

Estimating the Uncertainty in Emotion Attributes using Deep Evidential Regression

Anonymous ACL submission

Abstract

In automatic emotion recognition (AER), labels assigned by different human annotators to the same utterance are often inconsistent due to the inherent complexity of emotion and the subjectivity of perception. Though deterministic labels generated by averaging or voting are often used as the ground truth, it ignores the intrinsic uncertainty revealed by the inconsistent labels. This paper proposes a Bayesian approach, deep evidential emotion regression (DEER), to estimating the uncertainty in emotion attributes. Treating the emotion attribute labels of an utterance as samples drawn from an unknown Gaussian distribution, DEER places an utterance-specific normal-inverse gamma prior over the Gaussian likelihood and predicts its hyper-parameters using a deep neural network model. It enables a joint estimation of emotion attributes along with the aleatoric and epistemic uncertainties. AER experiments on the widely used MSP-Podcast and IEMOCAP datasets showed DEER produced state-of-the-art results for both the mean values and the distribution of emotion attributes¹.

1 Introduction

Automatic emotion recognition (AER) is the task that enables computers to predict human emotional states based on multimodal signals, such as audio, video and text. An emotional state is defined based on either categorical or dimensional theory. The categorical theory claims the existence of a small number of basic discrete emotions (*i.e.* anger and happy) that are inherent in our brain and universally recognised (Gunes et al., 2011; Plutchik, 2001). Dimensional emotion theory characterises emotional states by a small number of roughly orthogonal fundamental continuous-valued bipolar dimensions (Schlosberg, 1954; Nicolaou et al., 2011) such as valence-arousal and approach-avoidance (Russell and Mehrabian, 1977;

Russell, 1980; Grimm et al., 2007). These dimensions are also known as emotion attributes, which allow us to model more subtle and complex emotions and are thus more common in psychological studies. As a result, AER includes a classification approach based on emotion-class-based labels and a regression approach based on attribute-based labels. This paper focuses on attribute-based AER with speech input.

Emotion annotation is challenging due to the inherent ambiguity of mixed emotion, the personal variations in emotion expression, the subjectivity in emotion perception, *etc.* Most AER datasets use multiple human annotators to label each utterance, which often results in inconsistent labels, either as emotion categories or attributes. This is also a typical manifestation of the intrinsic data uncertainty, also referred to as aleatoric uncertainty (Matthies, 2007; Der Kiureghian and Ditlevsen, 2009), that arises from the natural complexity of emotion data. It is common to replace such inconsistent labels with deterministic labels obtained by majority voting (Busso et al., 2008, 2017) or (weighted) averages (Ringeval et al., 2013; Lotfian and Busso, 2019; Kossaifi et al., 2019; Grimm and Kroschel, 2005). However, this causes a loss of data samples when a majority agreed emotion class doesn't exist (Majumder et al., 2018; Poria et al., 2018; Wu et al., 2021) and also ignores the discrepancies between annotators and the aleatoric uncertainty in emotion data.

In this paper, we propose to model the uncertainty in emotion attributes with a Bayesian approach based on deep evidential regression (Amini et al., 2020), denoted deep evidential emotion regression (DEER). In DEER, the inconsistent human labels of each utterance are considered as observations drawn independently from an unknown Gaussian distribution. To probabilistically estimate the mean and variance of the Gaussian distribution, a normal inverse-gamma (NIG) prior is introduced,

¹Code will be publicly available upon acceptance.

which places a Gaussian prior over the mean and an inverse-gamma prior over the variance. The AER system is trained to predict the hyper-parameters of the NIG prior for each utterance by maximising the per-observation-based marginal likelihood of each observed label under this prior. As a result, DEER not only models the distribution of emotion attributes but also learns both the aleatoric uncertainty and the epistemic uncertainty (Der Kiureghian and Ditlevsen, 2009) without repeating the inference procedure for sampling. Epistemic uncertainty, also known as model uncertainty, is associated with uncertainty in model parameters that best explain the observed data. Aleatoric and epistemic uncertainty are combined to induce the total uncertainty, also called predictive uncertainty, that measures the confidence of attribute predictions. As a further improvement, a novel regulariser is proposed based on the mean and variance of the observed labels to better calibrate the uncertainty estimation. The proposed methods were evaluated on the MSP-Podcast and IEMOCAP datasets.

The rest of the paper is organised as follows. Section 2 summarises related work. Section 3 introduces the proposed DEER approach. Sections 4 and 5 present the experimental setup and results respectively, followed by the conclusion.

2 Related Work

There has been previous work by AER researchers to address the issue of inconsistent labels. For emotion categories, a single ground-truth label can be obtained as either a continuous-valued mean vector representing emotion intensities (Fayek et al., 2016; Ando et al., 2018), or as a multi-hot vector obtained based on the existence of emotions (Zhang et al., 2020; Ju et al., 2020). Recently, distribution-based approaches have been proposed, which consider the labels as samples drawn from emotion distributions (Chou et al., 2022; Wu et al., 2022b).

For emotion attributes, annotators often assign different values to the same attribute of each utterance. Davani et al. (2022) proposed a multi-annotator model which contains multiple heads to predict each annotator’s judgement. This approach is computationally viable only when the number of annotators is relatively small. The method requires sufficient annotations from each annotator to be effective. Deng et al. (2012) derived confidence measures based on annotator agreement to build emotion-scoring models. Han et al. (2017, 2021)

proposed predicting the standard deviation of the attribute label values as an extra task in the multi-task training framework. Dang et al. (2017, 2018) included annotator variability as a representation of uncertainty in a Gaussian mixture regression model. These techniques take the variance of human annotations either as an extra target or as an extra input. More recently, Bayesian deep learning has been introduced to the task, which models the uncertainty in emotion annotation without explicitly using the variance of human annotations. These include the use of Gaussian processes (Atcheson et al., 2018, 2019), variational auto-encoders (Sridhar et al., 2021), Bayesian neural networks (Prabhu et al., 2021), Monte-Carlo dropout (Sridhar and Busso, 2020b) and sequential Monte-Carlo methods (Markov et al., 2015; Wu et al., 2022a).

So far, these methods have not distinguished aleatoric uncertainty from epistemic uncertainty which are defined in the introduction. Our proposed DEER approach can simultaneously model these two uncertainties. In addition, our approach is more generic. It has no limits on the number of annotators, the number of annotators per utterance, and the number of annotations per annotator, and thus can cope with large crowd-sourced datasets.

3 Deep Evidential Emotion Regression

3.1 Problem setup

In contrast to Bayesian neural networks that place priors on model parameters (Blundell et al., 2015; Kendall and Gal, 2017), evidential deep learning (Sensoy et al., 2018; Malinin and Gales, 2018; Amini et al., 2020) places priors over the likelihood function. Every training sample adds support to a learned higher-order prior distribution called the evidential distribution. Sampling from this distribution gives instances of lower-order likelihood functions from which the data was drawn.

Consider an input utterance $\mathbf{x}$ with M emotion attribute labels $y^{(1)}, \dots, y^{(M)}$ provided by multiple annotators. Assuming $y^{(1)}, \dots, y^{(M)}$ are observations drawn *i.i.d.* from a Gaussian distribution with unknown mean μ and unknown variance σ^2 , where μ is drawn from a Gaussian prior and σ^2 is drawn from an inverse-gamma prior:

$$y^{(1)}, \dots, y^{(M)} \sim \mathcal{N}(\mu, \sigma^2)$$

$$\mu \sim \mathcal{N}(\gamma, \sigma^2 v^{-1}), \quad \sigma^2 \sim \Gamma^{-1}(\alpha, \beta)$$

where $\gamma \in \mathbb{R}$, $v > 0$, and $\Gamma(\cdot)$ is the gamma function with $\alpha > 1$ and $\beta > 0$.

