



Glaze: Protecting Artists from Style Mimicry by Text-to-Image Models

Shawn Shan, Jenna Cryan, Emily Wenger, Haitao Zheng,
Rana Hanocka, and Ben Y. Zhao, *University of Chicago*

<https://www.usenix.org/conference/usenixsecurity23/presentation/shan>

This paper is included in the Proceedings of the
32nd USENIX Security Symposium.

August 9–11, 2023 • Anaheim, CA, USA

978-1-939133-37-3

Open access to the Proceedings of the
32nd USENIX Security Symposium
is sponsored by USENIX.

Glaze: Protecting Artists from Style Mimicry by Text-to-Image Models

Shawn Shan, Jenna Cryan, Emily Wenger, Haitao Zheng, Rana Hanocka, Ben Y. Zhao

Department of Computer Science, University of Chicago

{shawnshan, jennacryan, ewillson, htzheng, ranahanocka, ravenben}@cs.uchicago.edu

Abstract

Recent text-to-image diffusion models such as MidJourney and Stable Diffusion threaten to displace many in the professional artist community. In particular, models can learn to mimic the artistic style of specific artists after “fine-tuning” on samples of their art. In this paper, we describe the design, implementation and evaluation of *Glaze*, a tool that enables artists to apply “style cloaks” to their art before sharing online. These cloaks apply barely perceptible perturbations to images, and when used as training data, mislead generative models that try to mimic a specific artist. In coordination with the professional artist community, we deploy user studies to more than 1000 artists, assessing their views of AI art, as well as the efficacy of our tool, its usability and tolerability of perturbations, and robustness across different scenarios and against adaptive countermeasures. Both surveyed artists and empirical CLIP-based scores show that even at low perturbation levels ($p=0.05$), *Glaze* is highly successful at disrupting mimicry under normal conditions ($>92\%$) and against adaptive countermeasures ($>85\%$).

1 Introduction

It is not an exaggeration to say that the arrival of text-to-image generator models has transformed, perhaps upended, the art industry. By sending simple text prompts like “A picture of a corgi on the moon” to diffusion models such as StableDiffusion or MidJourney, anyone can generate incredibly detailed, high resolution artwork that previously required many hours of work by professional artists. AI-art such as those in Figure 1 have won awards at established art conventions [69], served as cover images for magazines [47], and used to illustrate children’s books [61] and video games [77]. More powerful models continue to arrive [34, 57, 83], catalyzed by VC funding [59, 97, 98], technical research breakthroughs [4, 11, 37, 46, 52], and powered at their core by continuous training on a large volume of human-made art scraped from online art repositories such as ArtStation, Pinterest and DeviantArt.

Only months after their arrival, these models are rapidly growing in users and platforms. In September 2022, MidJourney reported over 2.7 million users and 275K AI art images generated *each day* [31]. Beyond simple prompts,

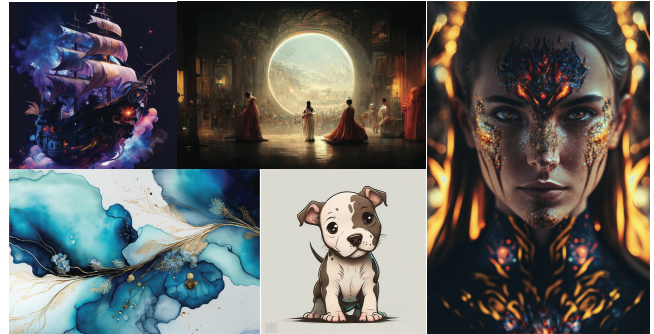


Figure 1. Sample AI-generated art pieces from the MidJourney community showcase [53, 69].

many have taken the open sourced StableDiffusion model, and “fine-tuned” it on additional samples from specific artists, allowing them to generate AI art that *mimics* the specific artistic styles of that artist [32]. In fact, entire platforms have sprung up where home users are posting and sharing their own customized diffusion models that specialize on mimicking specific artists, likeness of celebrities, and NSFW themes [14].

Beyond open questions of copyrights [6], ethics [27, 60], and consent [21, 26, 30], it is clear that these AI models have had significant negative impacts on independent artists. For the estimated hundreds of thousands of independent artists across the globe, most work on commissions, and attract customers by advertising and promoting samples of their artwork online. First, professional artists undergo years of training to develop their individual artistic styles. A model that mimics this style profits from that training without compensating the artist, effectively ending their ability to earn a living. Second, as synthetic art mimicry continues to grow for popular artists, they displace original art in search results, further disrupting the artist’s ability to advertise and promote work to potential customers [32, 75]. Finally, these mimicry attacks are demoralizing art students training to be future artists. Art students see their future careers replaced by AI models even if they can successfully find and develop their own artistic styles [55].

