



Generative adversarial minority enlargement—A local linear over-sampling synthetic method

Ke Wang^{a,b}, Tongqing Zhou^a, Menghua Luo^c, Xionglve Li^a, Zhiping Cai^{a,*}

^a Computer College, National University of Defense Technology, Changsha, China

^b Information College, GuiZhou University of Finance and Economics, Guiyang, China

^c Big Data Application and Economics College, GuiZhou University of Finance and Economics, Guiyang, China

ARTICLE INFO

Keywords:

Imbalanced data problem
Generative adversarial model
Synthetic numerical samples
Local linear property of samples
Over-sampling strategy

ABSTRACT

The imbalanced data problem is widely recognized in real-world datasets. To avoid learning bias on imbalanced data, the over-sampling strategy is well studied and adopted for generating synthetic minority samples. Wherein, the Synthetic Minority Over-sampling Technology and its improved algorithms become standard baselines. In recent years, the popular Generative Adversarial Networks and its enhanced variants, introduced from the computer vision community, are believed to generate better samples, by approximating the true data distribution. Yet, we notice that the synthetic samples for the minority category in these existing methods is usually restrained in a limited samples space known in advance, which may mislead the classifiers trained on them to take data in the unsampled region of the minority category as from the majority category. Given such limitations, we propose a Generative Adversarial Minority Enlargement (GAME) method to intentionally extend the sampling margin during data generative and adversarial phases. This is accomplished by adjusting the parameters of a local linear model to approach the majority category. We conduct an extensive evaluation on 28 datasets of different domains, extracted from the UCI real-world datasets. The results show that GAME can achieve more balanced and stable results efficiently than 18 state-of-the-art methods.

1. Introduction

Dealing with imbalanced data is a well-known challenge in conducting data analysis based on machine learning for both academia and industry. In fact, for most of the available datasets obtained from real-world situations, the number of samples for different categories is hardly approximated to be equal. Requiring more real samples for the minority classes usually means enlarging the annotation cost, and is quite hard, if not impossible, for specific applications, such as fraud detection (Zhang et al., 2022), traffic identification (Mehranian, Bagi, Moshiri, & Al-Basir, 2021), and medical diagnosis (Orooji & Kermani, 2021).

The imbalanced data problem incurs uncertainty and inaccuracy for the training process of the classifier (Zhang et al., 2022). The classification performance for the majority class is always good with sufficient samples, while the machine learning models usually fail to detect the minority class. In order to relieve the pitfall, over-sampling approaches are widely adopted as an efficient technique to generate synthetic samples for the minority class, thus facilitating classification models with balanced inputs in terms of each category.

However, we point out that the existing efforts devoted to over-sampling construct new samples by merely looking at the minority group, which leads to large unsampled space between borders of different categories. On one hand, classic methods, including SMOTE (Chawla, Bowyer, Hall, & Kegelmeyer, 2002; Dablain, Krawczyk, & Chawla, 2022) and the improved versions (e.g., SMOTE borderline 1, SMOTE-borderline 2 (Han, Wang, & Mao, 2005), ADASYN (He, Bai, Garcia, & Li, 2008), SMOTE-SVM (Wang, Luo, Huang, Feng, & Liu, 2017), Polynom-fit SMOTE (Gazzah & Amara, 2008), ProWSyn (Barua, Islam, & Murase, 2013), ProWRAS (Bej, Schulz, Srivastava, Wolfien, & Wolkenhauer, 2021; Kovács, 2019), and CURE-SMOTE (Ma & Fan, 2017).), propose to generate new samples by randomly selecting data points on the line segments between known labeled data in the minority class. Obviously, such a generation process cannot get rid of the possible limitation of the original data distribution. On the other hand, the popular Generative Adversarial Networks (GAN) (Goodfellow, 2017; Goodfellow et al., 2014) and its enhanced variants (e.g., DC-GAN (Radford, Metz, & Chintala, 2015), CGAN (Mirza & Osindero, 2014), WGAN (Arjovsky, Chintala, & Bottou, 2017), SGAN (Huang,

* Corresponding author.

E-mail addresses: kewang@mail.gufe.edu.cn (K. Wang), zhoutongqing@nudt.edu.cn (T. Zhou), mh_luo@mail.gufe.edu.cn (M. Luo), lixionglve17@nudt.edu.cn (X. Li), zpcai@nudt.edu.cn (Z. Cai).

<https://doi.org/10.1016/j.eswa.2023.121696>

Received 30 March 2022; Received in revised form 28 April 2023; Accepted 16 September 2023

Available online 21 September 2023

0957-4174/© 2023 Elsevier Ltd. All rights reserved.

Li, Poursaeed, Hopcroft, & Belongie, 2017), CTGAN (Xu, Skoularidou, Cuesta-Infante, & Veeramachaneni, 2019) and CTABGAN (Zhao, Kunar, Birke and Chen, 2021)) are powerful in generating synthetic images based on the game between the generative and discriminative models. The given minority samples can hardly depict the true distribution comprehensively, so the synthetic data of GAN is restrained to a limited range. Considering the above limitations, existing methods will yield classifiers that show an unexpected preference for the majority category. A data synthetic method that can formally approximate the true distribution of the minority category is still in need.

In this paper, we propose a Generative Adversarial Minority Enlargement model (GAME) to construct synthetic data with an enlarged over-sampling space. A modified Jacobian-based (Papernot et al., 2016) data augmentation method is designed to obtain initial synthetic data in the generative phase. Then we bring forth the local linear property of samples based on the idea that data points falling in a sufficiently small space can be classified with a line. Given this insight, we adopt logistic regression as the building block to classify the updated samples in the adversarial phase and iteratively adjust the synthetic samples to the direction of the majority category by injecting the majority samples into the training data. In this way, our GAME can provide synthetic samples that are distributed in the undistinguishable space of the original dataset, facilitating truly balanced data in over-sampling. Note that the training cost can be shortened and the effect can be better guaranteed as we combine generative adversarial thinking and our local linear classification unit.

Contributions. The main contributions of this work are summarized as follows:

- We propose a novel over-sampling model, GAME, by exploiting the idea of local linear property of samples. GAME integrates logistic regression into a generative and adversarial process to generate synthetic data.
- We leverage the adversarial process to calibrate the generated data for better representing the minority category. By using the mixture of majority and minority samples for training, we can gradually adjust the data with an enlarged sampling space that is expected to better approximate the true minority samples distribution.
- We evaluate the classification performance of the proposal by 4 classifiers on 28 datasets of various imbalanced ratios. Overall, GAME can yield better F1 scores than the 18 baselines, and shows good execution efficiency at the same time.

2. Related work

We summarize the relevant work from two aspects: the SMOTE series algorithms (SMOTEs) and the GAN-based methods (GANs).