Denote $\{\mu, \sigma^2\}$ and $\{\gamma, v, \alpha, \beta\}$ as Ψ and Ω . The posterior $p(\Psi|\Omega)$ is a NIG distribution, which is the Gaussian conjugate prior:

$$\begin{aligned} p(\Psi|\Omega) &= p(\mu|\sigma^2, \Omega) p(\sigma^2|\Omega) \\ &= \mathcal{N}(\gamma, \sigma^2 v^{-1}) \Gamma^{-1}(\alpha, \beta) \\ &= \frac{\beta^\alpha \sqrt{v}}{\Gamma(\alpha) \sqrt{2\pi\sigma^2}} \left(\frac{1}{\sigma^2}\right)^{\alpha+1} \\ &\quad \cdot \exp\left\{-\frac{2\beta + v(\gamma - \mu)^2}{2\sigma^2}\right\} \end{aligned}$$

Drawing a sample Ψ_i from the NIG distribution yields a single instance of the likelihood function $\mathcal{N}(\mu_i, \sigma_i^2)$. The NIG distribution therefore serves as the higher-order, evidential distribution on top of the unknown lower-order likelihood distribution from which the observations are drawn. The NIG hyper-parameters Ω determine not only the location but also the uncertainty, associated with the inferred likelihood function.

By training a deep neural network model to output the hyper-parameters of the evidential distribution, evidential deep learning allows the uncertainties to be found by analytic computation of the maximum likelihood Gaussian without the need for repeated inference for sampling (Amini et al., 2020). Furthermore, it also allows an effective estimate of the aleatoric uncertainty computed as the expectation of the variance of the Gaussian distribution, as well as the epistemic uncertainty defined as the variance of the predicted Gaussian mean. Given an NIG distribution, the prediction, aleatoric, and epistemic uncertainty can be computed as:

$$\begin{aligned} \text{Prediction: } \mathbb{E}[\mu] &= \gamma \\ \text{Aleatoric: } \mathbb{E}[\sigma^2] &= \frac{\beta}{\alpha - 1}, \quad \forall \alpha > 1 \\ \text{Epistemic: } \text{Var}[\mu] &= \frac{\beta}{v(\alpha - 1)}, \quad \forall \alpha > 1 \end{aligned}$$

3.2 Training

The training of DEER is structured as fitting the model to the data while enforcing the prior to calibrate the uncertainty when the prediction is wrong.

3.2.1 Maximising the data fit

The likelihood of an observation y given the evidential distribution hyper-parameters Ω is computed by marginalising over the likelihood parameters Ψ :

$$\begin{aligned} p(y|\Omega) &= \int_{\Psi} p(y|\Psi) p(\Psi|\Omega) d\Psi \\ &= \mathbb{E}_{p(\Psi|\Omega)}[p(y|\Psi)] \end{aligned} \quad (1)$$

An analytical solution exists in the case of placing an NIG prior on the Gaussian likelihood function:

$$\begin{aligned} p(y|\Omega) &= \frac{\Gamma(1/2 + \alpha)}{\Gamma(\alpha)} \sqrt{\frac{v}{\pi}} (2\beta(1 + v))^\alpha \\ &\quad \cdot (v(y - \gamma)^2 + 2\beta(1 + v))^{-(\frac{1}{2} + \alpha)} \\ &= \text{St}_{2\alpha}\left(y|\gamma, \frac{\beta(1 + v)}{v\alpha}\right) \end{aligned} \quad (2)$$

where $\text{St}_\nu(t|r, s)$ is the Student's t-distribution evaluated at t with location parameter r , scale parameter s , and ν degrees of freedom. The predicted mean and variance can be computed analytically as

$$\mathbb{E}[y] = \gamma, \quad \text{Var}[y] = \frac{\beta(1 + v)}{v(\alpha - 1)} \quad (3)$$

$\text{Var}[y]$ represents the total uncertainty of model prediction, which is equal to the summation of the aleatoric uncertainty $\mathbb{E}[\sigma^2]$ and epistemic uncertainty $\text{Var}[\mu]$ according to the law of total variance:

$$\begin{aligned} \text{Var}[y] &= \mathbb{E}[\text{Var}[y|\Psi]] + \text{Var}[\mathbb{E}[y|\Psi]] \\ &= \mathbb{E}[\sigma^2] + \text{Var}[\mu] \end{aligned}$$

To fit the NIG distribution, the model is trained by maximising the sum of the marginal likelihoods of each human label $y^{(m)}$. The negative log likelihood (NLL) loss can be computed as

$$\begin{aligned} \mathcal{L}^{\text{NLL}}(\mathbf{x}; \Theta) &= -\frac{1}{M} \sum_{m=1}^M \log p(y^{(m)}|\Omega) \\ &= -\frac{1}{M} \sum_{m=1}^M \log \left[\text{St}_{2\alpha}\left(y^{(m)}|\gamma, \frac{\beta(1 + v)}{v\alpha}\right) \right] \end{aligned} \quad (4)$$

This is our proposed per-observation-based NLL loss, which takes each observed label into consideration for AER. This loss serves as the first part of the objective function for training a deep neural network model Θ to predict the hyper-parameters $\{\gamma, v, \alpha, \beta\}$ to fit all observed labels of $\mathbf{x}$.

3.2.2 Calibrating the uncertainty on errors

The second part of the objective function regularises training by calibrating the uncertainty based on the incorrect predictions. A novel regulariser is formulated which contains two terms: $\mathcal{L}^\mu$ and $\mathcal{L}^\sigma$ that respectively regularises the errors on the estimation of the mean μ and the variance σ^2 of the Gaussian likelihood.

The first term $\mathcal{L}^\mu$ is proportional to the error between the model prediction and the average of the observations:

$$\mathcal{L}^\mu(\mathbf{x}; \Theta) = \Phi |\bar{y} - \mathbb{E}[\mu]|$$

where $\|\cdot\|$ is L1 norm, $\bar{y} = \frac{1}{M} \sum_{m=1}^M y^{(m)}$ is the averaged label which is usually used as the ground truth in regression-based AER, and Φ is an uncertainty measure associated with the inferred posterior. The reciprocal of the total uncertainty is used as Φ in this paper, which can be calculated as

$$\Phi = \frac{1}{\text{Var}[y]} = \frac{v(\alpha - 1)}{\beta(1 + v)}$$

The regulariser imposes a penalty when there's an error in prediction and dynamically scales it by dividing by the total uncertainty of inferred posterior. It penalises the cases where the model produces an incorrect prediction with a small uncertainty, thus preventing the model from being over-confident. For instance, if the model produces an error with a small predicted variance, Φ is large, resulting in a large penalty. Minimising the regularisation term enforces the model to produce accurate prediction or increase uncertainty when the error is large.

In addition to imposing a penalty on the mean prediction as in (Amini et al., 2020), a second term $\mathcal{L}^\sigma$ is proposed in order to calibrate the estimation of the aleatoric uncertainty. As discussed in the introduction, aleatoric uncertainty in AER is shown by the different emotional labels given to the same utterance by different human annotators. This paper uses the variance of the observations to describe the aleatoric uncertainty in the emotion data. The second regularising term is defined as:

$$\mathcal{L}^\sigma(\mathbf{x}; \Theta) = \Phi |\bar{\sigma}^2 - \mathbb{E}[\sigma^2]|$$

where $\bar{\sigma}^2 = \frac{1}{M} \sum_{m=1}^M (y^{(m)} - \bar{y})^2$.

3.3 Summary and implementation details

For an AER task that consists of N emotion attributes, DEER trains a deep neural network model to simultaneously predict the hyperparameters $\{\Omega_1, \dots, \Omega_N\}$ associated with the N attribute-specific NIG distributions, where $\Omega_n = \{\gamma_n, v_n, \alpha_n, \beta_n\}$. A DEER model thus has $4N$ output units. The system is trained by minimising the total loss *w.r.t.* Θ as:

$$\mathcal{L}_{\text{total}}(\mathbf{x}; \Theta) = \sum_{n=1}^N \epsilon_n \mathcal{L}_n(\mathbf{x}; \Theta) \quad (5)$$

$$\mathcal{L}_n(\mathbf{x}; \Theta) = \mathcal{L}_n^{\text{NLL}}(\mathbf{x}; \Theta) + \lambda_n [\mathcal{L}_n^\mu(\mathbf{x}; \Theta) + \mathcal{L}_n^\sigma(\mathbf{x}; \Theta)] \quad (6)$$

where ϵ_n is the weight satisfying $\sum_{n=1}^N \epsilon_n = 1$, λ_n is the scale coefficient that trades off the training between data fit and uncertainty regulation.

