

Asynchronous Event Error-Minimizing Noise for Safeguarding Event Dataset

Ruofei Wang¹ Peiqi Duan^{2,3} Boxin Shi^{2,3} Renjie Wan^{1*}

¹Department of Computer Science, Hong Kong Baptist University

²State Key Laboratory of Multimedia Information Processing, School of Computer Science, Peking University

³National Engineering Research Center of Visual Technology, School of Computer Science, Peking University

ruofei@life.hkbu.edu.hk, {duanqi0001, shiboxin}@pku.edu.cn, renjiewan@hkbu.edu.hk

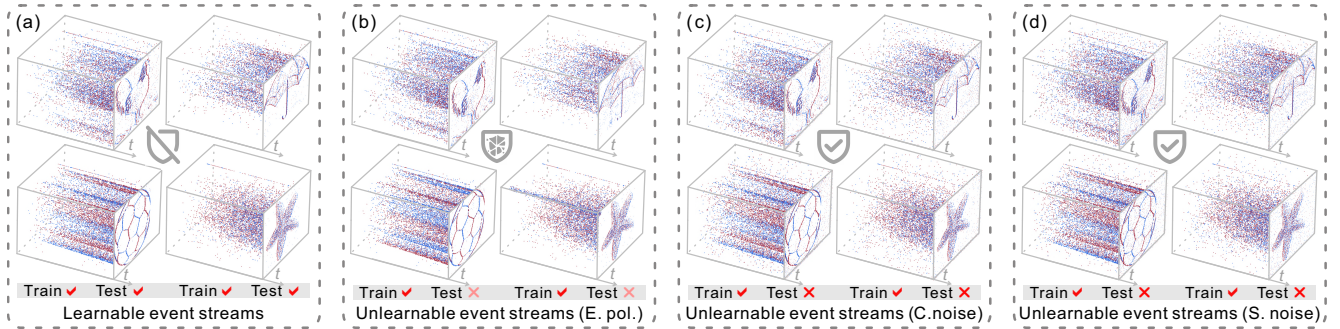


Figure 1. Comparison between the learnable event streams (*i.e.*, clean data), unlearnable event streams generated by event pollution (E. pol.), and our UEVs, including class-wise noise (C. noise) and sample-wise noise (S. noise). In case (b), we degrade the quality of these event streams by polluting the coordinates, timestamps, polarity, and adding new random events (reading order).

Abstract

With more event datasets being released online, safeguarding the event dataset against unauthorized usage has become a serious concern for data owners. Unlearnable Examples are proposed to prevent the unauthorized exploitation of image datasets. However, it's unclear how to create unlearnable asynchronous event streams to prevent event misuse. In this work, we propose the first unlearnable event stream generation method to prevent unauthorized training from event datasets. A new form of asynchronous event error-minimizing noise is proposed to perturb event streams, tricking the unauthorized model into learning embedded noise instead of realistic features. To be compatible with the sparse event, a projection strategy is presented to sparsify the noise to render our unlearnable event streams (UEVs). Extensive experiments demonstrate that our method effectively protects event data from unauthorized exploitation, while preserving their utility for legitimate use. We hope our UEVs contribute to the advancement of secure and trustworthy event dataset sharing. Code is available at: <https://github.com/rfww/uevs>.

1. Introduction

Recently, the easy availability of various event datasets [15, 27, 50, 54] has significantly accelerated the development of event vision tasks across various domains, including autonomous driving [39], pose estimation [24, 47], sign-language recognition [59], *etc.* However, a concerning fact is that some datasets are collected for restricted applications, without permission for unauthorized usage [1, 50]. The lack of protection mechanisms for event datasets leaves them susceptible to misuse for unauthorized purposes.

To prevent unauthorized dataset usage, a novel concept known as Unlearnable Examples (UEs) [22] is proposed. Different from data pollution that degrades the dataset quality to avoid misuse [58, 60], UEs incorporates an imperceptible adversarial noise in clean images to make them unlearnable for machine learning models. This technique perturbs the data in a way that disrupts the learning process of unauthorized models, effectively rendering the data unusable for unauthorized purposes. However, for authorized users, these data still retain the original utility and appearance. Therefore, this approach not only reduces the risk of data misuse but also preserves the dataset's fidelity and integrity, posing an effective way for safeguarding datasets.

Though UEs achieve great success in protecting image

*Corresponding author. This work was carried out at Renjie Group.

datasets [12, 46, 57], how to generate effective unlearnable asynchronous event dataset is unclear. As illustrated in Fig. 1 (a), event data comprises a sequence of asynchronous events (*i.e.*, (x, y, t, p)), where (x, y) represents the x - y coordinates, t denotes the timestamp, and p indicates the polarity. This special characteristic creates an event stream that resembles spatio-temporal point clouds rather than conventional 2D images [9]. Therefore, it is impossible to employ existing image UEs methods [22, 26, 57] to safeguard asynchronous event datasets.

Owing to the unique characteristics of event data, there are two main challenges in achieving the unlearnable functionality of event data. First, event data are distinguished by the binary polarity: $p = \pm 1$. It creates a highly discrete pattern distinct from the image samples used in previous UEs. Directly incorporating noise based on previous settings can diminish the effectiveness, as the noise contains values outside the binary pattern. This raises a significant challenge: generating effective error-minimizing noise that aligns with the discrete nature of event data. Second, current event vision tasks need to transfer the event data into the event stack and then feed it into the downstream DNN models. However, these stacks prevent noise from being injected directly into the event stream. Thus, effectively injecting noise into the event stream presents another significant challenge.

To this end, we propose the Unlearnable Event stream (UEVs) method with the Event Error-Minimizing Noise (E^2MN) to safeguard event datasets. **First**, we formulate an event-based optimization function to generate E^2MN . Specifically, our E^2MN includes two forms: class-wise noise (Fig. 1 (c)) and sample-wise noise (Fig. 1 (d)), with the former generated on the class-by-class basis and the latter on case-by-case basis. These two forms all aim to build a shortcut between the input sample and target labels to prevent the model from learning real semantic features. **Second**, we propose an adaptive projection mechanism to sparsify the generated noise, resulting in the projected noise that can be compatible with event stacks seamlessly. Thereby, the possibility of our E^2MN being detected by malicious users is reduced. **Third**, we propose a new retrieval strategy to reconstruct the event stream from its corresponding representation. This strategy ensures that E^2MN is effectively transferred from unlearnable event stacks to unlearnable event streams.

Our main contributions are as follows:

- We present the first Unlearnable Event stream (UEVs) with an event error-minimizing noise (E^2MN) to prevent unauthorized usage of the valuable event dataset.
- We propose an adaptive projection mechanism to sparsify the noise E^2MN , enabling the noise to be compatible with the original event stack. Subsequently, a retrieval strategy is introduced to reconstruct the unlearnable event stream from its corresponding event stack.

- We conduct extensive experiments to demonstrate that our E^2MN outperforms various event pollution operations (Fig. 1 (b)) in terms of effectiveness, imperceptibility, and robustness.

2. Related work

2.1. Event vision

Event vision [8, 10, 37] has gained increasing popularity recently due to the advantages brought by the event data [25, 33], which consists of a series of asynchronous events by recording the change of pixel-wise brightness [13]. Compared with conventional images, event data shows merits in terms of high dynamic range and temporal resolution, low time latency and power consumption [14, 50]. The easily available and diverse event datasets [1, 15, 27, 31, 42] have been crucial for the increasing attention to various event vision tasks. Researchers can effortlessly ensure their event data is compatible with current deep vision models via utilizing existing event representation methods [14, 23, 30], thereby effectively facilitating various explorations in event-based learning. Millerdurai *et al.* [41] propose employing an egocentric monocular event camera for 3D human motion capture, which addresses the failure of RGB cameras under low lighting and fast motions. Sun *et al.* [52] propose an unsupervised domain adaption method for semantic segmentation on event data, which motivates the segmentor to learn semantic information from labeled images to unlabeled events. Moreover, event-based studies achieve satisfactory performances in pose estimation [24], segmentation [5], deblurring [3, 28], denoising [9], optical flow [32], object recognition [61], object tracking [56], monitoring [18], *etc.*

Despite the significant interest and promising performance of event-based learning in various domains, there have been no investigations into preventing the unauthorized exploitation of event datasets. The absence of this exploration could result in the valuable event datasets being stolen for illegal purposes, leading to serious threat.

2.2. Unlearnable examples

Unlearnable Examples (UEs) are proposed to prevent the unauthorized training of Deep Neural Networks (DNNs) on some private or protected datasets [22, 48]. Generally, UEs are generated through a min-min bilevel optimization pipeline with a surrogate model [22]. The generated noise is called Error-Minimizing Noise (EMN) because it gradually reduces losses from the training data. This noise aims to deceive the target model into believing that correct predictions can be made merely based on perturbations, resulting in overlooking semantic features [22, 46]. To improve the robustness of vanilla UEs, Fu *et al.* [12] introduce adversarial noise into the EMN against adversarial training. He *et*

al. [19] propose to generate unlearnable examples based on unsupervised contrastive learning, which extends the supervised unlearnable noise generation into unsupervised learning. Ren *et al.* [46] propose Classwise Separability Discriminant, which aims to better transfer the unlearnable effects to other training settings and datasets by enhancing the linear separability. Meng *et al.* [40] propose a deep hiding strategy that adaptively hides semantic images into the latent domain of poisoned samples to generate UEs. Although current UEs have shown promising results in image area [34, 35, 51], this technique cannot be directly applied to protect asynchronous event streams due to the different data characteristics. To this end, our work focuses on exploring the unlearnable event stream to safeguard the event dataset against unauthorized usage.

3. Methodology

3.1. Preliminary of event data

Given an event data-based classification dataset $\mathcal{D} = \{(\mathcal{E}_i, l_i)\}_{i=1}^N$, where $\mathcal{E}$, l , and N indicate the event stream, the corresponding category label, and the dataset length, respectively. The event stream $\mathcal{E}$ consists of a variety of individual events, recorded as:

$$\mathcal{E} = \{e_k\}_{k=1}^K = \{(x_k, y_k, t_k, p_k)\}_{k=1}^K, \quad (1)$$

where (x_k, y_k, t_k, p_k) indicates the x and y direction coordinates, time stamp, and polarity of the k -th event, respectively. K is the length of the event stream $\mathcal{E}$ [13]. Specifically, an event, e_k , has occurred when the variation of the log brightness at each pixel exceeds the threshold σ , *i.e.*, $|\log(x_k, y_k, t_k) - \log(x_k, y_k, t_{k-1})| > \sigma$. The polarity $p_k = +1.0$ when the difference between bi-temporal brightness is higher than $+\sigma$. Otherwise, p_k is set to -1.0 .

