

Generation of Intracranial Hemorrhage CT Scan Images using GAN and VAE Frameworks for Improved Medical Imaging

Huda Mirza Saifuddin¹ and Vani Ashok²

^{1,2}JSS Science and Technology University, Mysuru, Karnataka, India

Email: hudasai91@gmail.com, vanisj@sjce.ac.in

Abstract—Intracranial hemorrhages are one of many neurological diseases that can be diagnosed and followed up on with the help of medical imaging. However, due to privacy concerns and need for more data, the availability of annotated medical imaging datasets, particularly those pertaining to severe illnesses like intracranial hemorrhages are frequently constrained. In order to solve the lack of medical data, synthetic data generation techniques like Generative Adversarial Networks (GAN) and Variational Autoencoders (VAE) have shown significant improvements. This research approaches the synthesis of intracranial hemorrhage computer tomography images through GAN and VAE constructs to advance deep-learning models for ICH identification. Performance evaluation of the proposed methodology relies on the dice coefficient measurements during synthetic dataset comparison against the small real-world dataset.

Index Terms—Intracranial Haemorrhage, Generative Adversarial Networks, Variational Autoencoders.

I. INTRODUCTION

Medical experts define Intracranial haemorrhage (ICH) as dangerous bleeding within the skull that develops anywhere within the brain region or adjacent areas [1-4]. This medical condition can sometimes become fatal. ICH presents itself in various forms as epidural, subdural, subarachnoid, intracerebral and intraventricular haemorrhages. Vein tears from falls and head trauma mainly produce subdural haemorrhages whereas epidural haemorrhages occur following head damage that breaks arteries. A subarachnoid haemorrhage manifests as a sudden and severe headache because the brain tissue bleeds between itself and the surrounding membrane because of aneurysm rupture. Adults commonly experience intracranial haemorrhages, since this form of bleeding occurs within brain tissue while trauma is caused especially when high blood pressure and vascular problems occur. Intraventricular haemorrhage predominantly affects the brain tissue in its ventricles and most often develops in premature infants and seriously ill adults. Intracerebral haemorrhage develops due to various catalysts which include head trauma along with blood pressure unsteadiness, vascular disorders, drug abuse, coagulation abnormality and brain tumors. When ICH occurs there is a series of crucial symptoms that include sudden severe headaches combined with nausea, vomiting, disorientation, seizures and neurological damage affecting speech, mobility or dexterity.

Current medical practice uses MRI scans or CT scans to diagnose this condition and it involves patient stabilization followed by intracranial pressure management along with surgical intervention for clot removal or vein repair when needed. The patient's survival chances together with the extent of brain damage depend heavily on the location of the bleeding and the speed of medical response after the haemorrhage occurs.

CT scans are indeed vital in diagnosing and managing ICH, and especially because they can give a quick diagnosis of brain haemorrhage [5-8]. Disorders suspected to be ICH, particularly in the emergency setting where the diagnosis should not wait for a long time, CT scanning is often the first investigative procedure to be done. A wide array of haemorrhages such as epidural, subdural, intracerebral and subarachnoid, can be easily identified with the help of it. Due to the fact that blood is denser than most other tissues on the CT scan, haemorrhages can be easily distinguished from the normal brain tissue since one will appear as a bright white compared to the other. CT can provide crucial information including the size, position, and volume of the bleed and is very useful in detecting acute haemorrhages, particularly in the initial hours after commencement. Determining the condition's severity, making surgical decisions, and estimating the chance of consequences such as brain herniation all depend on this information. For example, a subdural haematoma manifests as a crescent-shaped collection of blood, whereas an epidural haematoma usually manifests as a biconvex mass. CT scans are also beneficial for tracking the course of bleeding, especially while haemorrhages are changing or after surgery, in order to check for problems such as rebleeding. Even though CT is quite sensitive for identifying larger, more visible bleeding, it might not be as good at identifying smaller or earlier haemorrhages, including subarachnoid haemorrhage, where MRI might be a better option for more comprehensive imaging. Notwithstanding this drawback, CT is still the recommended diagnostic method in emergency situations because of its accessibility, quickness, and capacity to deliver the data required for prompt clinical judgement.

Due to deficiency of appropriate datasets in many sectors, the necessity for synthetic data has grown in importance. Real-world data can be scarce, costly, or difficult to get in fields like healthcare, autonomous driving, and rare event prediction, especially when specialized or labelled data is needed. Furthermore, the real-world data for training the machine learning models is relatively difficult because of the restrictions of privacies and rules involved for the particular data like personal or medical records. With this, synthetic data comes as a solution as it generates large, diverse, and clearly labelled datasets that can mimic the statistical properties of real data but do not offend privacy issues. This ability helps in overcoming data sparse, in making sure that in current datasets rare occurrences are captured and it makes it possible to build and test models faster and better. Synthetic data is now used to support AI and machine learning because real-world datasets can be many times scarce and small or unsuitable for the development of AI applications that require large volumes of data.

In reference to objectives of the study, this work largely deals with the generation of synthetic CT images of intracranial hemorrhages using GAN and VAE frameworks. The reality in this approach also helps to moderate the existence of inadequate data, not to mention, improving the training of deep learning models for accurate ICH diagnosis. The FID is considered as the primary valuation criterion since it assesses the quality of the synthetic data. FID measures the difference in distribution of generated and real images hence making it effective in determining how realistic the generated dataset is. Moreover, approaches for training state-of-the-art ICH detection models using the proposed synthetic dataset are evaluated and it shows promising results than the ICH detection models trained on real scenario data. This paper demonstrates how synthetic data can be used to aid such developments towards improving the availability of medical imaging and diagnosis.