At test-time, the predictive posteriors are N separate Student's t-distributions $p(y|\Omega_1), p(y|\Omega_2), \dots, p(y|\Omega_N)$, each of the same form as derived in Eqn. (2)². Apart from obtaining a distribution over the emotion attribute of the speaker, DEER also allows analytic computation of the uncertainty terms, as summarised in Table 1.

Term	Expression
Predicted mean	$\mathbb{E}[y] = \mathbb{E}[\mu] = \gamma$
Predicted variance (Total uncertainty)	$\text{Var}[y] = \frac{\beta(1+v)}{v(\alpha-1)}$
Aleatoric uncertainty	$\mathbb{E}[\sigma^2] = \frac{\beta}{\alpha-1}$
Epistemic uncertainty	$\text{Var}[\mu] = \frac{\beta}{v(\alpha-1)}$

Table 1: Summary of the uncertainty terms.

4 Experimental Setup

4.1 Dataset

The MSP-Podcast (Lotfian and Busso, 2019) and IEMOCAP datasets (Busso et al., 2008) were used in this paper. The annotations of both datasets use $N = 3$ with valence, arousal (also called activation), and dominance as the emotion attributes. MSP-Podcast contains natural English speech from podcast recordings and is one of the largest publicly available datasets in speech emotion recognition. A seven-point Likert scale was used to evaluate valence (1-negative vs 7-positive), arousal (1-calm vs 7-active), and dominance (1-weak vs 7-strong). The corpus was annotated using crowd-sourcing. Each utterance was labelled by at least 5 human annotators and has an average of 6.7 annotations per utterance. Ground-truth labels were defined by the average value. Release 1.8 was used in the experiments, which contains 73,042 utterances from 1,285 speakers amounting to more than 110 hours of speech. The average variance of the labels assigned to each sentence is 0.975, 1.122, 0.889 for valence, arousal, and dominance respectively. The standard splits for training (44,879 segments), validation (7,800 segments) and testing (15,326 segments) were used in the experiments.

The IEMOCAP corpus is one of the most widely used AER datasets. It consists of approximately 12 hours of English speech including 5 dyadic conversational sessions performed by 10 professional

²Since NIG is the Gaussian conjugate prior, the posterior is in the same parametric family as the prior. Therefore, the predictive posterior has the same form as the marginal likelihood. Detailed derivations see Appendix A.

actors with a session being a conversation between two speakers. There are in total 151 dialogues including 10,039 utterances. Each utterance was annotated by three human annotators using a five-point Likert scale. Again, ground-truth labels were determined by taking the average. The average variance of the labels assigned to each sentence is 0.130, 0.225, 0.300 for valence, arousal, and dominance respectively. Unless otherwise mentioned, systems on IEMOCAP were evaluated by training on Session 1-4 and testing on Session 5.

4.2 Model structure

The model structure used in this paper follows the upstream-downstream framework (Yang et al., 2021), as illustrated in Figure 1. WavLM (Chen et al., 2022) was used as the upstream model, which is a speech foundation model pre-trained by self-supervised learning. The BASE+ version³ of the model was used in this paper which has 12 Transformer encoder blocks with 768-dimensional hidden states and 8 attention heads. The parameters of the pre-trained model were frozen and the weighted sum of the outputs of the 12 Transformer encoder blocks was used as the speech embeddings and fed into the downstream model.

The downstream model consists of two 128-dimensional Transformer encoder blocks with 4-head self-attention, followed by an evidential layer that contains four output units for each of the three attributes, which has a total of 12 output units. The model contains 0.3M trainable parameters. A Softplus activator⁴ was applied to $\{v, \alpha, \beta\}$ to ensure $v, \alpha, \beta > 0$ with an additional +1 added to α to ensure $\alpha > 1$. A linear activation was used for $\gamma \in \mathbb{R}$. The proposed DEER model was trained to simultaneously learn three evidential distributions for the three attributes. The weights in Eqn. (5) were set as $\epsilon_v = \epsilon_a = \epsilon_d = 1/3$. The scale coefficients were set to $\lambda_v = \lambda_a = \lambda_d = 0.1$ for Eqn. (6)⁵.

A dropout rate of 0.3 was applied to the transformer parameters. The system was implemented using PyTorch and the SpeechBrain toolkit (Ravanelli et al., 2021). The Adam optimizer was used with an initial learning rate set to 0.001. Training took ~ 8 hours on an NVIDIA A100 GPU.

³<https://huggingface.co/microsoft/wavlm-base-plus>

⁴Softplus(x) = $\ln(1 + \exp(x))$

⁵The values were manually selected from a small number of candidates.

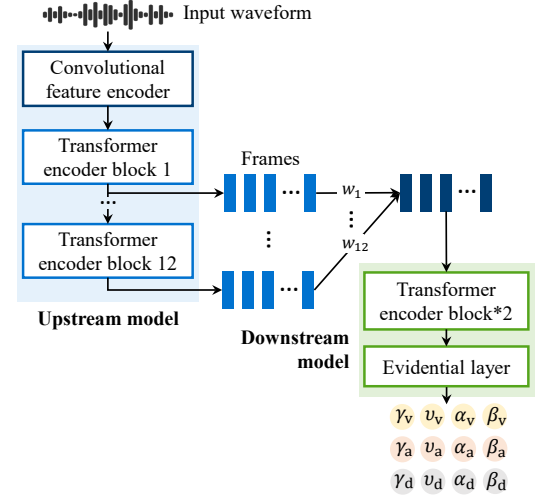


Figure 1: Illustration of the model structure. Weights $w_1, \dots, w_{12}$ for the weighted sum of the 12 Transformer encoder outputs are trainable and satisfy $\sum_{i=1}^{12} w_i = 1$.

4.3 Evaluation metrics

4.3.1 Mean prediction

Following prior work in continuous emotion recognition (Ringeval et al., 2015, 2017; Sridhar and Busso, 2020a; Leem et al., 2022), the concordance correlation coefficient (CCC) was used to evaluate the predicted mean. CCC combines the Pearson’s correlation coefficient with the square difference between the mean of the two compared sequences:

$$\rho_{ccc} = \frac{2\rho\sigma_{\text{ref}}\sigma_{\text{hyp}}}{\sigma_{\text{ref}}^2 + \sigma_{\text{hyp}}^2 + (\mu_{\text{ref}} - \mu_{\text{hyp}})^2},$$

where ρ is the Pearson correlation coefficient between a hypothesis sequence (system predictions) and a reference sequence, where μ_{hyp} and μ_{ref} are the mean values, and σ_{hyp}^2 and σ_{ref}^2 are the variance values of the two sequences. Hypotheses that are well correlated with the reference but shifted in value are penalised in proportion to the deviation. The value of CCC ranges from -1 (perfect disagreement) to 1 (perfect agreement).

The root mean square error (RMSE) averaged over the test set is also reported. Since the average of the human labels, $\bar{y}$, is defined as the ground truth in both datasets, $\bar{y}$ were used as the reference in computing the CCC and RMSE. However, using $\bar{y}$ also indicates that these metrics are less informative when the aleatoric uncertainty is large.

4.3.2 Uncertainty estimation

It is common to use NLL to measure the uncertainty estimation ability (Gal and Ghahramani, 2016; Amini et al., 2020). NLL is computed by fitting data to the predictive posterior $q(y)$.

