overview of the acoustic network. LSTM is chosen because the number of training samples is adequate (> 5000 samples) and it shows good results on the previous research (Schmitt and Schuller, 2018). Before entering the LSTM layers, the LLD features at the input layer are fed into a batch normalization layer to speed up the computation process. The three subsequent LSTM layers are stacked with 256 nodes in each layer, following one of the configurations in AbdelWahab and Busso (2018). Instead of returning the last output of the last LSTM layer, we designed the network to return the full sequence and flatten it before inputting it to three dense layers that represent

Table 2: The hyper-parameter used in experiments

Hyper-parameter	Acoustic network	Text network
network type	LSTM	LSTM
number of layers	3	3
number of units	256	256
fourth layer	Flatten	Dense
hidden activation	linear	linear
output activation	linear (LLD) / tanh (HSF)	linear
dropout rate	0.3 (LLD) / 0 (HSF)	0.3
learning rate	0.001	0.001
batch size	8	8
maximum epochs	50	50
optimizer	RMSprop	RMSprop

valence, arousal, and dominance. The outputs of these last dense layers are then the predictions for those emotion attributes, i.e., the degrees of valence, arousal, and dominance in the range $[-1, 1]$.

485 The tuning of hyper-parameters follows the previous research (Atmaja et al., 2019; Atmaja and Akagi, 2020). A batch size of 8 was used with a maximum of 50 epochs. An early stop criterion with ten patiences stops the training process if no improvement were made in 10 epochs (before the maximum epoch) and used the last highest-score model to predict the development data. An RMSprop
490 optimizer was used with its default learning rate, i.e., 0.001. Table 2 shows the setups on acoustic and text networks. These setups were obtained based on experiments with regard to the size of networks. For instance, the smaller acoustic networks with HSF features employed tanh output activation function and did not use the dropout rate while the larger acoustic networks (with LLD)
495 and text networks employed linear activation function and dropout rate.

For the HSF1 and HSF2 inputs on acoustic networks, the same setup applies.

These two feature sets are very small as compared to the LLDs: HSF1 has a size of 1×46 , while HSF2 has a size of 1×47 . This big difference in input size (1:1800) leads to faster computation on HSF1 and HSF2 than on the LLDs. Note that, although Figure 3 shows HSF2 as the input feature, the same architecture also applies for the LLDs and HSF1.

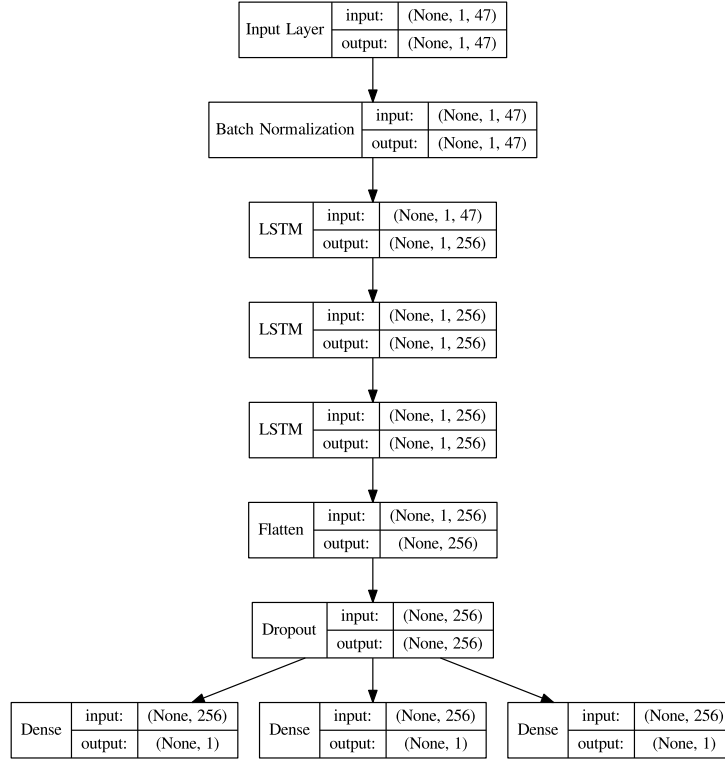


Figure 3: Structure of acoustic network to process acoustic features.

The idea of using LSTM is to hold the last output in memory and use the that output as successive step. For instance, LLD with (3409, 23) feature size will process the first time step 1 to the last time step 3409. For HSF1 and HSF2, which contains a single time stamp, the data is processed only once ([1, 46] and [1, 47] for HSF1 and HSF2). Here, the only difference, from multi time steps, is that the network performs three passes (forget gate, input gate, and output gate) instead of a single pass (see Eyben et al. (2010)). This information will

include all information from the networks’s memory.

510 4.2. Text emotion recognition system

The text network, shown in Figure 4 for the MSP-I+N dataset, uses the same input size for the three different text features. The WE, WE with pretrained Word2Vec, and WE with pretrained GloVe embedding have 300 dimensions for each word. The longest sequence in the IEMOCAP dataset is 100 sequences
515 (words), while for MSP-I+N, the longest is 300 sequences. Hence, the input feature sizes for the LSTM layers are 100×300 for IEMOCAP and 300×300 for MSP-I+N with its corresponding number of samples. The same three LSTM layers are stacked as in the acoustic network, but the last LSTM layer only returns the last output. A dense layer with a size of 128 nodes is added
520 after the LSTM layers and before the last three dense layers. Between the dense layers is a dropout layer with the same probability of 0.3 to avoid overfitting.

