

MULTIPLE OUTPUT SAMPLES PER INPUT IN A SINGLE-OUTPUT GAUSSIAN PROCESS

Jeremy H. M. Wong, Huayun Zhang, and Nancy F. Chen

Institute for Infocomm Research (I²R), A*STAR, Singapore

ABSTRACT

The standard Gaussian Process (GP) only considers a single output sample per input in the training set. Datasets for subjective tasks, such as spoken language assessment, may be annotated with output labels from multiple human raters per input. This paper proposes to generalise the GP to allow for these multiple output samples in the training set, and thus make use of available output uncertainty information. This differs from a multi-output GP, as all output samples are from the same task here. The output density function is formulated to be the joint likelihood of observing all output samples, and latent variables are not repeated to reduce computation cost. The test set predictions are inferred similarly to a standard GP, with a difference being in the optimised hyper-parameters. This is evaluated on speechocean762, showing that it allows the GP to compute a test set output distribution that is more similar to the collection of reference outputs from the multiple human raters.

Index Terms— Gaussian process, subjective, uncertainty, inter-rater agreement, spoken language assessment

1. INTRODUCTION

The Gaussian Process (GP) [1] expresses a distributional uncertainty [2] that naturally increases for inputs further away from the training data. It may also be interpreted as marginalising over functions [1], which reduces the influence of model uncertainty [3, 4]. However, a GP does not consider whether outputs are similar for proximate inputs, and thus may lack in modelling data uncertainty [3, 5]. Better uncertainty estimation may be especially desirable for subjective tasks, such as Spoken Language Assessment (SLA), where multiple human annotators may provide differing output labels for the same input [6]. A collection of human annotations for the same input may be interpreted as a reference of data uncertainty that an automatic model should also aim to compute. In such settings, the uncertainties expressed by the model and the human annotators can be explicitly compared.

The standard GP assumes that each input in the training set is paired with only one output, treated as the ground truth. This paper proposes to extend the GP formulation, to use available information of data uncertainty from multiple output samples for each input in the training set. A naive approach is to repeat the inputs for each output, but this greatly increases the computational cost. A proposed omission of repeated latent variables avoids this expense. All outputs here are from the same task, while the multi-output GP [7, 8] instead considers a multi-task framework [9].

The aim of SLA is to assign a score related to oral proficiency, to a speech input. Examples of proficiency aspects that are often assessed include pronunciation accuracy, intonation, fluency, prosody, task completion, and topic relevance. As previously mentioned, SLA

is a subjective task, where differing scores are often assigned by multiple expert raters for the same input. The diversity of scores can be reduced by increasing the coverage of rules in a rubric. However, the choice of rules may not agree with the needs of diverse users. This paper instead embraces subjectivity, by using the diversity of reference scores to better predict output uncertainty. This uncertainty information may allow the system to better calibrate its feedback to the user, by not penalising a student for a low score that raters would have disagreed upon, and instead seek clarification or human intervention. This also reduces the need for a strict rubric, thereby allowing better generalisation to diverse user needs.

2. GAUSSIAN PROCESS REGRESSION

When given a training set of N input feature vectors of dimension D , $\mathbf{X} \in \mathbb{R}^{N \times D}$, a GP places a jointly Gaussian prior over latent variables, $\mathbf{f} \in \mathbb{R}^N$, as $p(\mathbf{f}|\mathbf{X}) = \mathcal{N}(\mathbf{f}; \mathbf{0}, \mathbf{K}^{\mathbf{X}\mathbf{X}})$. The covariance of the latent variables is defined as a pair-wise similarity between the inputs, computed by the kernel, $\mathbf{K}$. Here, the squared exponential kernel is used, with elements defined as $k_{ij}^{\mathbf{X}\mathbf{X}'} = s^2 \exp[-\frac{1}{2l^2}(\mathbf{x}_i - \mathbf{x}_j)^\top (\mathbf{x}_i - \mathbf{x}_j)]$, where i and j are the data point indexes, l is a length hyper-parameter, and s is a scale hyper-parameter. The GP makes the assumption that the outputs, $\mathbf{y} \in \mathbb{R}^N$, are conditionally independent of the inputs when given the latent variables. The outputs are Gaussian distributed, $p(\mathbf{y}|\mathbf{f}) = \mathcal{N}(\mathbf{y}; \mathbf{f}, \sigma^2 \mathbf{I})$, about a mean of the latent variable, with an observation noise hyper-parameter, σ , where $\mathbf{I}$ is the identity matrix.

