

Generating Data using Monte Carlo Dropout

Kristian Miok

West University of Timisoara
Computer Science Department
Romania

Email: kristian.miok@e-uvt.ro

Dong Nguyen-Doan

West University of Timisoara
Computer Science Department
Romania

Email: dong.nguyen10@e-uvt.ro

Daniela Zaharie

West University of Timisoara
Computer Science Department
Romania

Email: daniela.zaharie@e-uvt.ro

Marko Robnik-Šikonja

University of Ljubljana
Faculty of Computer and Information Science
Slovenia

Email: marko.robnik@fri.uni-lj.si

Abstract—For many analytical problems the challenge is to handle huge amounts of available data. However, there are data science application areas where collecting information is difficult and costly, e.g., in the study of geological phenomena, rare diseases, faults in complex systems, insurance frauds, etc. In many such cases, generators of synthetic data with the same statistical and predictive properties as the actual data allow efficient simulations and development of tools and applications. In this work, we propose the incorporation of Monte Carlo Dropout method within Autoencoder (MCD-AE) and Variational Autoencoder (MCD-VAE) as efficient generators of synthetic data sets. As the Variational Autoencoder (VAE) is one of the most popular generator techniques, we explore its similarities and differences to the proposed methods. We compare the generated data sets with the original data based on statistical properties, structural similarity, and predictive similarity. The results obtained show a strong similarity between the results of VAE, MCD-VAE and MCD-AE; however, the proposed methods are faster and can generate values similar to specific selected initial instances.

I. INTRODUCTION

We live in times of big data; yet, there are many application areas that lack sufficient data for analyses, simulations, and development of analytical approaches. For example, many studies within bio-medical domain require strict and expensive experimental conditions and can produce only small samples within the allocated budget. Similar examples are domains for which data is difficult to obtain, such as rare diseases, private records, or rare grammatical structures [1]. Thus, there is a need for machine learning methods that can generate new data preserving the statistical and predictive characteristics of the original data set.

Since its introduction by Diederik et al. [2], Variational autoencoders (VAE) become one of the most used unsupervised learning methods within the family of autoencoder (AE)

techniques [3]. They are used in various problems: predicting dense trajectories of pixels in computer vision [4], anomaly detection [5], and conversion of molecular discrete representations to and from multidimensional continuous representations [6]. A short description of VAEs is provided in Section 3. Our interest in VAEs is due to their ability to generate new data [7, 8].

The main goal of this work is to introduce Monte Carlo dropout into (variational) autoencoder-based data generating methods that can provide comparable results to existing VAE generators in a shorter time. To show favorable properties of the new generators, we conduct comparisons among three groups of data sets:

- 1) original data sets,
- 2) data sets produced by the VAE generator,
- 3) data sets generated using the newly introduced MCD-VAE and MCD-AE approaches.

We compare statistics of individual attributes in each of the data sets, structures of the data sets as determined by clustering algorithms, predictive performance of machine learning algorithms trained and tested on data sets from each group, and times required for generation of new instances.

The outline of the paper is as follows. In Section 2, we shortly discuss related work. In Section 3, we introduce the methodology and architecture of our methods. Section 4 describes how the VAE, MCD-VAE and MCD-AE generators were compared followed by the results obtained in Section 5. We compare the computational performance of the three generators in Section 6 and derive conclusions in Section 7.

II. RELATED WORK

Methods that learn the distribution from existing data in order to generate new instances are of recent interest to

scientific community. Till recently, generative methods were based on models that provide a parametric specification of a probability distribution function and models that can estimate kernel density [9]. For example, [10] and Yang et al. [11] used kernel density estimation to generate new virtual instances. However, those methods work only for data sets with low dimensionality. An interesting method that generates new records using an evolutionary algorithm was proposed in [12]. This method does not take dependencies between attributes into account. The generator based on Radial Basis Function (RBF) networks [1] corrects this shortcoming but is less suitable for really high dimensional data sets (such as images and text). Two popular generators for images are VAEs [13] and Generative Adversarial Networks (GAN) [9]. Interesting combinations of those two methods were proposed by Larsen et al. [14] and Rosca et al. [15] suggesting that a GAN discriminator can be used in place of a VAEs decoder.