The SMOTE series algorithms (SMOTEs). In over-sampling, the enlargement of minority samples is the key task. The duplicated samples method is phased out because of poor accuracy and over-fitting problem. The synthetic minority over-sampling technique (SMOTE) (Chawla et al., 2002) is the cornerstone of the existing methods. This algorithm randomly generates synthetic samples on the selected line segment. On this basis, some improvements have been made in the past years. Borderline-SMOTE (Han et al., 2005) is a significant improvement of SMOTE. This method divides the minority samples into three groups, DANGER, SAFE, and NOISE, where DANGER is the borderline samples set. The method has two specific algorithms, *borderline1* generates synthetic samples between DANGER samples and *k* nearest minority neighbors. In *borderline2*, the *k* nearest neighbors include both minority and majority samples. ADASYN (He et al., 2008) is another milestone in the development of SMOTE. It proposes a density distribution of majority samples in *k* nearest neighbors, according to which assigns the number of the synthetic data samples. SMOTE-SVM (Wang et al., 2017) uses the SMOTE on support vectors to improve the classification effect. Polynom-fit SMOTE (Gazzah & Amara, 2008) oversamples

the minority samples using polynomial fitting functions. Proximity Weighted Synthetic algorithm (ProWSyn) (Barua et al., 2013) generates effective weights for the minority samples based on sample's proximity information as distance from boundary. ProWRAS (Bej et al., 2021; Kovács, 2019) integrates the Localized Random Affine Shadowsampling (LoRAS) algorithm and ProWSyn and improves performance by controlling the variance of the synthetic samples, as well as a proximity-weighted clustering system of the minority samples. CURE-SMOTE (Ma & Fan, 2017) uses the combination of Clustering Using Representatives (CURE) to enhance the original SMOTE algorithms. Although different strategies are adopted, the core approach of these methods to generate synthetic samples is not changed. The synthetic samples must be limited within the range of the given minority samples and cannot be further adjusted.

The GAN-based methods (GANs). GANs become more and more popular in the field of computer vision for image generation. The game of discriminative and generative models leads to many achievements in image generation. Although varying improved models and techniques (Salimans et al., 2016) are proposed, the non-convergence problem, collapse problem, and uncontrollability have become bottlenecks. Conditional Generative Adversarial Nets (CGAN) (Mirza & Osindero, 2014) introduces condition data on both generator and discriminator. Deep Convolutional Generative Adversarial Networks (DCGAN) (Radford et al., 2015) have certain architectural constraints good for unsupervised learning. Wasserstein GAN (WGAN) (Arjovsky et al., 2017) introduces a stable way of learning dealing with mode collapse problem. Stacked Generative Adversarial Networks (SGAN) (Huang et al., 2017) have a bottom-up discriminative network to invert the hierarchical representations. CTGAN (Xu et al., 2019) uses a conditional generative adversarial network to generate synthetic samples for tabular data. CTABGAN (Zhao, Kunar et al., 2021) is another GAN-based method for tabular data which introduces the information loss, classification loss and generator loss to the conditional GAN and encodes the mixed data type and skewed distribution of data variable by a conditional vector. The approximation of the true data distribution is also a restriction of generating synthetic samples from the limited minority samples. The synthetic samples are essentially limited to the interior of the given minority samples. From the perspective of imbalanced data, it only considers the given minority samples and does not consider the influence of the majority of samples, so there must be a lack of overall data information.

Some Special Approaches. Some effective special approaches focus on adjusting the algorithm itself or integrating domain knowledge to deal with given imbalanced data, without concerning the data distribution. A typical approach here is the cost-sensitive method (Mienye & Sun, 2021). The main point of this method is to design the cost matrix with domain knowledge. To a certain dataset, with the help of domain experts, we may get a satisfactory result, which is both an advantage and a limitation. Ensemble learning methods (Chen, Duan, Kang, & Qiu, 2021; Xu, Yu, Chen, & Liu, 2021) utilize multiple models and some adjustments to the dataset to improve performance. One-class classification method (Hayashi, Fujita, & Hernandez-Matamoros, 2021) is more suitable for the situation of only very few minority samples existing, which tries to shrink the boundary of the majority samples. The classification result is no majority or minority, but whether belonging to the majority. Feature selection (Agrawal, Abutarboush, Ganesh, & Mohamed, 2021) is an effective method to some imbalanced data problems. The imbalanced samples always come with imbalanced features, which means choosing the discriminate features can improve the minority samples' recognition rate. A contrastive variational autoencoder method (Dai et al., 2019) leverages both the majority and minority classes information on 2 clinic datasets with highly imbalanced outcomes. Experiment results show that their algorithm performs better when the number of minority class samples is very small. But the application field of this method is relatively limited. A conditional variational autoencoder-based self-transferred algorithm (Zhao, Hao,

Tang, Chen and Wei, 2021) is proposed to solve the highly imbalanced classification problem, where the training instances of the minority classes are rare. The Variational Autoencoder method is a well-designed data generation model, but there is a certain amount of information loss in its encoding process. As in the image dataset, it is directly calculating the mean square error of the generated image and the original image instead of learning by adversarial like GAN, which makes the generated image a bit blurry. An oversampling framework SWIM (Bellinger, Sharma, Japkowicz, & Zaiane, 2020) uses the density of the well-sampled majority class to guide the generation process, which is designed to address the extremely imbalanced problem without advantages for different imbalanced ratios. The introduction of some mathematical methods (Du, Zhou, Liu, Vong, & Wang, 2021; Xia, Zheng, Wang, Gao, & Wang, 2021) has also produced good results on certain datasets. These methods work well on specific datasets and problems. But they are not universal and difficult to generalize.

The previous approaches have some drawbacks: The generating method of SMOTEs has a limitation of data generating range and is not convenient for adjustment; The GANs methods are difficult to break through the limitations of given data, do not consider the global data information, and consume large computing resources; Some special methods are difficult to widespread use. Our method breaks the limitations mentioned above and achieves better results on 28 real-world datasets of various fields.

3. Local linear property of samples in imbalanced data problem

In the data space, the classification hyperplane is always an irregular surface. In machine learning, varying models and algorithms are proposed to deal with hyperplane fitting problems. In over-sampling problems, insufficient data makes it impossible to obtain real global data distribution of the minority category. So we choose a local perspective in our research. In the samples space which is a subspace of the data space, we just focus on the local top k points on both sides of the hyperplane, as shown in Fig. 1(a). The blue triangles stand for the minority samples, while the red points represent the majority samples. We select a tiny space around the hyperplane, where k is small enough, as the yellow space indicated by the cyan arrow in Fig. 1(a). Let us zoom the yellow space, as shown in Fig. 1(b). The several samples in this tiny space can always be linear separable. And a tiny classification line can be calculated as the blue line in Fig. 1(b). All these tiny local classification lines make up the accurate classification multiple dimensional curves, without the requirement of composing a global classification hyperplane. "Local linear property" means that in a very small part of the whole data, where are only several, or only one, minority sample(s), making the minority samples and the majority samples in this small region linear separable.

The traditional deep learning method is to construct the hyperplane globally and then adjust and optimize it through iterations. The local linear view is to optimize the local area. For data synthesis, the local samples' generation is the focus point, but not the global hyperplane. From the view of our local linear property of samples, local samples' characteristic is the decisive factor for the synthetic data generation, while the data points far away are of little significance for local data synthesis. Therefore, using the properties of local data points to adjust and optimize the generated samples is the core of our method. In addition, the local linear property of samples makes use of local computing to save computing resources and improve efficiency. This property brings a lot of convenience to machine learning processing. The linear method is simple and mature. For a small number of data, it can deal with linear and nonlinear classification problems in a reasonable range of errors and reduce the probability of the over-fitting problem. From the point of view of the model training process, the linear models are much faster than the DNN and perform as well as them in small training samples.

4. Generative adversarial minority enlargement model

According to our local linear property of samples, the whole dataset is made up of many local groups. There are only two kinds of groups, the inside group consists of only samples of one category, and the outside group is from both the majority and the minority samples. For the inside groups of majority samples, we maintain the existing state. For the inside groups of minority samples, we generate some synthetic samples by SMOTE to keep the data distribution stable. For the outside groups, we design and use the GAME to generate the synthetic samples.