Training a GP involves estimating the hyper-parameters s , l , and σ . One approach is to maximise the marginal log-likelihood of the training data, $\mathcal{F} = \log p(\mathbf{y}^{\text{ref}}|\mathbf{X})$, where $\mathbf{y}^{\text{ref}}$ are the reference outputs. The marginal likelihood is

$$p(\mathbf{y}|\mathbf{X}) = \int p(\mathbf{y}|\mathbf{f}) p(\mathbf{f}|\mathbf{X}) d\mathbf{f} = \mathcal{N}(\mathbf{y}; \mathbf{0}, \mathbf{K}^{\mathbf{X}\mathbf{X}} + \sigma^2 \mathbf{I}). \quad (1)$$

Gradient-based optimisation of the hyper-parameters requires inverting $\mathbf{K}^{\mathbf{X}\mathbf{X}} + \sigma^2 \mathbf{I}$ [1], with a number of operations scaling as $\mathcal{O}(N^3)$. Algorithms such as [10] reduce the polynomial power, but still require more than $\mathcal{O}(N^2)$ operations. For evaluation, outputs, $\hat{\mathbf{y}}$, are predicted from test inputs, $\hat{\mathbf{X}}$, by first computing a posterior over predicted latent variables,

$$p(\hat{\mathbf{f}}|\hat{\mathbf{x}}, \mathbf{y}, \mathbf{X}) = \frac{p(\hat{\mathbf{f}}, \mathbf{y}|\hat{\mathbf{x}}, \mathbf{X})}{p(\mathbf{y}|\mathbf{X})} = \mathcal{N}(\hat{\mathbf{f}}; \hat{\boldsymbol{\mu}}, \hat{\boldsymbol{\Sigma}}), \quad (2)$$

where $\hat{\boldsymbol{\mu}} = \mathbf{k}^{\mathbf{X}\hat{\mathbf{x}}}^\top [\mathbf{K}^{\mathbf{X}\mathbf{X}} + \sigma^2 \mathbf{I}]^{-1} \mathbf{y}$ and $\hat{\boldsymbol{\Sigma}} = \mathbf{k}^{\hat{\mathbf{x}}\hat{\mathbf{x}}} - \mathbf{k}^{\mathbf{X}\hat{\mathbf{x}}}^\top [\mathbf{K}^{\mathbf{X}\mathbf{X}} + \sigma^2 \mathbf{I}]^{-1} \mathbf{k}^{\mathbf{X}\hat{\mathbf{x}}}$, which similarly to training, also requires computing $[\mathbf{K}^{\mathbf{X}\mathbf{X}} + \sigma^2 \mathbf{I}]^{-1}$. The posterior for the test output is then $p(\hat{\mathbf{y}}|\hat{\mathbf{x}}, \mathbf{y}, \mathbf{X}) = \mathcal{N}(\hat{\mathbf{y}}; \hat{\boldsymbol{\mu}}, \hat{\boldsymbol{\Sigma}} + \sigma^2)$. Finally, the predicted scalar output for each test data point, i , can be inferred from this, using a decision rule of, for example, choosing the mean, $\hat{y}_i^* = \hat{\mu}_i$. The

This work was supported by the A*STAR Computational Resource Centre through the use of its high performance computing facilities.

mean, mode, and median are equivalent because of the Gaussian symmetry and uni-modality.