II. RELATED WORK

Based on our survey, the T1-MRI domain was represented by a modified T1-MRI domain VAE-GAN model developed by Stenzel Cackowski et al. [9]. ImUnity is another tool used in MRI harmonization with improved feature enhancement which is a modified VAE-GAN model. This design coincided with their proposed 2.5D model that was based on the VAE-GAN. With the objective of protecting the biological information, it provided meaningful outcomes and allowed to dispel an impact of peculiar data sets.

In case of medical picture synthesis, Feihong Li et al. [10] put forward a new model architecture based on VAE and GAN. In addition to the useful arterial spin labelling, a new model namely GAN-based model was introduced for synthesizing the labelling picture. Technically, it was integrated into the structure of GAN as its generator part, yet its role in the process of generation is significant. Several GANs suitable for data augmentation in colorectal cancer histopathology picture segmentation were evaluated by R Sujatha et al. [11]. For the purpose of applying, it to cancer prediction, they compared the performance of the two GAN models, named as deep convolutional GAN (DC-GAN) and variational autoencoder GAN (VAE-GAN) in generating realistic synthetic images.

To have a more disjointed interpretation of the attributes, later in the similar context of VAE, Irem Cetin et al. [12] developed the Attri-VAE that incorporated the attribute regularization term to ascertain the correspondence of clinical and medical imaging qualities to several regularized dimensions in the resulting latent space. In order to describe the distribution of 3D MR brain volumes, Anna Volokitin et al. [13] used a Gaussian model that captures the relationships between slices with a 2D slice VAE. Over the slice direction, they calculated the sample mean and covariance in the 2D model's latent space.

Recently introduced deep learning GANs were used by Maayan Frid-Adar et al. [14] to create artificial medical images. Additionally, they demonstrated how generated medical images can enhance CNN's performance for medical image classification and be used for synthetic data augmentation. To create high-quality liver lesion ROIs, they first took advantage of GAN architectures. then introduced a brand-new CNN-based liver lesion classification scheme. Lastly, they compared the CNN's performance after training it with both our synthetic data augmentation and traditional data augmentation. Shrinivas D. Desai et al. [15] demonstrated the use of GAN for limited labelled datasets in breast cancer detection. ImageJ, a visualization tool, was used to assess how similar the synthetic and actual photos were. Medical professionals assisted in conducting a visual Turing test to validate the suggested model.

In order to address the issue of gradient disappearance in GANs, Jiwei Lv, et al. [16] shown that the output layer of the generator employs the Tanh activation function, the remaining levels use ReLU, and all discriminator layers have LeakyReLU activation functions. Simple initial fraction (ISs) and complex initial fraction (ISc) indices were developed from the two elements of picture quality and image production diversity, respectively, after the enhanced DCGAN model was validated on the MNIST dataset.

S. Nikkath Bushra et al. [17] assessed DCGAN for COVID-19 detection using CT/X-ray chest scans. Using X-ray and CT scan images taken from various infected individuals, they sought to provide an overview of the details of the systems recently developed to diagnose novel COVID-19. Several studies have been done on DCGAN for Data Augmentation of Chest X-ray Images where, Sagar Kora Venu et al. [18] assessed the performance. For data augmentation in their study, they employed generative modelling such as the DCGAN and generate new chest X-ray images of the limited samples of the underrepresented class. To solve the diabetic retinopathy grading problem, Shuqiang Wang et al. [19] developed a multichannel-based GAN with semi-supervision. To achieve this, the generative model had to be a multichannel model to generate a set of sub-fundus images that match with the scattering properties of diabetic retinopathy. The proposed semi-supervised MGAN could be utilizing high-resolution fundus images without compression for identifying the features of the lesion and depend less on labelled data. Anuja Negi et al. [20] presented an approach based on GAN for segmenting the tumor in breast ultrasound image. Sharpening their focus on precise segmentation of colorectal tumor, Xiaoming Liu et al. [21] have embarked on applying the label assignment generative adversarial network. The works related to the present research deals with the use of GAN and VAE architectures to produce synthetic images used for different purposes.

III. PROPOSED METHOD

In the following sub-section, the techniques for creating synthetic data of intracranial haemorrhage from brain CT images using GAN and VAE are described. Diagnostic CT scan are made from high-definition images that can be used to identify cerebral haemorrhage. These images with better spectral and spatial resolution help the doctor in a better way to detect the blotting areas of the body.

A. Generative adversarial networks

One of the remarkable classes of the AI models termed as GANs has improved the areas of generative modelling and picture synthesis heavily. Generally known as GANs, the method was proposed by Ian Goodfellow and his team in 2014 has become one of the most efficient and famous deep learning types. It is important to note that the present day GANs fall under the category of unsupervised learning methodologies. They are very effective in generating high quality synthetic data which resembles the actual data and instances of a given world. They can create images, sounds, texts and other data kinds that resemble samples of those offered in real-world due to their unique adversarial game system which includes the Generator and Discriminator. Recently, GANs have been incorporated into many fields such as data augmentation, image synthesis and drug discovery. These are further expected in the study of generative modelling and overall advancement in artificial intelligence in a various application when the research of GANs happens.