In this section, we will first introduce the local linear function, then go through the generative part and discriminative part of our GAME model, and finally depict the whole data generation process.

4.1. The local linear function

This function selects the local linear region data and prepares it for the follow-up work, where the details are shown in Algorithm 1. First, it calculates the initial classification line $line_0$ for the given minority samples S_{min} and majority samples S_{maj} (Line Input), even if they are nonlinear separable. Then it makes S_{maj} , S_{min} , and $line_0$ as inputs for the local linear function. For each sample x_j in S_{maj} , it calculates D_{maj} , the set of distances between x_j and $line_0$ (Line 1–3). Next, the top n samples with shortest distances in D_{maj} are chosen as $Near_{maj}$, which contains the nearest n majority samples to the minority samples (Line 5). After that, for each x_i in S_{min} , it selects local k samples from the minority samples and $Near_{maj}$, requiring at least one of them is from $Near_{maj}$. The $local_i$ here is the data set of x_i and its k -nearest neighbors points. $Local_k$ is the set of all the $local_i$. D is the data set of distances from $Local_k$ (Line 6–11). At last, it returns $Local_k$ and D for the generative process.

Algorithm 1 Local_Linear_F

Input: S_{maj} , S_{min} , $line_0$

Output: $Local_k$, D

```

1: for each  $x_j \in S_{maj}$  do
2:    $d_j = distance(x_j, line_0)$ 
3:    $D_{maj} = D_{maj} + d_j$ 
4: end for
5: select the minimal  $n$  of  $D_{maj}$  as  $Near_{maj}$ 
6: for each  $x_i \in S_{min}$  do
7:    $[x_1, x_2, \dots, x_k] = KNN(x_i, (S_{min} + Near_{maj}))$ 
8:    $local_i = [x_i, x_1, x_2, \dots, x_k]$ 
9:    $Local_k = Local_k + local_i$ 
10:   $D = distance(KNN(Local_k))$ 
11: end for
12: return  $Local_k$ ,  $D$ 

```

4.2. The generative part

The generating method is the core of the over-sampling synthetic technology. The purpose of our method is to generate synthetic samples of minority samples, not to calculate a global piecewise classification boundary. The local repetitions generate the required samples, but do not consider whether the local classification line can form an overall classification boundary or hyperplane. In this case, the required local synthetic samples can be iterative generated. When the local generation of each minority sample is completed, the goal of balancing data is achieved, regardless of the distribution of the known samples or the global classification boundary. We compare the present methods and adjust them as the basis of our generative method.

In SMOTE and its improved algorithms, the generation process almost uses Eq. (1). It randomly picks up x_j , one of the k nearest neighbors of the selected minority sample x_i , and then generates synthetic sample x_g with Eq. (1), where $\lambda \in (0, 1)$. In simple words, x_g is

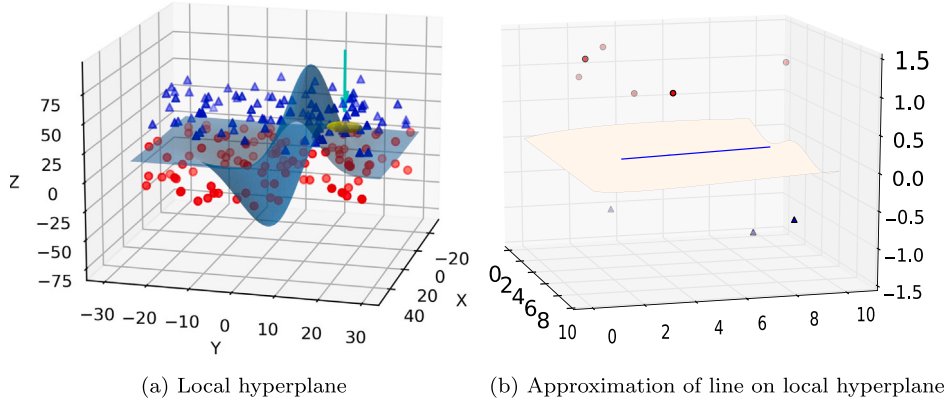


Fig. 1. Local linear property of samples in imbalanced data problem.

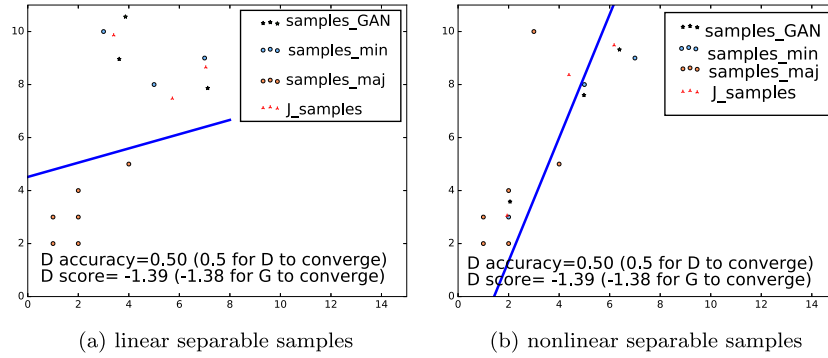


Fig. 2. Jacobian method VS. GAN in the local generating process.

selected randomly on the line segment (x_i, x_j) . In SMOTE-borderline2, the parameter λ is adjusted ($\lambda \in (0, 0.5)$) to make the generated samples close to the minority sample x_i .

$$\vec{x}_g = \vec{x}_i + \lambda * (\vec{x}_j - \vec{x}_i) \quad (1)$$

There is a space range between the majority and minority samples. The SMOTEs and GANs generate data within the given data distribution. These generating methods neglect the most space. The direction of our study is to enlarge the generating space and to make the synthetic samples adjustable in the required direction. In order to generalize the method of synthetic data in a more general way, we propose Eq. (2), where the $\Delta noise$ here is the key point of the algorithm design. Obviously, Eq. (1) is a specific form of Eq. (2).

$$\vec{x}_g = \vec{x}_i + \Delta noise \quad (2)$$

$$\vec{x}_g = \vec{x}_i + \lambda_i * \text{sgn}(J_F(O(\vec{x}_i))) \quad (3)$$

A Jacobian-based data augmentation technology (Papernot et al., 2016) is proposed as Eq. (3), where λ_i is the step size of the learning, J_F is the Jacobian Matrix, and $O(\vec{x}_i)$ is the label of x_i . To the demand of our work, we adjust the Jacobian method as Eq. (4). We change $O(\vec{x}_i)$ to $\vec{x}_i$ as for the Jacobian Matrix dimension requirement. We replace λ_i with $distance_i$ which is the distance between two of the top k points.

$$\vec{x}_g = \vec{x}_i + distance_i * \text{sgn}(J_F(\vec{x}_i)) \quad (4)$$

On the basis of full study of traditional over-sampling methods and GANs, in the GAME, the generating technology is the Jacobian method, as Eq. (4). The Jacobian Matrix represents the optimal linear approximation of a differentiable equation to a given point. Here is the sign of the Jacobian Matrix value of input $\vec{x}_i$, and a mapping direction