3.2. Our scenario

We propose UEVs to prevent unauthorized training on the valuable event datasets that effectively protect the data owner's interests. Generally, two parties are mentioned in unlearnable event streams, *i.e.*, the protector and the hacker. The protector has full accessibility to the original dataset and generates UEVs for his/her private data before the dataset is released. To ensure the unlearnability, the protector can use various surrogate models to generate the event error-minimizing noise. However, he/she cannot touch any information about the models or training skills that the hacker would use. As a hacker, he/she has no knowledge about the original datasets and surrogate models. He/she aims to steal these data to train their illegal models with any training tricks. In this paper, we are the protector to make the collected event dataset unlearnable with the imperceptible noise E²MN.

3.3. Unlearnable event stream

Overview. The overall pipeline of our UEVs is shown in Fig. 2. The framework takes the original event stream as input and outputs an unlearnable one. Based on the fact that an example with higher training loss contains more useful knowledge to be learned [12], we formulate the pipeline of our Unlearnable Event streams (UEVs) as:

$$\arg \min_{\theta} \mathbb{E}_{(\mathcal{E}, l) \in \mathcal{D}} [\min_{\delta} \mathcal{L}(f'_{\theta}(\mathcal{R}(\mathcal{E}) + \delta), l)] \text{ s.t. } \|\delta\|_{\infty} \leq \epsilon, \quad (2)$$

where f' denotes a surrogate model used for noise δ generation, $\mathcal{L}$ indicates the loss function, and $\mathcal{R}$ means event representation. Eq. (2) is a min-min bi-level optimization function, where both constrained optimization items all aim to minimize the loss of the surrogate model. Specifically, the inner minimization is employed to find the L_{∞} -norm bounded noise δ to improve the model's performance, while the outer optimization finds the optimal parameter θ to minimize the model's classification loss. Obviously, the core component of UEVs is to find the effective noise δ to minimize the loss of f' , thereby suppressing the model from learning real semantic features from input events.

To generate unlearnable event streams, first, we convert the asynchronous event data into a regular event stack; second, we feed this stack into a surrogate model to compute the noise δ ; third, we project δ to be compatible with the binary-polarity event stack; and fourth, we reconstruct the unlearnable event stream from the unlearnable event stack.

Event representation. To be compatible with existing vision models [20, 49, 53], the event stream needs to be transferred into image-like stacks and then input as images for the following predictions [2, 45]. To meet the regular input requirement of the surrogate model, we propose binning events into a C -channel event stack. We set C with a large value to prevent possible event corruption since the single or three-channel maps may easily lead to the coverage of previous events. Each channel of this stack consists of polarities of those events distributed in the divided time bin Δ_t . Hence, our event stack only contains three types of values at each coordinate: $\{0, 0.5, 1.0\}$, denoting an event with polarity = -1 , no event, and an event with polarity = $+1$, respectively. Note that the event data cannot be reconstructed from the event stack if any values other than these two specific values are used.

Noise generation. The event stack is then fed into the surrogate model f' used for the noise generation. The training steps of f' are limited since training the surrogate model with larger steps would lead to ineffective event error-minimizing noise optimization. Thus, we train the parameter θ of f' with only M iterations first and then op-

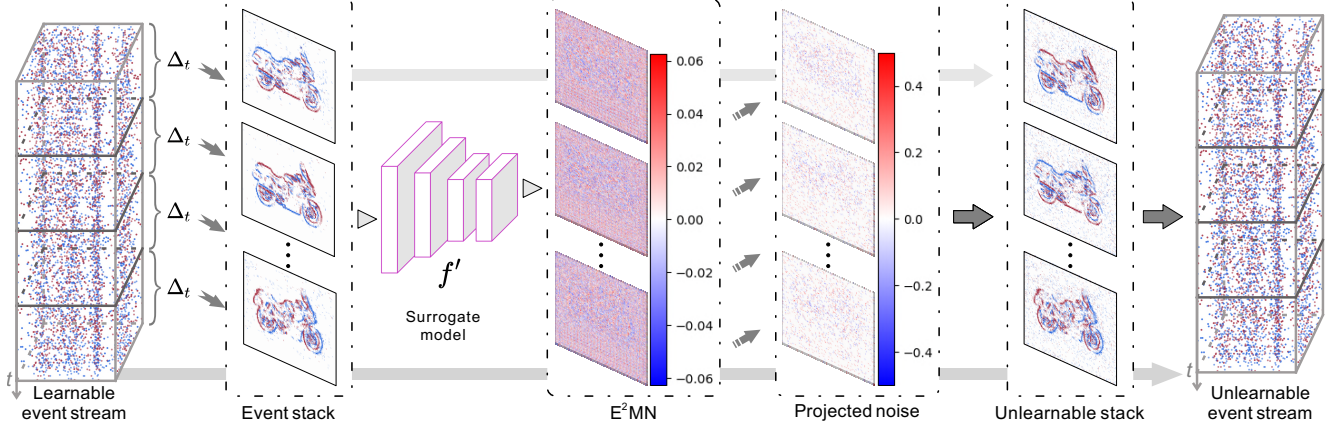


Figure 2. Framework of our UEVs. UEVs converts an event stream to the C -channel event stack ($t = C \times \Delta_t$) and then employs a surrogate model f' to generate the event error-minimizing noise (E^2MN). Subsequently, the E^2MN is projected into $\{-0.5, 0, 0.5\}$ and integrated with the target event stack to generate the unlearnable event stack. The final unlearnable event stream can be reconstructed from the unlearnable event stack and original event stream via the proposed retrieval strategy.

timize δ across the entire $\mathcal{D}_c$. This entire bi-level optimization process (Eq. (2)) is terminated once the classification accuracy is higher than γ . Essentially, the generated event error-minimizing noise is a kind of easy-to-learn feature (*i.e.*, the shortcut between input samples and target labels) that guides the model trained on E^2MN to ignore their realistic semantics. Our event error-minimizing noise consists of two types:

- **Sample-wise noise.** Sample-wise noise Δ_s is obtained by the first-order optimization method PGD [38] as follows:

$$\hat{\mathcal{E}}'_{s+1} = \Pi_{\epsilon}(\hat{\mathcal{E}}'_s - \alpha \cdot \text{sign}(\nabla_{\hat{\mathcal{E}}} \mathcal{L}(f'(\hat{\mathcal{E}}'_s), l))), \quad (3)$$

where, $\hat{\mathcal{E}} = \mathcal{R}(\mathcal{E})$ indicates the event stack, s denotes the current perturbation step, $\nabla_{\hat{\mathcal{E}}} \mathcal{L}(f'(\hat{\mathcal{E}}'_s), l)$ is the gradient of the loss with respect to the $\hat{\mathcal{E}}$, Π_{ϵ} is a projection function that clips the perturbed event stack back to the ϵ -ball around the event stack $\hat{\mathcal{E}}$ when it goes beyond, and α is the step size. $\mathcal{L}$ is the cross entropy loss function. The generated event error-minimizing noise is achieved by: $\delta_i = \hat{\mathcal{E}}'_i - \hat{\mathcal{E}}_i$ and calculates it case-by-case to get $\Delta_s = \{\delta_1, \delta_2, \dots, \delta_N\}$. The output $\hat{\mathcal{E}}'$ is not our final unlearnable event stack since it does not exhibit a binary-polarity characteristic, which is crucial for the reconstruction of unlearnable events. We need to project this noise into a sparse pattern that is compatible with the original event stack to generate our unlearnable event streams.

- **Class-wise noise.** Since sample-wise noise depends on each input event, memory consumption would increase sharply as the dataset scale grows. Therefore, we introduce the class-wise noise: $\Delta_c = \{\delta_{l_1}, \delta_{l_2}, \dots, \delta_{l_N}\}$, which applies the same noise to different events within the same class. For Δ_c , we firstly initialize the noise in the class-level and then employ Eq. (3) to optimize the

noise δ_l with all event streams sampled from the class l . Class-wise noise shows advantages over sample-wise noise in terms of practicality and efficiency. It can be easily applied to newly captured event streams and has a smaller noise scale. Same to the sample-wise noise, class-wise noise cannot be directly used to generate unlearnable event streams, which also needs to be projected into a sparse form for UEVs generation.

Noise projection. As demonstrated in Sec. 3.1, the event stream consists solely of polarities of -1 and $+1$, which are normalized to 0 and 1 in event stack, receptively. Therefore, we need to project the generated event error-minimizing noise into a sparse and event-friendly pattern to facilitate the subsequent generation of unlearnable event streams. Here, we propose projecting the generated noise into $\{-0.5, 0.0, +0.5\}$ as follows:

$$\mathbf{P}(\delta) = \begin{cases} -0.5, & \text{if } \delta_{i,j} < \mu - \tau \times \pi, \\ 0.0, & \text{if } \mu - \tau \times \pi \leq \delta_{i,j} \leq \mu + \tau \times \pi, \\ +0.5, & \text{if } \delta_{i,j} > \mu + \tau \times \pi, \end{cases} \quad (4)$$

where, the μ and π denote the mean and $1/2$ bound of noise δ , τ is a balancing parameter.

τ imposes great influence on two criteria of our UEVs, *i.e.*, effectiveness and imperceptibility. Effectiveness means that the projected noise still owns the unlearnable functionality for event streams. Imperceptibility denotes that the injected noise is imperceptible to users. A higher τ benefits better imperceptibility but corrupts the effectiveness. To ensure the effectiveness of the projected noise, we enforce the surrogate model to discriminate between the unlearnable features and clean features during model training. Thereby,

the model can generate a strong shortcut during noise optimization, preventing the model from learning real semantic features. We reformulate the loss function of the outer optimization in Eq. (2) as:

$$\mathcal{L}^* = \lambda_1 \mathcal{L} + \lambda_2 \mathcal{L}_s. \quad (5)$$

$\mathcal{L}_s$ is the similarity loss that aims to enlarge the difference between clean and unlearnable features.