3. MULTIPLE OUTPUT SAMPLES

In the standard GP, each input is paired with a single output sample. In SLA, using a GP with only a scalar output label as the reference [11] may forego information about the inter-rater uncertainty. This paper proposes extending the GP to use the multiple output samples in the training set. Let each training input be associated with R output reference samples, $\mathbf{y}_i^{\text{ref}} = [y_{i1}^{\text{ref}}, \dots, y_{iR}^{\text{ref}}]$, where y_{ir}^{ref} is the r th reference output sample for the i th data point. It is assumed for simplicity that the number of output samples is the same for all inputs, but the formulation can be extended for varying numbers of raters [12]. The outputs are stacked across all data points to yield a matrix $\mathbf{Y}^{\text{ref}} \in \mathbb{R}^{N \times R}$.

A naive approach to accommodate multiple output samples is to flatten $\mathbf{Y}^{\text{ref}}$ into a vector and repeat the inputs for each output sample. This will yield a $\mathbb{R}^{NR \times NR}$ kernel of $\tilde{\mathbf{K}} = \mathbf{K}(\mathbf{X}, \mathbf{X}) \otimes \mathbf{1}^{(R)}$, where $\mathbf{1}^{(R)}$ is a $\mathbb{R}^{R \times R}$ matrix of 1s and $\otimes$ represents a Kronecker (tensor) product. There is no need to explicitly invert this non-full-rank kernel, as only the inverse of the full-rank $\tilde{\mathbf{K}} + \sigma^2 \mathbf{I}$ is needed. Inversion requires a number of operations that scales as $\mathcal{O}(N^3 R^3)$, as opposed to $\mathcal{O}(N^3)$ when using only a single output sample. Computational constraints may then impose the need for sparse approximation [13], even for modest training set sizes.

When using this repeated kernel, the latent variables prior is $p(\tilde{\mathbf{f}}|\mathbf{X}) = \mathcal{N}(\tilde{\mathbf{f}}|\mathbf{0}, \tilde{\mathbf{K}})$. The repeated structure of $\tilde{\mathbf{K}}$ causes all latent variables associated with the same training data point to be perfectly correlated, and thus have the same value. It seems computationally redundant to separately model the repeated latent variables. This paper proposes that there is no need to separately compute each of the repetitions, and to instead only express a single instance of each latent variable per training data point. The joint density of the multiple output samples is then conditioned on these non-repeated latent variables, $p(\mathbf{Y}|\mathbf{f})$. The diagonal covariance in the output density function of a standard GP implies that the outputs from different training data points are conditionally independent of each other when given the latent variables. This conditional independence, applied to the case of repeated inputs in the naive approach, yields in this proposed approach a joint output density that can be factorised as

$$p(\mathbf{Y}|\mathbf{f}) = \prod_{r=1}^R \mathcal{N}(\mathbf{y}_r; \mathbf{f}, \sigma^2 \mathbf{I}) \quad (3)$$

$$= \mathcal{N}\left(\bar{\boldsymbol{\mu}}; \mathbf{f}, \frac{\sigma^2}{R} \mathbf{I}\right) \mathcal{N}\left(\bar{\boldsymbol{\eta}}; \mathbf{0}, \frac{\sigma^2}{R} \mathbf{I}\right) \frac{(2\pi\sigma^2)^{(2-R)\frac{N}{2}}}{R^N},$$