Today, all of these consequences have indeed occurred in the span of a few months. Art students are quitting the field; AI models that mimic specific artists are uploaded and shared for free; and professional artists are losing their livelihoods to models mimicking their unique styles. Artists are fighting

back via lawsuits [20, 35], online boycotts and petitions [22], but legal and regulatory action can take years, and are difficult to enforce internationally. Thus most artists are faced a choice to 1) do nothing, or 2) stop sharing samples of their art online to avoid training models, and in doing so cripple their main way to advertise and promote their work to customers.

In this paper, we present the design, implementation and evaluation of a technical alternative to protect artists against style mimicry by text-to-image diffusion models. We present Glaze, a system that allows an artist to apply carefully computed perturbations to their art, such that diffusion models will learn significantly altered versions of their style, and be ineffective in future attempts at style mimicry. We worked closely with members of the professional artist community to develop Glaze, and conduct multiple user studies with 1,156 participants from the artist community to evaluate its efficacy, usability, and robustness against a variety of active countermeasures.

Intuitively, *Glaze* works by taking a piece of artwork, and computing a minimal perturbation (a “style cloak”) which, when applied, shifts the artwork’s representation in the generator model’s feature space towards a chosen target art style. Training on multiple cloaked images teaches the generator model to shift the artistic style it associates with the artist, leading to mimicry art that fails to match the artist’s true style.

Our work makes several key contributions:

- We engage with top professional artists and the broader community, and conduct user studies to understand their views and concerns towards AI art and the impact on their careers and community.
- We propose *Glaze*, a system that protects artists from style mimicry by adding minimal perturbations to their artwork to mislead AI models to generate art different from the targeted artist. 92% of surveyed artists find the perturbations small enough not to disrupt the value of their art.
- Surveyed artists find that *Glaze* successfully disrupts style mimicry by AI models on protected artwork. 93% of artists rate the protection is successful under a variety of settings, including tests against real-world mimicry platforms.
- In challenging scenarios where an artist has already posted significant artworks online, we show *Glaze* protection remains high. 87.2% of surveyed artists rate the protection as successful when an artist is only able to cloak 1/4 of their online art (75% of art is uncloaked).
- We evaluate *Glaze* and show that it is robust (protection success > 85%) to a variety of adaptive countermeasures.
- We discuss *Glaze* deployment and post-deployment experiences, including countermeasures in the wild.

Ethics. Our user study was reviewed and approved by our institutional review board (IRB). All art samples used in experiments were used with explicit consent by their respective artists. All user study participants were compensated for their time, although many refused payment.

2 Background: AI Art and Style Mimicry

In this section, we provide critical context in the form of basic background on current AI art models and style mimicry.

2.1 Text-to-Image Generation

Since Text-to-image generation was first proposed in 2015 [50], a stream of research has proposed newer model architectures and training methods enabling generation of higher-quality images [45, 63, 99, 103, 107]. The high level design of recent models used for AI art generation [17, 65, 67] is shown in Figure 3. During training, the model takes in an image x and uses a feature extractor Φ to extract its features, producing $\Phi(x)$. Simultaneously, a conditional image generator G takes in a corresponding text caption (s) and outputs a predicted feature vector $G(s)$. Then the parameters of G are optimized so the text feature vector $G(s)$ matches the image feature vector $\Phi(x)$. At generation time, a user gives G a text prompt s_0 , and G outputs an image feature vector $G(s_0)$. A decoder D then decodes $G(s_0)$ to produce the final generated image.

Compared to earlier models based on generative adversarial networks (GANs) or variational autoencoders (VAE) [63, 86, 107], more recent models [66, 67] leveraging *diffusion* models produce significantly higher quality images. Feature extractor (Φ) is used to reduce the dimensionality of the input image to facilitate the generation process. The extractor Φ and decoder D are often a pair of variational autoencoder (VAE) [65, 67], *i.e.*, extractor (encoder) extracts image features and decoder map features back to images.

Training Data Sources. The training datasets of these models typically contain image/ALT text pairs scraped from the Internet. They are extremely large, e.g. LAION [78] contains 5 billion images collected from 3 billion webpages.

These datasets are subject to minimal curation and governance. Data collectors typically only filter out data with extremely short or incorrect text captions (based on an automated text/image alignment metric [78]). Since copyrighted images are not filtered [78], these datasets are rife with private, sensitive content, including copyrighted artworks.

2.2 Style Mimicry

In a *style mimicry* attack, a bad actor uses an AI art model to create art in a particular artist’s style without their consent. More than 67% of art pieces showcased on a popular AI-art-sharing website leverage style mimicry [53].

Style mimicry techniques. Today, a “mimic” can easily copy the style of a victim artist with only an open-source text-to-image model and a few samples of artwork from the artist. A naive mimicry attack directly queries a generic text-to-image model using the name of the victim artist. For example, the prompt “a painting in the style of Greg Rutkowski” would cause the model to generate images in the style of Polish



Figure 2. Real-world incident of AI plagiarizing the style of artist Hollie Mengert [3]. **Left:** original artwork by Hollie Mengert. **Right:** plagiarized artwork generated by a model trained to mimic Hollie’s style.

artist Greg Rutkowski. This is because many of Rutkowski’s artworks appear in training datasets of these generic models labeled with his name.