Event reconstruction. After getting the projected event error-minimizing noise, we incorporate this noise into $\hat{\mathcal{E}}$ to generate the unlearnable counterpart. Note that what we obtain here is only the unlearnable event stack. We need to reconstruct the unlearnable event stream from this stack. Therefore, a retrieval strategy is proposed to recover the compressed dense temporal information when converting event data into event stacks. For an event within the unlearnable event stack $\hat{\mathcal{E}}$, **1**) if it is recorded in original event stream $\mathcal{E}$ then search for the time stamp t in $\mathcal{E}$ to assign it directly; **2**) if it is a newly generated event, initialize a new adaptive time stamp based on the selected Δ_t to set it. Based on this strategy, the generated unlearnable events can be transferred from unlearnable event stacks into event streams, as shown in Fig. 1. The overall pipeline of our UEVs is illustrated in Algorithm 1. $\mathcal{R}(\cdot)$ and $\mathcal{R}'(\cdot)$ denote event representation and event reconstruction, respectively.

3.4. Implementation details

Following [22], we employ the ResNet18 [20] as the surrogate model. The SGD optimizer is employed with a learning rate of 10^{-4} and a momentum of 0.9. An exponential learning rate scheduler with a gamma of 0.9 is used. The iteration number M and epoch number are set to 10 and 30, respectively. The batch size is 16. The C in the event stack is set to 16. In noise generation, the step number S , parameter ϵ , and step size α are 10, 0.5, and 0.8/255, respectively. The balancing ratio τ is 3/4. The accuracy γ is 0.99. We train all models on a single NVIDIA V100 GPU.

4. Experiment

4.1. Setup

Dataset. Four event-based datasets are used in our experiments, including N-Caltech101 [42], CIFAR10-DVS [31], DVS128 Gesture [1], and N-ImageNet [27]. Due to the limited computing resources, the first 10 classes are sampled from N-ImageNet [27] for experiments. Detailed information is shown in Table 6 of the Supplementary Material.

Baseline. To evaluate the effectiveness, we implement several event pollution operations (see Fig. 1 (b)) as our baselines, including coordinate shifting (CS), time stamp shifting (TS), polarity inversion (PI), manual pattern injection (MP), and area shuffling (AS).

Algorithm 1: UEVs Perturbation Algorithm.

Input: Surrogate model $f'_\theta(\cdot)$, clean dataset $\mathcal{D}_c$, optimization number M , accuracy γ .

Output: Perturbation δ , Unlearnable dataset $\mathcal{D}_u$.

```

1  $\delta \leftarrow \text{randomization};$ 
2 Do
3   for  $i$  in  $\text{range}(M)$  do
4      $(\mathcal{E}, l) = \text{Next}(\mathcal{D}_c);$ 
5      $\theta \leftarrow \theta - \nabla_\theta(\mathcal{L}^*(f'_\theta(\mathcal{R}(\mathcal{E}) + [\delta_i / \delta_{l_i}]), l))$ 
         $\triangleright$  Sample or class-wise noise:  $\delta_i / \delta_{l_i}$ 
6   end
7   for  $(\mathcal{E}, l)$  in  $\mathcal{D}_c$  do
8      $\delta = \mathbb{P}(\mathcal{E}, l, f', \delta, \theta);$ 
         $\triangleright$  Perturbed by Eq. (3)
9      $\text{Clip}(\delta, -\epsilon, +\epsilon);$ 
10  end
11   $\text{acc} = \text{evaluation}(\mathcal{D}_c, f', \delta, \theta);$ 
12 Until  $\text{acc}$  higher than  $\gamma$ ;
13 for  $(\mathcal{E}_j, l_j)$  in  $\mathcal{D}_c$  do
14    $\delta' = \mathbb{P}(\delta \leftarrow [\delta_j \text{ or } \delta_{l_j}]); \triangleright \text{Eq. (4)}$ 
15    $\mathcal{E}' = \mathcal{R}'(\text{Clip}(\mathcal{R}(\mathcal{E}_j) + \delta', 0, 1));$ 
16    $\mathcal{D}_u.\text{append}((\mathcal{E}'_j, l_j)); \triangleright \text{Unlearnable event}$ 
17 end
```

To test the generalizability of UEVs, we introduce ResNet18 (RN18) [20], ResNet50 (RN50) [20], VGG16 (VG16) [49], DenseNet121 (DN121) [21], EfficientNet-B1 (EN-B1) [53], ViT_B [7], and Swin_B [36] as baselines in our experiment.

Metric. Following the existing UEs method [22], we use the test accuracy to evaluate the data protection ability of the proposed method. The lower the test accuracy, the stronger the protection ability. To assess the imperceptibility of the generated noise, we employ PSNR, SSIM, and MSE to measure the visual perception on the event stack [9].

4.2. Evaluation

Visualization comparison of aforementioned various noise forms is shown in Fig. 3. We implement five kinds of event distortion operations as baselines to evaluate our UEVs. The coordinate shifting (CS) only introduces an offset in x and y directions without adding any additional noise, as shown in the 2nd column. We put off the time stamps of each event by a fixed time bin to achieve the TS, which shows that some new events are introduced in the current representations. Reverse the polarity of all events (PI) is another noise form that shows a noticeable difference between the clean data and unlearnable counterpart in the 4th column. MP is implemented by injecting special event streams manually, which shows an obvious pattern in the

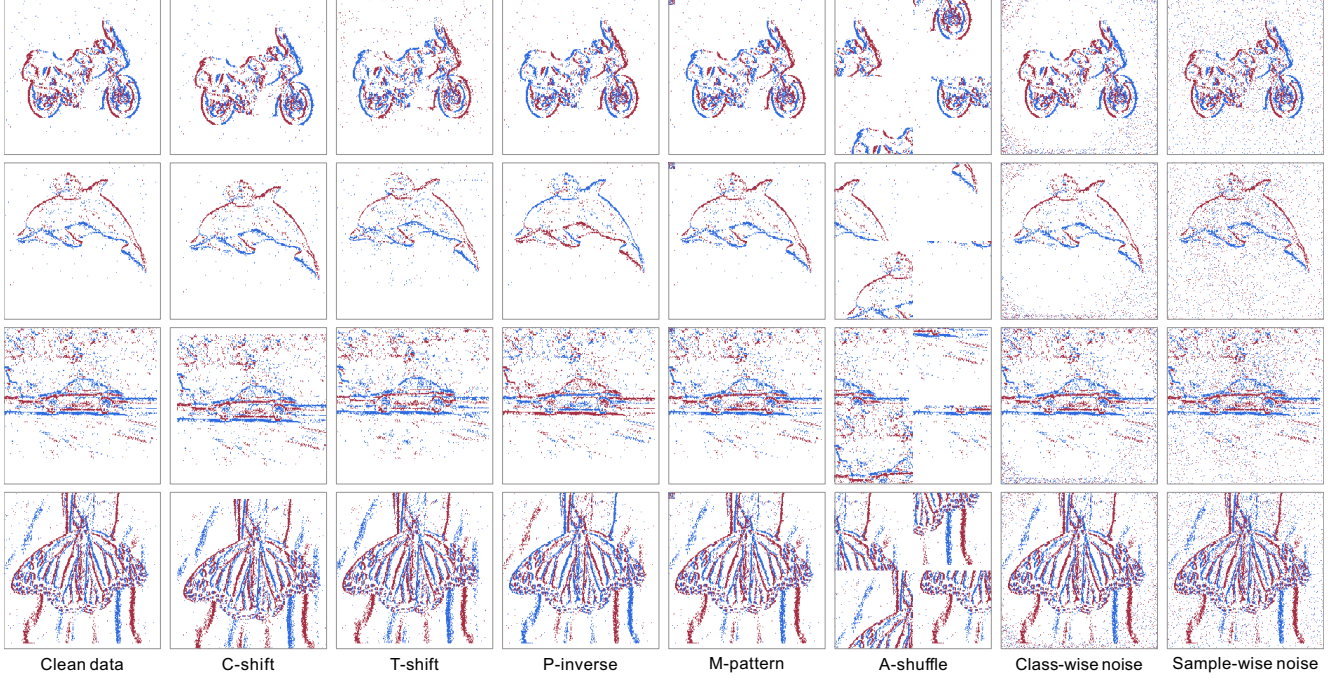


Figure 3. Visualization results of various noise forms on the N-Caltech101 dataset [42]. Blue/Red points denote the events with $p = +1/-1$. Our noise E^2MN does not corrupt the target objects, and the noise distributed in the background region appears quite realistic.

top-left corner of the representation. We also shuffle the event in block-level (AS) to confuse the model training, which shows a noticeable difference between the clean data and unlearnable one poisoned by AS. Our class-wise noise and sample-wise noise are shown in the last two columns of Fig. 3. Detailed imperceptibility evaluation of each kind of noise is shown in Table 1. Injecting a small manual pattern (MP) can achieve a good imperceptibility. But our Δ_c and Δ_s achieve the best overall performance in imperceptibility and unlearnability. The PSNR of Δ_c and Δ_s are 20.39 and 18.22, respectively.

Quantitative results of various unlearnable noise forms on the four datasets are shown in Table 2. We can find that our sample-wise and class-wise noise achieve a noticeable accuracy drop among the whole classifiers compared with the baselines. Shifting the coordinates or time stamps of the event cannot achieve a reliable unlearnable effect. For example, on the DVS128 Gesture dataset, ResNet18 achieves an improvement by 1.21% when train the classifier on $\mathcal{D}_u$ poisoned by CS than the clean dataset $\mathcal{D}_c$. Setting the polarity of each event to inverse can confuse the model to learn semantic features, where PI achieves the accuracy drop by 40.87% on the N-Caltech101 dataset. However, this noise can be removed easily via simple data augmentation skills. Injecting a manual pattern (MP) into the event stream to build a shortcut between the event data and target label is in-

Table 1. Imperceptibility evaluation of various noise forms.

	CS	TS	PI	MP	AS	Δ_c	Δ_s
PSNR	12.19	13.07	16.76	28.82	11.75	20.36	18.22
SSIM	0.316	0.418	0.996	0.997	0.245	0.791	0.571
MSE	0.182	0.143	0.067	0.069	0.212	0.028	0.040

spired by backdoor attacks [55]. We randomize patterns for each class to poison the event data that influences the accuracy of ResNet18 from 65.19% to 36.12% on the CIFAR10-DVS dataset. On N-ImageNet, the MP achieves an obvious accuracy drop by 45.40 since it is easier to be captured by deep models than the real complex features of input data. Shuffling the event in block-level (AS) (see the 6th column of Fig. 3) cannot protect the event stream effectively, which achieves the accuracy by 48.25%, 39.81%, 58.68%, and 28.00% on N-Caltech101, CIFAR10-DVS, DVS128 Gesture, and N-ImageNet, respectively.