MSP-Podcast	CCC $\uparrow$			RMSE $\downarrow$			NLL(avg) $\downarrow$			NLL(all) $\downarrow$		
	v	a	d	v	a	d	v	a	d	v	a	d
$\mathcal{L}$ in Eqn. (6)	0.506	0.698	0.613	0.772	0.680	0.576	1.334	1.285	1.156	1.696	1.692	1.577
$\mathcal{L}^\sigma = 0$	0.451	0.687	0.607	0.784	0.679	0.580	1.345	1.277	1.159	1.706	1.705	1.586
$\mathcal{L}^{\text{NLL}} = \bar{\mathcal{L}}^{\text{NLL}}$	0.473	0.682	0.609	0.808	0.673	0.566	1.290	1.060	0.899	2.027	2.089	1.969
IEMOCAP	v			a			d			v		
	v	a	d	v	a	d	v	a	d	v	a	d
$\mathcal{L}$ in Eqn. (6)	0.596	0.755	0.569	0.755	0.457	0.638	1.070	0.795	1.035	1.275	1.053	1.283
$\mathcal{L}^\sigma = 0$	0.582	0.752	0.553	0.772	0.466	0.655	1.180	0.773	1.061	1.408	1.069	1.294
$\mathcal{L}^{\text{NLL}} = \bar{\mathcal{L}}^{\text{NLL}}$	0.585	0.759	0.555	0.786	0.444	0.633	1.001	0.727	1.036	1.627	1.329	1.441

Table 2: DEER results variations of the loss in Eqn. (6). ‘v’, ‘a’, ‘d’ stands for valence, arousal, dominance. ‘ $\uparrow$ ’ denotes the higher the better, ‘ $\downarrow$ ’ denotes the lower the better. The ‘ $\mathcal{L}$ in Eqn. (6)’ row systems used the complete total loss of DEER. The ‘ $\mathcal{L}^\sigma = 0$ ’ row systems had no $\mathcal{L}^\sigma$ regularisation term in the total loss. The ‘ $\mathcal{L}^{\text{NLL}} = \bar{\mathcal{L}}^{\text{NLL}}$ ’ row systems replaced the individual human labels with $\bar{\mathcal{L}}^{\text{NLL}}$ in the total loss.

In this paper, NLL(avg) defined as $-\log q(\bar{y})$ and NLL(all) defined as $-\frac{1}{M} \sum_{m=1}^M \log q(y^{(m)})$ are both used. NLL(avg) measures how much the averaged label $\bar{y}$ fits into the predicted posterior distribution, and NLL(all) measures how much every single human label $y^{(m)}$ fits into the predicted posterior. A lower NLL indicates better uncertainty estimation.

5 Experiments and Results

5.1 Effect of the aleatoric regulariser $\mathcal{L}^\sigma$

First, by setting $\mathcal{L}^\sigma = 0$ in the total loss, an ablation study of the effect of the proposed extra regularising term $\mathcal{L}^\sigma$ is performed. The results are given in the ‘ $\mathcal{L}^\sigma = 0$ ’ rows in Table 2. In this case, only $\mathcal{L}^\mu$ is used to regularise $\mathcal{L}^{\text{NLL}}$ and the results are compared to those trained using the complete loss defined in Eqn. (6), which are shown in the ‘ $\mathcal{L}$ in Eqn. (6)’ rows. From the results, $\mathcal{L}^\sigma$ improves the performance in CCC and NLL(all), but not in NLL(ref), as expected.

5.2 Effect of the per-observation-based $\mathcal{L}^{\text{NLL}}$

Next, the effect of our proposed per-observation-based NLL loss defined in Eqn. (4), $\mathcal{L}^{\text{NLL}}$, is compared to an alternative. Instead of using $\mathcal{L}^{\text{NLL}}$,

$$\bar{\mathcal{L}}^{\text{NLL}} = -\log p(\bar{y}|\Omega)$$

is used to compute the total loss during training, and the results are given in the ‘ $\mathcal{L}^{\text{NLL}} = \bar{\mathcal{L}}^{\text{NLL}}$ ’ rows in Table 2. While $\mathcal{L}^{\text{NLL}}$ considers the likelihood of fitting each individual observation into the predicted posterior, $\bar{\mathcal{L}}^{\text{NLL}}$ only considers the averaged observation. Therefore, it is expected that using $\bar{\mathcal{L}}^{\text{NLL}}$ instead of $\mathcal{L}^{\text{NLL}}$ yields a smaller NLL(avg) but larger NLL(all), which have been validated by the results in the table.

5.3 Baseline comparisons

Three baseline systems were built:

- A Gaussian Process (GP) with a radial basis function kernel, trained by maximising the per-observation-based marginal likelihood.
- A Monte Carlo dropout (MCdp) system with a dropout rate of 0.4. During inference, the system was forwarded 50 times with different dropout random seeds to obtain 50 samples.
- An ensemble of 10 systems initialised and trained with 10 different random seeds.

The MCdp and ensemble baselines used the same model structure as the DEER system, expect that the evidential output layer was replaced by a standard fully-connected output layer with three output units to predict the values of valence, arousal and dominance respectively. Following prior work (Al-Badawy and Kim, 2018; Atmaja and Akagi, 2020b; Sridhar and Busso, 2020b), the CCC loss,

$$\mathcal{L}_{\text{ccc}} = 1 - \rho_{\text{ccc}}$$

was used for training the MCdp and ensemble baselines. The CCC loss was computed based on the sequence within each mini-batch of training data. The CCC loss has been shown by previous studies to improve the continuous emotion predictions compared to the RMSE loss (Povolny et al., 2016; Trigeorgis et al., 2016; Le et al., 2017). For MCdp and ensemble, the predicted distribution of the emotion attributes were estimated based on the obtained samples by kernel density estimation.

The results are listed in Table 3. The proposed DEER system outperforms the baselines on most of the attributes and the overall values. In particular, DEER outperforms all baselines consistently in the NLL(all) metric.

MSP-Podcast	CCC $\uparrow$			RMSE $\downarrow$			NLL_ref $\downarrow$			NLL_all $\downarrow$		
	v	a	d	v	a	d	v	a	d	v	a	d
DEER	0.506	0.698	0.613	0.772	0.680	0.576	1.334	1.285	1.156	1.696	1.692	1.577
GP	0.342	0.595	0.486	0.811	0.673	0.566	1.447	1.408	1.297	1.727	1.808	1.592
MCdp	0.476	0.667	0.594	0.874	0.702	0.623	1.680	1.300	1.071	2.050	2.027	1.776
Ensemble	0.511	0.679	0.608	0.855	0.692	0.615	1.864	1.384	1.112	2.096	2.066	1.795
IEMOCAP	v			a			d			v		
	v	a	d	v	a	d	v	a	d	v	a	d
DEER	0.596	0.756	0.569	0.755	0.457	0.638	1.070	0.795	1.035	1.275	1.053	1.283
GP	0.535	0.717	0.512	0.763	0.479	0.657	1.209	0.791	1.047	1.295	1.205	1.380
MCdp	0.539	0.724	0.568	0.786	0.561	0.702	1.291	0.849	1.133	1.549	1.325	1.747
Ensemble	0.580	0.754	0.560	0.778	0.476	0.686	1.296	0.864	1.110	1.584	1.218	1.749

Table 3: Comparison with the baselines. ‘v’, ‘a’, ‘d’ stands for valence, arousal, dominance. ‘ $\uparrow$ ’ denotes the higher the better, ‘ $\downarrow$ ’ denotes the lower the better. Best results in each column shown in bold.

5.4 Cross comparison of mean prediction

Table 4 compares results obtained with those previously published in terms of the CCC value. Previous papers have reported results on both version 1.6 and 1.8 of the MSP-Podcast dataset. For comparison, we also conducted experiments on version 1.6 for comparison. Version 1.6 of MSP-Podcast database is a subset of version 1.8 and contains 34,280 segments for training, 5,958 segments for validation and 10,124 segments for testing. For IEMOCAP, apart from training on Session 1-4 and testing on Session 5 (Ses05), we also evaluated the proposed system by a 5-fold cross-validation (5CV) based on a “leave-one-session-out” strategy. In each fold, one session was left out for testing and the others were used for training. The configuration is speaker-exclusive for both settings. As shown in Table 4, our DEER systems achieved state-of-the-art (SOTA) results both versions of MSP-Podcast and both test setting of IEMOCAP.

5.5 Analysis of uncertainty estimation

5.5.1 Visualisation

Based on a randomly selected subset test set of MSP-Podcast version 1.8, the aleatoric, epistemic and total uncertainty of the dominance attribute predicted by our proposed DEER system are shown in Figure 2.