Algorithm 2 Generator

Input: $G_n, Local_k, D, Num_{Max_iter}$

Output: Generated Samples: X_{gen}

```

1: for each  $x_i \in S_{min}$  do
2:    $local_i = Local_k[i]$ 
3:    $d_i = D[i]$ 
4:    $Maj_i, Min_i = Local_i$ 
5:    $X_{g0} = Min_i$ 
6:    $Num_{iter} = [G_n / Num(S_{min})] / Num(Min_i)$ 
7:   if  $Num_{iter} > Num_{Max\_iter}$  then
8:      $Num_{iter} = Num_{Max\_iter}$ 
9:   end if
10:  for  $j = 1; j \leq Num_{iter}; j++$  do
11:    if Discriminatorflag then
12:      break
13:    end if
14:    line( $w, b$ ) = LogisticRegression( $Maj_i, Min_i$ )
15:     $J_F(x) = \text{Jacobian}(\text{line}(w, b)(x))$ 
16:     $X_{g_j} = Min_i + d_i \times \text{sgn}(J_F(X_{g_{j-1}}))$ 
17:     $X_{gen} = X_{gen} + X_{g_j}$ 
18:     $Min_i = Min_i + X_{gen}$ 
19:  end for
20:   $Num_{inside} = G_n - Num_{X_{gen}}$ 
21:   $X_{g_{inside}} = \text{SMOTE}(Num_{inside})$ 
22:   $X_{gen} = X_{gen} + X_{g_{inside}}$ 
23: end for
24: return  $X_{gen}$ 

```

of the change from $\vec{x}_i$ to the straight line is obtained, so that the added noise can be near $\vec{x}_i$ and the classification line without being

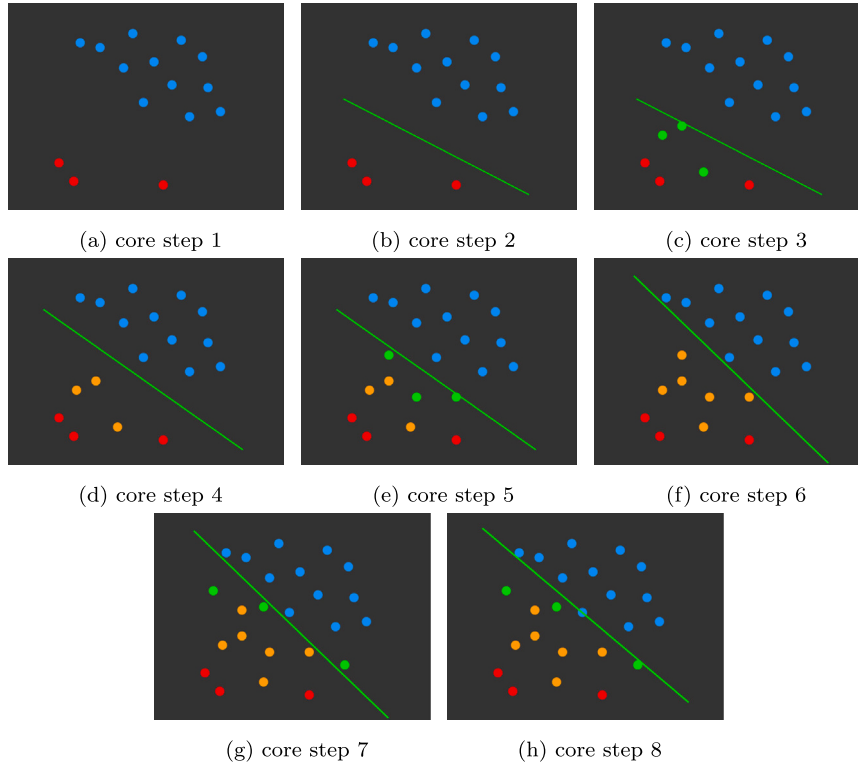


Fig. 3. The demonstration of core steps.

too random. And it is faster than simple random initialization and the GAN model.

In order to get efficient and fast initialized generated samples, we compare the Jacobian method with GAN, as shown in Figs. 2(a) and 2(b). After thousands of iterations of training, GAN gives the convergent results, while the Jacobian method directly generates the final results without training. From the results shown in Fig. 2, there are no obvious differences between the two approaches in both linear and nonlinear separable samples. So in the generative part, we choose the Jacobian method as the initialized generating approach. First, we get coefficients of the required generating number G_n , the maximum number of iterations Num_{Max_iter} , and $Local_k, D$ from the local linear function (Line Input). The next step is the core generation process. For each $x_i \in S_{min}$ with the $Local_k[i]$ and $D[i]$, the local majority samples Maj_i and minority samples Min_i , an initialized input data X_{g0} and the required number of iterations Num_{iter} are prepared for the iterative generation process (Line 2–6). If the required number of iterations is greater than the maximum value, the maximum value is taken as the actual number of iterations. During the loop, it determines whether the discriminator condition is met first (Line 11). After that, a new classification line is calculated by Maj_i and Min_i (Line 14). Then a Jacobian Matrix $J_F(x)$ is calculated by which the synthetic samples' matrix in this iteration X_{g_j} is generated (Line 15–16). Next, X_{g_j} is integrated into the generated dataset X_{gen} . And X_{gen} is added to Min_i for the next iteration (Line 17–18). At last, if there is an insufficient amount of synthetic samples, the rest are generated by the inside group with the SMOTE algorithm. This generative process is shown as Algorithm 2.

The demonstration of core steps is shown as Fig. 3. For each minority sample (the red point), its k-nearest neighbors $local_i$ (shown in the red and blue points) and the corresponding distances set d_i are loaded as Fig. 3(a). The classification line (green line) is calculated as Fig. 3(b). The synthetic samples (green points) are generated by Eq. (4) as Fig. 3(c). In the iterative process, the new minority samples (the red and orange points) and majority samples (the blue points) recalculate

the classification line as Figs. 3(d) to 3(h). This process proceeds until the classification line in the discriminator cannot tell the majority points (blue points) and the new synthetic samples (green points) as Fig. 3(h).

Comparing the classic generating method (Eq. (1)) in SMOTE and its improved algorithms, our generative approach is more efficient and controllable. The SMOTE method randomly generates synthetic samples in both inside group and outside group, which reflects the characteristics of the minority samples distribution in some degree but is uncontrollable. GANs care about the overall optimization process without focusing on local details, whose optimization effect does not have advantages for several local points. In our way, the generated synthetic samples are optimized locally and close to the real data distribution within an enlarged given data range. In the next section, we discuss the condition for the discriminator.

4.3. The discriminative part

In the training process of GANs, the minority samples are treated as real data, but the majority samples are not involved at all. The synthetic samples are trained to follow the real data distribution. In our method, the core target is to refine and make full use of data distribution information. We introduce the majority samples in the discriminator.