On the contrary, our sample-wise noise and class-wise noise all achieve credible unlearnable performance on four datasets. Since ResNet18 is selected as the surrogate model in our experiment, the best unlearnable effectiveness is obtained when a malicious user trains ResNet18 on the unlearnable datasets. To evaluate the generalizability of the UEVs, we have trained a series of popular DNN classifiers on the clean data and our unlearnable counterpart. As shown in Table 2, the class-wise noise can deceive almost

Table 2. The testing accuracy (%) of various DNN classifiers trained on the clean training sets ($\mathcal{D}_c$) and their unlearnable ones ($\mathcal{D}_u$) generated by five kinds of event pollutions, sample-wise noise (Δ_s), and class-wise noise (Δ_c).

Noise Form	Model	N-Caltech101		CIFAR10-DVS		DVS128 Gesture		N-ImageNet	
		$\mathcal{D}_c$	$\mathcal{D}_u$	$\mathcal{D}_c$	$\mathcal{D}_u$	$\mathcal{D}_c$	$\mathcal{D}_u$	$\mathcal{D}_c$	$\mathcal{D}_u$
CS	Res18	78.32	50.43 _(-27.79)	65.19	46.80 _(-18.39)	74.14	75.35 _(+01.21)	56.60	42.60 _(-14.00)
TS			37.54 _(-40.87)		35.34 _(-29.85)		72.92 _(-01.22)		41.20 _(-15.40)
PI			40.78 _(-37.45)		22.72 _(-42.47)		31.94 _(-42.20)		41.80 _(-4.80)
MP			50.26 _(-28.06)		36.12 _(-29.07)		68.06 _(-06.08)		11.20 _(-45.40)
AS			48.25 _(-30.07)		39.81 _(-25.38)		58.68 _(-15.46)		28.00 _(-28.00)
Δ_c	RN18	78.52	01.90 _(-76.62)	65.22	22.33 _(-42.89)	74.31	14.93 _(-59.38)	56.80	10.00 _(-46.80)
	RN50	80.64	04.88 _(-75.76)	67.67	18.83 _(-48.84)	78.47	24.31 _(-54.16)	61.60	10.00 _(-51.60)
	VG16	71.63	03.68 _(-67.95)	67.77	30.19 _(-37.58)	81.94	20.49 _(-61.45)	56.20	04.00 _(-52.20)
	DN121	75.59	02.18 _(-73.31)	62.72	24.37 _(-38.35)	74.31	22.22 _(-52.09)	59.20	08.80 _(-50.40)
	EN-B1	74.44	12.81 _(-61.63)	72.72	23.50 _(-49.22)	62.50	19.10 _(-43.40)	56.20	06.20 _(-50.00)
	ViT _B	44.69	01.55 _(-43.14)	45.44	09.51 _(-35.93)	66.67	26.74 _(-39.93)	41.00	10.00 _(-31.00)
	Swin _B	90.70	00.52 _(-90.18)	75.24	29.03 _(-46.21)	83.29	28.12 _(-55.17)	72.00	10.00 _(-62.00)
Δ_s	RN18	78.12	00.52 _(-77.60)	65.15	15.51 _(-49.64)	73.96	14.54 _(-59.42)	56.40	10.20 _(-46.20)
	RN50	80.30	06.09 _(-74.21)	67.86	12.62 _(-55.24)	79.17	20.03 _(-59.14)	61.20	16.80 _(-44.40)
	VG16	72.03	14.13 _(-57.90)	67.18	27.86 _(-39.32)	81.94	26.39 _(-55.55)	55.40	14.80 _(-40.60)
	DN121	75.30	11.49 _(-63.85)	64.27	15.73 _(-48.54)	74.31	20.83 _(-53.48)	59.20	11.40 _(-47.80)
	EN-B1	73.81	17.98 _(-55.83)	72.82	14.95 _(-57.87)	64.93	16.32 _(-48.61)	55.60	11.40 _(-44.20)
	ViT _B	44.73	01.09 _(-43.64)	47.38	26.02 _(-21.36)	67.01	30.21 _(-36.80)	40.40	08.80 _(-31.60)
	Swin _B	90.70	21.14 _(-69.56)	75.44	27.18 _(-48.26)	83.68	17.01 _(-66.67)	71.20	09.00 _(-62.20)

all of the classifiers to learn shortcuts on N-Caltech101, where the accuracy of these classifiers is less than 5% except EfficientNet-B1. Moreover, it also achieves a great accuracy drop higher than 35% on both CIFAR10-DVS and DVS128 Gesture datasets. For sample-wise noise, every data has a different low-error noise, which noise can reduce the loss during the optimization process of the model, which leads to the event data unlearnable. In Table 2, sample-wise noise Δ_s achieves the biggest accuracy drop by 77.60%, 57.87%, and 66.67% on four datasets, respectively.

4.3. Ablation study

Similarity loss. To ensure the event error-minimizing noise can still prevent model learning after being projected, we introduce the similarity loss to enlarge the distance between the clean and unlearnable stacks during the surrogate model training. Detailed results are shown in Table 3. Compared with **E1** in Table 3, the effectiveness of our sample-wise noise (**E2**) is decreased when we remove the similarity loss. Minimizing the similarity between the clean and unlearnable stacks is to improve the difference between the shortcut and real semantic features of input samples under the correct classification, avoiding possible feature learning. As shown in Fig. 4, this similarity supervision is crucial for ensuring the unlearnability of our UEVs due to the big difference between the original noise and projected noise.

Mixed noise. In the main experiment, we have conducted exploration about the sample-wise noise and class-wise noise, respectively. According to detailed results, we can

Table 3. Results of ablation experiments for UEVs (sample-wise noise) on N-Caltech101 dataset.

	RN18	RN50	VG16	DN121	EN-B1	ViT _B	Swin _B
E1	0.52	6.09	14.13	11.49	17.98	1.09	21.14
E2	3.85	15.22	38.31	10.17	24.18	37.51	27.51
E3	5.17	5.40	22.23	6.15	18.78	1.21	16.66
E4	5.34	5.17	8.16	5.17	5.86	1.03	0.57
E5	13.04	8.04	14.99	10.74	16.54	9.88	25.73
E6	27.74	33.08	44.28	34.92	27.17	8.85	24.41

find that sample-wise and class-wise noises achieve the best performance on different datasets among various classifiers. Here, we propose a mixed counterpart to evaluate the effectiveness of our UEVs, which is beneficial for personalized utilization and preventing noise exposure. Since it's hard to optimize these two kinds of noises simultaneously, we mix the obtained two noises as the mixed counterpart instead of training it from scratch. Specifically, we design two modes to generate the mixed noise. First, we randomly select δ from Δ_c and Δ_s to poison the input sample, denoting as $\Delta_c \vee \Delta_u$. Second, we employ the element-wise addition to generate the unlearnable examples, naming $\Delta_c + \Delta_u$.

As shown in Table 3, **E3** and **E4** represent the $\Delta_c \vee \Delta_u$ and $\Delta_c + \Delta_u$ respectively, where the proposed event error-minimizing noise remains highly effective. Although employing $\Delta_c + \Delta_u$ can achieve a better unlearnable performance, simple addition leads to the low stealthiness of the noise. Randomly selecting noise from $\Delta_c \vee \Delta_u$ provides another possible solution for dataset protection, making our

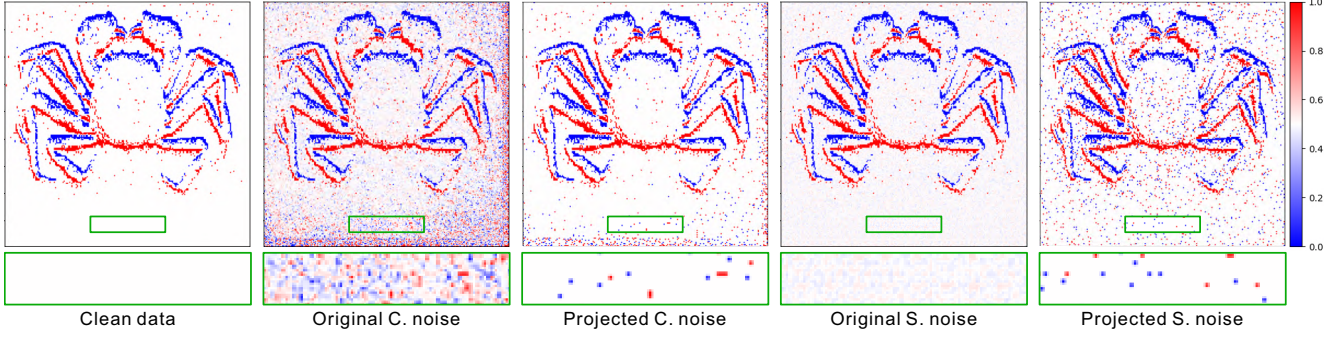


Figure 4. Visualization results of clean data representation, unlearnable representations generated by class-wise noise and sample-wise noise with and without projection, respectively. Details are zoomed in the green box.

UEVs more flexible in practical usage.

Robustness. In event-based learning, data augmentation is always adopted to enhance the model’s performance. Therefore, studying the robustness of our unlearnable event streams against data augmentation methods is of great importance. Here, we employ the common random shift, crop, flip, and event drop [17] as augmentation techniques in our experiments. As shown in Table 3 E5, these augmentation techniques slightly eliminate the effectiveness of our event error-minimizing noise. This evaluation demonstrates the high practicality of our UEVs in various training scenarios.

Different adversarial attacks. As shown in Eq. (3), our E^2MN is generated inspired by the PGD attack. Therefore, exploring the generalization ability of our method encountering other adversarial attack methods (*e.g.*, FGSM [16]) becomes natural. For example, we reformulate the Eq. (3) via FGSM as $\hat{\mathcal{E}}' = \hat{\mathcal{E}} - \alpha \cdot \text{sign}(\nabla_{\mathcal{E}} \mathcal{L}(f'_{\theta}(\hat{\mathcal{E}}), l))$ to induce the model focusing on easy-learned features. Here, we set α to 8/255 and use the same configurations with Eq. (3) to train the surrogate model. In Table 3, E6 shows the unlearnable performance of our UEVs via FGSM on N-Caltech101 dataset. The unlearnable ability of our method is indeed decreased by using FGSM. However, UEVs can still prevent unauthorized models from achieving reliable predictions under the protection provided by UEVs.