Figure 2 (a) shows the predicted mean $\pm$ square root of the predicted aleatoric uncertainty ($\mathbb{E}[\mu] \pm \sqrt{\mathbb{E}[\sigma^2]}$) and the average label $\pm$ the standard deviation of the human labels ($\bar{y} \pm \bar{\sigma}$). It can be seen that the predicted aleatoric uncertainty (blue) overlaps with the label standard deviation (grey) and the overlapping is more evident when the mean predictions are accurate (*i.e.* samples around index

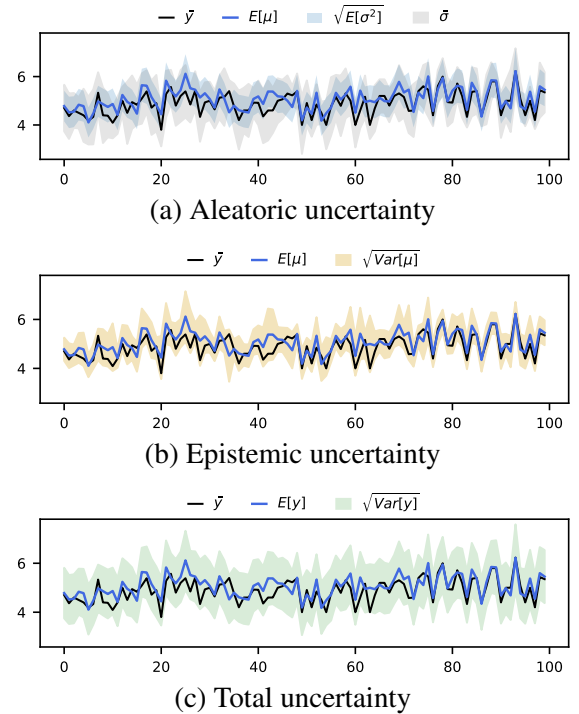


Figure 2: Visualisation of (a) aleatoric (b) epistemic (c) total uncertainty of dominance for MSP-Podcast. x -axis is the test utterance index.

80-100).

Figure 2 (b) shows the predicted mean $\pm$ square root of the predicted epistemic uncertainty ($\mathbb{E}[\mu] \pm \sqrt{\text{Var}[\mu]}$). The epistemic uncertainty is high when the predicted mean deviates from the target (*i.e.* samples around index 40-50) while low then the predicted mean matches the target (*i.e.* samples around index 80-100).

Figure 2 (c) shows the predicted mean $\pm$ square root of the total epistemic uncertainty ($\mathbb{E}[y] \pm \sqrt{\text{Var}[y]}$) which combines the aleatoric and epistemic uncertainty. The total uncertainty is high

	Paper	Version	v	a	d	Average
MSP-podcast	Ghriss et al. (2022)	1.6	0.412	0.679	0.564	0.552
	Mitra et al. (2022)	1.6	0.57	0.75	0.67	0.663
	Srinivasan et al. (2022)	1.6	0.627	0.757	0.671	0.685
	DEER	1.6	0.629	0.777	0.684	0.697
	Leem et al. (2022)	1.8	0.212	0.572	0.505	0.430
	DEER	1.8	0.506	0.698	0.613	0.606
	Paper	Setting	v	a	d	Average
IEMOCAP	Atmaja and Akagi (2020a)	Ses05	0.421	0.590	0.484	0.498
	Atmaja and Akagi (2021)	Ses05	0.553	0.579	0.465	0.532
	DEER	Ses05	0.596	0.756	0.569	0.640
	Srinivasan et al. (2022)	5CV	0.582	0.667	0.545	0.598
	DEER	5CV	0.625	0.720	0.548	0.631

Table 4: Cross comparison of the CCC value on MSP-Podcast and IEMOCAP. ‘v’, ‘a’, ‘d’ stands for valence, arousal, dominance. ‘Version’ of MSP-Podcast denotes the release version of the dataset., and only the results from the same dataset version are comparable. ‘Test set’ of IEMOCAP denotes the train/set split. ‘Ses05’ denotes training on Session 1-4 and tested on Session 5. ‘5CV’ denotes leave-one-session-out 5-fold cross validation.

either when the input utterance is complex or the model is not confident.

5.5.2 Reject option

A reject option was applied to analyse the uncertainty estimation performance, where the system has the option to accept or decline a test sample based on the uncertainty prediction. Since the evaluation of CCC is based on the whole sequence rather than individual samples, its computation would be affected when the sequence is modified by rejection (Wu et al., 2022a). Therefore, the reject option is performed based on RMSE.

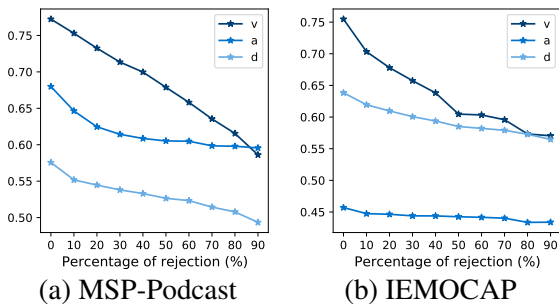


Figure 3: Reject Option of RMSE based on predicted variance for (a) MSP-Podcast and (b) IEMOCAP.

Confidence is measured by the total uncertainty given in Eqn. (3). Figure 3 shows the performance of the proposed DEER system with a reject option on MSP-Podcast and IEMOCAP. A percentage of utterances with the largest predicted variance were rejected. The results at 0% rejection corresponds to the RMSE achieved on the entire test data. As the percentage of rejection increases, test coverage

decreases and the average RMSE decreases showing the predicted variance succeeded in confidence estimation. The system then trades off between the test coverage and performance.

6 Conclusions

Two types of uncertainty exist in AER: (i) aleatoric uncertainty arising from the inherent ambiguity of emotion and personal variations in emotion expression; (ii) epistemic uncertainty associated with the estimated network parameters given the observed data. This paper proposes DEER for estimating those uncertainties in emotion attributes. Treating observed attribute-based annotations as samples drawn from a Gaussian distribution, DEER places a normal-inverse gamma (NIG) prior over the Gaussian likelihood. A novel training loss is proposed which combines a per-observation-based NLL loss with a regulariser on both the mean and the variance of the Gaussian likelihood. Experiments on the MSP-Podcast and IEMOCAP datasets show that DEER can produce SOTA results in estimating both the mean value and the distribution of emotion attributes. The use of NIG, the conjugate prior to the Gaussian distribution, leads to tractable analytic computation of the marginal likelihood as well as aleatoric and epistemic uncertainty associated with attribute prediction. Uncertainty estimation is analysed by visualisation and a reject option. Beyond the scope of AER, DEER could also be applied to other tasks with subjective evaluations yielding inconsistent labels.

Limitations

The proposed approach (along with other methods for estimating uncertainty in inconsistent annotations) is only viable when the raw labels from different human annotators for each sentence are provided by the datasets. However, some multiple-annotated datasets only released the majority vote or averaged label for each sentence (*i.e.* Poria et al., 2019).

The proposed method made a Gaussian assumption on the likelihood function for the analytic computation of the uncertainties. The results show that this modelling approach is effective. Despite the effectiveness of the proposed method, other distributions could also be considered.

Apart from aleatoric and epistemic uncertainty, distributional uncertainty can also occur when testing on unseen data. Distributional uncertainty (also called dataset shift) results from the mismatch between the training and test distributions. Aleatoric uncertainty is sometimes considered as “known-unknown” – the knowledge the model learnt from the training set can be well-generalised to the test data and the model knows how difficult it is to predict from the data. Distributional uncertainty is an “unknown-unknown” – the model is unfamiliar with the test data and cannot make predictions confidently. Unseen data is not discussed in this paper but can be an interesting extension to investigate and the DEER method should have the potential to address this issue.