4.3. Multitask learning

The task here in dimensional emotion recognition is to simultaneously predict the degrees of three emotion attributes, i.e., the degrees of valence, arousal, and dominance, for any given utterance. As the main target metric is the concordance correlation coefficient (CCC), the loss function is the CCC loss (CCCL). CCC also measures the performance of regression analyses by comparing gold standard labels and the predictions. CCCL computes the score difference between the labels and predicted values for the three attributes. The CCC and CCCL are formulated as the following:

$$CCC = \frac{2\rho\sigma_x\sigma_y}{\sigma_x^2 + \sigma_y^2 + (\mu_x - \mu_y)^2}, \quad (4)$$

$$CCCL = 1 - CCC, \quad (5)$$

where ρ is the Pearson correlation coefficient between the predicted emotion degree x and the true emotion degree y , σ^2 is the variance, and μ is the mean. As the learning process minimizes three variables, we use the following multitask

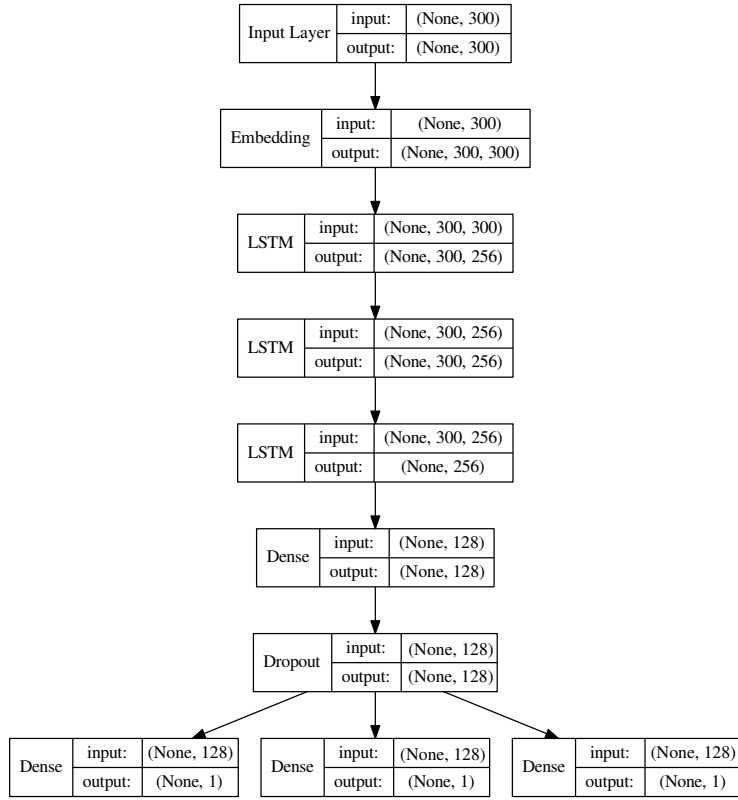


Figure 4: Structure of text network to process word embeddings/vectors.

learning approach to optimize the CCC score:

$$CCCL_{tot} = \alpha CCCL_V + \beta CCCL_A + (1 - \alpha - \beta)CCCL_D, \quad (6)$$

where α and β are respective CCCL parameters for valence (V) and arousal (A). The parameter for dominance (D) is obtained by subtracting α and β from 1. The same parameter range $[0, 1]$ with 0.1 step is investigated for α and β for both the acoustic and text networks, resulting in different optimal parameters, which are obtained by using linear search. Note that only positive values of $CCCL_D$'s parameters are used to investigate the optimal parameters.

4.4. SVM-based late fusion

We choose an SVM as the final classifier to fuse the outputs of the acoustic and text networks because of its effectiveness in handling smaller data (as compared to a DNN) and its computation speed. The datapoints produced by LSTM processing as the input of SVM is small; i.e., 1600, 1538, 1147, and 1148 for IEMOCAP-SD, IEMOCAP-LOSO, MSPIN-SD and MSPIN-LOSO, respectively. The SVM then applies a regression analysis to map them to the given labels. Figure 5 shows the architecture of this two-stage emotion recognition system using DNNs and an SVM. Each prediction from the acoustic and text networks is fed into the SVM. From two values (e.g., valence predictions from the acoustic and text networks), the SVM learns to generate a final predicted degree (e.g., for valence). The concept of using the SVM as the final classifier can be summarized as follows.

Suppose that two valence prediction outputs from the acoustic and text networks, $x_i = [x_{ser}[i], x_{ter}[i]]$, are generated by the DNNs, and that y_i is the corresponding valence label. The problem in dimensional SER fusing acoustic

545 and text results is to minimize the following:

$$\begin{aligned}
& \min_{w, b, \zeta, \zeta^*} \quad \frac{1}{2} w^T w + C \sum_{i=1}^n \zeta_i + C \sum_{i=1}^n \zeta_i^* \\
& \text{subject to} \quad w^T \phi(x_i) + b - y_i \leq \epsilon + \zeta_i, \\
& \quad y_i - w^T \phi(x_i) - b \leq \epsilon + \zeta_i^*, \\
& \quad \zeta_i, \zeta_i^* \geq 0, i = 1, \dots, n,
\end{aligned} \tag{7}$$

where w is a weighting vector, C is a penalty parameter, ζ and ζ^* is the distance between misclassified points and the corresponding marginal boundary (above or below). Here, ϕ is the kernel function. We choose a radial basis function (RBF) kernel because of its flexibility to model a nonlinear process with
550 a dimensional emotion model close to this kernel. The function ϕ for the RBF kernel is formulated as

$$K(x_i, x_j) = e^{\gamma(x_i - x_j)^2}, \tag{8}$$

where γ defines how much influence a single training has on the model. All parameters in this SVM are obtained empirically via linear search in a specific range. Although the explanation above uses valence, the same also applies for
555 arousal and dominance.

4.5. Reproducibility

The experimental code was written in Python and, for the sake of research reproducibility, it is available in the following repository: <https://github.com/bagustris/two-stage-ser>. The DNN part was implemented using Keras
560 by Chollet and Others (2015) and Tensorflow, while the SVM-based fusion was implemented using the scikit-learn toolkit by Pedregosa et al. (2011). To obtain consistent results for each run, some fixed numbers are initialized at the beginning, as can be found in the repository above.

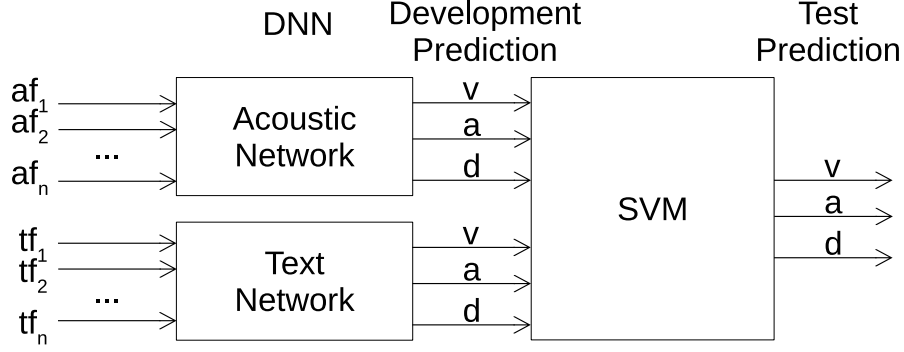


Figure 5: Proposed two-stage dimensional emotion recognition method using DNNs and an SVM. The inputs are acoustic features (af) and text features (tf); the outputs are valence (v), arousal (a), and dominance (d).