A GAN is made up of two neural networks: the Discriminator and Generator networks which play a adversarial game as illustrated below in Fig.1. As a matter of fact, a GAN's discriminator is simply a classifier. He will try

to distinguish between raw data and fake data from the generator. Using the discriminator feedback, the generator part which is a part of GAN provides the fake data. This makes it possible to improve the ability of the discriminator in giving the label 'real' to its output. Thus, high-quality synthetic data is produced by the interaction of these two networks. Such GANs are trained through what is described as the min-max optimization process. It is to reach the point for the Generator to generate realistic samples as much as possible as well as for the Discriminator to fail to distinguish between the real and fake data.

Some random noise vectors which are generated by the Generator at the start of training are provided to the network. These vectors are converted into synthetic data samples and given to Discriminator as well actual data samples is given. In the case of each sample, there is a score which is provided by the discriminator telling it to what extent it is right in the classification as true or bogus. The Discriminator parameters are then updated by back-propagation of the gradients of these probabilities through it. This is because the Generator also aims at generating samples that are as close as possible to the Discriminator's ability to detect as genuinely fake. The values of the Generator parameters are then adjusted in a manner of backpropagating the derivative of these probabilities. The generator is trained to produce samples as real as possible and the discriminator's objective is to distinguish between real and fake samples in an attempt.

Discriminator: As the name implies, the discriminator is the component that is designed to identify between actual target distribution data samples and artificial data samples produced by the generator. Its work is to distinguish between target distribution data samples and fake product samples generated by the discriminator. The discriminator receives input in the form of data samples and returns a probability score of the input belonging to the real world. The discriminator in the same way also comprises a deep neural network, that can be for instance a CNN for image data, or other structures for other types of data. During the training process, the parameters that appear are adjusted to enhance its ability to classify actual and real data from fake one. Based on the above two modules, the main role of the discriminator is to identify such data as realistic and measure the discriminant ability at the maximum level. The training process improves in differentiating real data from other data apart from the sample data created by the developing generator.

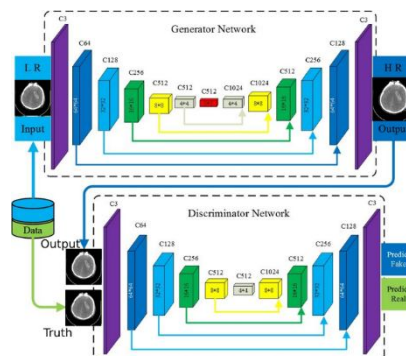


Fig 1. Architecture of GAN

Generator: The primary innovation of GANs is the competitive game between the discriminator and the generator. The discriminator and generator are trained in a cycle, with each iteration improving the performance of one while challenging the other. The goal of the generator is to provide fake data that is convincing enough for the discriminator to believe it to be authentic. The discriminator, on the other hand, seeks to accurately separate real data from the generator's fake data. As the discriminator becomes stronger at accurately distinguishing between actual and bogus data, the generator gets better at creating more realistic samples. Through this continuous adversarial process, the generator is compelled to generate increasingly realistic samples, ultimately settling on data that is almost exactly the same as the distribution of real data. When the discriminator is unable to safely discern between actual and fake data and the generator produces high-quality synthetic data, a competitive equilibrium is reached.

Implementation of GAN

Implementation of Generator: It accepts an input of a normal distribution-based random vector. This random vector is routed through convolution layers after passing through a dense layer and being reshaped. Here, convolution layers downscale our latent vector. Our downscaled latent vector is then upsampled using Conv2DTranspose following a sequence of convolution batch normalization and LeakyRelu layers. Generator generates an image with a final output layer that is 128 by 128 by 3. Utilizing hyperbolic tangent as activation,

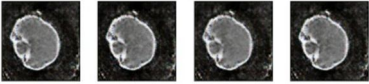
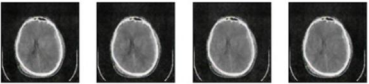
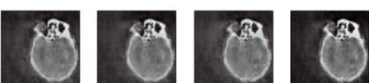
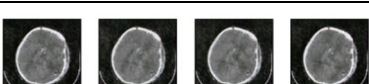
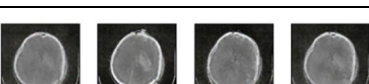
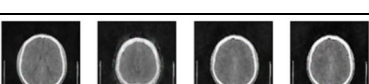
the final layer of the generator compresses the value to fall between -1 and 1. The generator model resembles a straightforward autoencoder paradigm in which input data is first down sampled and then up sampled.

Implementation of Discriminator: Discriminator models use images that are 128 by 128 by 3 and can be either real or fake. A single neuron is fed this input image after it has been flattened and downsampled using a Convolution layer so that it can tell the difference between a real and fake image. Since the final layer activates using the sigmoid function, its output value ranges from 0 to 1. Here, a true image has a value larger than 0.5 while a fraudulent image has a value less than 0.5. The output of the discriminator is used to train the generator. Table 1. displays a few samples of synthetic data generation images using DCGAN.

The implemented architecture represents a Generative Adversarial Network (GAN), comprising two primary components: the Generator and the Discriminator. The Generator is responsible for synthesizing realistic images from a random noise vector of size $\text{latent_dim} = 100$. It begins with a dense layer that transforms the input noise into an initial feature map, which is reshaped into an image-like structure. This structure is progressively refined through a series of convolutional layers and transposed convolutional (upsampling) layers, interleaved with LeakyReLU activations and BatchNormalization to ensure stable training and effective gradient flow. The final layer employs a tanh activation function, ensuring the generated pixel values are normalized to the range $[-1, 1]$, which is standard in GAN training.