Data collection processes for AER datasets vary in terms of recording conditions, emotional elicitation scheme, and annotation procedure *etc.* This work was tested on two typical datasets: IEMOCAP and MSP-Podcast. The two datasets are both publicly available and differ in various aspects:

- IEMOCAP contains emotion acted by professional actors while MSP-Podcast contains natural emotion.
- IEMOCAP contains dyadic conversations while MSP-Podcast contains Podcast recordings.
- IEMOCAP contains 10 speakers and MSP-Podcast contains 1285 speakers.
- IEMOCAP contains about 12 hours of speech and MSP-Podcast contains more than 110 hours of speech.
- IEMOCAP was annotated by six professional evaluators with each sentence being annotated

by three evaluators. MSP-Podcast was annotated by crowd-sourcing where a total of 11,799 workers were involved and each work annotated 41.5 sentences on average.

The proposed approach has been shown effective over both datasets. We believe the proposed technique should be generic. Furthermore, although validated only for AER, the proposed method could also be applied to other tasks with disagreements in subjective annotations such as hate speech detection and language assessment.

Ethics Statement

For subjective tasks such as emotion recognition, it is common to employ multiple human annotators to give multiple annotations to each data instance. In face of annotator disagreements, majority voting and averaging are commonly used to derive single ground truth labels for training supervised machine learning models. However, in many subjective tasks, there is usually no single “correct” answer. And a potential risk of enforcing a single ground truth is ignoring the valuable nuance in each annotator’s evaluation and their disagreements. This can cause minority views to be under-represented. The DEER approach proposed in this work could be beneficial to this concern as it models uncertainty in annotator disagreements and provides some explainability of the predictions.

While our method helps preserve minority perspectives, misuse of this technique might lead to ethical concerns. Emotion recognition is at risk of exposing a person’s inner state to others and this information could be abused. Furthermore, since the proposed approach takes each annotation into consideration, it is important to protect the anonymity of annotators.

References

- Ehab A AlBadawy and Yelin Kim. 2018. Joint discrete and continuous emotion prediction using ensemble and end-to-end approaches. In *Proc. ICMI*, Boulder.
- Alexander Amini, Wilko Schwarting, Ava Soleimany, and Daniela Rus. 2020. Deep evidential regression. In *Proc. NeurIPS*, Vancouver.
- Atsushi Ando, Satoshi Kobashikawa, Hosana Kamiyama, Ryo Masumura, Yusuke Ijima, and Yushi Aono. 2018. Soft-target training with ambiguous emotional utterances for DNN-based speech emotion classification. In *Proc. ICASSP*, Brighton.

683	Mia Atcheson, Vidhyasaharan Sethu, and Julien Epps.	Ting Dang, Vidhyasaharan Sethu, Julien Epps, and	738
684	2018. Demonstrating and modelling systematic time-	Eliathamby Ambikairajah. 2017. An investigation of	739
685	varying annotator disagreement in continuous emo-	emotion prediction uncertainty using gaussian mix-	740
686	tion annotation. In <i>Proc. Interspeech</i> , Hyderabad.	ture regression. In <i>Proc. Interspeech</i> , Stockholm.	741
687	Mia Atcheson, Vidhyasaharan Sethu, and Julien Epps.	Aida Mostafazadeh Davani, Mark Díaz, and Vinodku-	742
688	2019. Using Gaussian processes with LSTM neu-	mar Prabhakaran. 2022. Dealing with disagreements:	743
689	ral networks to predict continuous-time, dimensional	Looking beyond the majority vote in subjective an-	744
690	emotion in ambiguous speech. In <i>Proc. ACII</i> , Cam-	notations. <i>Transactions of the Association for Com-</i>	745
691	bridge.	<i>putational Linguistics</i> , 10:92–110.	746
692	Bagus Tris Atmaja and Masato Akagi. 2020a. Improv-	Jun Deng, Wenjing Han, and Björn Schuller. 2012. Con-	747
693	ing valence prediction in dimensional speech emotion	confidence measures for speech emotion recognition: A	748
694	recognition using linguistic information. In <i>Proc. O-</i>	start. In <i>Speech Communication; 10. ITG Sympo-</i>	749
695	<i>COCOSDA</i> , Yangon.	<i>sium</i> , pages 1–4. VDE.	750
696	Bagus Tris Atmaja and Masato Akagi. 2020b. Multi-	Armen Der Kiureghian and Ove Ditlevsen. 2009.	751
697	task learning and multistage fusion for dimensional	Aleatory or epistemic? does it matter? <i>Structural</i>	752
698	audiovisual emotion recognition. In <i>Proc. ICASSP</i> ,	<i>Safety</i> , 31(2):105–112.	753
699	Conference held virtually.	H.M. Fayek, M. Lech, and L. Cavedon. 2016. Model-	754
700	Bagus Tris Atmaja and Masato Akagi. 2021. Two-	ing subjectiveness in emotion recognition with deep	755
701	stage dimensional emotion recognition by fusing pre-	neural networks: Ensembles vs soft labels. In <i>Proc.</i>	756
702	dictions of acoustic and text networks using svm.	<i>IJCNN</i> , Vancouver.	757
703	<i>Speech Communication</i> , 126:9–21.	Yarin Gal and Zoubin Ghahramani. 2016. Dropout	758
704	Alexei Baevski, Yuhao Zhou, Abdelrahman Mohamed,	as a bayesian approximation: Representing model	759
705	and Michael Auli. 2020. Wav2Vec 2.0: A framework	uncertainty in deep learning. In <i>Proc. ICML</i> , New	760
706	for self-supervised learning of speech representations.	York City.	761
707	In <i>Proc. NeurIPS</i> , Conference held virtually.	Ayoub Ghriss, Bo Yang, Viktor Rozgic, Elizabeth	762
708	Charles Blundell, Julien Cornebise, Koray	Shriberg, and Chao Wang. 2022. Sentiment-aware	763
709	Kavukcuoglu, and Daan Wierstra. 2015. Weight	automatic speech recognition pre-training for en-	764
710	uncertainty in neural network. In <i>Proc. ICML</i> , Lille.	hanced speech emotion recognition. In <i>Proc.</i>	765
711	C. Busso, M. Bulut, C.-C. Lee, A. Kazemzadeh, E.M.	<i>ICASSP</i> , Singapore.	766
712	Provost, S. Kim, J.N. Chang, S. Lee, and S.S.	Michael Grimm and Kristian Kroschel. 2005. Eval-	767
713	Narayanan. 2008. IEMOCAP: Interactive emotional	uation of natural emotions using self assessment	768
714	dyadic motion capture database. <i>Language Re-</i>	manikins. In <i>Proc. ASRU</i> , Cancun.	769
715	<i>sources and Evaluation</i> , 42:335–359.	Michael Grimm, Kristian Kroschel, Emily Mower, and	770
716	Carlos Busso, Srinivas Parthasarathy, Alec Burmania,	Shrikanth Narayanan. 2007. Primitives-based evalu-	771
717	Mohammed AbdelWahab, Najmeh Sadoughi, and	ation and estimation of emotions in speech. <i>Speech</i>	772
718	Emily Mower Provost. 2017. MSP-IMPROV: An	<i>Communication</i> , 49(10-11):787–800.	773
719	acted corpus of dyadic interactions to study emotion	Hatice Gunes, Björn Schuller, Maja Pantic, and Roddy	774
720	perception . <i>IEEE Transactions on Affective Comput-</i>	Cowie. 2011. Emotion representation, analysis and	775
721	<i>ing</i> , 8(1):67–80.	synthesis in continuous space: A survey. In <i>Proc.</i>	776
722	Sanyuan Chen, Chengyi Wang, Zhengyang Chen,	<i>FG</i> , Santa Barbara.	777
723	Yu Wu, Shujie Liu, Zhuo Chen, Jinyu Li, Naoyuki	Jing Han, Zixing Zhang, Zhao Ren, and Björn Schuller.	778
724	Kanda, Takuya Yoshioka, Xiong Xiao, et al. 2022.	2021. Exploring perception uncertainty for emotion	779
725	WavLM: Large-scale self-supervised pre-training for	recognition in dyadic conversation and music listen-	780
726	full stack speech processing. <i>IEEE Journal of Se-</i>	<i>ing</i> . <i>Cognitive Computation</i> , 13(2):231–240.	781
727	<i>lected Topics in Signal Processing</i> .	Jing Han, Zixing Zhang, Maximilian Schmitt, Maja	782
728	Huang-Cheng Chou, Wei-Cheng Lin, Chi-Chun Lee,	Pantic, and Björn Schuller. 2017. From hard to soft:	783
729	and Carlos Busso. 2022. Exploiting annotators’	Towards more human-like emotion recognition by	784
730	typed description of emotion perception to maximize	modelling the perception uncertainty. In <i>Proc. ACM</i>	785
731	utilization of ratings for speech emotion recognition.	<i>MM</i> , Mountain View.	786
732	In <i>Proc. ICASSP</i> , Singapore.	Xincheng Ju, Dong Zhang, Junhui Li, and Guodong	787
733	Ting Dang, Vidhyasaharan Sethu, and Eliathamby Am-	Zhou. 2020. Transformer-based label set generation	788
734	bikairajah. 2018. Dynamic multi-rater Gaussian mix-	for multi-modal multi-label emotion detection. In	789
735	ture regression incorporating temporal dependencies	<i>Proc. ACM MM</i> , Seattle.	790
736	of emotion uncertainty using Kalman filters. In <i>Proc.</i>		
737	<i>ICASSP</i> , Calgary.		