5. Results and discussions

5.1. Results from single modality

Before presenting the bimodal feature-fusion results, it is important to show the results of unimodal emotion recognition. The goals here are (1) to observe the (relative) improvement of bimodal feature fusion over using a single modality, and (2) to observe the effects of different features on different emotion attributes.

Tables 3 and 4 list the single-modality results of dimensional emotion recognition from the acoustic and text networks, respectively. In general, acoustic-based SER gave better results than text-based SER in terms of the average CCC score. For particular emotion attributes, the text network gave a higher CCC score for valence prediction than those obtained by the acoustic network, except on the MSPIN datasets. This confirms the previous finding by Karadogan and Larsen (2012) that valence is better estimated by semantic features, while arousal is better predicted by acoustic features. In addition, we found that the dominance dimension was better predicted by acoustic features than by text features. This finding can be inferred from both tables, in which the CCC scores for the dominance dimension are frequently higher from the acoustic

network than from the text network.

The exception of a higher valence score on the MSPIN-SD dataset by the acoustic network can be seen as the effect of either the DNN architecture or the dataset’s characteristics. In Chen et al. (2017), the obtained score was higher for valence than for arousal or liking (the third dimension, instead of dominance) with their strategy on acoustic features. In contrast, AbdelWahab and Busso (2018) obtained a lower score for valence than for arousal and dominance by using their proposed DANN method on the same MSP-IMPROV dataset (whole data, all four scenarios). Given this comparison, we conclude that the higher valence score obtained here was an effect of the DNN architecture, because of the multitask learning. Our result on a single modality (acoustic network) outperformed the DANN result on MSP-IMPROV, where their highest CCC scores were (0.303, 0.176, 0.476) as compared to our scores of (0.404, 0.605, 0.517) for valence, arousal, and dominance, respectively.

To find the optimal parameter values for α and β , a linear search was performed on the scale [0.0, 1.0] with a step of 0.1. Using this conventional technique, we found four sets of optimal parameters for the acoustic and text networks. Note that, while only the improvised and natural scenarios (MSP-I+N) were used to find the optimal text-network parameters for the MSP-IMPROV dataset, the whole dataset was used to find the optimal acoustic-network parameters. Table 5 lists the optimal parameter values for α and β .

To summarize the single-modality results, average CCC scores from three emotion dimensions can be used to justify which features perform better among others. The results show that HSF2 was the most useful of the acoustic feature sets (in two of four datasets), while the word embedding (WE) with pretrained GloVe embedding was the most useful of the text feature sets. The performance of dimensional emotion recognition in the speaker-independent (LOSO) case was lower than in the speaker-dependent (SD) case, as predicted. Note that both acoustic and text emotion networks used a fixed seed number to achieve the same result for each run; however, text network resulted in different scores. Hence, standard deviations were given to measure fluctuation in 20 runs.

Table 3: CCC score results of dimensional emotion recognition using an acoustic network. The best results on the test set are in bold. LLDs: low-level descriptors from GeMAPS (Eyben et al., 2016a); HSF1: mean + std of LLDs; HSF2: mean + std + silence.

Feature set	V	A	D	Mean
IEMOCAP-SD				
LLD	0.153	0.522	0.534	0.403
HSF1	0.186	0.535	0.466	0.396
HSF2	0.192	0.539	0.469	0.400
MSPIN-SD				
LLD	0.299	0.545	0.441	0.428
HSF1	0.400	0.603	0.506	0.503
HSF2	0.404	0.605	0.517	0.508
IEMOCAP-LOSO				
LLD	0.168	0.486	0.442	0.365
HSF1	0.206	0.526	0.442	0.391
HSF2	0.204	0.543	0.442	0.396
MSPIN-LOSO				
LLD	0.176	0.454	0.369	0.333
HSF1	0.201	0.506	0.357	0.355
HSF2	0.206	0.503	0.346	0.352

Table 4: CCC score results of dimensional emotion recognition using text networks; each score is an averaged score of 20 runs with its standard deviation. WE: word embedding; Word2Vec: WE weighted by pretrained word vector (Mikolov et al., 2013); GloVe: WE weighted by pretrained global vector (Pennington et al., 2014).

Feature set	V	A	D	Mean
IEMOCAP-SD				
WE	0.389 ± 0.008	0.373 ± 0.010	0.398 ± 0.017	0.387 ± 0.010
Word2Vec	0.393 ± 0.012	0.371 ± 0.018	0.366 ± 0.024	0.377 ± 0.016
GloVe	0.410 ± 0.007	0.381 ± 0.013	0.393 ± 0.016	0.395 ± 0.010
MSPIN-SD				
WE	0.120 ± 0.047	0.148 ± 0.023	0.084 ± 0.024	0.105 ± 0.026
Word2Vec	0.138 ± 0.031	0.108 ± 0.024	0.101 ± 0.024	0.116 ± 0.017
GloVe	0.147 ± 0.043	0.141 ± 0.019	0.098 ± 0.017	0.128 ± 0.015
IEMOCAP-LOSO				
WE	0.376 ± 0.008	0.359 ± 0.018	0.370 ± 0.020	0.368 ± 0.013
Word2Vec	0.375 ± 0.058	0.357 ± 0.058	0.365 ± 0.065	0.366 ± 0.059
GloVe	0.405 ± 0.009	0.382 ± 0.020	0.378 ± 0.021	0.389 ± 0.014
MSPIN-LOSO				
WE	0.076 ± 0.013	0.196 ± 0.011	0.136 ± 0.015	0.136 ± 0.009
Word2Vec	0.162 ± 0.008	0.202 ± 0.005	0.147 ± 0.003	0.170 ± 0.000
GloVe	0.192 ± 0.004	0.189 ± 0.007	0.129 ± 0.004	0.170 ± 0.003

Table 5: Optimal parameters for multitask learning.

Dataset	Modality	α	β
IEMOCAP	acoustic	0.1	0.5
	text	0.7	0.2
MSP-IMPROV	acoustic	0.3	0.6
	text	0.1	0.6

5.2. Results from SVM-based fusion

The main proposal of this research is the late-fusion approach combining the
615 results from acoustic and text networks for dimensional emotion recognition.
This subsection presents the results for the late-fusion approach, including the
obtained performances, comparison with the single-modality results, which pairs
of acoustic-text results performed better, and our overall findings.