While the Discriminator is binary in nature and is trained to determine the capability of a given image as a genuine image or an image that has been synthesized artificially. The input of this model is image of size (SIZE, SIZE, 3) and using a series of convolutional layers, it extracts hierarchal features while downsampling the input. It is to note that each layer of the convolution has been completed with LeakyReLU activations and a normalization step using BatchNormalization in order to stabilize the training process. The extracted features are flattened and are later on fed through a dense layer having one neuron which outputs the probability that the input image is real using the sigmoid activation function.

TABLE I. SYNTHETIC DATA GENERATED USING DCGAN

Type of ICH	Synthetic data samples
Epidural	
Subdural	
Subarachnoid	
Intraventricular	
Intraparenchymal	
No haemorrhage	

Both the Generator and Discriminator get trained in an adversarial way, which means that the Generator will continually try to produce images that are so good that Discriminator will not be able to distinguish them from the actual images while Discriminator also works in such a way that, it will try to identify the images produced by the generator as fake as possible. This feedback mechanism allows the generator, over the course of their training, to create more realistic images as a result of the discriminator unable to distinguish their generated data. This architecture makes good use of the interaction or combination of the two networks in the production of the synthetic images. In a GAN, the Discriminator and the Generator are both trained simultaneously and

adversarially to create high-quality fake data or in general, data that has not been witnessed before, perhaps images. The Generator is designed to generate a data similar as real life data while the Discriminator is designed to distinguish between fake and original data.

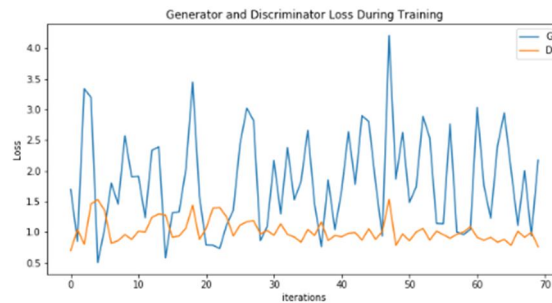


Fig 2. GAN Generator and Discriminator Loss

In a GAN, the Discriminator and the Generator are both trained simultaneously and adversarially to create high-quality fake data or in general, data that has not been witnessed before, perhaps images. The Generator is designed to generate a data similar as real life data while the Discriminator is designed to distinguish between fake and original data. The changes in losses over time can be used to gauge how well the training is going. Ideally, the Generator loss and Discriminator loss both decrease, demonstrating that both networks are becoming more adept at their respective duties.

B. Variational Autoencoder

A revolutionary development in the fields of generative modelling and unsupervised learning is the use of variational autoencoders. VAEs, which Kingma and Welling [23] proposed in 2013, combine components of probabilistic modelling and autoencoders, offering a strong foundation for learning detailed data representations unsupervised. VAEs are a flexible tool for tasks like data generation, dimensionality reduction, and data imputation because they provide a novel method for producing fresh data while concurrently learning significant latent structures. The encoder and the decoder are the two fundamental parts of a variational autoencoder as seen in Fig 3. Together, these elements provide an autoencoding structure with a distinctive probabilistic twist that sets VAEs apart from traditional autoencoders.

Encoder: The encoder transfers an input data point to a latent space, where each point in the space represents a condensed and abstract representation of the input data. The encoder is frequently realized as a neural network. However, the encoder creates a probability distribution across all potential points rather than a single point in the latent space. The mean and variance parameters of this distribution, which is often a multivariate Gaussian, define it.

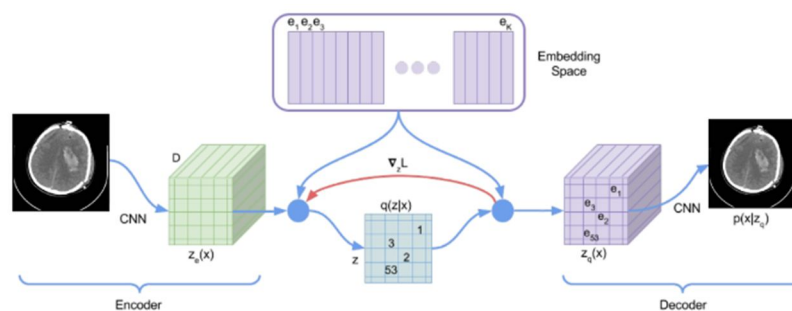


Fig 3. Architecture of VAE

Decoder: A point from the latent space is used by the decoder, a different neural network, to reassemble the initial input data. In order to convey the inherent uncertainty in the reconstruction process, the decoder's output, like that of the encoder, is not a deterministic reconstruction but rather a distribution across possible reconstructions.

KL-divergence serves as the loss function for the VAE, whose objective is to reduce the discrepancy between a dataset's original distribution and its assumed distribution. As a result, the total loss is composed of two terms: reconstruction error and KL-divergence loss as shown in equation 1.

$$\text{Loss} = L(x, x^{\wedge}) + \sum_j KL(q_j(z|x) || p(z)) \quad (1)$$

where $L(x, x^{\wedge})$ is the reconstruction loss that measures how well the reconstructed output $x^{\wedge}$ resembles the original input x . $\sum_j KL(q_j(z|x)||p(z))$ is the KL divergence that regularizes the learned latent distribution $q_j(z|x)$ to be close to the prior $p(z)$ and ensures the latent space is structured and prevents overfitting.

The main advancement of VAEs is their probabilistic basis. Traditional autoencoders provide deterministic mappings by concentrating on reducing the reconstruction error between the input and the output. VAEs, in contrast, include a probabilistic component that enables them to capture the underlying variability and uncertainty in data.