As the GAN generated data that can be very different from the original data set its outputs cannot be used to simulate the original data. On the other hand, the shortcoming of VAE is that the newly generated values strongly depend on the distribution of the whole training set. Hence, in case we want to generate instances similar to specific instances, e.g., outliers, this is impossible. The proposed method addresses the mentioned shortcomings of VAEs and improves upon it in terms of flexibility of the generated instances and speed of generation.

III. METHODS

We first present the background information on AE, VAE and Monte Carlo Dropout method and then explain how we can harness the power of both to produce flexible and efficient data generators. Finally, we visually demonstrate the differences between different generators on a digit recognition data set.

A. (Variational) Autoencoders

A typical AE is made of two neural networks called an encoder and a decoder. The encoder compresses the data into an internal representation and the decoder tries to decompress from this compressed representation (or latent vector) back into the original data using a reconstruction loss function [16]. VAEs inherit the architecture of classical AEs introduced by Rumelhart et al. [17]; however, their learning process uses the data to explicitly estimate *the distribution* from which the latent space is sampled [3]. Hence, VAEs store the latent variables in the form of probability distributions. As depicted in Fig. 1, VAEs resample latent values z from the generated distribution that are further transformed using the decoder network. From the Bayesian perspective the encoder is doing an approximation of the posterior distribution $p(z|x)$:

$$p(z|x) = \frac{p(x|z)p(z)}{p(x)},$$

where z denotes the hidden variable values and x the input data. As this distribution usually does not have analytical

closed form solution, we have to approximate it. In order to avoid computationally expensive sampling procedure like Markov Chain Monte Carlo (MCMC) sampling, the Variational Inference (VI) method is applied. The VI method [18] samples from the distribution for which the Kullback-Leibler divergence to the posterior distribution is minimal.

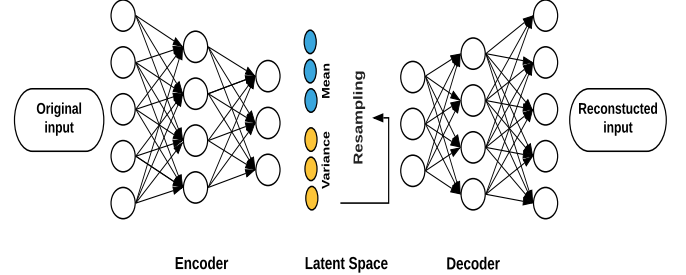


Fig. 1: Variational Autoencoder Diagram.

B. Monte Carlo Dropout Method

Deep learning is the state-of-the-art approach for many problems where machine learning is applied. However, standard deep neural networks do not provide information on reliability of predictions. Bayesian neural networks (BNN) can overcome this issue by probabilistic interpretation of model parameters. Apart from prediction uncertainty estimation, BNNs offer robustness to overfitting and can be efficiently trained even on small data sets [19]. While there exist several BNN variants and implementations, our work is based on Monte Carlo Dropout (MCD) method proposed by Gal and Ghahramani [20]. The idea of this approach is to capture prediction uncertainty using the dropout as a regularization technique. Authors prove that the use of dropout in NNs can be seen as a Bayesian approximation of the Gaussian process probabilistic models. Generating new values can be seen as the uncertainty estimation process of predicting the original instance for which generation is done [21]. The generated values shall reflect the distributional properties of the original instances.

The bias in the prediction accuracy can come from different sources. Based on where uncertainty is coming from, we distinguish: model uncertainty, data uncertainty, and distributional uncertainty. Model uncertainty describes how well the model fits the data and it can be reduced using larger training set. The data uncertainty is caused by the nature of the data set used and is irreducible by current techniques. Distributional uncertainty arises from the distributional incompatibility between the training and testing data sets. In case of the Bayesian inference, the overall uncertainty is captured with the data and model uncertainty [22]. The prediction uncertainties within the Bayesian framework can be summarized with the posterior predictive distribution (PPD) [23]. Once the posterior distribution is estimated, the PPD can be calculated using the formula:

$$p(y^*|x^*, X, Y) = \int p(y^*|f^\omega(x^*)) p(\omega|X, Y) d\omega$$

where the $p(y^*|f^\omega(x^*))$ is the likelihood function that contains the data uncertainty while the $p(\omega|X, Y)$ is the posterior distribution of the model parameters ω presenting uncertainty of the model.