For each dataset (IEMOCAP-SD, MSPIN-SD, IEMOCAP-LOSO, MSPIN-
620 LOSO), nine combinations of acoustic-text result pairs could be fed to the SVM
system. Tables 6, 7, 8, and 9 list the respective CCC results for these datasets.
Generally, our proposed two-stage dimensional emotion recognition improved
the CCC score from single-modality emotion recognition. The pair of results
from HSF2 (acoustic) and Word2Vec (text) gave the highest CCC score on
625 speaker-dependent scenarios.

On the speaker-independent IEMOCAP dataset (IEMOCAP-LOSO), the
result from the pair of HSF2 and GloVe gave the highest CCC score. This result
linearly correlated with the single-modality results for that dataset, in which
HSF2 obtained the highest CCC score among the acoustic features, and GloVe
630 was the best among the text features. On the four datasets, the results from
HSF2 obtained the highest CCC score for two out of four datasets while GloVe
obtained the highest CCC score for all four datasets. Hence, we conclude that
the highest result from a single modality, when paired with the highest result
from another modality, will achieve the highest performance among possible
635 pairs.

To evaluate the improvement obtained by SVM-based late fusion, average
CCC scores again can be used as a single metric. The rightmost column in Ta-
bles 6, 7, 8, and 9 shows the average CCC scores obtained from the nine pairs
of acoustic and text results on the four different datasets. Comparing these
640 bimodal results to unimodal results (Table 3 and 4) shows the difference. All
results from SVM improved unimodal results. In speaker-independent (LOSO)
results (which are more appropriate for real-life analysis), the scores resulted by
pairs of HSF with any word vector obtain remarkable improvements, particularly

Table 6: CCC score results after late fusion using an SVM on the IEMOCAP-SD test set.

Inputs	V	A	D	Mean
LLD + WE	0.520	0.602	0.519	0.547
LLD + Word2Vec	0.552	0.613	0.524	0.563
LLD + GloVe	0.546	0.606	0.520	0.557
HSF1 + WE	0.578	0.575	0.490	0.548
HSF1 + Word2Vec	0.599	0.590	0.491	0.560
HSF1 + GloVe	0.595	0.582	0.495	0.557
HSF2 + WE	0.598	0.591	0.502	0.564
HSF2 + Word2Vec	0.595	0.601	0.499	0.565
HSF2 + GloVe	0.598	0.591	0.502	0.564

Table 7: CCC score results after late fusion using an SVM on the MSPIN-SD dataset.

Inputs	V	A	D	Mean
LLD + WE	0.344	0.591	0.447	0.461
LLD + Word2Vec	0.326	0.586	0.439	0.450
LLD + GloVe	0.344	0.585	0.439	0.456
HSF1 + WE	0.461	0.637	0.517	0.538
HSF1 + Word2Vec	0.464	0.634	0.518	0.539
HSF1 + GloVe	0.466	0.630	0.510	0.535
HSF2 + WE	0.475	0.640	0.522	0.546
HSF2 + Word2Vec	0.486	0.641	0.524	0.550
HSF2 + GloVe	0.485	0.638	0.523	0.549

Table 8: CCC score results after late fusion using an SVM on the IEMOCAP-LOSO test set.

Inputs	V	A	D	Mean
LLD + WE	0.537	0.583	0.431	0.517
LLD + Word2Vec	0.528	0.580	0.421	0.510
LLD + GloVe	0.539	0.587	0.430	0.518
HSF1 + WE	0.565	0.565	0.453	0.528
HSF1 + Word2Vec	0.536	0.559	0.434	0.510
HSF1 + GloVe	0.559	0.570	0.452	0.527
HSF2 + WE	0.524	0.566	0.452	0.514
HSF2 + Word2Vec	0.531	0.571	0.445	0.516
HSF2 + GloVe	0.553	0.579	0.465	0.532

Table 9: CCC score results after late fusion using an SVM on the MSPIN-LOSO test set.

Inputs	V	A	D	Mean
LLD + WE	0.204	0.485	0.387	0.358
LLD + Word2Vec	0.267	0.487	0.386	0.380
LLD + GloVe	0.269	0.482	0.375	0.376
HSF1 + WE	0.224	0.565	0.410	0.400
HSF1 + Word2Vec	0.286	0.558	0.411	0.418
HSF1 + GloVe	0.282	0.555	0.409	0.415
HSF2 + WE	0.232	0.566	0.421	0.406
HSF2 + Word2Vec	0.287	0.562	0.411	0.420
HSF2 + GloVe	0.291	0.570	0.405	0.422

in the MSPIN-LOSO dataset. For any other pair involving LLDs, the obtained
645 score was also lower as compared to other pairs. Considering all low scores
involved LLD results, improving the performance of dimensional emotion recog-
nition by using LLDs is more complicated than by using HSF1 and HSF2, apart
from the larger feature size and the longer training time. The large network size
created by an LLD input as a result of its much bigger feature dimension (e.g.,
650 3409×23 on IEMOCAP) did not help either the single-modality or late-fusion
performance. In contrast, the small sizes of the functional features (HSF1 and
HSF2) enabled better performance on a single modality, which led to better
performance for the late-fusion score. To obtain functional features, however, a
set of LLD features must be obtained first. This problem is a challenging future
655 research direction, especially for implementing dimensional emotion recognition
with real-time processing.

Aside from the fact that a speaker-independent dataset is usually more dif-
ficult than a speaker-dependent dataset, the low score on MSPIN-LOSO was
due to its low scores on a single modality. In other words, lower pair perfor-
660 mance from a single modality will result in low performance in late fusion. In
particular, these low results derive from low CCC scores from the text modality
(Table 4). The average CCC score for the text modality on the MSPIN-LOSO
dataset was less than 0.16, compared to an average score higher than 0.34 for
the acoustic modality. All nine pairs in late-fusion approaches improved on
665 the single-modality results because of the two-stage DNN and SVM regression
analysis. Thus, out of 36 trials ($9 \text{ pairs} \times 4 \text{ datasets}$), our proposed two-stage
dimensional emotion recognition outperforms any single modality result (used
in a pair).