Noise projection. To ensure the generated noise can be compatible with event streams, the noise projection has been proposed to process the generated noise. Here, we conduct the ablation study to showcase the necessary of this operation. Note that the polarity of events is normalized into $\{0.0, 0.5, 1.0\}$, where 0.0/1.0 denote the polarity of an event is -1/+1, 0.5 means there is no event. As shown in Fig. 4, there are many noises in the background of the representation frames while removing the projection operation (original noise). Hence, it’s impossible to reconstruct an unlearn-

able event stream from non-projected representations. We propose Eq. (4) to project our noise into $\{-0.5, 0, +0.5\}$. In Eq. (4), τ is introduced to balance the effectiveness and stealthiness of the injected noise in our UEVs. A larger τ can achieve higher stealthiness while corrupting the noise’s unlearnable ability. Minimizing the τ can achieve good effectiveness but the stealthiness would be corrupted.

5. Conclusion

This paper proposes the first unlearnable framework for asynchronous event streams (UEVs), which provides the possibility to protect the event data from being learned by unauthorized models. Specifically, a new form of event error-minimizing noise (E^2MN) is presented to integrate with event data, preventing the model from learning its real semantic features. To ensure the learned noise can be compatible with sparse event data, we propose a projection mechanism to project the E^2MN into event friendly counterpart that facility the generation of unlearnable event streams. Additionally, a retrieval strategy is proposed to reconstruct the unlearnable event streams from their corresponding unlearnable stacks. Extensive ablation studies and analysis experiments conducted on four datasets with seven popular DNN models demonstrate the effectiveness, imperceptibility, and robustness of UEVs.

Acknowledgement. Renjie Group is supported by the National Natural Science Foundation of China under Grant No. 62302415, Guangdong Basic and Applied Basic Research Foundation under Grant No. 2022A1515110692, 2024A1515012822, and the Blue Sky Research Fund of HKBU under Grant No. BSRF/21-22/16. Boxin Shi and Peiqi Duan are supported by National Natural Science Foundation of China (Grant No. 62088102, 62402014), Beijing Natural Science Foundation (Grant No. L233024), and Beijing Municipal Science & Technology Commission, Administrative Commission of Zhongguancun Science Park (Grant No. Z241100003524012).

References

- [1] Arnon Amir, Brian Taba, David Berg, Timothy Melano, Jeffrey McKinstry, Carmelo Di Nolfo, Tapan Nayak, Alexander Andreopoulos, Guillaume Garreau, Marcela Mendoza, et al. A low power, fully event-based gesture recognition system. In *Proc. CVPR*, pages 7243–7252, 2017. 1, 2, 5, 4
- [2] Lorenzo Berlincioni, Luca Cultrera, Chiara Albisani, Lisa Cresti, Andrea Leonardo, Sara Picchioni, Federico Becattini, and Alberto Del Bimbo. Neuromorphic event-based facial expression recognition. In *Proc. CVPR*, 2023. 3
- [3] Marco Cannici and Davide Scaramuzza. Mitigating motion blur in neural radiance fields with events and frames. In *Proc. CVPR*, pages 9286–9296, 2024. 2
- [4] Nicholas Carlini and David Wagner. Towards evaluating the robustness of neural networks. In *S&P*, pages 39–57, 2017. 3
- [5] Zhiwen Chen, Zhiyu Zhu, Yifan Zhang, Junhui Hou, Guangming Shi, and Jinjian Wu. Segment any event streams via weighted adaptation of pivotal tokens. In *Proc. CVPR*, pages 3890–3900, 2024. 2
- [6] Yinpeng Dong, Fangzhou Liao, Tianyu Pang, Hang Su, Jun Zhu, Xiaolin Hu, and Jianguo Li. Boosting adversarial attacks with momentum. In *Proc. CVPR*, pages 9185–9193, 2018. 3
- [7] Alexey Dosovitskiy, Lucas Beyer, Alexander Kolesnikov, Dirk Weissenborn, Xiaohua Zhai, Thomas Unterthiner, Mostafa Dehghani, Matthias Minderer, Georg Heigold, Sylvain Gelly, Jakob Uszkoreit, and Neil Houlsby. An image is worth 16x16 words: Transformers for image recognition at scale. In *ICLR*, 2021. 5
- [8] Peiqi Duan, Zihao W Wang, Boxin Shi, Oliver Cossairt, Tiejun Huang, and Aggelos K Katsaggelos. Guided event filtering: Synergy between intensity images and neuromorphic events for high performance imaging. *TPAMI*, 44(11): 8261–8275, 2021. 2
- [9] Peiqi Duan, Zihao W Wang, Xinyu Zhou, Yi Ma, and Boxin Shi. EventZoom: Learning to denoise and super resolve neuromorphic events. In *Proc. CVPR*, pages 12824–12833, 2021. 2, 5, 3
- [10] Peiqi Duan, Boyu Li, Yixin Yang, Hanyue Lou, Minggui Teng, Xinyu Zhou, Yi Ma, and Boxin Shi. Eventaid: Benchmarking event-aided image/video enhancement algorithms with real-captured hybrid dataset. *TPAMI*, 2025. 2
- [11] Li Fei-Fei, Rob Fergus, and Pietro Perona. Learning generative visual models from few training examples: An incremental bayesian approach tested on 101 object categories. In *Proc. CVPRW*, pages 178–178, 2004. 4
- [12] Shaopeng Fu, Fengxiang He, Yang Liu, Li Shen, and Dacheng Tao. Robust unlearnable examples: Protecting data privacy against adversarial learning. In *ICLR*, 2022. 2, 3
- [13] Guillermo Gallego, Tobi Delbrück, Garrick Orchard, Chiara Bartolozzi, Brian Taba, Andrea Censi, Stefan Leutenegger, Andrew J Davison, Jörg Conradt, Kostas Daniilidis, et al. Event-based vision: A survey. *TPAMI*, 44(1):154–180, 2020. 2, 3
- [14] Daniel Gehrig, Antonio Loquercio, Konstantinos G Derpanis, and Davide Scaramuzza. End-to-end learning of representations for asynchronous event-based data. In *Proc. ICCV*, pages 5633–5643, 2019. 2
- [15] Mathias Gehrig, Willem Aarents, Daniel Gehrig, and Davide Scaramuzza. Dsec: A stereo event camera dataset for driving scenarios. *RAL*, 6(3):4947–4954, 2021. 1, 2
- [16] Ian J Goodfellow, Jonathon Shlens, and Christian Szegedy. Explaining and harnessing adversarial examples. In *ICLR*, 2015. 8, 3
- [17] Fuqiang Gu, Weicong Sng, Xuke Hu, and Fangwen Yu. Eventdrop: Data augmentation for event-based learning. In *Proc. IJCAI*, 2021. 8, 3
- [18] Friedhelm Hamann, Suman Ghosh, Ignacio Juarez Martinez, Tom Hart, Alex Kacelnik, and Guillermo Gallego. Low-power continuous remote behavioral localization with event cameras. In *Proc. CVPR*, pages 18612–18621, 2024. 2
- [19] Hao He, Kaiwen Zha, and Dina Katabi. Indiscriminate poisoning attacks on unsupervised contrastive learning. In *ICLR*, 2023. 3
- [20] Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. Deep residual learning for image recognition. In *Proc. CVPR*, pages 770–778, 2016. 3, 5
- [21] Gao Huang, Zhuang Liu, Laurens Van Der Maaten, and Kilian Q Weinberger. Densely connected convolutional networks. In *Proc. CVPR*, pages 4700–4708, 2017. 5
- [22] Hanxun Huang, Xingjun Ma, Sarah Monazam Erfani, James Bailey, and Yisen Wang. Unlearnable examples: Making personal data unexploitable. In *ICLR*, 2021. 1, 2, 5, 3
- [23] Ze Huang, Li Sun, Cheng Zhao, Song Li, and Songzhi Su. Eventpoint: Self-supervised interest point detection and description for event-based camera. In *Proc. WACV*, pages 5396–5405, 2023. 2
- [24] Jianping Jiang, Jiahe Li, Baowen Zhang, Xiaoming Deng, and Boxin Shi. Evhandpose: Event-based 3D hand pose estimation with sparse supervision. *TPAMI*, 2024. 1, 2
- [25] Jianping Jiang, Xinyu Zhou, Bingxuan Wang, Xiaoming Deng, Chao Xu, and Boxin Shi. Complementing event streams and rgb frames for hand mesh reconstruction. In *Proc. CVPR*, pages 24944–24954, 2024. 2
- [26] Wan Jiang, Yunfeng Diao, He Wang, Jianxin Sun, Meng Wang, and Richang Hong. Unlearnable examples give a false sense of security: Piercing through unexploitable data with learnable examples. In *Pro. ACM MM*, pages 8910–8921, 2023. 2
- [27] Junho Kim, Jaehyeok Bae, Gangin Park, Dongsu Zhang, and Young Min Kim. N-imagenet: Towards robust, fine-grained object recognition with event cameras. In *Proc. ICCV*, pages 2146–2156, 2021. 1, 2, 5, 4
- [28] Taewoo Kim, Hoonhee Cho, and Kuk-Jin Yoon. Frequency-aware event-based video deblurring for real-world motion blur. In *Proc. CVPR*, pages 24966–24976, 2024. 2
- [29] Alex Krizhevsky, Geoffrey Hinton, et al. Learning multiple layers of features from tiny images. 2009. 4
- [30] Xavier Lagorce, Garrick Orchard, Francesco Galluppi, Bertram E Shi, and Ryad B Benosman. Hots: a hierarchy of event-based time-surfaces for pattern recognition. *TPAMI*, 39(7):1346–1359, 2016. 2