Naive mimicry can succeed when the artist is well-known and has a significant amount of art online, but fail on other artists. In more recent mimicry attacks, a mimic *fine-tunes* a generic text-to-image model on samples of a target artist’s work (as few as 20 unique pieces) downloaded from online sources. This calibrates the model to the victim artist’s style, identifying important features related to the victim style and associating these regions in the feature space with the victim artist’s name [28, 70]. This enables style mimicry with impressive accuracy. The entire fine-tuning process takes less than 20 minutes on a low-end consumer GPU¹.

Real-work mimicry incidents. The first well-known incident of mimicry was when a Reddit user stole American artist Hollie Mengert’s style and open-sourced the style-specific model on Reddit [3]. Figure 2 has a side-by-side comparison of Hollie’s original artwork and plagiarized artwork generated via style mimicry. Later, famous cartoonist Sarah Andersen reported that AI art models can mimic her cartoon drawings [2], and other similar incidents abound [54, 100].

Several companies [77] have even hosted style mimicry as a service, allowing users to upload a few art pieces painted by victim artists and producing new art in the victim styles. CivitAI [14] built a large online marketplace where people share their customized stable diffusion models, fine-tuned on certain artwork.

3 Collaborating with Artists

Next, we explain our collaborative relationship with professional artists, and its significant impact on our key evaluation metrics in this paper. We also summarize key results from our first user study on views of AI art and mimicry by members of the artist community.

¹It takes an average of 18.3 minutes on a GTX 1080 GPU

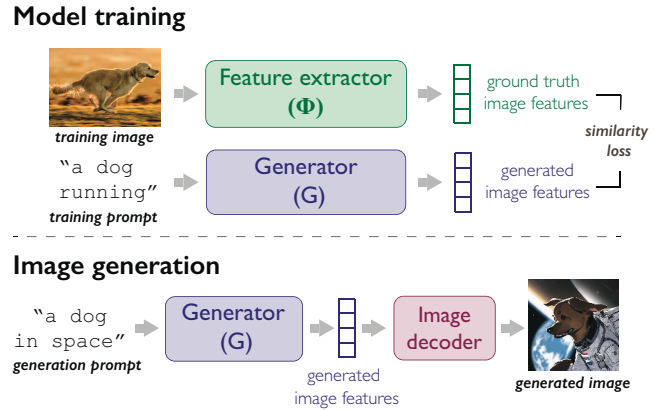


Figure 3. High level model architecture of text-to-image models.

Artists have spoken out against style mimicry in numerous venues, focusing particularly on how it violates their intellectual property rights and threatens their livelihoods [15, 88, 92, 93]. Others have taken direct action. The Concept Art Association raised over \$200K to fight AI art, and filed a class action lawsuit in the US against AI art companies [35]. In November 2022, artists organized a large protest against ArtStation [92], the large digital art sharing platform that allowed users to post AI artwork without identification. Anti-AI images flooded the site for several weeks, until ArtStation banned the protest images [23]. Recently, the Writers Guild of America (WGA) went on strike demanding contractual changes to ban generative AI [12].

Members of the professional art community reached out to us in Sept 2022. We joined online town halls and meetings alongside hundreds of professionals, including Emmy winners and artists at major film studios. After learning more, we began an active collaboration with multiple professional artists, including award-winning artist Karla Ortiz, who leads efforts defending artists and is lead plaintiff in the class action suit. The artists helped this project in multiple ways, by 1) sharing experiences about specific ways AI-art has impacted them and their colleagues; 2) sharing domain knowledge about what is acceptable to artists in terms of perturbations on their art; and 3) helping to widely disseminate our user study to members of their professional organizations, including the Concept Art Association and the Animation Guild (TAG839).

Evaluation via Direct Feedback from Artists. Our goal is to help artists disrupt AI models trying to mimic their artistic style, without adversely impacting their own artwork. Because “success” in this context is highly subjective (“Did this AI-art successfully mimic Karla’s painting style?”), we believe the only reliable evaluation metric is direct feedback by professional artists themselves. Therefore, wherever possible, the evaluation of *Glaze* is done via detailed user studies engaging members of the professional artist community, augmented by an empirical score we develop based on genre prediction using CLIP models.

Survey	# of artists	Content
Survey 1	1156	1) Broad views of AI art and style mimicry (§3.1) 2) Glaze's usability, i.e. acceptable levels of cloaking (§6.3) 3) Glaze performance in disrupting style mimicry (§6.3)
Survey 2 (Extension to Survey 1)	151	1) Additional performance tests (§6.3) 2) Robustness to advanced scenarios (§6.4) and countermeasures (§7) 3) Additional system evaluation (Appendix A)

Table 1. Information on our user studies: the number of artist participants and where we report the results of the studies. We sent Survey 2 to some specific participants from survey 1 who volunteered to participate in a followup study.