The idea of MCD method is to replace the complex Bayesian process of seizing those uncertainties during the regularization using dropout. Practically, the dropout is equivalent to several forward passes through the network and recalculation of the results. At each backward pass, the model ends-up with new optimization results of the model weights. Keeping all this information, the method mimics the Bayesian inference and is equivalent to the Bayesian posterior distribution estimation [24].

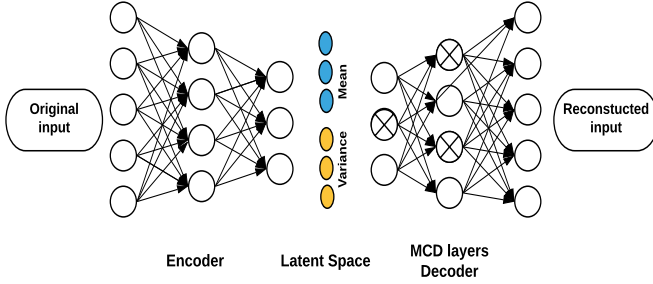


Fig. 2: Variational Autoencoder with MCD Decoder. Note the difference to Figure 1.

C. VAEs for Data Generation

For the VAE architecture (Figure 1), we use two intermediate layers (fully-connected layers) with size of M (e.g. 512) and N (e.g. 256) in the encoder. Similarly, the decoder contains two fully-connected layers with N and M neurons. To take into account various types of data sets used in our experiments, we choose the number of latent variables L to be equal to one-half of attributes present in each data set. This value is chosen in order to keep an important part of the information from which the new data can be generated.

There are two approaches to generate the data from the VAE, once the model is trained. The first approach is to generate the sampled latent vectors from the estimated normal distribution (μ, Σ) where the Σ is a diagonal covariance matrix. The sampled values are then sent through the decoder part to get the final generated instances. The second approach is to send existing instances through the trained encoder and decoder layers. In this paper, we are interested to generate new values similar to existing values present in the training set, therefore we focus on the second approach.

The process of generating data using the VAE method can be described as follows.

- 1) Obtain the distribution of latent vectors (μ_i, σ_i) with $i = 1, \dots, L$ from each value in the seeding data set by using the encoder.
- 2) Resample t times from the obtained latent space distribution, where t is the number of new instances we want

to generate for a single seeding instance as in following equation:

$$z_i = \mu_i + \sigma_i \cdot \epsilon, \text{ where } \epsilon \sim N(0, 1).$$

- 3) Decode the resampled values by the decoder.

D. MCD-VAE for Data Generation

The MCD-VAE architecture (Figure 2) has a similar structure to the VAE generator, with the exception that the MCD regularization is used within the decoder layers.

The process of generating data with MCD-VAE can be described as follows.

- 1) Obtain the distribution of latent vectors (μ_i, σ_i) with $i = 1, \dots, L$ from each value in the seeding data set.
- 2) Send the means $\mu_1, \dots, \mu_L$ through the MCD decoder t times, where t is the number of new instances we want to generate for a single seeding instance.

As evident from the above description, MCD-VAE utilizes MCD within the decoder part to get additional fine grained control over the generated instances. Namely, once the MCD-VAE is prepared for a single seeding instance, due to dropout, it can produce many different outputs by going forward through the network. This increases the speed of generation and gives the user of the generator much finer control on the generated instances.

E. MCD-AE for Data Generation

We can apply the MC dropout method also in the decoder part of AE and get the generator called MCD-AE. The structure of MCD-AE in our experiments is similar to VAE and MCD-VAE described before. The process to generate data in MCD-AE is outlined below.

- 1) For each value in the seeding data set, obtain latent vectors of size L .
- 2) Send the latent vectors through the MCD decoder. The decoder samples a new dropout mask in each of the t forward passes through the network and generates t values for a single input.

The decoder part of the MCD-AE generator is identical to the decoder in MCD-VAE. The difference between the two generators is that MCD-AE does not assume any distributional constraints for the latent space representation.