where $\bar{\boldsymbol{\mu}} = \frac{1}{R} \sum_r \mathbf{y}_r$ and $\bar{\boldsymbol{\eta}}_i = \sqrt{\frac{1}{R} \sum_r (y_{ir} - \bar{\mu}_i)^2}$, and $\mathbf{y}_r \in \mathbb{R}^N$ is a vector of the r th output sample for all inputs. The factorisation and diagonal covariance in (3) assume that the multiple output samples for the same input and across different inputs are independent of each other, when given the latent variable. As a result, the output density does not depend on the order of the multiple output samples for each input. This can be explicitly seen in the re-parameterisation of the multiple outputs as the empirical mean $\bar{\boldsymbol{\mu}}$ and empirical biased standard deviation $\bar{\boldsymbol{\eta}}$, both of which are independent of the ordering of the output samples. Also, (3) may seem unnormalised over $\bar{\boldsymbol{\mu}}$ and $\bar{\boldsymbol{\eta}}$. However, it should be noted that $p(\mathbf{Y}|\mathbf{f})$ is expected to sum to one when integrated over $\mathbf{Y}$, and not over $\bar{\boldsymbol{\mu}}$ and $\bar{\boldsymbol{\eta}}$.

The marginal likelihood can be computed from the non-redundant prior and the output density of (3), as

$$p(\mathbf{Y}|\mathbf{X}) = \int p(\mathbf{Y}|\mathbf{f}) p(\mathbf{f}|\mathbf{X}) d\mathbf{f} = \mathcal{N}\left(\bar{\boldsymbol{\mu}}; \mathbf{0}, \mathbf{K}^{\text{xx}} + \frac{\sigma^2}{R} \mathbf{I}\right) g(\mathbf{Y}), \quad (4)$$

where $g(\mathbf{Y}) = \mathcal{N}(\bar{\boldsymbol{\eta}}; \mathbf{0}, \frac{\sigma^2}{R} \mathbf{I}) (2\pi\sigma^2)^{(2-R)\frac{N}{2}} R^{-N}$. Analogously to the standard GP, and equivalently to the naive approach, the hyper-parameters here can be optimised by maximising the joint marginal log-likelihood of all output samples, $\mathcal{F}_{\text{joint}} = \log p(\mathbf{Y}^{\text{ref}}|\mathbf{X})$. Unlike a naive implementation, the computational cost of gradient-based optimisation of the proposed method scales as $\mathcal{O}(N^3)$, similarly to a standard GP with a single output sample.

When performing evaluation, the predicted output can be inferred from the posterior. First, the joint density between the training set outputs and the test latent variables needs to be computed. When the training set comprises multiple output samples per input, this paper again proposes that there is no need to separately model redundant latent variables, by expressing this joint density as

$$p(\hat{\mathbf{f}}, \mathbf{Y}|\hat{\mathbf{x}}, \mathbf{X}) = \int p(\hat{\mathbf{f}}, \mathbf{f}|\hat{\mathbf{x}}, \mathbf{X}) p(\mathbf{Y}|\mathbf{f}) d\mathbf{f} \quad (5)$$

$$= \mathcal{N}\left(\begin{bmatrix} \hat{\mathbf{f}} \\ \bar{\boldsymbol{\mu}} \end{bmatrix}; \mathbf{0}, \begin{bmatrix} k^{\hat{\mathbf{x}}\hat{\mathbf{x}}} & \mathbf{k}^{\hat{\mathbf{x}}\mathbf{x}} \\ \mathbf{k}^{\mathbf{x}\hat{\mathbf{x}}} & \mathbf{K}^{\mathbf{xx}} + \frac{\sigma^2}{R} \mathbf{I} \end{bmatrix}\right) g(\mathbf{Y}).$$

The test set latent variables posterior is

$$p(\hat{\mathbf{f}}|\hat{\mathbf{x}}, \mathbf{Y}, \mathbf{X}) = \frac{p(\hat{\mathbf{f}}, \mathbf{Y}|\hat{\mathbf{x}}, \mathbf{X})}{p(\mathbf{Y}|\mathbf{X})} = \mathcal{N}(\hat{\mathbf{f}}; \check{\boldsymbol{\mu}}, \check{\boldsymbol{\nu}}), \quad (6)$$