The low score on MSPIN for the text modality can be tracked to the origin
670 of the dataset, that is, there may have been a number of sentences semantically
identical to the target sentences in the dataset we used. Although we chose
sentences only from the improvised dialogues (minus the target sentence) and
from the natural interactions (those sentences produced by the experimenters
and subjects during the breaks), some of the sentences in this corpus were se-

675 mantically identical to that of the target sentences in the “Target-Improvised” data set. This was confirmed retroactively by manually checking the provided transcription and our automatic transcription. Given the nature of the elicitation task in a dialogue framework, this is not surprising. A similarly low result for the text modality on this MSPIN dataset was also shown in Zhang et al.
 680 (2019). In general, compared to the IEMOCAP dataset, the MSPIN dataset suffers from low accuracy in recognizing the valence category by using acoustic and lexical properties. Interestingly, however, those authors also did not show improvement on the IEMOCAP scripted dataset, another text-based session in which lexical/text features do not contribute significantly.

685 To measure the improvement by our proposed two-stage late fusion, we calculated the relative improvement obtained by late fusion from the highest CCC scores for a single modality. For example, the pair of LLD + WE used the results from the LLDs in the acoustic network and the WE in the text network. We compared the result for LLD + WE with that for the LLDs, as it had a
 690 higher score than the WE did. Figure 6 thus shows the relative improvement for all nine pairs. All of 36 trials showed improvements ranging from 5.11% to 40.32%. Table 10 lists the statistics for the obtained relative improvement. Our results show higher relative accuracy improvement as compared to those obtained by Zhang et al. (2019) for valence prediction, which ranged from 6% to
 695 9%. Nevertheless, their multistage fusion method also showed benefits over the multimodal and single-modality approaches. These findings confirm the benefits of using bimodal/multimodal fusion instead of single-modality processing for valence, arousal, and dominance prediction.

5.3. *Speaker-dependent vs. speaker-independent text emotion recognition*

700 While speech-based emotion recognition is performed with a fixed random seed to generate the same result for each run, text-based emotion recognition resulted in different scores for each run. The different results on text emotion recognition probably is caused by initiation of weightings on embedding layers. In this case, statistical tests can be performed on text emotion results to observe

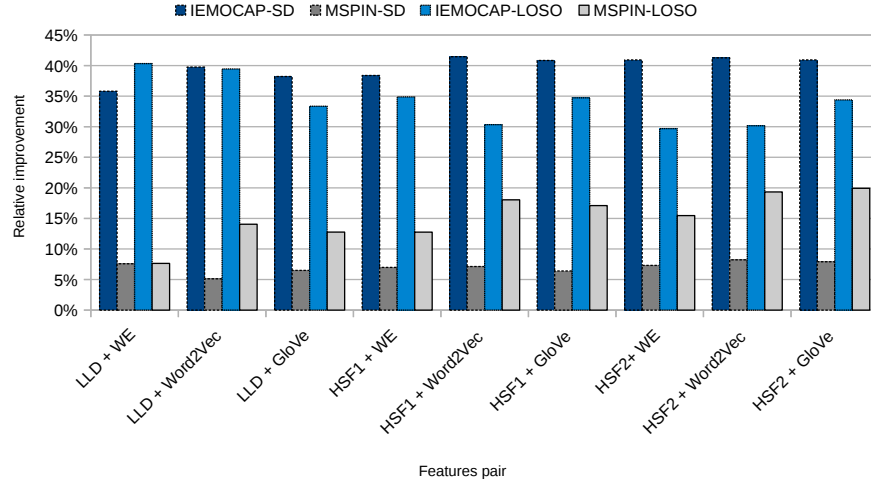


Figure 6: Relative improvement in average CCC scores from late fusion using an SVM as compared to the highest average CCC scores from a single modality.

Table 10: Statistics of relative improvement by late fusion using an SVM as compared to the highest scores for a single modality across datasets; the scores were extracted from the data shown in Figure 6.

Statistic	IEMOCAP-SD	MSPIN-SD	IEMOCAP-LOSO	MSPIN-LOSO
Average	39.73%	7.01%	34.15%	15.23%
Max	41.45%	8.22%	40.32%	19.93%
Min	35.80%	5.11%	29.69%	7.64%
Std	1.90%	0.93%	3.84%	3.90%

Table 11: Significant difference between speaker-dependent and speaker-independent scenario on the same text feature set; statistical tests were performed using two-tail paired t -test with p -value = 0.005.

Feature	IEMOCAP	MSPIN
WE	Yes	Yes
Word2Vec	No	Yes
GloVe	Yes	Yes

705 the difference between speaker-dependent and speaker-independent scenario. In contrast, statistical tests cannot be performed between acoustic results and bimodal acoustic-text results due to differences of the data (deterministic vs non-deterministic).

Table 11 shows if there is a significant difference between speaker-dependent
710 and speaker-independent results on the same feature set. We set p -value = 0.05 with two-tail paired t -test between mean scores of speaker-dependent and speaker-independent results. This paired t -test was based on the assumption that there are no outliers (after pre-processing) and two different inputs are fed into the same system. Only one result from text emotion recognition shows
715 no significant difference on IEMOCAP dataset while all results on MSPIN dataset shows a significant difference between speaker-dependent and speaker-independent results. This result reveals a tendency for a difference on evaluating speaker-dependent and speaker-independent data. The results from speaker-dependent data were different from those of the speaker-independent data. In
720 other words, results from speaker-dependent data cannot be used to justify speaker-independent or whole data.

5.4. Effect of removing target sentence from MSPIN dataset

Since the goal of this research is to evaluate the contribution of both acoustic and linguistic information in affective expressions, it is necessary to have
725 sentences in the dataset that are free from any stimuli control. However, the original MSP-IMPROV dataset contains 20 “target” sentences; a sentence with

the same linguistic content, but produced with different emotions. These parts of MSP-IMPROV dataset are irrelevant to this study; hence, we remove it from the dataset (*Target - improvised* and *Target - read* parts). However, we found
730 that the results show low CCC scores (Table 4, particularly on valence) indicating influence from target sentences. These results may be explained, as mentioned in section 5.2, that some utterances in the data analyzed in this study also inadvertently included sentences semantically the same as those in the improvised target sentences.