As an input parameter in the discriminative part, X_{gen} from the Generator iteration is integrated with S_{min} , of which the result and S_{maj} determine a new classification line $line_{new}$ (Line 1). Then it selects samples from S_{maj} with k nearest distances to $line_{new}$ as S'_{maj} , where k equals the number of X_{gen} (Line 2–4). We compute and record the classification results of $prob_{maj}$ and $prob_{min}$ for Loss, as Eq. (5) (Line 5–7). In detail, first, we make sure that the synthetic samples and the given minority samples are on the same side of the classification line. Then, each sample in X_{gen} is moved towards $line_{new}$ by a value which is the smallest value in D_{new} multiplied by a random number between 0 and 1. After that, it updates X_{gen} from the next round of iteration in the Generator (Line 8). So that the new classification line gradually

Algorithm 3 Discriminator

Input: $X_{gen}, S_{maj}, S_{min}$
Output: Synthetic Samples: X_{syn}

- 1: $line_{new} = Logistic_Regression((X_{gen}, S_{min}), S_{maj})$
- 2: $D_{new} = distance(S_{maj}, line_{new})$
- 3: $k = Num_{X_{gen}}$
- 4: $S'_{maj} = KNN(S_{maj})$
- 5: $prob_{min} = line_{new}(X_{gen})$
- 6: $prob_{maj} = line_{new}(S'_{maj})$
- 7: $Loss = \nabla \frac{1}{m} \sum_{i=1}^m [\log(1 - prob_{maj}) + \log(prob_{min})]$
- 8: Adjust X_{gen} and Update the Generator with $line_{new}$ and X_{gen} by ascending the Loss.
- 9: **if** Loss converges or iteration stops **then**
- 10: $X_{syn} = X_{gen}$
- 11: **end if**
- 12: **return** X_{syn}

closes to the majority samples. In this way, we expand the local data distribution reasonably and simplify the training model effectively. The details are shown in Algorithm 3.

$$Loss = \nabla \frac{1}{m} \sum_{i=1}^m [\log(1 - prob_{maj}) + \log(prob_{min})] \quad (5)$$

4.4. More details on GAME

After specifying the local linear function, the Generator and the Discriminator above, we summarize the whole GAME model in Algorithm 4. First, we compute G_n (Line 1), the amount of the generated samples and set the maximum iteration number Num_{Max_iter} . Then, the initial classification line of S_{maj} and S_{min} is calculated (Line 2). After that, we put these parameters into the local linear function (Line 3). Next, the Generator and the Discriminator perform data generation and update the loss according to their algorithms (Line 4–5). At last, the synthetic samples, the given minority and majority samples form balanced samples (Line 6). We update the local linear function and the Generator by ascending the stochastic gradient as Eq. (5) in the Discriminator. The iterative adjustment process of classification lines and synthetic samples gradually enlarges the data generating space.

Algorithm 4 the whole GAME algorithm

Input: $S_{maj}, S_{min}, Num_{Max_iter}$
Output: S_{syn}

- 1: $G_n = Num(S_{maj}) - Num(S_{min})$
- 2: $line_0 = Logistic_Regression(S_{maj}, S_{min})$
- 3: $Local_k, D = Local_Linear_F(S_{maj}, S_{min}, line_0)$
- 4: $X_{gen} = Generator(G_n, Local_k, D, Num_{Max_iter})$
- 5: $X_{syn} = Discriminator(X_{gen}, S_{maj}, S_{min})$
- 6: $S_{syn} = X_{syn} + S_{min} + S_{maj}$
- 7: **return** S_{syn}

In the imbalanced data problem, the minority samples can rarely stand for the whole minority sample distribution. So the synthetic samples generated by SMOTEs and GANs are not representative of the minority samples. In our view, the SMOTEs are uncontrollable and not adjustable, while the GANs are a fine-grained learning process with a coarse-grained generation. The main function of SMOTEs and GANs is to generate synthetic samples that satisfy the given data distribution. GAME is not only to maintain the original minority samples characteristics but also to break through the given data distribution to enlarge the data generating space. The local data generation complements

Table 1

The majority-minority ratio description.

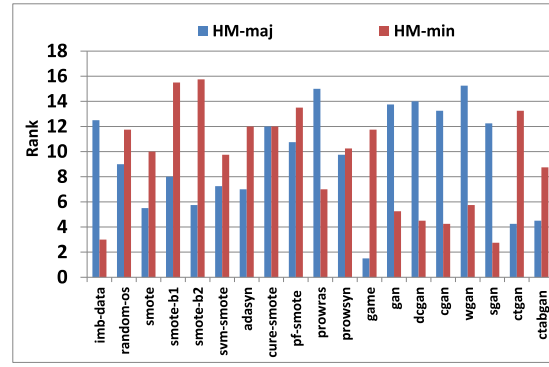
Dataset	original-maj:min	test-maj:min	train-maj:min
HB	255:81	67:67	65:13(5:1)
BCW1			200:38(5:1)
BCW2	444:239	200:200	200:20(10:1)
BCW3			210:14(15:1)
BCW4			244:12(20:1)
PID1		220:220	220:44(5:1)
PID2	500:268		240:24(10:1)
PID3		240:240	240:16(15:1)
PID4			240:12(20:1)
GM1			250:50(5:1)
GM2	700:300	250:250	250:25(10:1)
GM3			225:15(15:1)
GM4			260:13(20:1)
DCCC1		5500:5500	5680:1136(5:1)
DCCC2		6036:6036	6000:600(10:1)
DCCC3		6222:6222	6210:414(15:1)
DCCC4	23 364:6636	6320:6320	6320:316(20:1)
DCCC5		6506:6506	6500:130(50:1)
DCCC6		6571:6571	6500:65(100:1)
DCCC7		6623:6623	6500:13(500:1)
Shuttle1		10 345:10 345	10 345:2069(5:1)
Shuttle2		11 268:11 268	11 280:1128(10:1)
Shuttle3		11 639:11 639	11 625:775(15:1)
Shuttle4		11 823:11 823	11 820:591(20:1)
Shuttle5	45 586:12 414	12 171:12 171	12 150:243(50:1)
Shuttle6		12 292:12 292	12 200:122(100:1)
Shuttle7		12 390:12 390	12 000:24(500:1)
Shuttle8		12 402:12 402	12 000:12(1000:1)

the overall data shortage without calculating the overall classification hyperplane. A large number of hyperparameter calculations in the iteration are avoided, and computational resources are saved.

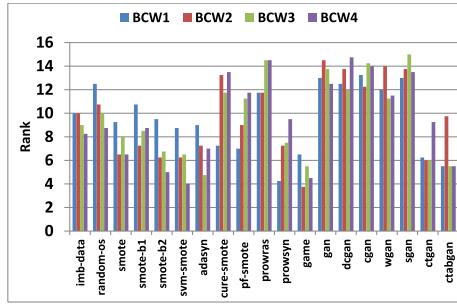
5. Experiment results and discussions**5.1. Datasets and experimental settings**

We use 28 tabular datasets with various imbalanced ratios designed from six UCI real-world datasets of different domains to examine our ideas. Haberman's Survival (HM) dataset contains cases on the survival of breast cancer patients with 306 instances and 3 attributes. Breast Cancer Wisconsin Original (BCW) dataset records the clinical cases of Dr. Wolberg with 699 instances and 10 attributes, from which we remove the missing values in our experiment. The Pima Indians Diabetes (PID) dataset draws from the National Institute of Diabetes and Digestive and Kidney Diseases with 768 instances and 8 attributes, where the patients are all female Pima Indian heritage. German Credit Data (GM) classifies people described by a set of attributes as good or bad credit risks with 1000 instances and 20 attributes. Default of Credit Card Clients (DCCC) data set describes the cases of customers' default payments in Taiwan with 30,000 instances and 24 attributes. Statlog Dataset (Shuttle) collects shuttle logs with 58,000 instances and 9 attributes.