F. Visual Comparison of the Generators

We visually demonstrate the differences between the three generators (VAE, MCD-VAE, and MCD-AE). For this we have chosen a well-known MNIST data set of hand-written digits¹ and used it to train the three generators. The architecture for VAE and MCD-VAE generators contains a fully-connected layer with 1024 units and a latent layer with the size 10. The generated images are presented in Figure 3. The original seeding image is always given in the first column. In this experiment, we investigated generation of digits 9, 5, and 1

¹<http://yann.lecun.com/exdb/mnist/>

(see the three blocks of images). Digits 9 and 5 were generated from seeding instances that are written in nonstandard way, with the shape that differs from the rest of digits in their class. The digit 1 that was used as a seeding instance is written in a standard way.

The five images generated for the digit 9 using VAE (top group, first row) have the same structures as the seeding digit 9 but do not reflect much specifics of the seeding image. Contrarily, the images generated using MCD-VAE and MCD-AE (top group, second and third row) tend to better reflect the actual structure of the seeding images. The digit 5, used as a seeding instance in the middle group of images is a complete outlier - on the first sight one can not be sure if it is 5 or 6. The five generated images for digit 5 using VAE (middle group, first row) reflect all the training instances and do not take specifics of the seeding instance into account; hence, VAE generates images a bit similar to the digit 8. On the other hand, the images generated using MCD-VAE and MCD-AE better mimic the seeding image. The images generated from the seeding digit 1, written in the standard way, do not seem to differ much between the three generators (bottom group).

IV. EXPERIMENTAL SETTING

In this section, we first describe the methodology used to compare original and generated data in Section IV-A. We compare statistical, structural, and prediction properties of two data sets presented in Section IV-B. In Section IV-C, we present the data sets which served as original data in our evaluation.

A. Data Generation Experiment

To prepare a training data set for generators, the original data set is randomly split into two equal parts as shown in Figure 4. The first part is further split into the equal-sized training and generator seeding parts, while the second part of the original data set is left for evaluation. The training part is used to train the generators, while the generator seeding part is used in data generation. From each instance in the generator seeding set, two new instances were generated. Thus, the newly generated data sets are of the same size as the evaluation data sets.

In order to deal with multi-valued categorical attributes, we encode them with several binary substitute attributes, where the presence of a given categorical value in the original attribute sets the substitute variable corresponding to that value to 1. For example, for a multi-valued attribute X with three values $\{red, green, blue\}$ we form three substitute binary variables $X_{red}, X_{green}, X_{blue}$. If the original attribute contains value $X = blue$, the values of the substitute attributes are $X_{red} = 0, X_{green} = 0, X_{blue} = 1$. After the data is generated, we perform the reverse operation and decode the substitute variables into one multi-valued attribute.

B. Data Set Comparison

In evaluation, presented in Section V, we take an existing data set and based on it we generate three synthetic data

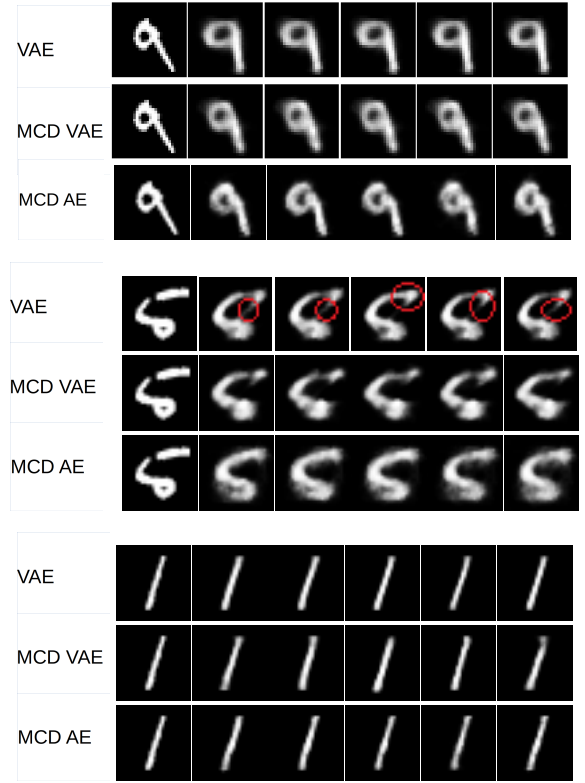


Fig. 3: The generated numbers 9, 5, and 1 are grouped in the top, middle and bottom, respectively. Each block of images contains the original seeding image (in the first column) and five generated images using VAE (the first row), MCD VAE (the second row), MCD AE (the third row).