791	Alex Kendall and Yarin Gal. 2017. What uncertainties	847
792	do we need in bayesian deep learning for computer	848
793	vision? In <i>Proc. NeurIPS</i> , Long Beach.	849
794	Jean Kossaifi, Robert Walecki, Yannis Panagakis, Jie	850
795	Shen, Maximilian Schmitt, Fabien Ringeval, Jing	
796	Han, Vedhas Pandit, Antoine Toisoul, Björn Schuller,	851
797	et al. 2019. SEWA DB: A rich database for audio-	852
798	visual emotion and sentiment research in the wild.	853
799	<i>IEEE Transactions on Pattern Analysis and Machine</i>	854
800	<i>Intelligence</i> , 43(3):1022–1040.	
801	Duc Le, Zakaria Aldeneh, and Emily Mower Provost.	
802	2017. Discretized continuous speech emotion recog-	
803	nition with multi-task deep recurrent neural network.	
804	In <i>Proc. Interspeech</i> , Stockholm.	
805	Seong-Gyun Leem, Daniel Fulford, Jukka-Pekka On-	
806	nela, David Gard, and Carlos Busso. 2022. Not all	
807	features are equal: Selection of robust features for	
808	speech emotion recognition in noisy environments.	
809	In <i>Proc. ICASSP</i> , Singapore.	
810	Yinhan Liu, Myle Ott, Naman Goyal, Jingfei Du, Man-	
811	dar Joshi, Danqi Chen, Omer Levy, Mike Lewis,	
812	Luke Zettlemoyer, and Veselin Stoyanov. 2019.	
813	RoBERTa: A robustly optimized BERT pretraining	
814	approach. <i>arXiv preprint arXiv:1907.11692</i> .	
815	R. Lotfian and C. Busso. 2019. Building naturalis-	
816	tic emotionally balanced speech corpus by retriev-	
817	ing emotional speech from existing podcast record-	
818	ings. <i>IEEE Transactions on Affective Computing</i> ,	
819	10(4):471–483.	
820	N. Majumder, D. Hazarika, A. Gelbukh, E. Cambria,	
821	and S. Poria. 2018. Multimodal sentiment analy-	
822	sis using hierarchical fusion with context modeling.	
823	<i>Knowledge-Based Systems</i> , 161:124–133.	
824	Andrey Malinin and Mark Gales. 2018. Predictive un-	
825	certainty estimation via prior networks. In <i>Proc.</i>	
826	<i>NeurIPS</i> , Montréal.	
827	Konstantin Markov, Tomoko Matsui, Francois Septier,	
828	and Gareth Peters. 2015. Dynamic speech emotion	
829	recognition with state-space models. In <i>Proc. EU-</i>	
830	<i>SIPCO</i> , Nice.	
831	Hermann G Matthies. 2007. Quantifying uncertainty:	
832	Modern computational representation of probability	
833	and applications. In <i>Extreme man-made and natural</i>	
834	<i>hazards in dynamics of structures</i> , pages 105–135.	
835	Springer.	
836	Vikramjit Mitra, Hsiang-Yun Sherry Chien, Vasudha	
837	Kowtha, Joseph Yitan Cheng, and Erdrin Azemi.	
838	2022. Speech Emotion: Investigating Model Rep-	
839	resentations, Multi-Task Learning and Knowledge	
840	Distillation . In <i>Proc. Interspeech 2022</i> , pages 4715–	
841	4719.	
842	Mihalis A Nicolaou, Hatice Gunes, and Maja Pantic.	
843	2011. Continuous prediction of spontaneous affect	
844	from multiple cues and modalities in valence-arousal	
845	space. <i>IEEE Transactions on Affective Computing</i> ,	
846	2(2):92–105.	
	Vassil Panayotov, Guoguo Chen, Daniel Povey, and	
	Sanjeev Khudanpur. 2015. Librispeech: An ASR	
	corpus based on public domain audio books. In <i>Proc.</i>	
	<i>ICASSP</i> , South Brisbane.	
	Robert Plutchik. 2001. The nature of emotions: Human	
	emotions have deep evolutionary roots, a fact that	
	may explain their complexity and provide tools for	
	clinical practice. <i>American Scientist</i> , 89(4):344–350.	
	Soujanya Poria, Devamanyu Hazarika, Navonil Ma-	
	jumder, Gautam Naik, Erik Cambria, and Rada Mi-	
	halcea. 2019. MELD: A multimodal multi-party	
	dataset for emotion recognition in conversations. In	
	<i>Proc. ACL</i> , Florence.	
	Soujanya Poria, Navonil Majumder, Devamanyu Haz-	
	arika, Erik Cambria, Alexander Gelbukh, and Amir	
	Hussain. 2018. Multimodal sentiment analysis: Ad-	
	dressing key issues and setting up the baselines.	
	<i>IEEE Intelligent Systems</i> , 33(6):17–25.	
	Filip Povolny, Pavel Matejka, Michal Hradis, Anna Pop-	
	ková, Lubomír Otrusina, Pavel Smrz, Ian Wood, Ce-	
	cile Robin, and Lori Lamel. 2016. Multimodal emo-	
	tion recognition for avec 2016 challenge. In <i>Proc.</i>	
	<i>ACM MM</i> , Amsterdam.	
	Navin Raj Prabhu, Guillaume Carbajal, Nale Lehmann-	
	Willenbrock, and Timo Gerkmann. 2021. End-to-	
	end label uncertainty modeling for speech emotion	
	recognition using bayesian neural networks. <i>arXiv</i>	
	<i>preprint arXiv:2110.03299</i> .	
	Mirco Ravanelli, Titouan Parcollet, Peter Plantinga,	
	Aku Rouhe, Samuele Cornell, Loren Lugosch, Cem	
	Subakan, Nauman Dawalatabad, Abdelwahab Heba,	
	Jianyuan Zhong, Ju-Chieh Chou, Sung-Lin Yeh,	
	Szu-Wei Fu, Chien-Feng Liao, Elena Rastorgueva,	
	François Grondin, William Aris, Hwidong Na, Yan	
	Gao, Renato De Mori, and Yoshua Bengio. 2021.	
	SpeechBrain: A general-purpose speech toolkit .	
	ArXiv:2106.04624.	
	Fabien Ringeval, Björn Schuller, Michel Valstar, Roddy	
	Cowie, and Maja Pantic. 2015. AVEC 2015: The	
	5th international audio/visual emotion challenge and	
	workshop. In <i>Proc. ACM MM</i> , Brisbane.	
	Fabien Ringeval, Björn Schuller, Michel Valstar,	
	Jonathan Gratch, Roddy Cowie, Stefan Scherer,	
	Sharon Mozgai, Nicholas Cummins, Maximilian	
	Schmitt, and Maja Pantic. 2017. AVEC 2017: Real-	
	life depression, and affect recognition workshop and	
	challenge. In <i>Proc. ACM MM</i> , Mountain View.	
	Fabien Ringeval, Andreas Sonderegger, Jürgen Sauer,	
	and Denis Lalanne. 2013. Introducing the RECOLA	
	multimodal corpus of remote collaborative and affec-	
	tive interactions. In <i>Proc. FG</i> , Shanghai.	
	James A Russell. 1980. A circumplex model of af-	
	fect. <i>Journal of Personality and Social Psychology</i> ,	
	39(6):1161.	