We deployed two user studies during the course of this project (see Table 1). Both are IRB-approved by our institution. Both draw participants from professional artists informed via their social circles and professional networks. The first (Survey 1, §3.1, §6.3), asked participants about their broad views of AI style mimicry, and then presented them with a number of inputs and outputs of our tool, and asked them to give ratings corresponding to key metrics we wanted to evaluate. We select a subset of participants from the first study to participate in a longer and more in-depth study (Survey 2) where they were asked to evaluate the performance of *Glaze* in additional settings (§6.3, §6.4, §7, and Appendix A).

3.1 Artists' Opinions on Style Mimicry

While we expected artists to view style mimicry negatively, we wanted to better understand how much individual artists understood this topic and how many perceived it as a threat. Here we describe results from Survey 1 to gather perceptions of the potential impact of AI art on existing artists.

Survey Design. Our survey consisted of both multiple choice and free response questions to understand how well people understand the concept of AI art, and how well the models successfully imitate the style of artists. Additionally, we asked artists about the extent to which they anticipate the emergence of AI art to impact their artistic activities, such as posting their art online and their job security. A handful of professional artists helped disseminate our survey to their respective artist community groups. Overall, we collected responses from 1,207 participants, consisting primarily of professional artists (both full-time (46%) and part-time/freelancer (50%)) and some non-artist members of the art community who felt invested in the impact of AI art (4%). Of the participants who consider themselves artists, their experience varied: <1 year (13%), 1-5 years (49%), 5-10 years (19%), 10+ years (19%). Participants' primary art style varied widely, including: animation, concept art, abstract, anime, game art, digital 2D/3D, illustration, character artwork, storyboarding, traditional painting/drawing, graphic design, and others.

Key Results. Our study found that 91% of the artists have read about AI art extensively, and either know of or worry about their art being used to train the models. Artists expect AI mimicry to have a significant impact on artist community: 97% artists state it will decrease some artists' job security; 88% artists state it will discourage new students from studying

art; and 70% artists state it will diminish creativity. "Junior positions will become extinct," as stated by one participant.

Many artists (> 89% artists) have already or plan to take actions because of AI mimicry. Over 95% of artists post their artwork online. Out of these artists, 53% of them anticipate reducing or removing their online artwork, if they haven't already. Out of these artists, 55% of them believe reducing their online presence will significantly impact their careers. One participant stated "AI art has unmotivated myself from uploading more art and made me think about all the years I spent learning art." 78% of artists anticipate AI mimicry would impact their job security, and this percentage increases to 94% for the job security of newer artists. Further, 24% of artists believe AI art has *already* impacted their job security, and an additional 53% expect to be affected within the next 3 years. Over 51% of artists expressed interest in proactive measures, such as personally joining class action lawsuits against AI companies.

Professional artists thought AI mimicry was very successful at mimicking the style of specific artists. We showed the artists examples of original artwork from 23 artists, and the artwork generated by a model attempting to mimic their styles (detailed mimicry setup in §6). 77% of artists found the AI model *successfully* or *very successfully* mimic the styles of victim artists, with one stating "it's shocking how well AI can mimic the original artwork." Additionally, 19% of participants thought the AI mimicry is somewhat successful, leaving only < 5% of artists rating the mimicry as unsuccessful. Several artists also pointed out that, as artists, upon close inspection they could spot differences between the AI art and originals, but were skeptical the general public would notice them.

A significant concern of most participants, surprisingly, is not just the existence of AI art, but rather scraping of existing artworks without permission or compensation. As one participant stated: "If artists are paid to have their pieces be used and asked permission, and if people had to pay to use that AI software with those pieces in it, I would have no problem." However, without consent to use their artwork to train the models, "it's incredibly disrespectful to the artist to have their work 'eaten' by a machine [after] many years to grow our skills and develop our styles."

4 Preliminaries

We propose *Glaze*, a tool that protects artists against AI style mimicry. An artist uses *Glaze* to add small digital perturbations ("cloak") to images of their own art before sharing online (Figure 5). A text-to-image model that trains on cloaked images of artwork will learn an incorrect representation of the artist's style in feature space *i.e.*, the model's internal understanding of artistic styles. When asked to generate art pieces in victim's style, the model will fail to mimic the style of the victim, and instead output art pieces in a recognizably different style.