Fig. 4: Splits of each original data set used in the experimental evaluation: the generator training set (25%), generator seeding set (25%), and evaluation set (50%).

sets, using VAE, MCD-VAE, and MCD-AE. The original and the three generated data sets are compared using a general data set comparison framework [25] which consist of three components, statistical evaluation of differences between attributes, structural comparison of data sets based on clustering, and predictive comparison based on classification models. We describe the three components below.

- 1) *Statistical evaluation of attributes.* test the mean, standard deviation, and differences in distributions between matching attributes in two compared data sets. In order to make comparison sensible for all statistics, the attributes are normalized to $[0, 1]$ scale. The value that summarizes the difference between the two data sets is calculated as the median value of pairwise attribute differences. For example, to compare mean across the

whole data set, we compute the differences in means for each of the attributes and then average these values and report it as the final measure. We therefore report Δmean and Δstd .

- 2) *Clustering performance evaluation* is performed based on the structured based distance comparing two data sets using the adjusted Rand index (ARI) [26]. The ARI value is in range of $[0, 1]$ having 0 in the case of random distributions of clusters and 1 for ideally matching clusters. The clusters of two data sets are separately computed and the process obtains the medoids for each of the clusterings. The instances in the second data set are assigned to the nearest clusters in the first data set based on the medoids computed for the first data set. The same assignment is repeated with the first data set, as instances of the first data set are assigned to clusters computed on the second data set based on the medoids from these clusters. In this way, we obtain two clusterings that contain instances from both data sets. Finally, we use ARI to summarize the clustering similarity between the two clusterings and report it as the data sets topological similarity value.
- 3) *Classification performance* based evaluation measures the predictive similarity of two data sets by comparing random forest classification accuracies on the two data sets. Let us assume that the original data set is denoted as d_1 and the generated data sets are labeled with d_2 . Both d_1 and d_2 are split into two parts, where the first parts are used to train the random forest models m_1 and m_2 , while the second parts are used for testing. Four accuracy values are computed: m_1d_1 - model computed on the first data set and evaluated on the first data set; m_1d_2 - model computed on the first data set and evaluated on the second data set; m_2d_1 - model computed on the second data set and evaluated on the first data set; and m_2d_2 - model computed on the second data set and evaluated on the second data set. If those four values are similar (in particular if accuracies on the original data set are close, i.e. accuracies of m_1d_1 and m_2d_1), one can conclude that the first and the second data set have similar predictive performance. We report only the difference of $m_2d_1 - m_1d_1$ as the predictive similarity Δacc .

C. Data Sets

To evaluate the difference between results of the three generators, we use data sets from UCI (University of California Irvine) repository [27]. The R package readMLDATA [28] was used for data manipulation. We selected classification data sets with between 500 and 1000 instances. The characteristics of the used data sets are provided in Table I.

TABLE I: The characteristics of the used data sets. The columns are: n - number of instances, a - number of attributes, num - number of numeric attributes, disc - number of discrete attributes, v/a - average number of values per discrete attribute, C - number of class values, $\text{majority \%}$ - proportion of majority class in percentages, $\text{missing \%}$ - percentage of missing values.