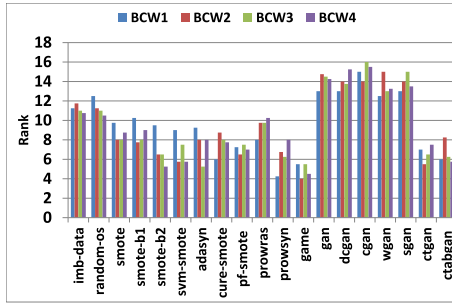
To fully compare the data generated by the algorithms, we design the datasets based on three principles. The first is to maximize the utilization of minority samples, which means to keep minority samples as many as possible when constructing different datasets. The second is to use a balanced dataset as the testing set, which is used to maximize the effect of distinguishing between majority and minority class tests. And the last is to keep the same number of majority samples in both training and testing sets, which is to minimize the problems caused by the differences in the amount of majority samples. From the perspective of the imbalanced ratio, our manuscript does not target a specific ratio but rather aims to accommodate as many ratios as possible. In this paper, we construct datasets with imbalanced ratios starting from 5:1



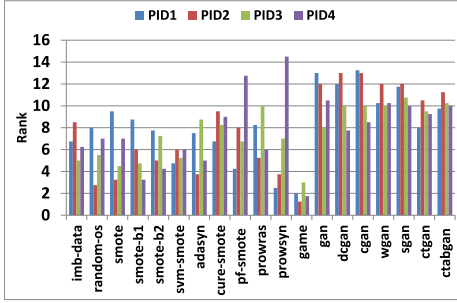
(a) Average Ranks in HM Majority & Minority



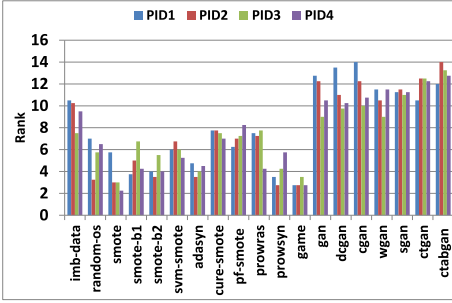
(b) Average Ranks of BCW Majority



(c) Average Ranks of BCW Minority



(d) Average Ranks of PID Majority



(e) Average Ranks of PID Minority

Fig. 4. F1 scores average ranks of different algorithms by datasets.

to 1000:1 as long as the total number of minority class samples is more than 10, so as to verify the actual effect of each algorithm. The specific description is shown in Table 1.

In order to measure the adaptability of the synthetic data generated by the algorithms to classifiers, we choose four classifiers of different design principles, Decision Tree(DT), Logistic Regression (LR), Support Vector Machine (SVM), and Naive Bayes (NB).

We select 18 methods as the baselines which are directly using the original imbalanced data (imb-data), random over-sampling (random-os), SMOTE, SMOTE-borderline 1 (Somte-b1), SMOTE-borderline 2 (Smote-b2), SMOTE-SVM (Smote-SVM), ADASYN, polynom-fit-SMOT (pf-Smote), CURE-SMOTE, ProWSyn, ProWRAS, GAN, DCGAN, CGAN, WGAN, SGAN, CTGAN, and CTABGAN. In evaluation, we calculate the F1 scores of the majority and the minority separately, to compare the balanced effects of the algorithms for imbalanced data. For the convenience of statistics, the experimental results are the average values of 50 cross-validations for each dataset. The standard deviations of all 19 algorithms are measured in 10^{-2} and retained to 2 decimal places.

The parameters in our algorithm are set as follows. The maximum number of iterations $Num_{Max_{iter}}$ is set to 50. The k is set to 5 in the nearest neighbors of each minority sample, which contains at least 1 known sample from the majority samples. The structure of the neural

network is not used in this method and there is no need to calculate the overall hyperplane, so no global hyper parameters are needed. The local data generation complements the overall data shortage without calculating the overall classification hyperplane. A large number of hyper parameters calculations in iteration are avoided, and computational resources are saved.

5.2. Analysis and discussion

The detailed experimental results are shown in Table 3 to Table 8. For the convenience of aggregating results in various datasets with different classifiers, we sorted out the ordinal ranking of each f1 score in Table 3 to Table 8 as the average rank, which is the average ranking position of the F1 scores from the highest to the lowest in all the result tables. The ranking details for all datasets are shown in Figs. 4a to 4k, and the final summary results are summarized in Table 2. We analyze the results and summarize the advantages of our algorithm below.

In this experiment, the minority samples have been generated in the enlarged space. The majority samples keep the original. If the enlargement space of the generated minority samples is too large and takes up the actual distribution space of the majority samples, the F1 Score of the majority samples will be affected. If the generating space

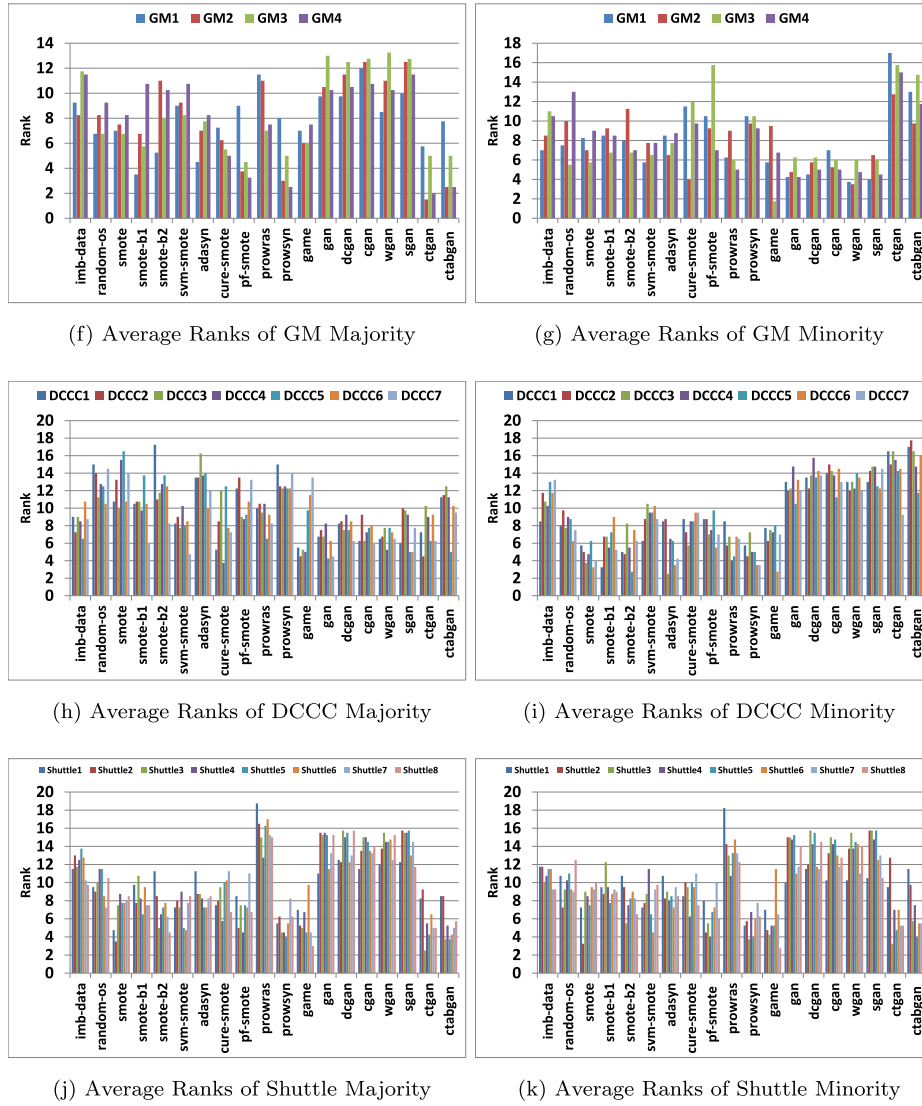


Fig. 4. (continued).