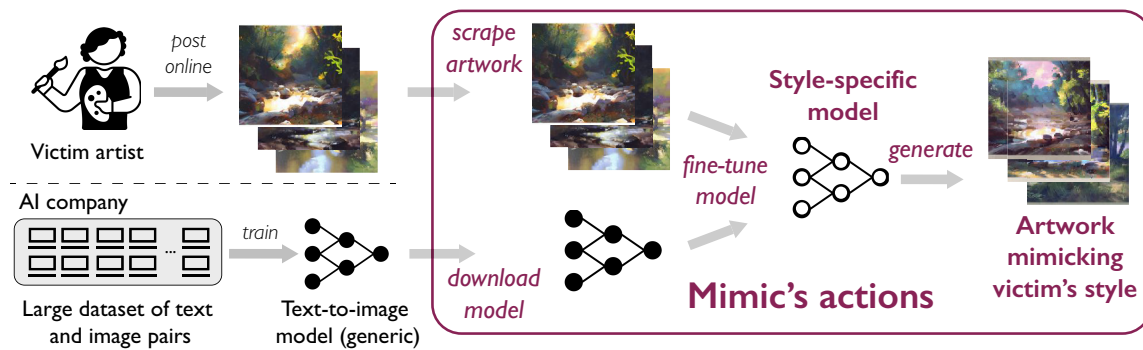


Figure 4. High level overview of the mimicry attack scenario. The mimic scrapes copyrighted artwork from the victim artist and uses these to fine-tune a pre-trained, generic text-to-image model. The generic model is trained and open-sourced by an AI company. The mimic then uses the fine-tuned model to generate artwork in the style of the victim artist.

Here, we first introduce the threat model, then discuss existing alternatives to the AI style mimicry problem. We present the intuition behind *Glaze* and detailed design in §5.

4.1 Threat Model

Here we state assumptions for both the artists protecting their own art and the users training models to replicate their artistic style. We refer to these AI art model trainers as “mimics.”

Artists. Artists want to share and promote their artwork online without allowing mimics to train models that replicate their art styles. Sharing art online enables artists to sell their work and attract commissioned work, fueling their livelihoods (§3). Artists protect themselves by adding imperceptible perturbations to their artwork before sharing them as shown in Figure 5. The goal of the *Glaze* cloak is to disrupt the style mimicry process, while only introducing minimal perturbation on images of the artwork.

We assume the artists have access to moderate computing resources (e.g., a laptop) and add perturbation to images of their artwork locally before posting online. We also assume artists have access to some public feature extractor (e.g., open-source models such as Stable Diffusion). We begin with assumption that artists use the same feature extractor as mimics (large majority of mimics use the open-source Stable Diffusion model). We later relax this assumption.

Mimics. The mimic’s goal is to train a text-to-image model that generates *high-quality* art pieces of any subject in the *victim’s style*. A mimic could be a well-funded AI company, e.g., Stability AI or OpenAI, or an individual interested in the style of victim artist. We assume the mimic has:

- access to the weights of generic text-to-image models well-trained on large datasets;
- access to art pieces from the target artist;
- significant computational power.

We assume the attack scenario where the mimic fine-tunes its model on images of the artist’s artwork (as shown in Figure 4). This is stronger than the naive mimic attack without

fine tuning. Finally, we assume the mimic is aware of our protection tool and can deploy adaptive countermeasures (§7).

4.2 Potential Alternatives and Challenges

A number of related prior works target protection against invasive and unauthorized facial recognition models. They proposed “image cloaking” as a tool to prevent a user’s images from being used to train a facial recognition model of them [9, 13, 24, 81, 95]. They share a similar high level approach, by using optimized perturbations that cause cloaked images to have drastically different feature representations from original user images. It is possible to adapt existing cloaking-based systems to protect artists against AI style mimicry. Protection system would compute a cloak on each artwork in order to perturb its feature space representation to be different from its unperturbed representation. This can succeed if the cloak significantly shifts the artwork’s feature representation, making resulting models generate dramatically different artwork.

We found that in practice, however, existing solutions are unable to introduce large-enough feature space shifts to achieve the desired protection. This is due to the properties of feature spaces in text-to-image models. Face recognition models *classify identities*, so their feature spaces mainly represent identity-related information. On the other hand, text-to-image models *reconstruct original images from extracted features*, so their feature spaces retain more information about the original image (objects, locations, color, style, etc.). Thus, producing the same shift in feature representation in a text-to-image model is much harder (requires more perturbation budget) than in a classification model. This observation is validated by prior work showing that adversarial perturbations are much less effective at attacking generative models [29, 39, 85]. Specifically, [39, 85] found that adversarial attack methods that are effective at attacking classifiers are significantly less effective at attacking autoencoders. We empirically confirm that existing cloaking methods cannot prevent AI mimicry (§A.1 in Appendix). We show that Fawkes [81]

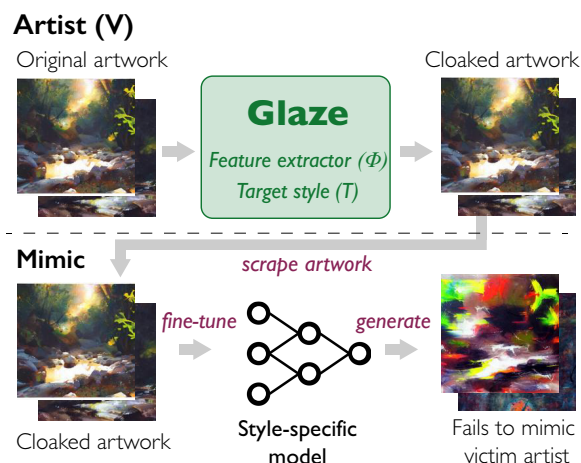


Figure 5. Overview of *Glaze*, a system that protects victim artists from AI style mimicry by cloaking their online artwork. (Top) An artist V applies the cloaking algorithm (uses a feature extractor Φ and a target style T) to generate cloaked versions of V 's art pieces. Each cloak is a small perturbation unnoticeable to human eye. (Bottom) A mimic scrapes the cloaked art pieces from online and uses them to fine-tune a model to mimic V 's style. When prompted to generate artwork in the style of V , mimic's model will generate artwork in the target style T , rather than V 's true style.

and LowKey [13] perform poorly in this setting, even when artists add highly visible cloaks to their artwork.