Data set	n	a	num	disc	v/a	C	majority (%)	missing (%)
Brest-WDBC	569	30	30	0	0.0	2	62.7	0.00
Brest-WISC	699	9	9	0	0.0	2	65.5	0.25
Credit-screening	690	15	6	9	4.4	2	55.5	0.64
PIMA-diabetes	768	8	8	0	0.0	2	65.1	0.00
Statlog-German	1000	20	7	13	4.2	2	70.0	0.00
Tic-tac-toe	958	9	0	9	3.0	2	65.3	0.00

V. EVALUATION AND RESULTS

Using the above described data sets we evaluated the quality of data generators. In Table II we compare the original data set with the generated data set using VAE architecture. The results comparing the original data set with the MCD-VAE and MCD-AE generators are presented in Tables III and IV, respectively. For comparison we use the statistical, structural, and predictive criteria, described in Section IV-B, i.e. the average difference in means (Δmean) and standard deviation (Δstd), similarity of produced clusters expressed with Adjusted Rand Index (ARI), and differences in predictive accuracy Δacc ($m_2d_1 - m_1d_1$).

TABLE II: Comparison between the original data and VAE generator.

Data set	Δmean	Δstd	ARI	Δacc
Brest-WDBC	-0.161	-0.089	0.909	-0.024
Brest-WISC	-0.069	0.001	0.970	-0.045
Credit-screening	-0.078	-0.041	0.474	-0.068
PIMA-diabetes	-0.171	-0.047	0.446	-0.015
Statlog-German	-0.040	0.040	0.167	-0.000
Tic-tac-toe	-	-	0.133	-0.092

TABLE III: Comparison between the original data and MCD-VAE generator.

Data set	Δmean	Δstd	ARI	Δacc
Brest-WDBC	-0.045	-0.044	0.876	-0.008
Brest-WISC	0.011	0.013	0.916	-0.011
Credit-screening	-0.028	-0.038	0.447	-0.061
PIMA-diabetes	-0.022	-0.028	0.715	-0.007
Statlog-German	-0.016	0.028	0.243	-0.001
Tic-tac-toe	-	-	0.122	-0.017

TABLE IV: Comparison between the original data and MCD-AE generator.

Data set	Δmean	Δstd	ARI	Δacc
Brest-WDBC	-0.059	-0.048	0.746	-0.014
Brest-WISC	0.004	0.021	0.994	-0.020
Credit-screening	-0.046	-0.036	0.393	-0.072
PIMA-diabetes	-0.077	-0.037	0.551	-0.012
Statlog-German	-0.030	0.021	0.235	0.000
Tic-tac-toe	-	-	0.224	-0.158

Comparing the results in Tables II, III, and IV, we can see that differences between the original and generated data are small. There is no clear pattern which of the three generators is better. We can conclude that all of them are useful, while minor differences in the quality of the generated data may depend on the structure of a data set. However, it can be

observed that MCD-VAE provide slightly better classification performance than VAE and MCD-AE based on the compared Δ acc value. On the other hand, for *Breast-WDBC*, *Breast-WISC* and *Credit-screening* datasets VAE generator has the better clustering performance than the two newly introduced generators.

VI. COMPARING EFFICIENCY OF GENERATORS

In order to compare the data generation time (in seconds) of VAE, MCD-VAE, and MCD-AE, we measure the time for 100 repetitions of the data generating process using the above described data sets. To get reliable measurements, we resample each seeding instance 1000 times (instead of 2 times as in the previous experiments). Table V reports the mean and standard deviation of the measured times. We generate data sets as described in Section III. For VAE, the instances in seeding data sets are encoded to obtain the latent values, then the latent values are resampled and decoded. For MCD-VAE and MCD-AE, we obtain the mean values with the seeding instances and obtain the generated data using the MCD decoder.

TABLE V: Comparison of time required for data generation in seconds.

Datasets/Models	VAE [s.d.]	MCD-VAE [s.d.]	MCD-AE [s.d.]
Breast-WDBC	1.04 [0.018]	0.89 [0.022]	0.89 [0.020]
Breast-WISC	0.90 [0.019]	0.85 [0.037]	0.89 [0.011]
Credit-screening	1.00 [0.030]	0.93 [0.025]	0.94 [0.010]
PIMA-diabetes	0.91 [0.034]	0.85 [0.021]	0.85 [0.016]
Statlog-German	1.07 [0.018]	1.03 [0.018]	1.10 [0.045]
Tic-tac-toe	0.99 [0.026]	0.93 [0.039]	0.94 [0.012]

The MCD-VAE and MCD-AE generators are consistently slightly faster than the VAE generator (between 5-10%). Although the MCD-AE generator is architecturally simpler, it is not faster than the MCD-VAE generator. The datasets used are relatively small, hence, for the larger datasets, we expect larger differences.