- [31] Hongmin Li, Hanchao Liu, Xiangyang Ji, Guoqi Li, and Luping Shi. Cifar10-dvs: an event-stream dataset for object classification. *Front. Neurosci.*, 11:309, 2017. 2, 5, 4
- [32] Haotian Liu, Guang Chen, Sanqing Qu, Yanping Zhang, Zhi-jun Li, Alois Knoll, and Changjun Jiang. Tma: Temporal motion aggregation for event-based optical flow. In *Proc. ICCV*, pages 9685–9694, 2023. 2
- [33] Haoyue Liu, Shihan Peng, Lin Zhu, Yi Chang, Hanyu Zhou, and Luxin Yan. Seeing motion at nighttime with an event camera. In *Proc. CVPR*, pages 25648–25658, 2024. 2
- [34] Xinwei Liu, Xiaojun Jia, Yuan Xun, Siyuan Liang, and Xiaochun Cao. Multimodal unlearnable examples: Protecting data against multimodal contrastive learning. In *Proc. ACM MM*, pages 8024–8033, 2024. 3
- [35] Yixin Liu, Kaidi Xu, Xun Chen, and Lichao Sun. Stable unlearnable example: Enhancing the robustness of unlearnable examples via stable error-minimizing noise. In *Proc. AAAI*, pages 3783–3791, 2024. 3
- [36] Ze Liu, Yutong Lin, Yue Cao, Han Hu, Yixuan Wei, Zheng Zhang, Stephen Lin, and Baining Guo. Swin transformer: Hierarchical vision transformer using shifted windows. In *Proc. ICCV*, pages 10012–10022, 2021. 5
- [37] Qi Ma, Danda Pani Paudel, Ajad Chhatkuli, and Luc Van Gool. Deformable neural radiance fields using rgb and event cameras. In *Proc. ICCV*, pages 3590–3600, 2023. 2
- [38] Aleksander Madry, Aleksandar Makelov, Ludwig Schmidt, Dimitris Tsipras, and Adrian Vladu. Towards deep learning models resistant to adversarial attacks. In *ICLR*, 2018. 4, 3
- [39] Ana I Maqueda, Antonio Loquercio, Guillermo Gallego, Narciso García, and Davide Scaramuzza. Event-based vision meets deep learning on steering prediction for self-driving cars. In *Proc. CVPR*, pages 5419–5427, 2018. 1
- [40] Ruohan Meng, Chenyu Yi, Yi Yu, Siyuan Yang, Bingquan Shen, and Alex C Kot. Semantic deep hiding for robust unlearnable examples. *TIFS*, 2024. 3
- [41] Christen Millerdurai, Hiroyasu Akada, Jian Wang, Diogo Luvizon, Christian Theobalt, and Vladislav Golyanik. EventEgo3D: 3D human motion capture from egocentric event streams. In *Proc. CVPR*, pages 1186–1195, 2024. 2
- [42] Garrick Orchard, Ajinkya Jayawant, Gregory K Cohen, and Nitish Thakor. Converting static image datasets to spiking neuromorphic datasets using saccades. *Front. Neurosci.*, 9: 437, 2015. 2, 5, 6, 4
- [43] Henri Rebecq, Timo Horstschaefer, and Davide Scaramuzza. Real-time visual-inertial odometry for event cameras using keyframe-based nonlinear optimization. In *Proc. BMVC*, pages 16–1, 2017. 3
- [44] Henri Rebecq, René Ranftl, Vladlen Koltun, and Davide Scaramuzza. Events-to-video: Bringing modern computer vision to event cameras. In *Proc. CVPR*, pages 3857–3866, 2019. 4
- [45] Henri Rebecq, René Ranftl, Vladlen Koltun, and Davide Scaramuzza. High speed and high dynamic range video with an event camera. *TPAMI*, 43(6):1964–1980, 2019. 3
- [46] Jie Ren, Han Xu, Yuxuan Wan, Xingjun Ma, Lichao Sun, and Jiliang Tang. Transferable unlearnable examples. In *ICLR*, 2023. 2, 3
- [47] Viktor Rudnev, Vladislav Golyanik, Jiayi Wang, Hans-Peter Seidel, Franziska Mueller, Mohamed Elgharib, and Christian Theobalt. Eventhands: Real-time neural 3D hand pose estimation from an event stream. In *Proc. ICCV*, pages 12385–12395, 2021. 1
- [48] Shawn Shan, Emily Wenger, Jiayun Zhang, Huiying Li, Haitao Zheng, and Ben Y Zhao. Fawkes: Protecting privacy against unauthorized deep learning models. In *USENIX Security*, pages 1589–1604, 2020. 2
- [49] K Simonyan and A Zisserman. Very deep convolutional networks for large-scale image recognition. In *ICLR*, 2015. 3, 5
- [50] Amos Sironi, Manuele Brambilla, Nicolas Bourdis, Xavier Lagorce, and Ryad Benosman. HATS: Histograms of averaged time surfaces for robust event-based object classification. In *Proc. CVPR*, pages 1731–1740, 2018. 1, 2, 3
- [51] Ye Sun, Hao Zhang, Tiehua Zhang, Xingjun Ma, and Yungang Jiang. Unseg: One universal unlearnable example generator is enough against all image segmentation. In *NeurIPS*, 2024. 3
- [52] Zhaoning Sun, Nico Messikommer, Daniel Gehrig, and Davide Scaramuzza. Ess: Learning event-based semantic segmentation from still images. In *Proc. ECCV*, pages 341–357, 2022. 2
- [53] Mingxing Tan and Quoc Le. EfficientNet: Rethinking model scaling for convolutional neural networks. In *ICML*, 2019. 3, 5
- [54] Ajay Vasudevan, Pablo Negri, Bernabe Linares-Barranco, and Teresa Serrano-Gotarredona. Introduction and analysis of an event-based sign language dataset. In *FG 2020*, pages 675–682, 2020. 1
- [55] Ruofei Wang, Qing Guo, Haoliang Li, and Renjie Wan. Event trojan: Asynchronous event-based backdoor attacks. In *Proc. ECCV*, pages 315–332, 2024. 6, 3
- [56] Xiao Wang, Shiao Wang, Chuanming Tang, Lin Zhu, Bo Jiang, Yonghong Tian, and Jin Tang. Event stream-based visual object tracking: A high-resolution benchmark dataset and a novel baseline. In *Proc. CVPR*, pages 19248–19257, 2024. 2
- [57] Jingwen Ye and Xinchao Wang. Ungeneralizable examples. In *Proc. CVPR*, pages 11944–11953, 2024. 2
- [58] Pei Zhang, Shuo Zhu, and Edmund Y Lam. Event encryption: Rethinking privacy exposure for neuromorphic imaging. *NCE*, 4(1):014002, 2024. 1
- [59] Pengyu Zhang, Hao Yin, Zeren Wang, Wenyue Chen, Shengming Li, Dong Wang, Huchuan Lu, and Xu Jia. Evsign: Sign language recognition and translation with streaming events. In *Proc. ECCV*, pages 335–351, 2025. 1
- [60] Xian Zhang, Yong Wang, Qing Yang, Yiran Shen, and Hongkai Wen. Ev-perturb: event-stream perturbation for privacy-preserving classification with dynamic vision sensors. *MULTIMED TOOLS APPL*, 83(6):16823–16847, 2024. 1
- [61] Nikola Zubić, Daniel Gehrig, Mathias Gehrig, and Davide Scaramuzza. From chaos comes order: Ordering event representations for object recognition and detection. In *Proc. ICCV*, pages 12846–12856, 2023. 2

Asynchronous Event Error-Minimizing Noise for Safeguarding Event Dataset

Supplementary Material

Ruofei Wang¹ Peiqi Duan^{2,3} Boxin Shi^{2,3} Renjie Wan^{1*}

¹Department of Computer Science, Hong Kong Baptist University

²State Key Laboratory for Multimedia Information Processing, School of Computer Science, Peking University

³National Engineering Research Center of Visual Technology, School of Computer Science, Peking University
ruofei@life.hkbu.edu.hk, {duanqi0001, shiboxin}@pku.edu.cn, renjiewan@hkbu.edu.hk

6. Overview

- Sec. 7 discusses our **application scenario** again and shows the **difference** between image baselines and our unlearnable event streams.
- Sec. 8 illustrates the detailed explanation of our **algorithm**.
- Sec. 9 shows more details of our E²MN and the corresponding **mixed** counterparts.
- Sec. 10 shows the exploration about the proposed **projection** strategy.
- Sec. 11 illustrates the details of five kinds of event pollution operations used in our experiments.
- Sec. 12 add more experiments on **event representations**, **time bins**, **adversarial attack strategies**, **naive baselines**, *etc.*
- Sec. 13 lists the **dataset** details for N-Caltech101, CIFAR10-DVS, DVS128 Gesture, and N-ImageNet.
- Sec. 14 depicts more **visualization** results to evaluate the imperceptibility of our E²MN.
- Sec. 15 discusses the **social impact** of our UEVs.
- Sec. 16 lists our **future work**, including transferability evaluation, generation efficiency, and defense mechanism.

7. Our UEVs

We propose UEVs mainly focusing on preventing **unauthorized** event data usage. As shown in Fig. 5, the protector, *i.e.*, data owner, releases the unlearnable dataset for users, while only the authorized models can learn real semantic features from these data. The hacker’s unauthorized models are prevented from learning. This mechanism effectively protects the interests of data owners and avoids the privacy leakage caused by data misuse. Fig. 5 shows the working scenario of our UEVs.

Unlearnable Examples (UEs) are proposed to prevent image dataset from unauthorized usage. Compared with UEs, UEVs show great challenges. 1) The image perturbation is directly optimized by deep models to ensure its effectiveness and imperceptibility. However, the event data cannot be directly input into deep models to generate the unlearnable version. 2) Event data shows the binary polar-

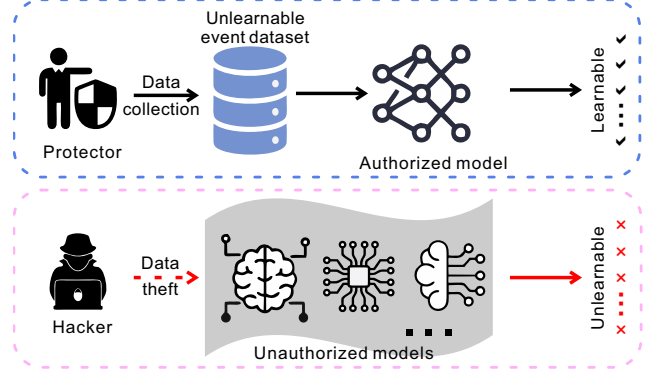


Figure 5. Our scenario of the unlearnable event streams (UEVs). For authorized training, the UEVs can be effectively used to train downstream models and achieve correct predictions. However, if hackers train their authorized networks without our authority that they cannot achieve the reliable performance.