For generative models, concurrent work [76] proposes PhotoGuard, a method to cloak images to prevent unauthorized image edits (inpainting) on cloaked images. Similar to existing cloaking systems, PhotoGuard tries to indiscriminately minimize all information contained in an image (*i.e.*, the norm of the feature vector) to prevent models from editing the image. Thus, it is also not effective at mimicry prevention.

Design Challenges. The main reason that existing cloaking methods fail to prevent AI mimicry is because they indiscriminately shift all features in an image, wasting the cloak perturbation budget on shifting unnecessary features (*e.g.*, object shape, location, etc.). Protecting artist's style requires only *shifting features related to the artistic style of victim*. This can be achieved if a text-to-image model learns to draw objects similar to those drawn by the victim artist *as long as the model cannot mimic the artist's unique style*. Thus, optimal protection from mimicry requires concentrating the cloak on style-specific features.

Unfortunately, identifying and separating out these style-specific features is difficult. Even assuming the existence of interpretability methods that perfectly explain the feature space of a text-to-image model, there is no clear way to mathematically define and calculate "artistic styles." In all likelihood, any definition would change across different styles. For example, "impressionist" likely correlates more strongly with color features, whereas "cubism" correlates with shape features. Even across multiple art pieces in the same style, the style may manifest differently.

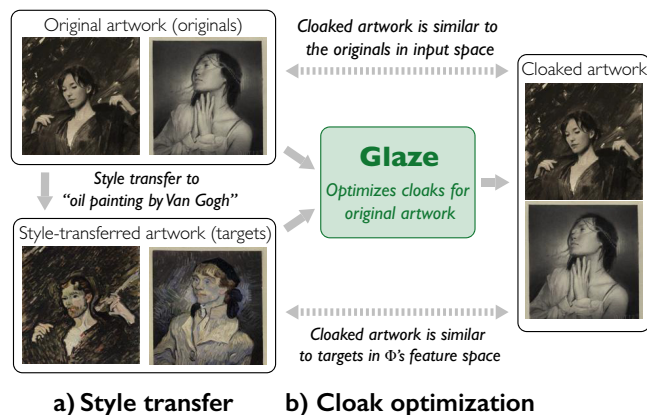


Figure 6. High level overview of how *Glaze* perturbs the style-specific features of the artwork. **a)** *Glaze* style transfers the original artwork to a different style, which changes its style but leaves other features unaltered. **b)** *Glaze* optimizes a cloak that makes the artwork's features representation match that of the style-transferred art, while constraining the amount of visible changes to the artwork.



Figure 7. Example style-transferred artwork with different target styles.

5 Disrupting Style Mimicry with Glaze

In this section, we introduce *Glaze*, its design intuition followed by the detailed algorithm.

5.1 Design Intuition

Our key intuition is to identify and isolate *style-specific features* of an artist's original artwork, *i.e.*, the set of image features that correspond to artistic style. Then *Glaze* computes cloaks while focusing the perturbation budget on these style-specific features to maximize impact on stylistic features.

As discussed, identifying and calculating style-specific features in model's feature space is difficult due to the poor interpretability of model features and how art style manifests differently across artworks. We overcome these two challenges by designing a style-dependent and artwork-dependent method that operates at image space. Given an artwork, we leverage "style transfer," an end-to-end computer vision technique, to modify and isolate its style components. "Style transfer" transforms an image into a new image with a differ-

ent style (*e.g.*, from impressionist style to cubist style) while keeping other aspects of the image similar (*e.g.*, subject matter and location).

We leverage style transfer in our protection technique as follows. Given an original artwork from the victim artist, we apply style-transfer to produce a similar piece of art with a different style, *e.g.*, in style of “an oil painting by Van Gogh” (Figure 6 a). The new version has similar content to the original, but its style mirrors that of Van Gogh. We show more style-transfer examples with different target styles in Figure 7. Now, we can use the style-transferred artwork as projection target to guide the perturbation computation. This perturbs the original artwork’s style-specific features towards that of the style-transferred version. We do this by optimizing a cloak that, when added to the original artwork, makes its feature representation similar to the style-transferred image. Since the content is identical between the pair of images, cloak optimization will focus its perturbation budget on style features.