VII. CONCLUSIONS AND FURTHER WORK

We constructed and compared three generators of semi-artificial data. The VAE generator is based on the variational autoencoder architecture while the MCD-AE and MCD-VAE employ Monte Carlo dropout within autoencoders and variational autoencoders. The comparison of the generated data sets based on statistical, structural, and predictive properties shows that the three generators produce similar data sets which are highly similar to the original data.

The advantages of the proposed Monte Carlo dropout employed within VAE and AE over the existing VAE method can be summarized with the following two points:

- Improved speed. Based on the results presented in Table V we can conclude that generating data using MCD-VAE and MCD-AE is slightly faster than using the VAE generator.
- Greater flexibility. The MCD-VAE and MCD-AE methods generate data similar to specific selected seeding instances. This can be very useful if the provided seeding instances are outliers or instances of special interest.

For example, in image generation, the newly generated images will be closer to the original one even when the original image is different from the rest of the images in the training set.

The advantage of the MCD-AE method over MCD-VAE method is that it does not make any distributional assumptions during the latent space generation. The information received from the encoder part is directly introduced into the MCD decoder. The time required for data generation using MCD-AE is similar to MCD-VAE. The more detailed differences between these generators are left for further investigation.

With methodological development of deep learning, the models that can estimate the distributions, e.g., the variational autoencoders, are becoming increasingly important. Hence, our further work will focus on investigating new and improving existing architectures that can generate new data efficiently and reliably. Further, we aim to test those architectures within different application contexts. As biomedical imaging is expensive and limited by the budget, our goal is to investigate data generation within this field. The Python code of the proposed generators is publicly available².

Acknowledgement

The work was supported by the Slovenian Research Agency (ARRS) core research programme P6-0411 (Marko Robnik-Šikonja). The research was carried out in the frame of the project Bioeconomic approach to antimicrobial agents - use and resistance financed by UEFISCDI by contract no. 7PC-CDI / 2018, cod PN-III-P1-1.2-PCCDI-2017-0361 (Kristian Miok and Daniela Zaharie). This project has also received funding from the European Unions Horizon 2020 research and innovation programme under grant agreement No 825153 (EMBEDDIA) (Kristian Miok and Marko Robnik-Šikonja).

REFERENCES

- [1] Marko Robnik-Šikonja. Data generators for learning systems based on RBF networks. *IEEE transactions on neural networks and learning systems*, 27(5):926–938, 2015.
- [2] P Kingma Diederik, Max Welling, et al. Auto-encoding variational bayes. In *Proceedings of the International Conference on Learning Representations (ICLR)*, 2014.
- [3] Carl Doersch. Tutorial on variational autoencoders. *arXiv preprint arXiv:1606.05908*, 2016.
- [4] Jacob Walker, Carl Doersch, Abhinav Gupta, and Martial Hebert. An uncertain future: Forecasting from static images using variational autoencoders. In *European Conference on Computer Vision*, pages 835–851. Springer, 2016.
- [5] Jinwon An and Sungzoon Cho. Variational autoencoder based anomaly detection using reconstruction probability. *Special Lecture on IE*, 2:1–18, 2015.
- [6] Rafael Gómez-Bombarelli, Jennifer N Wei, David Duvenaud, José Miguel Hernández-Lobato, Benjamín