analogously to (2), where here, $\check{\boldsymbol{\mu}} = \mathbf{k}^{\hat{\mathbf{x}}\hat{\mathbf{x}}} [\mathbf{K}^{\mathbf{xx}} + \frac{\sigma^2}{R} \mathbf{I}]^{-1} \bar{\boldsymbol{\mu}}$ and $\check{\boldsymbol{\nu}} = k^{\hat{\mathbf{x}}\hat{\mathbf{x}}} - \mathbf{k}^{\hat{\mathbf{x}}\mathbf{x}} [\mathbf{K}^{\mathbf{xx}} + \frac{\sigma^2}{R} \mathbf{I}]^{-1} \mathbf{k}^{\mathbf{x}\hat{\mathbf{x}}}$. The test output posterior is then $p(\hat{\mathbf{y}}|\hat{\mathbf{x}}, \mathbf{Y}, \mathbf{X}) = \mathcal{N}(\hat{\mathbf{y}}; \check{\boldsymbol{\mu}}, \check{\boldsymbol{\nu}} + \sigma^2)$, from which an output can be chosen. The matrix inverted here is of size $\mathbb{R}^{N \times N}$, as opposed to $\mathbb{R}^{NR \times NR}$ in the naive repetition approach. The number of operations thus scales as $\mathcal{O}(N^3)$, similarly to a standard GP with a single output sample.

The covariance of the standard and proposed posteriors are independent of the training set outputs, and only depend on the training set inputs and the hyper-parameters. As such, information about the proximity between the multiple output samples for the same input in the training set does not influence the resulting posterior covariance. Furthermore, whether the training set output values are near or far from each other, for two inputs that are near as measured by the kernel, does not influence the resulting posterior covariance for test inputs near to these training inputs. The posterior covariance is only affected by the distance between the test and training set inputs. This implies that a GP may be more suited for computing distributional uncertainty, and less so for the data uncertainty [5] that may be more relevant to the uncertainty expressed between the raters. Work in [12] extends this framework to allow the GP to take training set outputs into account for the posterior covariance.

The predictive mean, $\check{\boldsymbol{\mu}}$, is dependent on the mean of the multiple output samples, $\bar{\boldsymbol{\mu}}$, and independent of $\bar{\boldsymbol{\eta}}$. This is similar to a standard GP, when using a mean combination of the multiple rater scores to get the reference. As mentioned previously, the predictive covariance is independent of the training set outputs. Thus, during inference, the remaining difference between the standard GP and the proposal is in the hyper-parameters. The proposal optimises the hyper-parameters by maximising the joint marginal log-likelihood of all training set outputs, while the standard approach maximises the marginal log-likelihood of the combined output reference.

4. EXPERIMENTS

Experiments were performed on speechocean762 [6], comprising training and test sets, each with 2500 sentences and 125 disjoint speakers. The native Mandarin speakers read the sentences in English. Each sentence is annotated by 5 raters with a variety of score types at different linguistic levels, but only the sentence-level pronunciation accuracy was used, with an integer score ranging between 0 and 10. The scalar reference scores were computed as a mean of the multiple rater scores, differing from the median used in [6]. The mean may yield a baseline that better matches (3). Following [6], the predicted and combined reference scores were rounded to integers before computing the Pearson’s Correlation Coefficient (PCC) and Mean Squared Error (MSE). These metrics only compare scalar references and hypotheses. A model with a probabilistic output can also be assessed on how well it matches the distribution of scores from multiple raters, and thus its ability to predict the inter-rater uncertainty. The reference distribution over integer scores, c , can be expressed as a mixture of Kronecker δ -functions at each rater’s score, $P^{\text{ref}}(c|\hat{\mathbf{x}}_i) = \frac{1}{R} \sum_r \delta(c, y_{ir}^{\text{ref}})$, showing the fraction of raters who assigned each score. This Monte Carlo approximation to a Bayesian NN [4] marginalises over the human equivalent of model uncertainty, yielding a purer reference of data uncertainty. The model’s output posterior can be discretised as $P(c|\hat{\mathbf{x}}_i, \mathbf{Y}, \mathbf{X}) = \frac{\int_{c-0.5}^{c+0.5} p(\hat{y}_i|\hat{\mathbf{x}}_i, \mathbf{Y}, \mathbf{X}) d\hat{y}_i}{\sum_{c'} \int_{c'-0.5}^{c'+0.5} p(\hat{y}_i|\hat{\mathbf{x}}_i, \mathbf{Y}, \mathbf{X}) d\hat{y}_i}$, to express the probability of each integer score. The reference and hypothesised distributions can then be compared using a discrete Kullback-Leibler (KL) divergence between $P^{\text{ref}}(c|\hat{\mathbf{x}}_i)$ and $P(c|\hat{\mathbf{x}}_i, \mathbf{Y}, \mathbf{X})$. Unlike a continuous KL divergence, this is lower-bounded by 0, easing interpretability, and follows [6] by rounding scores before evaluating.