5.2 Computing Style Cloaks

Using this approach, we compute style cloaks to disrupt style mimicry as follows. Given an artwork (x), we use an existing feature extractor to compute the style-transferred version of x into target style T : $\Omega(x, T)$. We then compute a style cloak δ_x , such that δ_x moves x ’s style-specific feature representation to match that of $\Omega(x, T)$ while minimizing visual impact. The cloak generation optimization is:

$$\begin{aligned} \min_{\delta_x} & \text{Dist}(\Phi(x + \delta_x), \Phi(\Omega(x, T))), \\ \text{subject to } & |\delta_x| < p, \end{aligned} \quad (1)$$

where Φ is a generic image feature extractor commonly used in text-to-image generation tasks, $\text{Dist}(\cdot)$ computes the distance of two feature representations, $|\delta_x|$ measures the perceptual perturbation caused by cloaking, and p is the perceptual perturbation budget.

As discussed in §5.1, the use of the style-transferred image $\Omega(x, T)$ guides the cloak optimization in Eq (1) to focus on changing style-specific image features. To maximize cloak efficacy, the target style T should be dissimilar from artist’s original style in the feature space. We discuss our heuristic for selecting target styles in §5.

5.3 Detailed System Design

Now we present the detailed design of *Glaze*. Given a victim artist V , *Glaze* takes as input the set of V ’s artwork to be shared online X_V , an image feature extractor Φ , a style-transfer model Ω , and perturbation budget p . Note that in many cases, a single model (*e.g.* Stable Diffusion) provides both Φ and Ω .

Step 1: Choose Target Style. The selected target style T should be sufficiently different from V ’s style in model feature space to maximize chances of disrupting style mimicry. For

example, Fauvism and Impressionism are distinct art styles that often look visually similar to the untrained eye. Image of an impressionist painting style cloaked to Fauvism might not produce a visually discernible effect on model-generated paintings. Note that an artist can maximize their ability to avoid mimicry if they consistently style cloak all their artwork towards the same target T .

For a new user, *Glaze* uses the following algorithm to randomly select T from a set of candidate styles reasonably different from V ’s style. The algorithm first inspects a public dataset of artists, each with a specific style (*e.g.*, Monet, Van Gogh, Picasso). For each candidate target artist/style, it selects a few images in that style and calculates their feature space centroid using Φ . It also computes V ’s centroid in Φ using V ’s artwork. Then, it locates the set of candidate styles whose centroid distance to V ’s centroid is between the 50 to 75 percentile of all candidates. Finally, it randomly selects T from the candidate set.

Step 2: Style transfer. *Glaze* then leverages a pre-trained style-transfer model Ω [67] to generate the style-transferred artwork for optimization. Given each art piece $x \in X_V$ and target style T , it style transfers x to target style T to produce style-transferred image $\Omega(x, T)$.

Step 3: Compute cloak perturbation. Then, *Glaze* computes the cloak perturbation, δ_x for x , following the optimization defined by eq. (1), subject to $|\delta_x| < p$. Our implementation uses LPIPS (Learned Perceptual Image Patch Similarity) [106] to bound the perturbation. Different from the L_p distance used in previous work [8, 40, 71], LPIPS has gained popularity as a measure of user-perceived image distortion [13, 42, 68]. Bounding cloak generation with this metric ensures that cloaked versions of images are visually similar to the originals. We apply the *penalty method* [56] to solve the optimization in eq.(1) as follows:

$$\min_{\delta_x} \|\Phi(\Omega(x, T)) - \Phi(x + \delta_x)\|_2^2 + \alpha \cdot \max(\text{LPIPS}(\delta_x) - p, 0) \quad (2)$$

where α controls the impact of the input perturbation. L_2 distance is used to calculate feature space distance.

Upload artwork online. Finally, the artist posts the cloaked artwork online. For artists already with a large online presence, they can cloak and re-upload artwork on their online portfolio. While updating online images is not always possible, *Glaze* can be effective even when the mimic’s model has significant amount of uncloaked art (§6.4).

5.4 On the Efficacy of Style Cloaks

Glaze’s style cloaks work by shifting feature representation of artwork in the generator model. But how much shift do we need in order to have a noticeable impact on mimicked art?

Two reasons suggest that even small shifts in style will have a meaningful impact in disrupting style mimicry. First, generative models used for style mimicry have *continuous*

output spaces, *i.e.*, any shift in image feature representation results in changes in the generated image. Because generative models are trained to interpolate their continuous feature spaces [89, 96], any shift in the model’s representation of art style results in a new style, a “blend” between the artist and the chosen target style. Second, mimicked artwork must achieve reasonable quality and similarity in style to the artist to be useful. Small shifts in the style space often produce incoherent blends of conflicting styles that are enough to disrupt style mimicry, *e.g.*, thick oil brushstrokes of Van Gogh’s style mixed into a realism portrait.

These two factors contribute to *Glaze*’s success in more challenging scenarios (§6.4), and its robustness against countermeasures (*e.g.* adversarial training) that succeed against cloaking tools for facial recognition (§7).