One of VAE's key procedures is to set variational lower bound on the data probability. VAEs also learn reasonable representations in the latent space as well as the efficient generative skills by working with this lower bound.

The former is the KL divergence term that makes the distribution of the encoder force a prior distribution usually a Gaussian distribution while the latter one makes the output data close to the input data. During the training of VAEs the variational lower limit is estimated with the help of stochastic gradient descent or one of its modifications. For each iteration of the training, a batch of the input data is used to decide what the mean and variance of the latent distribution are. A point is then taken from this distribution and sent into the decoder to reconstruct the data. The optimization objective is to minimize the negative variational lower limit, which maximizes the likelihood of the data while minimizing the divergence between the prior and the encoder's distribution.

Implementation of VAE

An encoder, a latent space, and a decoder make up a VAE. While the decoder creates data samples from points in this space, the encoder maps input data to a distribution in the latent space. Such architectures can be created more easily thanks to Keras' layers and model-building features. The VAE has mainly 3 layers:

- Encoder: The encoder takes the input data and maps it to the parameters of the latent space distribution. It is often implemented as a neural network. The architecture of the encoder can be specified using Keras' Sequential or Functional API. It frequently consists of numerous layers, such as dense or convolutional layers, followed by the mean and log variance of the latent distribution.
- Sampling layer: A sample layer is added to create points from the latent distribution. The encoder's computed mean and log variance are used to sample from a normal distribution.
- Decoder: Data samples are produced by the decoder from points in the latent space. The Keras API can be used to define the decoder architecture, just like the encoder. Its architecture is an inverse mirror of the encoders.

For training the VAE model, in this study plain vanilla VAE model is used with a smaller encoding of 512 which produces substantially better results compared to the 2048 larger encoding. The model training is performed with the following hyperparameters and optimizers:

- Learning rate schedule: The learning rate at each epoch is determined using an exponential decay formula by the function `exponential_decay_fn(epoch)`. This function returns the matching learning rate after receiving the epoch number as input. Every 10 epochs, the learning rate is lowered by a factor of 0.1 from its initial value of 0.001. As training advances, this method gradually lowers the learning rate, which can assist in making the optimization process more reliable and effective. The function then initializes a Keras `LearningRateScheduler` callback. Prior to each epoch, this callback modifies the learning rate in accordance with the values supplied by the `exponential_decay_fn`.
- Early Stopping: The callback with the identifier `early_stopping_cb` defines for early stopping. Monitoring the validation loss during training is done using this callback. The training will end if the validation loss does not decrease after the patient parameter-specified number of epochs (in this case, 5 epochs) has passed. When the early stopping requirement is satisfied, the model's weights would be reset to the best weights discovered during training because the `restore_best_weights` option was set to `True`. The training was performed for 100 epochs with early stopping.

The presented architecture implements a convolutional autoencoder, consisting of two main components: an encoder and a decoder. This design facilitates the encoding of high-dimensional input images into a compact, low-dimensional representation of size `ENCODING_SIZE = 1024` and their subsequent reconstruction. The model is particularly suited for tasks such as dimensionality reduction, image denoising, and anomaly detection.

Encoder: The encoder compresses the input images, which have dimensions `[IMG_WIDTH, IMG_HEIGHT, 3]` (with 3 representing RGB channels), into a latent representation. It achieves this through a series of layers:

- Spatial Features: These extract the spatial feature from the input images by using 3×3 filters of the same type as Convolutional Layers. This would introduce non linearity with the ReLU activation function which allows the model to learn complex patterns.
- These layers are called MaxPooling Layers, which reduce spatial dimensions of the feature maps making hierarchical feature extraction easier but at the same time reduction in the computational complexity.
- In this case, a Flatten Layer is employed, which converts the spatial feature maps into a one-dimensional vector to be fed into a dense layer.
- Flattened vector (dense layer): maps the flattened vector into a latent space of size ENCODING_SIZE, as the compact representation of the input.

This special structured design of the encoder allows significant reduction in the dimensionality of the input as it captures the salient features of the input data in a more effective way.

Decoder: The decoder is responsible for reconstructing the original image from the latent representation. It performs the reverse operations of the encoder using the following layers:

- Dense Layer: Expands the latent vector back into a higher-dimensional space suitable for reshaping into 2D feature maps.
- Reshape Layer: Transforms the dense layer's output into feature maps of shape [24, 24, 128].
- Transposed Convolutional Layers: These layers progressively increase the spatial dimensions of the feature maps while applying learnable filters to refine the reconstructed image. ReLU activations are applied to introduce non-linearity.
- Output Transposed Convolutional Layer: Produces the final reconstructed image with the same dimensions as the input. A sigmoid activation function is applied to constrain the pixel values to the range [0, 1].

Autoencoder: The autoencoder is a sequential model combining the encoder and decoder into a single pipeline. The model is trained end-to-end using a reconstruction loss function, such as mean squared error (MSE) or binary cross-entropy, to minimize the difference between the input image and its reconstruction. This architecture provides an efficient framework for learning compact, meaningful representations of high-dimensional data and reconstructing it with high fidelity.