735 5.5. Final remarks

We tried to perform a benchmark between our results and others on the same datasets, scenarios, and metrics. Unfortunately, to the best of our knowledge, the only reference is the one reported by Atmaja and Akagi (2020), which reports an early fusion method on IEMOCAP dataset. We improve the average CCC
740 score from 0.508 to 0.532. This higher result suggests that late fusion is better than early fusion to model how humans fuse multimodal information, which is inline with neuropsychological research. This late-fusion approach can be embedded with current speech technology, i.e., ASR, in which the text output can be processed to weigh emotion prediction from acoustic features.

745 AbdelWahab and Busso (2018) used MSP-Podcast (Lotfian and Busso, 2019) as a target corpora, which is not available for the public yet, and IEMOCAP with MSP-IMPROV as source corpus to implement their DANN for cross-corpus speech emotion recognition. Although the goal is different, we observed similar patterns between theirs and our acoustic-only speech emotion recognition.
750 First, we observed that the order of highest to lowest CCC scores is arousal, dominance, and valence. This pattern is also consistent when IEMOCAP is mixed with MSP-IMPROV as reported by Parthasarathy and Busso (2017) (in Table 2). Second, we observed that the CCC scores obtained in IEMOCAP are higher than those obtained in MSP-IMPROV; we believe that this lower score
755 in MSP-IMPROV was due to the smaller size of the dataset.

Along with our SVM architecture, we also explored the parameters C and

γ , because both parameters are important for an RBF-kernel-based SVM architecture (Pedregosa et al., 2011). Linear search was used in the ranges of $[10^{-2}, 1, 10^2, 2 \times 10^2, 3 \times 10^2]$ for C and $[10^{-2}, 10^{-1}, 1, 10, 10^2]$ for γ with a fixed value of ϵ , i.e., 0.01. The best parameter values were $C = 200$ and $\gamma = 0.1$. The repository includes the detailed implementation of the SVM architecture.

Per our stated objective, we applied two-stage processing by using DNNs and an SVM for dimensional emotion recognition from acoustic and text features on four different datasets. We found that the combination of mean + std + silence from the acoustic features and word embeddings weighted by pretrained GloVe embeddings achieved the highest result among the nine pairs of acoustic-text results from DNNs trained with multitask learning. When the performance in obtaining one input to the SVM is very low, the resulting relative improvement due to the SVM is also low. For instance, the lowest improvement on MSPIN-LOSO was from LLD + WE features, in which WE obtained a low score ($\overline{CCC} = 0.136$) on text network. This phenomenon suggests a challenging future research direction for dealing with minimal linguistic information in the fusion strategy. One strategy applied in this research was to use a pretrained GloVe embedding on text features with HSF2 on acoustic features, which improved the $\overline{CCC}$ score from 0.358 (relative improvement = 7.64%) to 0.422 (relative improvement = 19.93%). Other strategies should also be proposed, such as on how to handle the data differently when the linguistically identical sentences elicits different emotions (i.e., whole MSP-IMPROV dataset). In contrast, the current evaluated word embeddings treat the same words to have the same representations, even when it conveys different emotions.

6. Conclusions

In conclusion, we summarize several findings. First, we found a linear correlation between the single-modality and late-fusion methods in dimensional emotion recognition. The best results from each modality, when they were paired, gave the best fusion result. In the same way, the worst results obtained from

each network, when they were paired, gave the worst fusion results for bimodal emotion recognition. This finding differs from that reported in Atmaja et al. (2019), which used an early-fusion approach for categorical emotion recognition. In their work, the best pair differs from the best methods in single modalities.

790 Second, text features strongly influenced the score of dimensional SER on the valence dimension, while acoustic features strongly influenced arousal and dominance scores. Accordingly, the proposed two-stage processing can take advantage of text features which are commonly used in predicting sentiment (valence) for the dimensional emotion recognition task. The proposed fusion
795 method improves all three emotion dimensions without attenuating the performance of any dimension. That is, the proposed method elevates the scores for valence, arousal, and dominance subsequently from the highest to the lowest gain.

Third, the combination of input pairs does not matter in the proposed fusion
800 method, as indicated by the low deviation in relative improvement across the nine possible input pairs. What does matter is the performance on the input in the DNN stage. If the performance of a feature set in the DNN stage is low ($\overline{CCC} \leq 0.2$), it will also result in low performance when paired with another low-performance input in the SVM stage.

805 Finally, this bimodal approach can be extended to a multimodal approach. Both acoustic and text features can be combined with visual and motion-capture measurements that have advantages in specific emotion dimensions (liking or naturalness). The results can be benchmarked with current results to observe such improvements by adding more modalities. The SVM stage itself can be
810 performed many times to obtain such improvements. These broad research directions are open challenges for researchers in human-computer interaction.

References

AbdelWahab, M., Busso, C., 2018. Domain Adversarial for Acoustic Emotion Recognition. *IEEE/ACM Transactions on Audio Speech and*

- 815 Language Processing 26, 2423–2435. doi:10.1109/TASLP.2018.2867099,
arXiv:1804.07690.
- Albornoz, E.M., Milone, D.H., Rufiner, H.L., 2011. Spoken emotion recognition
using hierarchical classifiers. Computer Speech & Language 25, 556–570.
doi:10.1016/j.cs1.2010.10.001.
- 820 Aldeneh, Z., Khorram, S., Dimitriadis, D., Provost, E.M., 2017. Pooling acoustic
and lexical features for the prediction of valence, in: ICMI 2017 - Proceedings
of the 19th ACM International Conference on Multimodal Interaction, ACM.
pp. 68–72. doi:10.1145/3136755.3136760.
- Alm, C.O., Roth, D., Sproat, R., 2005. Emotions from text: machine learn-
825 ing for text-based emotion prediction, in: Proceedings of the conference
on Human Language Technology and Empirical Methods in Natural Lan-
guage Processing - HLT '05, pp. 579–586. doi:10.3115/1220575.1220648,
arXiv:arXiv:1011.1669v3.
- Atmaja, B.T., Akagi, M., 2019. Speech Emotion Recognition Based on Speech
830 Segment Using LSTM with Attention Model, in: 2019 IEEE International
Conference on Signals and Systems (ICSigSys), IEEE. pp. 40–44. doi:10.
1109/ICSIGSYS.2019.8811080.
- Atmaja, B.T., Akagi, M., 2020. Dimensional speech emotion recognition from
speech features and word embeddings by using multitask learning. APSIPA
835 Transactions on Signal and Information Processing 9, e17. doi:10.1017/
ATSIP.2020.14.
- Atmaja, B.T., Shirai, K., Akagi, M., 2019. Speech Emotion Recognition Us-
ing Speech Feature and Word Embedding, in: 2019 Asia-Pacific Signal and
Information Processing Association Annual Summit and Conference (AP-
840 SIPA ASC), Lanzhou. pp. 519—523. doi:10.1109/APSIPAASC47483.2019.
9023098.