James A Russell and Albert Mehrabian. 1977. Evidence for a three-factor theory of emotions. *Journal of Research in Personality*, 11(3):273–294.

Harold Schlosberg. 1954. Three dimensions of emotion. *Psychological Review*, 61(2):81.

Murat Sensoy, Lance Kaplan, and Melih Kandemir. 2018. Evidential deep learning to quantify classification uncertainty. In *Proc. NeurIPS*, Montréal.

Kusha Sridhar and Carlos Busso. 2020a. Ensemble of students taught by probabilistic teachers to improve speech emotion recognition. In *Proc. Interspeech*, Shanghai.

Kusha Sridhar and Carlos Busso. 2020b. Modeling uncertainty in predicting emotional attributes from spontaneous speech. In *Proc. ICASSP*, Conference held virtually.

Kusha Sridhar, Wei-Cheng Lin, and Carlos Busso. 2021. Generative approach using soft-labels to learn uncertainty in predicting emotional attributes. In *Proc. ACII*, Chicago.

Sundararajan Srinivasan, Zhaocheng Huang, and Katrin Kirchhoff. 2022. Representation learning through cross-modal conditional teacher-student training for speech emotion recognition. In *Proc. ICASSP*, Singapore.

Andreas Triantafyllopoulos, Johannes Wagner, Hagen Wierstorf, Maximilian Schmitt, Uwe Reichel, Florian Eyben, Felix Burkhardt, and Björn W. Schuller. 2022. *Probing speech emotion recognition transformers for linguistic knowledge*. In *Proc. Interspeech 2022*, pages 146–150.

George Trigeorgis, Fabien Ringeval, Raymond Brueckner, Erik Marchi, Mihalis A Nicolaou, Björn Schuller, and Stefanos Zafeiriou. 2016. Adieu features? end-to-end speech emotion recognition using a deep convolutional recurrent network. In *Proc. ICASSP*, Shanghai.

Jingyao Wu, Ting Dang, Vidhyasaharan Sethu, and Eliathamby Ambikairajah. 2022a. A novel sequential Monte Carlo framework for predicting ambiguous emotion states. In *Proc. ICASSP*, Singapore.

Wen Wu, Chao Zhang, and Philip C. Woodland. 2021. Emotion recognition by fusing time synchronous and time asynchronous representations. In *Proc. ICASSP*, Toronto.

Wen Wu, Chao Zhang, Xixin Wu, and Philip C Woodland. 2022b. Estimating the uncertainty in emotion class labels with utterance-specific dirichlet priors. *arXiv preprint arXiv:2203.04443*.

Shu-wen Yang, Po-Han Chi, Yung-Sung Chuang, Cheng-I Jeff Lai, Kushal Lakhota, Yist Y Lin, Andy T Liu, Jiatong Shi, Xuankai Chang, Guan-Ting Lin, et al. 2021. SUPERB: Speech processing universal performance benchmark. *arXiv preprint arXiv:2105.01051*.

Dong Zhang, Xincheng Ju, Junhui Li, Shoushan Li, Qiaoming Zhu, and Guodong Zhou. 2020. Multi-modal multi-label emotion detection with modality and label dependence. In *Proc. EMNLP*, Conference held virtually.

A Derivation of the predictive posterior

Since NIG is the Gaussian conjugate prior,

$$\begin{aligned} p(\Psi|\Omega) &= \mathcal{N}(\gamma, \sigma^2 v^{-1}) \Gamma^{-1}(\alpha, \beta) \\ &= \frac{\beta^\alpha \sqrt{v}}{\Gamma(\alpha) \sqrt{2\pi\sigma^2}} \left(\frac{1}{\sigma^2}\right)^{\alpha+1} \\ &\quad \cdot \exp\left\{-\frac{2\beta + v(\gamma - \mu)^2}{2\sigma^2}\right\} \end{aligned}$$

its posterior $p(\Psi|\mathcal{D})$ is in the same parametric family as the prior $p(\Psi|\Omega)$. Therefore, given a test utterance $\mathbf{x}_*$, the predictive posterior $p(y_*|\mathcal{D})$ has the same form as the marginal likelihood $p(y|\Omega)$, where $\mathcal{D}$ denotes the training set.

$$p(y_*|\mathcal{D}) = \int p(y_*|\Psi)p(\Psi|\mathcal{D}) d\Psi \quad (7)$$

$$p(y|\Omega) = \int p(y|\Psi)p(\Psi|\Omega) d\Psi \quad (8)$$

In DEER, the predictive posterior and posterior are both conditioned on Ω , written as $p(y_*|\mathcal{D}, \Omega)$ and $p(\Psi|\mathcal{D}, \Omega)$ to be precise. Also, the information of $\mathcal{D}$ is contained in Ω_* since $\Omega_* = f_{\hat{\Theta}}(\mathbf{x}_*)$ and $\hat{\Theta}$ is the optimal model parameters obtained by training on $\mathcal{D}$. Then the predictive posterior can be written as $p(y_*|\Omega_*)$. Given the conjugate prior, the predictive posterior in DEER can be computed by directly substituting the predicted Ω_* into the expression of marginal likelihood derived in Eqn. (2), skipping the step of calculating the posterior.

B Fusion with text modality

This appendix presents bi-modal experiments that incorporate text information into the DEER model. Transcriptions were obtained from a publicly available automatic speech recognition (ASR) model “wav2vec2-base-960h”⁶ which fine-tuned the wav2vec 2.0 (Baevski et al., 2020) model on 960 hours Librispeech data (Panayotov et al., 2015). Transcriptions were first encoded by a RoBERTa model (Liu et al., 2019) and fed into another two-layer Transformer encoder. As shown in Figure 4, outputs from the text Transformer were concatenated with the outputs from the audio Transformer

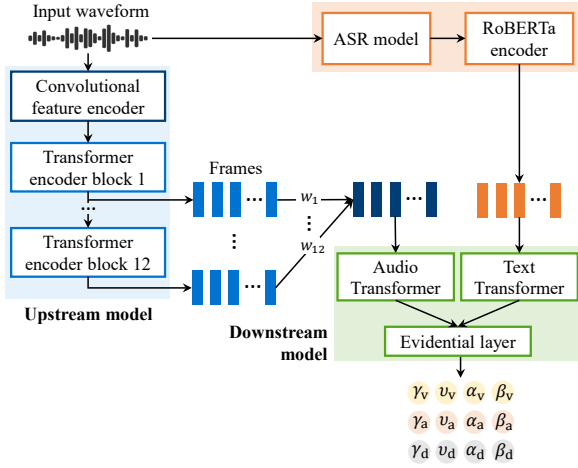


Figure 4: Model structure for bi-modal experiments.

	Modality	v	a	d
MSP-podcast	A	0.506	0.698	0.613
	A+T	0.559	0.699	0.614
	Modality	v	a	d
IEMOCAP	A	0.596	0.756	0.569
	A+T	0.609	0.754	0.575

Table 5: CCC value for bi-modal experiments. ‘A’ and ‘T’ stands for audio and text. ‘v’, ‘a’, and ‘d’ stand for valence, arousal, and dominance. Release 1.8 is used for MSP-Podcast. ‘Ses05’ setup used for IEMOCAP that trains on Session 1-4 and tested on Session 5.

encoder and fed into the evidential output layer. Results are shown in Table 5. Incorporating text information improves the estimation of valence but not necessarily for arousal and dominance. Similar phenomena were observed by (Triantafyllopoulos et al., 2022). A possible explanation is that text is effective for sentiment analysis (positive or negative) but may not be as informative as audio to determine a speaker’s level of excitement. CCC for dominance improves more for IEMOCAP than MSP-Podcast possibly because IEMOCAP is an acted dataset and the emotion may be exaggerated compared with MSP-Podcast which contains natural emotion.

⁶Available at: <https://huggingface.co/facebook/wav2vec2-base-960h>