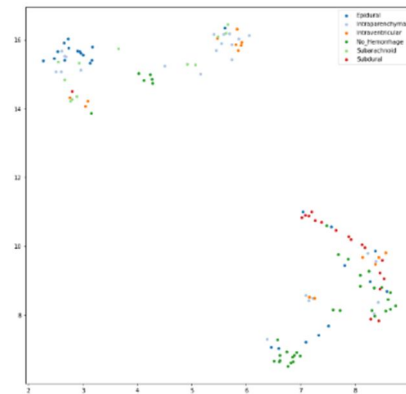


Fig 4. 2D UMAP projection

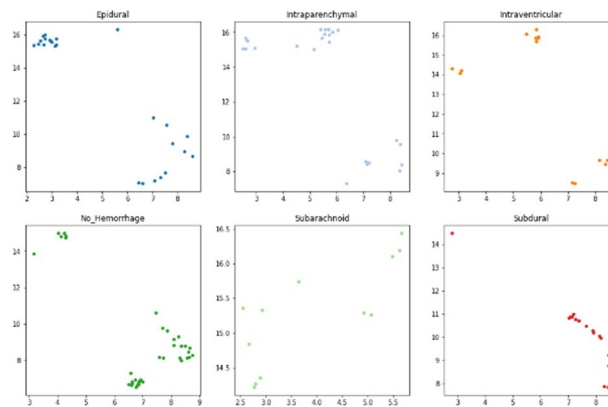


Fig 5. Embedding Interpretation of ICH types

The embeddings display some ICH type clustering when projected to 2 dimensions using Uniform Manifold Approximation and Projection (UMAP), although there is a substantial amount of overlap in different locations of the embedding space as seen in Fig 4. This may be partially due to the projection from 1024 dimensions to two dimensions. No matter the size of the autoencoder's embedding, findings showed a noisy clustering over a range of embedding dimensions (e.g., 512, 256, 32, 2). It is simpler to see which ICH are mixed and which are more clustered by looking at the Fig 5, which shows the UMAP projected embeddings for each type of ICH.

TABLE II. VAE RESULT ON CT SCANS

Type of ICH	Synthetic data sample 1	Synthetic data sample 2	Synthetic data sample 3
Epidural			
Subdural			
Subarachnoid			
Intraventricular			
Intraparenchymal			
No haemorrhage			

The encoder has a 1024-dimensional embedding layer, and it would be interesting to see how the flower images are represented in this central layer. To visualize the embeddings, a 1024-dimensional vector for each ICH image will be generated with the encoder (without decoding it). This will then be visualized using a UMAP projection with each point coloured by the ICH type. Table 2 displays the VAE result on CT scans with the original being the CT scan image from the dataset and the reconstructed image is the one obtained from the VAE result.

IV. RESULTS AND DISCUSSION

A. Dataset

The dataset utilized for the work includes CT pictures of the brain and bones. The dataset collection includes 5000 images, including 2500 brain window and 2500 bone window images from 82 brain CT scans. 30 slices, each around 5 mm thick, typically make up a patient's CT scan in the dataset [24]. The haemorrhagic areas in the dataset are contained in intracranial image masks that are linked to the CT scan pictures. The dataset collection includes CSV files with patient information and haemorrhage diagnosis information.

B. Evaluation Metrics

Fréchet Inception Distance calculates the separation between generated and real samples activation distributions. FID has been found to be more compatible with human evaluation in determining the realism and variance of the generated samples. It is a more comprehensive and principled metric.

C. Results

DCGAN won the comparison between VAE by demonstrating its superior ability to produce high-quality synthetic images. The Fréchet Inception Distance (FID) metric, which measures how different produced samples are from actual images, was used in the evaluation. The visual excellence and variety of created images are objectively assessed using the FID score. With an astoundingly low FID score of 11.34 in this investigation,

DCGAN demonstrated remarkable performance, demonstrating its capacity to create synthetic images that closely mimic real data distributions. The smaller the FIDs, the better the performance as seen in Table 3.

TABLE III. FIDS SCORE OF VAE AND DCGAN TRAINED ON THE DATASET

Method	FID on Data
VAE	67.4
DCGAN	11.34

On the contrary, VAE model exhibited very bad performance with an FID score of 67.4. The fact that the VAE model generates images that are farther away from true images raises the possibility that the data distribution in real images are intricate to the point that the generated images will deviate more from the real ones, hence this result. The higher FID score of the VAE indicates that the diversity and integrity in generated samples are hard to maintain, and even in the encoding decoding architecture there are limits. A large difference in the raw data of the game scores between DCGAN and VAE shows the advantages that DCGAN will bear on the task of game generation picture. As a result of the fact that DCGAN was only created for producing images, DCGAN's architecture, as the only one created just for generating images, use adversarial training so that its ability to produce images with finer features and more diversity may be improved. The complex interplay between the generator and discriminator networks enables DCGAN to obtain low FID score—close match between the produced and real picture distributions.

Deep learning models such as GANs and VAE's are scalable and applicable for healthcare domains only when the computational requirements and the efficiency of training these models are considered. To achieve this, we trained this model with the use of free GPU resources from Kaggle – an NVIDIA Tesla P100 GPU with 16GB of memory to be precise. The training process took place over the period of 12 hours, affording ample number of iterations to arrive at the point of convergence. This resource constrained environment illustrates the feasibility of training complex generative models using public platforms; however the extended training time shows the shortcoming of computational scalability. Access to high performance computing resources would then help in handling larger datasets and real time deployment in healthcare domain. Furthermore, employing efficient training strategies like hyperparameter tuning, data augmentation, and architecture refinement are crucial for training efficiency and minimizing computational overhead to maintain satisfactory model performance.