²<https://github.com/KristianMiok/MCD-VAE>

- Sánchez-Lengeling, Dennis Sheberla, Jorge Aguilera-Iparraguirre, Timothy D Hirzel, Ryan P Adams, and Alán Aspuru-Guzik. Automatic chemical design using a data-driven continuous representation of molecules. *ACS central science*, 4(2):268–276, 2018.
- [7] Daniel Jiwoong Im, Chris Dongjoo Kim, Hui Jiang, and Roland Memisevic. Generating images with recurrent adversarial networks. *arXiv preprint arXiv:1602.05110*, 2016.
- [8] Zhiting Hu, Zichao Yang, Xiaodan Liang, Ruslan Salakhutdinov, and Eric P Xing. Toward controlled generation of text. In *Proceedings of the 34th International Conference on Machine Learning-Volume 70*, pages 1587–1596. JMLR. org, 2017.
- [9] Ian Goodfellow, Jean Pouget-Abadie, Mehdi Mirza, Bing Xu, David Warde-Farley, Sherjil Ozair, Aaron Courville, and Yoshua Bengio. Generative adversarial nets. In *Advances in neural information processing systems*, pages 2672–2680, 2014.
- [10] Der-Chang Li and Yao-San Lin. Using virtual sample generation to build up management knowledge in the early manufacturing stages. *European Journal of Operational Research*, 175(1):413–434, 2006.
- [11] Jing Yang, Xu Yu, Zhi-Qiang Xie, and Jian-Pei Zhang. A novel virtual sample generation method based on gaussian distribution. *Knowledge-Based Systems*, 24(6):740–748, 2011.
- [12] Cinzia Meraviglia, Giulia Massini, Daria Croce, and Massimo Buscema. Gend an evolutionary system for resampling in survey research. *Quality and Quantity*, 40(5):825–859, 2006.
- [13] Karol Gregor, Ivo Danihelka, Alex Graves, Danilo Jimenez Rezende, and Daan Wierstra. Draw: A recurrent neural network for image generation. *arXiv preprint arXiv:1502.04623*, 2015.
- [14] Anders Boesen Lindbo Larsen, Søren Kaae Sønderby, Hugo Larochelle, and Ole Winther. Autoencoding beyond pixels using a learned similarity metric. *arXiv preprint arXiv:1512.09300*, 2015.
- [15] Mihaela Rosca, Balaji Lakshminarayanan, David Warde-Farley, and Shakir Mohamed. Variational approaches for auto-encoding generative adversarial networks. *arXiv preprint arXiv:1706.04987*, 2017.
- [16] Pierre Baldi. Autoencoders, unsupervised learning, and deep architectures. In *Proceedings of ICML workshop on unsupervised and transfer learning*, pages 37–49, 2012.
- [17] David E Rumelhart, Geoffrey E Hinton, and Ronald J Williams. Learning internal representations by error propagation. Technical report, California Univ San Diego La Jolla Inst for Cognitive Science, 1985.
- [18] Michael I Jordan, Zoubin Ghahramani, Tommi S Jaakkola, and Lawrence K Saul. An introduction to variational methods for graphical models. *Machine learning*, 37(2):183–233, 1999.
- [19] Alp Kucukelbir, Dustin Tran, Rajesh Ranganath, Andrew Gelman, and David M Blei. Automatic differentiation variational inference. *The Journal of Machine Learning Research*, 18(1):430–474, 2017.
- [20] Yarin Gal and Zoubin Ghahramani. Dropout as a bayesian approximation: Representing model uncertainty in deep learning. In *International conference on machine learning*, pages 1050–1059, 2016.
- [21] Kristian Miok. Estimation of prediction intervals in neural network-based regression models. In *2018 20th International Symposium on Symbolic and Numeric Algorithms for Scientific Computing (SYNAS)*, pages 463–468. IEEE, 2018.
- [22] Andrey Malinin and Mark Gales. Predictive uncertainty estimation via prior networks. In *Advances in Neural Information Processing Systems*, pages 7047–7058, 2018.
- [23] Pavel Myshkov and Simon Julier. Posterior distribution analysis for bayesian inference in neural networks. In *Workshop on Bayesian Deep Learning, NIPS*, 2016.
- [24] Yarin Gal. Uncertainty in deep learning. *University of Cambridge*, 2016.
- [25] Marko Robnik-Šikonja. Dataset comparison workflows. *International Journal of Data Science*, 3(2):126–145, 2018.
- [26] Lawrence Hubert and Phipps Arabie. Comparing partitions. *Journal of classification*, 2(1):193–218, 1985.
- [27] Kevin Bache and Moshe Lichman. Uci machine learning repository, 2013. URL <http://archive.ics.uci.edu/ml>, 5, 2013.
- [28] Petr Savicky. readMLData: Reading machine learning benchmark data sets from different sources in their original format, 2012.