- Batliner, A., Steidl, S., Schuller, B., Seppi, D., Vogt, T., Wagner, J., Devillers, L., Vidrascu, L., Aharonson, V., Kessous, L., Amir, N., 2011. Whodunnit - Searching for the most important feature types signalling emotion-related user states in speech. *Computer Speech and Language* 25, 4–28. doi:10.1016/j.cs1.2009.12.003.
- Berckmoes, C., Vingerhoets, G., 2004. Neural foundations of emotional speech processing. doi:10.1111/j.0963-7214.2004.00303.x.
- Busso, C., Bulut, M., Lee, C.C.C., Kazemzadeh, A., Mower, E., Kim, S., Chang, J.N., Lee, S., Narayanan, S.S., 2008. IEMOCAP: Interactive emotional dyadic motion capture database. *Language Resources and Evaluation* 42, 335–359. doi:10.1007/s10579-008-9076-6, arXiv:0507464v2.
- Charles, D., Paul, E., Phillip, P., 1872. The expression of the emotions in man and animals. Electronic Text Center, University of Virginia Library .
- Chen, S., Jin, Q., Zhao, J., Wang, S., 2017. Multimodal multi-task learning for dimensional and continuous emotion recognition, in: *Proceedings of the 7th Annual Workshop on Audio/Visual Emotion Challenge*, ACM. pp. 19–26.
- Chollet, F., Others, 2015. Keras. <https://keras.io>.
- El Ayadi, M., Kamel, M.S., Karray, F., 2011. Survey on speech emotion recognition: Features, classification schemes, and databases. *Pattern Recognition* 44, 572–587. doi:10.1016/j.patcog.2010.09.020.
- Elbarougy, R., Akagi, M., 2014. Improving speech emotion dimensions estimation using a three-layer model of human perception. *Acoustical Science and Technology* 35, 86–98. doi:10.1250/ast.35.86.
- Eyben, F., Scherer, K.R., Schuller, B.W., Sundberg, J., Andre, E., Busso, C., Devillers, L.Y., Epps, J., Laukka, P., Narayanan, S.S., Truong, K.P., 2016a. The Geneva Minimalistic Acoustic Parameter Set (GeMAPS) for Voice Research and Affective Computing. *IEEE Transactions on Affective Computing* 7, 190–202. doi:10.1109/TAFFC.2015.2457417.

- 870 Eyben, F., Wöllmer, M., Graves, A., Schuller, B., Douglas-Cowie, E., Cowie, R.,
2010. On-line emotion recognition in a 3-D activation-valence-time continuum
using acoustic and linguistic cues. *Journal on Multimodal User Interfaces* 3,
7–19. doi:10.1007/s12193-009-0032-6.
- Eyben, F., Wöllmer, M., Schuller, B., 2016b. *Opensmile*. November, ACM
875 Press, New York, New York, USA. doi:10.1145/1873951.1874246.
- Fabien Ringeval, Schuller, B., Valstar, M., Cowie, R., Kaya, H., Schmitt, M.,
Amiriparian, S., Cummins, N., Lalanne, D., Michaud, A., Salah, E.Ç., Ali,
H.G.A., Pantic, M., 2018. AVEC 2018 Workshop and Challenge: Bipolar
Disorder and Cross-Cultural Affect Recognition, in: *Proceedings of the 8th*
880 *International Workshop on Audio/Visual Emotion Challenge, AVEC’18*, co-
located with the 26th ACM International Conference on Multimedia, MM
2018, pp. 3–13.
- Fayek, H.M., Lech, M., Cavedon, L., 2017. Evaluating deep learning archi-
tectures for Speech Emotion Recognition. *Neural Networks* 92, 60–68.
885 doi:10.1016/j.neunet.2017.02.013.
- Fontaine, J.R.J., Scherer, K.R., Roesch, E.B., Phoebe, C., Fontaine, J.R.J.,
Scherer, K.R., Roesch, E.B., Ellsworth, P.C., 2017. The World of Emotions
Is Not Two-Dimensional. *Psychological science* 18, 1050–1057.
- Giannakopoulos, T., Pikrakis, A., Theodoridis, S., 2009. A dimensional ap-
890 proach to emotion recognition of speech from movies, in: *2009 IEEE Inter-
national Conference on Acoustics, Speech and Signal Processing, IEEE*. pp.
65–68. doi:10.1109/ICASSP.2009.4959521.
- Grice, H.P., 2002. Logic and Conversation, in: *Foundations of Cognitive Psy-
chology*. The MIT Press, pp. 41–58. doi:10.7551/mitpress/3080.003.0049.
- 895 Grimm, M., Kroschel, K., Mower, E., Narayanan, S., 2007. Primitives-based
evaluation and estimation of emotions in speech. *Speech Communication* 49,
787–800. doi:10.1016/j.specom.2007.01.010.

- Halevy, A., Norvig, P., Pereira, F., 2009. The unreasonable effectiveness of data. *IEEE Intelligent Systems* 24, 8–12. doi:10.1109/MIS.2009.36.
- 900 Hu, A., Flaxman, S., 2018. Multimodal sentiment analysis to explore the structure of emotions, in: *Proceedings of the ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, Association for Computing Machinery. pp. 350–358. doi:10.1145/3219819.3219853.
- Jin, Q., Li, C., Chen, S., Wu, H., 2015. Speech emotion recognition with acoustic and lexical features, in: *2015 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, IEEE. pp. 4749–4753. doi:10.1109/ICASSP.2015.7178872, arXiv:arXiv:1011.1669v3.
- Jin, X., Wang, Z., 2005. An Emotion Space Model for Recognition of Emotions in Spoken Chinese, in: *Lecture Notes in Computer Science (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)*, pp. 397–402. doi:10.1007/11573548_51.
- 910 Jurafsky, D., Martin, J.H., 2017. Lexicons for Sentiment and Affect Extraction, in: *Speech and Language Processing, An Introduction to Natural Language Processing, Computational Linguistics, and Speech Recognition*. 3 ed., pp. 326–345.
- Karadogan, S.G., Larsen, J., 2012. Combining semantic and acoustic features for valence and arousal recognition in speech, in: *2012 3rd International Workshop on Cognitive Information Processing (CIP)*, IEEE. pp. 1–6. doi:10.1109/CIP.2012.6232924.
- 920 Kim, S.M., Valitutti, A., Calvo, R.A., 2010. Evaluation of Unsupervised Emotion Models to Textual Affect Recognition, in: *Workshop on Computational Approaches to Analysis and Generation of Emotion in Text*, pp. 62–70.
- Kleine-cosack, C., 2006. Recognition and Simulation of Emotions. Technical Report.