We can conclude that in terms of producing high quality images, DCGAN has outperformed VAE much more than what FID metric comparison indicates. While the VAE scored 67.4, DCGAN scored even better by achieving an FID score of 11.34 during the generation of the images, showing that DCGAN is capable of generating synthetic, visually appealing and variable images. These results indicate, on the one hand, that we must circumspectly select the most suitable generative models for image creation in some cases, and, on the other hand, demonstrate the strong capability of DCGAN to deliver extraordinary results in the scenario.

D. Comparison with Discrete Wavelet Transform

We apply a segmentation framework which includes both spatial transformation and Discrete Wavelet Transform (DWT), in our experiments on real world medical images [25] as well as on the synthesized data from a GAN. However, in Fig 6, it is observed that the model performed reasonably well to the actual medical images in which we have performed segmentation with DWT. Figures 6 depicts that each image is divided into 3 images - 1st image is the original CT scan image 2nd image is the segmentation mask obtained by the application of the DWT and 3rd image is the resultant image with the segmented region of haemorrhage after DWT.

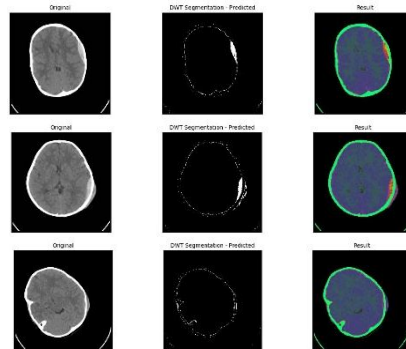


Fig 6. Segmentation results of DWT on dataset [25]

The generated data using GAN was not suboptimal. In turn, segmentation in the synthetic images neither reproduced the complexities and nuances of genuine medical scans nor resulted in outcomes that were either complete or consistent at a clinical level as shown in Fig 7. Each of the image is divided into 3 images – 1st image: synthetic data generated CT image, 2nd image: segmented mask obtained using DWT on the synthetic data image, 3rd image: The resultant image with the segmented particular region of haemorrhage after applying DWT on the synthetic data image.

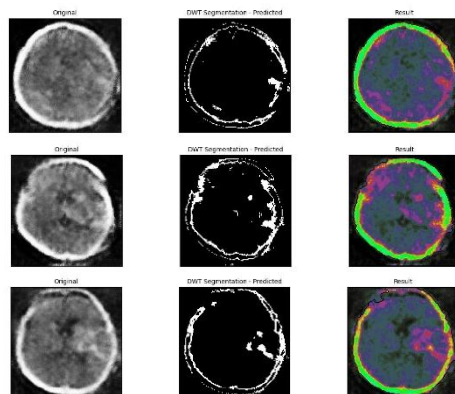


Fig 7. Segmentation results on generated synthetic data

In a medical context, complications or misleading predictions resulting from such errors carry a great risk to the patient safety and diagnostic accuracy. Thus, these findings indicate that even in the case of GANs data augmentation, not all the GAN generated data is ready to be used in real world medical applications. Diagnostics and research based on such data cannot lead to erroneous conclusions or patient harm if synthetic images do not sufficiently mimic the critical details of true medical images and more rigorous process is needed to ensure so.

V. CONCLUSION

As a final result, capabilities of DCGAN and VAE for producing pictures are thoroughly examined. To do so, we have tightly compared DCGAN against VAE using the Fréchet Inception Distance (FID) metric, and we have found that DCGAN invariably produces qualitatively slightly better synthesized images. As clear as day, there is emphasis to the role of architecture design in generative modelling as DCGAN achieves 11.34 whereas VAE lagged behind with 67.4. For image synthesis, DCGAN, designed especially for this task, makes use of adversarial training to generate images very much like the ones from true data distributions. This architectural benefit is evidenced by the fact that compared to those generated by DCGAN, DCGAN images display better fidelity, fidelity, and diversity, and a much-reduced FID score.

On the contrary, VAE's very high FID score gives rise to challenges in ensuring the image quality, especially when recording complex data pattern. Despite that, VAE is a useful model for representing the latent space and compressing data, the encoding-decoding structure of VAE may not always provide the best results for the image generation task as VAE lags behind from DCGAN on the quality aspect of produced realistic images. The importance of selecting a suitable generative model for a given task is our message. DCGAN is the best option when image quality and diversity are not negotiable. Considering its remarkable FID score, a FID score that is almost unrivalled in the field of image generation — DCGAN stands out as best option and its possibility to be applied in a various business related to art, entertainment and design is good.

This study concludes that in order to proceed the game, it is important to understand the specific advantages and disadvantages of a variety of generative models. This work also provides valuable lessons for future image synthesis advances and constitutes an invaluable resource for researchers and practitioners in the fields of deep learning and computer vision as we further study and improve generative models. Further research can be made into other GAN architectures and hyperparameters to understand better about data augmentation for the medical image analysis. In addition, transfer learning strategies and multi task learning can further improve the performance of the model and allow the prediction of more elements of ICH diagnosis. Moreover, standards can be created and standardized evaluation measures created for medical image analysis activities that may facilitate comparing, and validating, different models and methodologies. Finally, using these methods on larger and more

diverse datasets would likely enhance the efficacy and influence of how these methods can be used for early identification and treatment of ICH.