Statistical significance, ρ , for MSE and KL divergence was computed using a two-tailed paired t -test. The PCC is not easily expressed as a sum over data points, hindering application of the central limit theorem. Instead, the Z_1^* approach was used [14]. This first computes an approximately normally distributed transformation [15] from the two PCCs being compared, then the two-tailed cumulative density of this transformed variable yields the significance.

The feature extraction followed [16]. First, a hybrid speech recognition model force-aligned the audio with the transcriptions. This was trained using Kaldi [17] on 960 hours Librispeech [18] using cross-entropy, following [6]. From the forced alignment, goodness of pronunciation [19], log phone posterior [20], log posterior ratio [20], tempo [16], and pitch [21] features were extracted. A continuous skip-gram model [22], with a 32-node recurrent NN [23] hidden layer, was trained on the training set non-silence phone sequences, and used to extract phone embeddings [16]. These were concatenated to form one feature vector per phone.

It may not be trivial to allow a GP to use sequential inputs. Work in [24] designs kernels to operate on sequences, while [11] extracts hand-crafted sentence-level features for SLA. This paper computed sentence-level inputs using a NN feature extractor. The sequence of features, with length equal to the number of phones in the sentence, was fed into a Bidirectional Long Short-Term Memory (BLSTM) layer [25], with 32 nodes per direction. The output was then pooled across all phones in the sequence with equal weights, and fed through linear and sigmoid layers. The output was then scaled to the bounds of the output score range. The NN was trained by minimising the MSE toward the combined reference score. The activations after the pooling layer formed inputs for the GP [26]. The GP and NN feature extractor were not jointly fine-tuned, as doing so may overfit easily [27]. Principle component analysis whitening,

Table 1. Using training set scores from multiple raters in a GP

Model	PCC $\uparrow$	MSE $\downarrow$	KL $\downarrow$	Inference time (s) $\downarrow$
GP _{base}	0.710	1.149	3.10	136 $\pm$ 3
GP _{repeat}	0.713	1.136	0.85	8476 $\pm$ 510
GP _{joint}				138 $\pm$ 4

estimated on the training set, was applied to these features to better abide by the diagonal covariance assumption of the scalar length hyper-parameter in the kernel. Dropout [28] with an omission probability of 60% was used before the BLSTM and linear layers. For NN training, a 10% validation set was held out without speaker disjointment.

The experiment assesses the proposed extension of allowing the GP to use the separate output scores from multiple raters in the training set, referred to as GP_{joint}. This is compared against using only a single output reference in GP_{base}, and also against a naive multiple output samples extension of repeating the training set input features in GP_{repeat}. The GP_{joint} and GP_{repeat} hyper-parameters were optimised by maximising $\mathcal{F}_{\text{joint}}$. In table 1, the PCC and MSE results may suggest that the proposed GP_{joint} approach yields performance gains over GP_{base}, but these may not be significant, with $\rho_{\text{PCC}} = 0.065$ and $\rho_{\text{MSE}} = 0.036$. The PCC and MSE compare the prediction against the mean combined reference, assuming that this scalar reference is correct, and do not consider how well the model’s output uncertainty matches the uncertainty between the raters. The KL divergence is used to assess this uncertainty matching. The results show that using the separate scores from multiple training set raters improves the KL divergence of the GP, with $\rho_{\text{KL}} < 0.001$. This suggests that the proposed extension allows the GP’s output distribution to better match that represented by the collection of scores from the multiple raters. Such a more accurate uncertainty prediction may better inform the type of feedback that should be given to the user.