of the minority samples is not big enough, the F1 Score of the minority samples will be low. Therefore, while doing the whole data experiment, we separately count the F1 Scores of the majority and the minority samples, as Figs. 4a to 4k. From the Figures, almost every algorithm has a time when it performs optimally on a particular dataset, such as sgan in HM-min, svm-smote in BCW4-maj, smote in PID4-min, ctgan in GM2-maj, gan in DCCC5-maj and pf-smote in shuttle4-min. From Table 2a, it can be seen that our GAME algorithm has the best F1 scores in the majority samples. From Table 2b, the SMOTEs do well in DCCCs but get bad performance in HM, while the results are just the opposite for the GANs, where GAME keeps its stability in all the datasets. In Table 2c, the total average ranks of all the algorithms are listed, from which we can conclude that the GAME algorithm wins the most stable advantage ranking. From all 28 datasets and 4 classifiers, though not always perform best, the GAME method achieves the best-balanced effects for imbalanced data problems.

Then, the results should be discussed by imbalanced ratio and the amount of training data, as in Tables 3 to 8. The imbalanced ratios are considered to be high and low, where the boundary is 15:1. According to the amount, the datasets are divided into two groups here. The small group contains HM, BCWs, PIDs, and GMs, 13 datasets in all. The big group includes 15 datasets, DCCCs, and Shuttles. In the small group with low imbalanced ratios, the SMOTEs show good results, due to the limited data-generating space. In the big group with high

imbalanced ratios, the imb-data method and random-os approach gain their advantages, which benefit from a sufficient sample number. As for GANs, they play a role in some low imbalanced ratios situations because of enough given data points. Our GAME method is stable in all these cases.

Also there is a limitation of the method. It is that for problems with low imbalanced ratios and a large number of known minority samples, the results should not have a significant advantage. Theoretically, the amount of known samples in such problems is large and representative, and the sample space enlargement has a certain effect but with limitations. For the problems with high imbalanced ratios and a few known minority samples, the effect will be better.

After that, it comes to the adaptation to classifiers. To the four classifiers, all the algorithms perform stably on DT and LR but fluctuate on SVM and NB. As shown in Tables 5 to 7, the SVM is relatively weak in classification effects dealing with the PIDs, GMs, and DCCCs, which is irrelevant to the algorithm and the imbalanced ratio. In other words, the four classifiers are more sensitive to the distribution of the dataset itself rather than the algorithm principles, which is suitable to judge the quality of the data generated by the algorithms.

At last, from the perspective of executive efficiency, because of the local linear perspective in GAME, the number of iterations is set as 50, within which the DNN model is difficult to converge. For the GANs, the epochs parameter depends on the amount of training data,

Table 2
Average rank in all datasets.

(a) majority average rank							
Maj-Rank	HM	BCW	PID	GM	DCCC	Shuttle	Average Rank
imb-data	12.5	9.3125	6.625	10.1875	8.535714	11.90625	9.844494
random-os	9	10.5	5.8125	7.75	12.92857	9.71875	9.28497
smote	5.5	7.5625	6.0625	7.375	12.96429	7.0625	7.754464
smote-b1	8	8.8125	5.6875	6.6875	10.28571	8.4375	7.985119
smote-b2	5.75	6.875	6.0625	8.625	12.46429	7.125	7.816964
svm-smote	7.25	6.375	5.5	9.3125	8.071429	7.1875	7.282738
adasyn	7	7	6.25	6.875	13.28571	8.53125	8.156994
cure-smote	12	11.4375	8.375	6	8.142857	8.625	9.096726
pf-smote	10.75	9.75	7.9375	5.125	10.96429	7.25	8.629464
prowras	15	13.125	7.375	9.25	9.214286	15.8125	11.62946
prowsyn	9.75	7.125	6.9375	4.625	12.96429	5.59375	7.832589
game	1.5	5.0625	2	6.625	7.857143	5.71875	4.7939
gan	13.75	13.4375	10.875	10.875	6.32143	14.0625	11.55357
dcgan	14	13.25	10.6875	11.0625	7.964286	14	11.82738
cgan	13.25	13.4375	11.1875	12	7.25	13.78125	11.81771
wgan	15.25	12.1875	10.625	10.75	6.821429	14.09375	11.62128
sgan	12.25	13.8125	11.125	11.6875	7.535714	14.25	11.77679
ctgan	4.25	6.875	9.3125	3.5625	7.535714	5.78125	6.219494
ctabgan	4.5	6.5625	10.3125	4.4375	10.17857	5.59375	6.930804
(b) minority average rank							
Min-Rank	HM	BCW	PID	GM	DCCC	Shuttle	Average Rank
imb-data	3	11.1875	9.4375	9.25	11.32143	10.71875	9.15253
random-os	11.75	11.3125	5.625	9	8.142857	9.90625	9.289435
smote	10	8.625	3.5	7.5	4.67857	8.03125	7.055804
smote-b1	15.5	8.75	4.9375	8.25	6.25	9.34375	8.838542
smote-b2	15.75	6.9375	4.25	8.25	5.714286	8.15625	8.176339
svm-smote	9.75	7	6	6.9375	9.071429	8.15625	7.819196
adasyn	12	7.625	4.1875	7.875	5.75	8.71875	7.692708
cure-smote	12	7.625	7.5	9.3125	8.25	9.03125	8.953125
pf-smote	13.5	7.0625	7.1875	10.625	7.75	6.5	8.770833
prowras	7	9.4375	6.6875	6.5625	6.107143	13.71875	8.252232
prowsyn	10.25	6.3125	4.0625	10	4.928571	5.6875	6.873512
game	11.75	4.875	2.9375	5.9375	6.642857	5.90625	6.34152
gan	5.25	14.125	11.125	4.875	12.53571	13.34375	10.20908
dcgan	4.5	14	11.125	5.375	13.82143	13.34375	10.36086
cgan	4.25	15.125	11.75	5.8125	13.67857	13.125	10.62351
wgan	5.75	13.4375	10.625	4.5	12.82143	13.375	10.08482
sgan	2.75	13.875	11.25	5.25	13.71429	13.5625	10.06696
ctgan	13.25	6.625	11.9375	15.125	14.5	6.84375	11.38021
ctabgan	8.75	6.5625	13	12.3125	15.25	7.0625	10.48958
(c) all average rank							
Rank	HM	BCW	PID	GM	DCCC	Shuttle	Average Rank
imb-data	7.75	10.25	8.03125	9.71875	9.928571	11.3125	9.498512
random-os	10.375	10.90625	5.71875	8.375	10.51786	9.8125	9.284226
smote	7.75	8.09375	4.78125	7.4375	8.803571	7.546875	7.402158
smote-b1	11.75	8.78125	5.3125	7.46875	8.25	8.890625	8.408854
smote-b2	10.75	6.90625	5.15625	8.4375	9.071429	7.640625	7.993676
svm-smote	8.5	6.6875	5.75	8.125	8.553571	7.671875	7.547991
adasyn	9.5	7.3125	5.21875	7.375	9.5	8.625	7.921875
cure-smote	12	9.53125	7.9375	7.65625	8.357143	8.828125	9.051711
pf-smote	12.125	8.40625	7.5625	7.875	9.339286	6.875	8.697173
prowras	11	11.28125	7.03125	7.90625	7.642857	14.76563	9.937872
prowsyn	10	6.71875	5.5	7.3125	8.946429	5.64063	7.353051
game	6.625	4.96875	2.46875	6.28125	7.23214	5.8125	5.56473
gan	9.5	13.78125	11	7.875	9.428571	13.70313	10.88132
dcgan	9.25	13.625	10.90625	8.21875	10.89286	13.67188	11.09412
cgan	8.75	14.28125	11.46875	8.90625	10.46429	13.45313	11.22061
wgan	10.5	12.8125	10.625	7.625	9.821429	13.73438	10.85305
sgan	7.5	13.84375	11.1875	8.46875	10.625	13.90625	10.92188
ctgan	8.75	6.75	10.625	9.34375	11.01786	6.3125	8.799851
ctabgan	6.625	6.5625	11.65625	8.375	12.71429	6.328125	8.710193