REFERENCES

- [1] Caceres JA, Goldstein JN. Intracranial hemorrhage. *Emergency medicine clinics of North America*. 2012 Aug 1;30(3):771-94.
- [2] Naidech AM. Intracranial hemorrhage. *American journal of respiratory and critical care medicine*. 2011 Nov 1;184(9):998-1006.
- [3] Heit JJ, Iv M, Wintermark M. Imaging of intracranial hemorrhage. *Journal of stroke*. 2016 Dec 12;19(1):11.
- [4] Nishijima DK, Offerman SR, Ballard DW, Vinson DR, Chettipally UK, Rauchwerger AS, Reed ME, Holmes JF. Immediate and delayed traumatic intracranial hemorrhage in patients with head trauma and preinjury warfarin or clopidogrel use. *Annals of emergency medicine*. 2012 Jun 1;59(6):460-8.
- [5] Young RJ, Destian S. Imaging of traumatic intracranial hemorrhage. *Neuroimaging Clinics*. 2002 May 1;12(2):189-204.
- [6] Powers AY, Pinto MB, Tang OY, Chen JS, Doberstein C, Asaad WF. Predicting mortality in traumatic intracranial hemorrhage. *Journal of neurosurgery*. 2019 Feb 22;132(2):552-9.
- [7] Weisberg LA. Computerized tomography in intracranial hemorrhage. *Archives of neurology*. 1979 Jul 1;36(7):422-6.
- [8] Parizel P, Makkat S, Van Miert E, Van Goethem J, Van den Hauwe L, De Schepper A. Intracranial hemorrhage: principles of CT and MRI interpretation. *European radiology*. 2001 Sep;11:1770-83.
- [9] Cackowski S, Barbier EL, Dojat M, Christen T. ImUnity: A generalizable VAE-GAN solution for multicenter MR image harmonization. *Medical Image Analysis*. 2023 Aug 1;88:102799.
- [10] Li F, Huang W, Luo M, Zhang P, Zha Y. A new VAE-GAN model to synthesize arterial spin labeling images from structural MRI. *Displays*. 2021 Dec 1;70:102079.
- [11] Sujatha R, Mahalakshmi K, Yoosuf MS. Colorectal cancer prediction via histopathology segmentation using DC-GAN and VAE-GAN. *EAI Endorsed Transactions on Pervasive Health and Technology*. 2024;10.
- [12] Cetin I, Stephens M, Camara O, Ballester MA. Attri-VAE: Attribute-based interpretable representations of medical images with variational autoencoders. *Computerized Medical Imaging and Graphics*. 2023 Mar 1;104:102158.
- [13] Volokitin A, Erdil E, Karani N, Tezcan KC, Chen X, Van Gool L, Konukoglu E. Modelling the distribution of 3D brain MRI using a 2D slice VAE. In *Medical Image Computing and Computer Assisted Intervention—MICCAI 2020: 23rd International Conference, Lima, Peru, October 4–8, 2020, Proceedings, Part VII 23 2020* (pp. 657-666). Springer International Publishing.
- [14] Frid-Adar M, Diamant I, Klang E, Amitai M, Goldberger J, Greenspan H. GAN-based synthetic medical image augmentation for increased CNN performance in liver lesion classification. *Neurocomputing*. 2018 Dec 10;321:321-31.
- [15] Desai SD, Giraddi S, Verma N, Gupta P, Ramya S. Breast cancer detection using gan for limited labeled dataset. In *2020 12th International Conference on Computational Intelligence and Communication Networks (CICN) 2020 Sep 25* (pp. 34-39). IEEE.
- [16] Liu B, Lv J, Fan X, Luo J, Zou T. Application of an improved DCGAN for image generation. *Mobile information systems*. 2022;2022(1):9005552.
- [17] Bushra SN, Shobana G. A survey on deep convolutional generative adversarial neural network (dcgan) for detection of covid-19 using chest x-ray/ct-scan. In *2020 3rd International Conference on Intelligent Sustainable Systems (ICISS) 2020 Dec 3* (pp. 702-708). IEEE.
- [18] Kora Venu S, Ravula S. Evaluation of deep convolutional generative adversarial networks for data augmentation of chest x-ray images. *Future Internet*. 2020 Dec 31;13(1):8.
- [19] Wang S, Wang X, Hu Y, Shen Y, Yang Z, Gan M, Lei B. Diabetic retinopathy diagnosis using multichannel generative adversarial network with semisupervision. *IEEE Transactions on Automation Science and Engineering*. 2020 Apr 9;18(2):574-85.
- [20] Negi A, Raj AN, Nersisson R, Zhuang Z, Murugappan M. RDA-UNET-WGAN: an accurate breast ultrasound lesion segmentation using wasserstein generative adversarial networks. *Arabian Journal for Science and Engineering*. 2020 Aug;45:6399-410.
- [21] Liu X, Guo S, Zhang H, He K, Mu S, Guo Y, Li X. Accurate colorectal tumor segmentation for CT scans based on the label assignment generative adversarial network. *Medical physics*. 2019 Aug;46(8):3532-42.
- [22] Goodfellow I, Pouget-Abadie J, Mirza M, Xu B, Warde-Farley D, Ozair S, Courville A, Bengio Y. Generative adversarial nets. *Advances in neural information processing systems*. 2014;27.
- [23] Kingma, D. P. (2013). Auto-encoding variational bayes. *arXiv preprint arXiv:1312.6114*.
- [24] <https://www.kaggle.com/datasets/vbookshelf/computed-tomography-ct-images>
- [25] Saifuddin HM, Vijayalakshmi HC, Ramya J. Discrete Wavelet Transform: A breakthrough in segmentation of CT scans for Intracranial Hemorrhages. In *2023 International Conference on Artificial Intelligence and Smart Communication (AISC) 2023 Jan 27* (pp. 843-848). IEEE.