- 925 Li, X., Akagi, M., 2019. Improving multilingual speech emotion recognition by combining acoustic features in a three-layer model. *Speech Communication* 110, 1–12. doi:10.1016/j.specom.2019.04.004.
- Lotfian, R., Busso, C., 2019. Building Naturalistic Emotionally Balanced Speech Corpus by Retrieving Emotional Speech from Existing Podcast Recordings. *IEEE Transactions on Affective Computing* 10, 471–483. doi:10.1109/TAFFC.2017.2736999.
- 930 Mäntylä, M., Adams, B., Destefanis, G., Graziotin, D., Ortu, M., 2016. Mining valence, arousal, and dominance, in: *Proceedings of the 13th International Workshop on Mining Software Repositories - MSR '16*, ACM Press, New York, New York, USA. pp. 247–258. doi:10.1145/2901739.2901752, arXiv:1603.04287.
- 935 Mehrabian, A., Russell, J.A., 1974. *An approach to environmental psychology*. the MIT Press.
- Mikolov, T., Chen, K., Corrado, G., Dean, J., 2013. Efficient Estimation of Word Representations in Vector Space arXiv:1301.3781.
- 940 Mohammad, S.M., 2016. Sentiment Analysis, in: *Emotion Measurement*. Elsevier, pp. 201–237. doi:10.1016/B978-0-08-100508-8.00009-6, arXiv:0-387-31073-8.
- Moore, J.D., Tian, L., Lai, C., 2014. Word-level emotion recognition using high-level features, in: *International Conference on Intelligent Text Processing and Computational Linguistics*, Springer. pp. 17–31.
- 945 Mozilla, 2019. Project DeepSpeech: A TensorFlow implementation of Baidu’s DeepSpeech architecture, <https://github.com/mozilla/DeepSpeech>.
- Parthasarathy, S., Busso, C., 2017. Jointly Predicting Arousal, Valence and Dominance with Multi-Task Learning, in: *Interspeech*, pp. 1103–1107. doi:10.21437/Interspeech.2017-1494.
- 950

- Pedregosa, F., Varoquaux, G., Gramfort, A., Michel, V., Thirion, B., Grisel, O., Blondel, M., Prettenhofer, P., Weiss, R., Dubourg, V., Vanderplas, J., Passos, A., Cournapeau, D., Brucher, M., Perrot, M., Duchesnay, E., 2011. Scikit-learn: Machine Learning in Python. *Journal of Machine Learning Research* 12, 2825—2830.
- Pennington, J., Socher, R., Manning, C.D., 2014. GloVe: Global Vectors for Word Representation, in: *Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pp. 1532–1543.
- Petrushin, V.A., 1999. Emotion In Speech: Recognition And Application To Call Centers. *Proceedings of artificial neural networks in engineering* 710, 22–30.
- Russell, J.A., 1980. A circumplex model of affect. *Journal of Personality and Social Psychology* 39, 1161–1178. doi:10.1037/h0077714.
- Schmitt, M., Schuller, B., 2018. Deep Recurrent Neural Networks for Emotion Recognition in Speech, in: *DAGA*, pp. 1537–1540. doi:10.3390/s17112556.
- Sridhar, K., Parthasarathy, S., Busso, C., 2018. Role of Regularization in the Prediction of Valence from Speech, in: *Interspeech 2018, ISCA, ISCA*. pp. 941–945. doi:10.21437/Interspeech.2018-2508.
- Tripathi, S., Beigi, H., 2018. Multi-Modal Emotion recognition on IEMOCAP Dataset using Deep Learning. *CoRR arXiv:1804.05788*.
- Tris Atmaja, B., Akagi, M., 2020. The Effect of Silence Feature in Dimensional Speech Emotion Recognition, in: *10th International Conference on Speech Prosody 2020, ISCA, ISCA*. pp. 26–30. doi:10.21437/SpeechProsody.2020-6.
- Yang, Z., Hirschberg, J., 2018. Predicting Arousal and Valence from Waveforms and Spectrograms Using Deep Neural Networks, in: *Interspeech 2018, ISCA, ISCA*. pp. 3092–3096. doi:10.21437/Interspeech.2018-2397.

- 980 Yang, Z., Yang, D., Dyer, C., He, X., Smola, A., Hovy, E., 2016. Hierarchical attention networks for document classification, in: 2016 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, NAACL HLT 2016 - Proceedings of the Conference, Association for Computational Linguistics, Stroudsburg, PA, USA. pp. 1480–1489. doi:10.18653/v1/N16-1174.
- 985 Ye, W., Fan, X., 2014. Bimodal Emotion Recognition from Speech and Text. International Journal of Advanced Computer Science and Applications 5. doi:10.14569/ijacsa.2014.050204.
- Yoon, S., Byun, S., Jung, K., 2018. Multimodal Speech Emotion Recognition Using Audio and Text. Spoken Language Technology Workshop , 112–118arXiv:1810.04635.
- 990 Zadeh, A., Chen, M., Poria, S., Cambria, E., Morency, L.P., 2017. Tensor Fusion Network for Multimodal Sentiment Analysis, in: Proceedings of the 2017 Conference on Empirical Methods in Natural Language Processing, Association for Computational Linguistics, Stroudsburg, PA, USA. pp. 1103–1114. doi:10.18653/v1/D17-1115, arXiv:1707.07250.
- 995 Zhang, B., Khorram, S., Provost, E.M., 2019. Exploiting Acoustic and Lexical Properties of Phonemes to Recognize Valence from Speech, in: ICASSP, IEEE International Conference on Acoustics, Speech and Signal Processing - Proceedings, Institute of Electrical and Electronics Engineers Inc.. pp. 5871–5875. doi:10.1109/ICASSP.2019.8683190.
- 1000