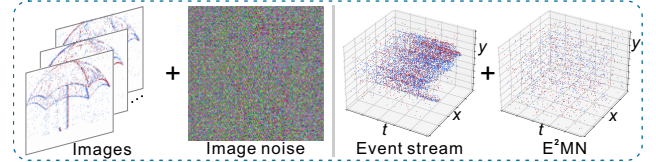


Figure 6. Comparison between the image noise and our E²MN (sample-wise noise).

ity and asynchronous nature that hinders the noise injection. This sparse property denotes that although we can represent the event stream via image-like features, the generated noise cannot be compatible with event data. 3) Different from image perturbation, event noise should include time stamps that ensure the unlearnable noise is closely aligned with the event data, enhancing imperceptibility.

Although UEVs and UEs adopt the same core idea: error minimizing loss function, to generate unlearnable noise, using UEs to protect event data is impossible, as shown in Fig. 6. The noise is only injected into event representations while the original event data is not protected. Our UEVs perturb the original event data that shows better practicality.

8. UEVs algorithm

The algorithm of our UEVs is shown in Algorithm 1. Here, more detailed explanations about our algorithm are provided. To generate unlearnable event streams, the protector needs to provide the surrogate model f' , target dataset $\mathcal{D}_c$, training step M , and classification accuracy γ . f' is employed to calculate the event error-minimizing noise on the dataset $\mathcal{D}_c$. M denotes the training iteration of f' in every epoch of noise generation, which is limited since more efforts should be attached on the δ optimization in Eq. (2). Specifically, the training process of f' is listed in Lines#3-6. We randomly sample M batches of event streams from $\mathcal{D}_c$ and incorporate with the noise (sample-wise noise or class-wise noise) to train the surrogate model. $\mathcal{L}^*$ is our loss function (Eq. (5)), consisting of the cosine similarity loss function and cross-entropy loss function. The final loss is calculated as:

$$\text{loss} = \lambda_1 \left(\left[1 + \frac{f'(\mathcal{R}(\mathcal{E}))_{\text{conv}} \cdot f'(\mathcal{R}(\mathcal{E}) + \delta)_{\text{conv}}}{\|f'(\mathcal{R}(\mathcal{E}))_{\text{conv}}\| \times \|f'(\mathcal{R}(\mathcal{E}) + \delta)_{\text{conv}}\|} \right] / 2 \right) + \lambda_2 \left(-[l_i \log f'(\mathcal{R}(\mathcal{E}) + \delta) + (1 - l_i) \log(1 - (f'(\mathcal{R}(\mathcal{E}) + \delta)))] \right), \quad (6)$$

where $(\cdot)_{\text{conv}}$ denotes the convolution features extracted by the last convolution layer of surrogate model f' , B indicates the batch size, $\mathcal{R}$ means event representation, converting an event stream into the event stack. Eq. (6) illustrates the training pipeline of f' on a batch of event streams. It is not required to calculate the cost of $f'(\mathcal{R}(\mathcal{E}))$ because this would cause the surrogate model to focus on learning the real semantic features, thereby preventing the optimization of our unlearnable noise.

After training, we generate the event error-minimizing noise for entire $\mathcal{D}_c$ according to Eq. (3), as shown in Lines#7-10. $\text{Clip}(\cdot)$ denotes clipping those noise that exceed $-\epsilon$ or $+\epsilon$ back to this region. The noise generation and surrogate model training would be terminated once the classification accuracy tested under δ is higher than γ . This termination demonstrates that the generated noise can effectively guide the model to conduct predictions without relying on image semantics. Therefore, the noise δ is able to prevent the unauthorized model from learning informative knowledge from our data.

Due to the special characteristic of event data, our noise δ is generated based on the event stack. It's necessary to conduct event reconstruction to generate the unlearnable event stream from its corresponding unlearnable event stack (Lines#13-17). To ensure the δ can be compatible with event data, we propose a projection strategy $\mathbf{P}(\cdot)$ to sparsify the noise into $\{-0.5, 0, +0.5\}^2$. Then, we integrate the projected noise with an event stack and clip it

²In image area, the generated noise is directly added on the images, rendering the unlearnable examples via modifying the pixel values. However, for event data, which consists solely of binary events, we can only trans-

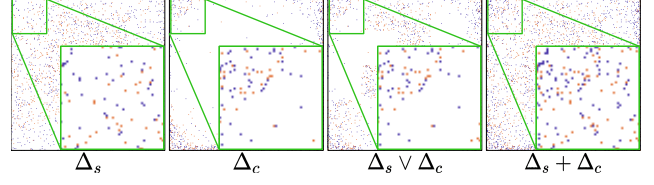


Figure 7. Illustrative examples of our E^2MN .

into $[0, 1]$ to generate the unlearnable event stack. Detailed noise embedding process is shown in Table 4. Finally, a retrieval strategy $\mathcal{R}'$ is proposed to search the compressed time stamps from the original event streams to achieve the reconstruction. Based on this algorithm, our valuable event data can be protected well that prevents unauthorized data exploitation, as shown in Fig. 5.

9. Class-wise and sample-wise noise

Our E^2MN consists of two kinds of noise: class-wise noise and sample-wise noise, which are all generated based on Eq. (2). The sample-wise noise is generated case by case, which leads to every generated noise being only workable for the single event stream. This limits the practicality of sample-wise noise, especially in the dataset scale grows or new event streams are captured. Hence, an alternative way is class-wise noise, which is generated class by class. It means that a kind of noise can be injected into different event streams sampled from the sample class. Another advantage is that class-wise noise consumes less memory than sample-wise noise in noise optimization. Although two kinds of noises achieve similar performance in Table 2, class-wise noise achieves better stealthiness than sample-wise noise as shown in Fig. 3 and Table 1.

Considering their respective advantages, we combine the two noises to explore whether it can bring more benefits. We have proposed union and addition operations in Sec. 4.3 to evaluate. For $\Delta_s \vee \Delta_c$, we randomly choose Δ_s or Δ_c to protect the event. The form of this kind of noise resembles both of them because only a simple random sampling operation is employed. For $\Delta_s + \Delta_c$, we fuse two kinds of noise by element-wise addition to perturb the event stream. As shown in Fig. 7, this configuration increases the number of perturbations introduced into the event data, resulting in higher effectiveness (see E4 of Table 3).

10. Projection discussion

Our projection strategy is designed to sparsify the noise δ to ensure compatibility with event stacks. We illustrate a detailed confusion matrix of the noise embedding in Ta-

form the event stream into an unlearnable one on the event level via deleting events or inserting new events. Therefore, we define the projected values as $\{-0.5, 0, +0.5\}$ to be compatible with event stacks ($\{0, 0.5, 1.0\}$).

Table 4. Confusion matrix of embedding the noise (E^2MN) into an event stack (E. stack). The final unlearnable event stack would be clipped into $[0, 1]$.

		δ		
		-0.5	0	+0.5
E. stack	0	original event	original event	event deletion
	0.5	event generation	no event	event generation
	1.0	event deletion	original event	original event

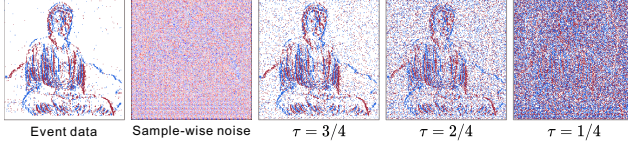


Figure 8. Visualization of event data, sample-wise noise, and the corresponding unlearnable event streams with different projection parameters τ .

ble 4. If $\delta = -0.5$ is added to a negative event (0), the event keeps its original value. If added to a positive event (1), the event is deleted. A new event with negative polarity is created when adding -0.5 to the pixel value 0.5 sampled from the event stack. In this section, we showcase the importance of the parameter τ in Eq. (4). As discussed in [22], the added noise should be imperceptible to human eyes and does not affect the normal data utility. Hence, we introduce the parameter τ to balance the imperceptibility and unlearnability of our E^2MN . As shown in Fig. 8, the larger τ can lead to better imperceptibility but the less unlearnable noise has been introduced, which harms the unlearnability. We have tested our sample-wise noise with $\tau = 7/8$ on the N-Caltech101 dataset and obtained the accuracy of ResNet18 by 4.88, which is higher than the accuracy (0.52) tested by $\tau = 3/4$. This demonstrates the great challenges in balancing the imperceptibility and unlearnability of our E^2MN .

11. Baseline setting

To further evaluate the effectiveness of our UEVs, we introduce the straightforward event distortions as our baselines, which are inspired from event data augmentation [17] and backdoor attacks [55]. Data augmentation is proposed to enrich the training data for improving the model’s performance, which usually employs data distortion operations to augment the sample. Generally, the quality of these augmented samples is lower than the original ones. Therefore, we propose simple coordinate shifting (CS), timestamp shifting (CS), polarity inversion (PI), and area shuffling (AS) based on [17] to corrupt the event streams for preventing unauthorized usage. Additionally, we also propose the manual pattern (MP) based on backdoor attacks [55] to perturb our event streams. We inject a pre-defined pat-

tern for those event streams sampled from the same class to prevent unauthorized data usage, which can be viewed as a class-wise noise. According to Table 2, we can find that compromising the quality of our event datasets can degrade the performance of downstream models. However, the unlearnability is rather limited and does not prevent the downstream models from learning informative knowledge.

12. Additional experiments

Various event representations. In our main experiments, we adopt the voxel-grid event stack as our event representation to evaluate the effectiveness. To test the generalizability of our UEVs among different representations, the event frame (EF) [43] and Time surface (TS) [50] are adopted. As shown in the E1 of Table 5, our UEVs can still prevent unauthorized event data usage under EF and TS representations, showing high robustness and generalizability.

Generalizability. According to [9], we set the time bin to 16 to represent the event stream. To evaluate the generalizability of our UEVs on different time intervals, we change the size of Δ_t to $0.5\times$ and $2\times$ to conduct ablation studies. as shown in E2 of Table 5, our UEVs still shows high protection ability to prevent the unauthorized event data usage.

Diverse Adversarial attacks. Apart from the PGD [38] and FGSM [16] attacks, we add new adversarial attack methods: C&W [4] and MIFGSM (MIF) [6]. CW attack is an optimization-based method that crafts minimal perturbations to mislead neural networks while remaining imperceptible. MIFGSM is an iterative adversarial attack that enhances the basic FGSM by incorporating the momentum. Compared with MIFGSM, CW attack achieves better unlearnable functionality. As shown in E3 of Table 5, our method can still achieve reliable unlearnable performance while adopting different adversarial attacks.

Table 5. Quantitative results tested by Res50 on N-Caltech101.