The computational saving of GP_{joint} over GP_{repeat} was assessed by repeating inference 10 times on a 48 core Intel Xeon Platinum 8268 2.90GHz CPU using the Numpy pseudo-inverse implementation. The mean and standard deviation of the total time to infer the whole test set accumulated over all threads are shown in the right-side column in table 1. Inference using a GP can be run efficiently by pre-computing the matrix inverse before any test data points are seen. It therefore seems more interpretable to view the total time across the whole test set, rather than the time per data point or real-time factor, assuming that the computation time is dominated by matrix inversion. The results show that GP_{joint} is indeed faster to infer from than GP_{repeat}. GP_{joint} and GP_{base} require comparable durations to infer from. Although not shown here, this computational saving also benefits the matrix inversion computation during training.

5. CONCLUSION

This work extends the GP to allow for multiple output samples per input in the training set with computational efficiency. The hyper-parameters can be optimised by maximising the joint marginal likelihood of all output samples. Inference is performed in a fairly similar way to a standard GP, with the main deviation being in the hyper-parameters that are optimised differently. This allows the posterior to better match the uncertainty expressed by the multiple raters. Such information may better inform the type of feedback given to the user.

6. REFERENCES

- [1] C. E. Rasmussen and C. K. I. Williams, *Gaussian processes for machine learning*, MIT Press, 2006.
- [2] A. Malinin and M. Gales, “Predictive uncertainty estimation via prior networks,” in *NeurIPS*, Montréal, Canada, Dec 2018, pp. 7047–7058.
- [3] A. Kendall and Y. Gal, “What uncertainties do we need in Bayesian deep learning for computer vision?,” in *NIPS*, Long Beach, USA, Dec 2017, pp. 5574–5584.
- [4] D. J. C. MacKay, “Bayesian interpolation,” *Neural Computation*, vol. 4, no. 3, pp. 415–447, May 1992.
- [5] B. Xu, R. Kuplicki, S. Sen, and M. P. Paulus, “The pitfalls of using Gaussian process regression for normative modeling,” *PLoS ONE*, vol. 16, no. 9, Sep 2021.
- [6] J. Zhang, Z. Zhang, Y. Wang, Z. Yan, Q. Song, Y. Huang, K. Li, D. Povey, and Y. Wang, “speechocean762: an open-source non-native English speech corpus for pronunciation assessment,” in *Interspeech*, Brno, Czechia, Aug 2021, pp. 3710–3714.
- [7] K. Yu, V. Tresp, and A. Schwaighofer, “Learning Gaussian processes from multiple tasks,” in *ICML*, Bonn, Germany, Aug 2005, pp. 1012–1019.
- [8] E. V. Bonilla, K. M. A. Chai, and C. K. I. Williams, “Multi-task Gaussian process prediction,” in *NIPS*, Vancouver, Canada, Dec 2007, pp. 153–160.
- [9] R. Caruana, “Multitask learning,” *Machine Learning*, vol. 28, no. 1, pp. 41–75, Jul 1997.
- [10] M. D. Petković and P. S. Stanimirović, “Generalized matrix inversion is not harder than matrix multiplication,” *Journal of Computational and Applied Mathematics*, vol. 230, no. 1, pp. 270–282, Aug 2009.