which is 5000 in the small group and 30,000 in the big group. The time consumptions are shown in the last column of each subtable of Tables 3 to 8. With the same hardware configuration, the SMOTE series completed data generation in one second, while the GANs take minutes even hours. Our GAME needs seconds to minutes, which is slower than SMOTEs but much faster than GANs.

In summary, we achieve more effective and more balanced results than the 18 mainstream algorithms.

6. Conclusion

In this paper, we propose a Generative Adversarial Minority Enlargement model to extend the data generating space. Based on the local linear property of samples, our GAME utilizes the adjustment of the generated samples and the parameters of the classification line, gradually approaching the majority category, enlarging the generating space, and getting a wide range of synthetic data. The target of our

- Du, J., Zhou, Y., Liu, P., Vong, C.-M., & Wang, T. (2021). Parameter-free loss for class-imbalanced deep learning in image classification. *IEEE Transactions on Neural Networks and Learning Systems*, 1–7. <http://dx.doi.org/10.1109/TNNLS.2021.3110885>.
- Gazzah, S., & Amara, N. E. B. (2008). New oversampling approaches based on polynomial fitting for imbalanced data sets. In *2008 the eighth iapr international workshop on document analysis systems* (pp. 677–684). IEEE.
- Goodfellow, I. (2017). NIPS 2016 tutorial: Generative adversarial networks. arxiv e-prints, arXiv:1701.00160.
- Goodfellow, I. J., Pouget-Abadie, J., Mirza, M., Xu, B., Warde-Farley, D., Ozair, S., et al. (2014). Generative adversarial nets. In *International conference on neural information processing systems* (pp. 2672–2680).
- Han, H., Wang, W. Y., & Mao, B. H. (2005). Borderline-SMOTE: a new over-sampling method in imbalanced data sets learning. In *International conference on advances in intelligent computing* (pp. 878–887).
- Hayashi, T., Fujita, H., & Hernandez-Matamoros, A. (2021). Less complexity one-class classification approach using construction error of convolutional image transformation network. *Information Sciences*, [ISSN: 0020-0255] 560, 217–234.
- He, H., Bai, Y., Garcia, E. A., & Li, S. (2008). ADASYN: Adaptive synthetic sampling approach for imbalanced learning. In *IEEE international joint conference on neural networks (IEEE world congress on computational intelligence), IJCNN 2008* (pp. 1322–1328).
- Huang, X., Li, Y., Poursaeed, O., Hopcroft, J., & Belongie, S. (2017). Stacked generative adversarial networks. In *2017 IEEE conference on computer vision and pattern recognition (CVPR)* (pp. 1866–1875).
- Kovács, G. (2019). An empirical comparison and evaluation of minority oversampling techniques on a large number of imbalanced datasets. *Applied Soft Computing*, [ISSN: 1568-4946] 83, Article 105662.
- Ma, L., & Fan, S. (2017). CURE-SMOTE algorithm and hybrid algorithm for feature selection and parameter optimization based on random forests. *BMC Bioinformatics*, 18(1), 1–18.
- Mehrannia, P., Bagi, S. S. G., Moshiri, B., & Al-Basir, O. A. (2021). Deep representation of imbalanced spatio-temporal traffic flow data for traffic accident detection. CoRR, arXiv:2108.09506.
- Mienye, I. D., & Sun, Y. (2021). Performance analysis of cost-sensitive learning methods with application to imbalanced medical data. *Informatics in Medicine Unlocked*, [ISSN: 2352-9148] 25, Article 100690.
- Mirza, M., & Osindero, S. (2014). Conditional generative adversarial nets. arXiv e-prints arXiv:1411.1784.
- Orooji, A., & Kermani, F. (2021). Machine learning based methods for handling imbalanced data in hepatitis diagnosis. *Frontiers in Health Informatics*, 10(1), 57.
- Papernot, N., McDaniel, P. D., Goodfellow, I. J., Jha, S., Celik, Z. B., & Swami, A. (2016). Practical black-box attacks against deep learning systems using adversarial examples. CoRR, arXiv:1602.02697.
- Radford, A., Metz, L., & Chintala, S. (2015). Unsupervised Representation Learning with Deep Convolutional Generative Adversarial Networks. arXiv e-prints, arXiv:1511.06434.
- Salimans, T., Goodfellow, I., Zaremba, W., Cheung, V., Radford, A., Chen, X., et al. (2016). Improved techniques for training GANs. In *Advances in neural information processing systems* 29 (pp. 2234–2242). Curran Associates, Inc..
- Wang, Q., Luo, Z., Huang, J., Feng, Y., & Liu, Z. (2017). A novel ensemble method for imbalanced data learning: Bagging of extrapolation-SMOTE SVM. *Computational Intelligence and Neuroscience*, 2017, 1827016:1–1827016:11.
- Xia, S., Zheng, S., Wang, G., Gao, X., & Wang, B. (2021). Granular ball sampling for noisy label classification or imbalanced classification. *IEEE Transactions on Neural Networks and Learning Systems*, 1–12. <http://dx.doi.org/10.1109/TNNLS.2021.3105984>.
- Xu, L., Skoularidou, M., Cuesta-Infante, A., & Veeramachaneni, K. (2019). Modeling tabular data using conditional GAN. vol. 32, In *Advances in neural information processing systems* (pp. 7335–7345).
- Xu, Y., Yu, Z., Chen, C. L. P., & Liu, Z. (2021). Adaptive subspace optimization ensemble method for high-dimensional imbalanced data classification. *IEEE Transactions on Neural Networks and Learning Systems*, 1–14. <http://dx.doi.org/10.1109/TNNLS.2021.3106306>.
- Zhang, T., Chen, J., Li, F., Zhang, K., Lv, H., He, S., et al. (2022). Intelligent fault diagnosis of machines with small & imbalanced data: A state-of-the-art review and possible extensions. *ISA Transactions*, 119, 152–171.
- Zhao, Y., Hao, K., Tang, X.-S., Chen, L., & Wei, B. (2021). A conditional variational autoencoder based self-transferred algorithm for imbalanced classification. *Knowledge-Based Systems*, 218, Article 106756.
- Zhao, Z., Kunar, A., Birke, R., & Chen, L. Y. (2021). Ctab-gan: Effective table data synthesizing. In *Asian conference on machine learning* (pp. 97–112). PMLR.