	E1		E2		E3		E4
	EF	TS	$0.5\Delta_t$	$2\Delta_t$	C&W	MIF	UEs
Δ_c	5.17	10.09	5.46	4.94	8.73	15.43	5.51
Δ_s	8.85	16.48	5.17	5.05	12.15	14.47	18.38

Naive image baselines. Image-based methods are unable to directly secure event data due to data differences, which can only safeguard the corresponding event representation. In E4 of Table 5, although the image approach, UEs [22], performs well in event representations, it cannot prevent malicious users from misusing the **original event**.

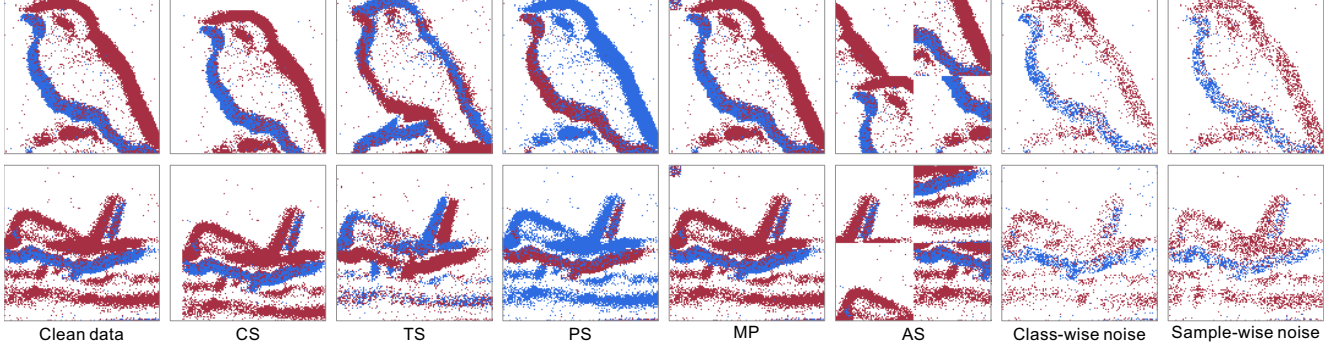


Figure 9. Visualization results of various noise forms on the CIFAR10-DVS dataset [31]. Blue/Red points denote the events with $p = +1/-1$. Our noise E^2MN does not introduce much noise in the background region, which maintains good imperceptibility.

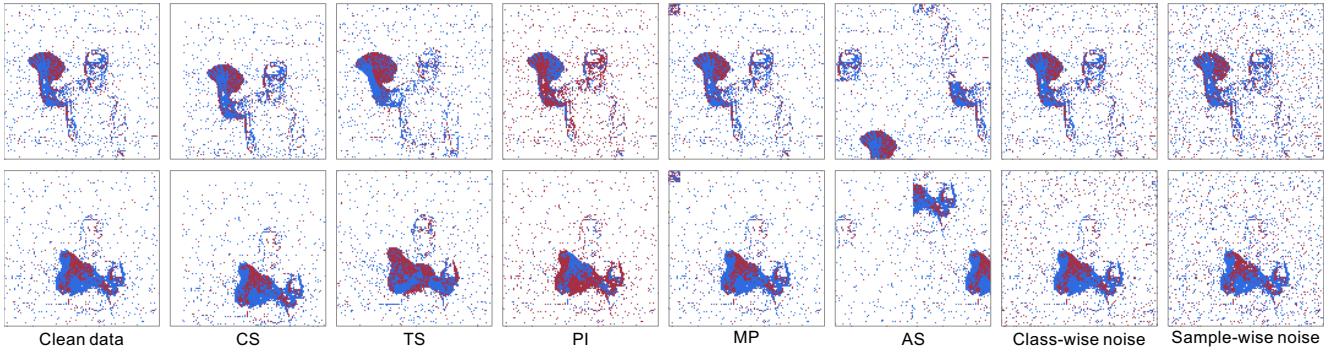


Figure 10. Visualization results of various noise forms on the DVS128 Gesture dataset [1]. Blue/Red points denote the events with $p = +1/-1$. Our noise E^2MN does not corrupt the target objects, and the noise distributed in the background region appears quite realistic.

Event to image reconstruction. We employ an event-to-image method (E2VID [44]) to evaluate the generazibility. As shown in Fig. 11, our method prevents E2VID [44] from reconstructing details from the protected event data, thereby providing solutions for privacy preserving.



Figure 11. Visualization frames reconstructed by E2VID [44].

Unlearnable cluster. We extend our class-wise noise into a cluster-wise one. We ❶ employ K-Means+ResNet50 to cluster the N-Caltech101 into 10 classes; ❷ train a surrogate model on 10 classes to generate **cluster-wise noise**; ❸ train ResNet18 on the whole classes (101) with cluster-wise noise. The cluster version of our method can reduce the classification accuracy from 0.787 to 0.189.

13. Dataset details

To evaluate the effectiveness of our method, we employ four popular event-based datasets in our experiments, includ-

ing N-Caltech101 [42], CIFAR10-DVS [31], DVS128 Gesture [1], and N-ImageNet [27]. N-Caltech101 is the neuromorphic version of the image dataset, Caltech101 [11], which has 101 classes and 4356, 2612, and 1741 samples for training, validation, and testing, respectively. CIFAR10-DVS is generated based on image datasets CIFAR-10 [29], where the training set, validation set and testing set contain 7000, 1000, and 2000 samples, respectively. DVS128 Gesture contains 11 classes from 29 subjects under 3 illumination conditions, which has 1176 training samples and 288 testing samples. The N-ImageNet (mini) dataset is derived from the ImageNet dataset. It utilizes an event camera to capture RGB images shown on a monitor. This dataset includes 100 object classes, with each class having 1,300 streams for training and 50 streams for validation. Details of each dataset are shown in Table 6.

14. Visualization comparison

To evaluate the imperceptibility of our E^2MN , we show more visualization results in Fig. 9 and Fig. 10. In Fig. 9, we sample event streams from the CIFAR10-DVS dataset [31] to generate the unlearnable ones via five straightforward distortions and our two kinds of noise. It's clear that there

Table 6. Details of four used datasets in our experiments.

	N-Caltech101	CIFAR10-DVS	DVS128 Gesture	N-ImageNet
Type	Simulated	Simulated	Real	Simulated
Calsses	101	10	11	1000
Resolution	180 × 240	128 × 128	128 × 128	480 × 640
Event Camera	ATIS camera	DVS128	DVS128	Samsung DVS Gen3
Train	4536	7000	1176	130000
Val	2612	1000	288	50000
Test	1741	2000	288	50000

is an x - y offset in the unlearnable event streams generated by CS compared to the clean data. TS alters time stamps to pollute the input event streams, resulting in a noticeable difference between the distorted samples and the clean data. PI reverses the polarity of the event stream to degrade the quality, thereby reducing the quality of the training data. AS rearranges the event data at the block level to produce unlearnable data, which exhibits low imperceptibility. Our class-wise noise and sample-wise noise perturb the event stream with imperceptible noise that shows better invisibility than other comparison methods. The visualization results sampled from DVS128 Gesture dataset [1] are shown in Fig. 10. It’s clear that our E²MN achieves the best unlearnability while maintaining good imperceptibility.

To visualize the influence caused by our E²MN, we employ GradCAM to highlight several unlearnable event streams in Fig. 12. Figures (A), (B), (C), (D), and (E) are rendered by A-shuffle, M-pattern, P-inverse, Class-wise noise, and sample-wise noise, respectively. Our method builds a shortcut between the input samples and labels that suppresses the model learning semantic features, resulting in a lower response on the foreground regions.

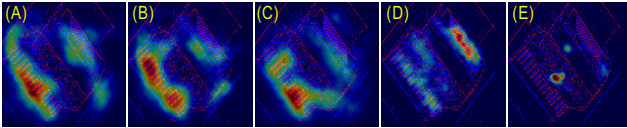


Figure 12. GradCAM of event distortions and our methods.

15. Social impact

The social impact of UEVs is multifaceted, addressing key issues around event data security, privacy, and ethics. By making event datasets unlearnable, our method helps protect individuals’ data from being used without authorization. This is particularly important in an era where event data privacy concerns are paramount, and unauthorized data usage can lead to significant privacy breaches. The method provides a robust protection way for event data owners, ensuring that their event streams cannot be exploited by unauthorized entities. This fosters greater trust in event data sharing. With the application of UEVs, there is a push towards more ethical event data practices. Entities will need to obtain proper authorization and consent before us-

ing event data, promoting a culture of respect for data ownership and user rights.

Overall, our UEVs contributes significantly to the advancement of secure and trustworthy data sharing, promoting a safer and more ethical event data ecosystem.

16. Future work

UEVs is the first method designed for generating unlearnable event streams, which provides a possible solution to prevent the unauthorized usage of our valuable event data. We mainly focus on studying the unlearnability of event streams in the main paper, while causing some limitations in terms of transferability evaluation, generation efficiency, and defense mechanism. To address these issues beyond our research topic in this work, we will explore the following directions in the future:

- **Transferability evaluation:** We plan to extend the evaluation of our method from classification task to other event datasets and event vision tasks, enhancing the transferability. This comprehensive testing is helpful to demonstrate the effectiveness of UEVs that promote the trustworthy event data sharing. However, the effectiveness of UEVs may be decreased across different vision tasks. A possible solution is that adopt the foundation model as our surrogate model to calculate the event error-minimizing noise, which is trained with a large amount of data that has strong generalization capabilities. It can be applied to a variety of tasks without training independent models for each specific task.
- **Generation efficiency:** According to our experiments, we find that the efficiency of the sample-wise noise generation depends on the scale of the event dataset. The larger the scale of the event dataset, the lower the efficiency of the sample-wise noise generation. We propose addressing this issue via a noise generator. We train this generator with the surrogate model jointly, which aims to enable the generated noise to minimize the cost of the surrogate model. This training pipeline avoids storing the generated medium noise, which saves efficiency significantly. Once the optimization has been finished, we can employ this generator to generate sample-wise noise for each input sample with high efficiency.
- **Defense mechanism:** It’s crucial to investigate potential defense mechanisms against the generation of malicious unlearnable event streams. If we release our unlearnable dataset online, hackers could manipulate these samples to force their models to learn the information. By understanding possible defense mechanisms, we can develop more reliable methods for creating unlearnable event datasets, thereby preventing unauthorized usage. Furthermore, elucidating these defense mechanisms can help users improve the dataset quality, ultimately saving training time and computing resources.