- [11] R. C. van Dalen, K. M. Knill, and M. J. F. Gales, “Automatically grading learners’ English using a Gaussian process,” in *SLATE*, Leipzig, Germany, Sep 2015, pp. 7–12.
- [12] J. H. M. Wong, H. Zhang, and N. F. Chen, “Variational Gaussian process data uncertainty,” in *ASRU*, Taipei, Taiwan, Dec 2023.
- [13] J. Quiñero-Candela and C. E. Rasmussen, “A unifying view of sparse approximate Gaussian process regression,” *JMLR*, vol. 6, no. 65, pp. 1939–1959, Dec 2005.
- [14] J. H. Steiger, “Tests for comparing elements of a correlation matrix,” *Psychological Bulletin*, vol. 87, no. 2, pp. 245–251, 1980.
- [15] O. J. Dunn and V. Clark, “Correlation coefficients measured on the same individuals,” *Journal of the American Statistical Association*, vol. 64, no. 325, pp. 366–377, Mar 1969.
- [16] H. Zhang, K. Shi, and N. F. Chen, “Multilingual speech evaluation: case studies on English, Malay and Tamil,” in *Interspeech*, Brno, Czechia, Aug 2021, pp. 4443–4447.
- [17] D. Povey, A. Ghoshal, G. Boulianne, L. Burget, O. Glembek, N. Goel, M. Hannemann, P. Motlíček, Y. Qian, P. Schwarz, J. Silovský, G. Stemmer, and K. Veselý, “The Kaldi speech recognition toolkit,” in *ASRU*, Hawaii, USA, Dec 2011.
- [18] V. Panayotov, G. Chen, D. Povey, and S. Khudanpur, “Librispeech: an ASR corpus based on public domain audio books,” in *ICASSP*, Brisbane, Australia, Apr 2015, pp. 5206–5210.
- [19] S. M. Witt and S. J. Young, “Phone-level pronunciation scoring and assessment for interactive language learning,” *Speech Communication*, vol. 30, no. 2-3, pp. 95–108, Feb 2000.
- [20] W. Hu, Y. Qian, F. K. Soong, and Y. Wang, “Improved mispronunciation detection with deep neural network trained acoustic models and transfer learning based logistic regression classifiers,” *Speech Communication*, vol. 67, pp. 154–166, Mar 2015.
- [21] P. Ghahremani, B. BabaAli, D. Povey, K. Riedhammer, J. Trmal, and S. Khudanpur, “A pitch extraction algorithm tuned for automatic speech recognition,” in *ICASSP*, Florence, Italy, May 2014, pp. 2494–2498.
- [22] T. Mikolov, K. Chen, G. Corrado, and J. Dean, “Efficient estimation of word representations in vector space,” in *ICLR*, Scottsdale, USA, May 2013.
- [23] J. L. Elman, “Finding structure in time,” *Cognitive Science*, vol. 14, no. 2, pp. 179–211, Apr 1990.
- [24] H. Lodhi, J. Shawe-Taylor, N. Cristianini, and C. Watkins, “Text classification using string kernels,” in *NIPS*, Denver, USA, Nov 2000, pp. 563–569.
- [25] A. Graves and J. Schmidhuber, “Framewise phoneme classification with bidirectional LSTM networks,” in *IJCNN*, Montréal, Canada, Jul 2005, pp. 2047–2052.
- [26] R. Salakhutdinov and G. E. Hinton, “Using deep belief nets to learn covariance kernels for Gaussian processes,” in *NIPS*, Vancouver, Canada, Dec 2007, pp. 1249–1256.
- [27] S. W. Ober, C. E. Rasmussen, and M. van der Wilk, “The promises and pitfalls of deep kernel learning,” in *UAI*, Jul 2021, pp. 1206–1216.
- [28] N. Srivastava, G. E. Hinton, A. Krizhevsky, I. Sutskever, and R. Salakhutdinov, “Dropout: a simple way to prevent neural networks from overfitting,” *Journal of Machine Learning Research*, vol. 15, no. 56, pp. 1929–1958, Jun 2014.