6 Evaluation

In this section, we evaluate *Glaze*’s efficacy in protecting artists from style mimicry. We first describe the datasets, models, and experimental configurations used in our tests. Then we present the results of *Glaze*’s protection in a variety of settings. Due to *Glaze*’s highly visual nature, we evaluate its performance using both direct visual assessment by **human artists** in a user study, and **automated metrics** (see §6.2 for details).

Summary of results. Over 93% of artists surveyed believe *Glaze* effectively protects artists’ styles from AI style mimicry attacks. Protection efficacy remains high in challenging settings, like when the mimic has access to unprotected artwork. *Glaze* also achieves high protection performance against a real-world mimicry-as-a-service platform. Of our 1156 artist participants, over 92% found the perturbations introduced by cloaking small enough not to disrupt the value of their art, and over 88% would like to use *Glaze* to protect their own artwork from mimicry attacks.

6.1 Experiment Setup

Mimicry dataset. We evaluate *Glaze*’s performance in protecting the styles of the following two groups of artists:

- *Current artists*: 4 professional artists let us use their artwork in our experiments. These artists have different styles and backgrounds (*e.g.*, full-time/freelancers, watercolor painters/digital artists, well-known/independent). Each provided us with between 26 to 34 *private* original art pieces for our experiments. We use perceptual hashing [38] to verify that none of these are included in existing public datasets used to train generic text-to-image models (*e.g.* [10, 78]).
- *Historical artists*: We also evaluate *Glaze*’s protection on 195 historical artists (*e.g.*, van Gogh, Monet) from the WikiArt dataset [73]. The WikiArt dataset contains 42,129 art pieces from 195 artists. Each art piece is labeled with

its genre (*e.g.*, impressionism, cubism). We randomly sampled 30 art pieces from each artist to use in style mimicry attacks. Generic text-to-image models found online have been trained on some artwork from these artists. Using this art simulates a more challenging scenario in which a famous artist attempts to disrupt a model that already understands their style.

Mimicry attack setup. We recreate the strongest-possible mimicry attack scenario, based on techniques used in real-world mimicry incidents [3, 70, 100], that works as follows. First, we take art pieces from the victim artist *V* and generate a text caption for each piece using an image captioning model [49]. The pretrained image captioning model generates a short sentence to describe the image. We found that this model can correctly caption protected images (examples in Figure 18), likely because *Glaze* focuses on perturbing style features while the captioning models focus on image content. Then, we append the artist’s name to each caption, *e.g.*, “mountain range by Vincent van Gogh”. Finally, we fine-tune a pre-trained generic text-to-image model (details below) on the caption/image pairs.

We use 80% of the art pieces from the victim artists to fine-tune models that mimic each artist’s style, reserving the rest for testing. We fine-tune for 3000 optimization steps, which we find achieves the best mimicry performance (Figure 19 in Appendix). We then use the fine-tuned, style-specific model to generate mimicked artwork in style of each victim artist. We query the model using the generated captions (which include *V*’s name) from the held-out test artwork set. We generate 5 pieces of mimicked art for each text caption using different random seeds and compare these to the real victim art pieces with this caption. Additional details on training and generation parameters, as well as its sensitivity to random seed selection and the number of training art pieces are in Appendix A.2.

Text-to-image models. We use two state-of-the-art, public, generic text-to-image models in our experiments:

- *Stable Diffusion (SD)*: Stable Diffusion is a popular and high-performing open-source text-to-image model [83], trained on 11.5 million images from the LAION dataset [78]. SD training takes over 277 GPU months (on A100 GPU) and costs around \$600K [83]. SD uses diffusion methods to generate images and achieves state-of-the-art performance on several benchmarks [67]. Viewed as one of the best open-source models, SD has powered many recent developments in text-to-image generation [1, 43, 57, 77]. We use SD version 2.1 in the paper [83], the most up-to-date version as of December 2022.
- *DALL·E-mega (DALL·E-m)*: DALL·E-m-mega, an updated version of the more well-known DALL·E-m-mini, is an open-source model based on OpenAI’s DALL·E-m 1 [65]. The model leverages a VAE for image generation and is trained on 17 million images from three different datasets [10, 82, 87]. Training takes 2 months on 256 TPUs [16]. While

DALL·E-m performs worse than diffusion-based models like SD, we use it to evaluate how *Glaze* generalizes to different model architectures.

Glaze configuration. We generate cloaks for each of victim V 's art pieces following the methodology of §5.3. First, we use the target selection algorithm to select a target style T . We choose from a set of 1119 candidate target styles, collected by querying the WikiArt dataset with artist and genre names, e.g., "Impressionism painting by Monet"². We then style transfer each victim art piece into the target style leveraging the style transfer functionality of stable diffusion model (stable diffusion model has both text-to-image and style transfer functionality). A style transfer model takes in an original image and a target prompt as input. Leveraging a similar diffusion process, the model modifies the original image to a style similar to that described in the target prompt. More information on style transfer can be found in [72]. Finally, we optimize a cloak for each art piece using Eq. 2 by running the Adam optimizer for 500 steps. We benchmark *Glaze*'s runtime on artwork with resolution ranging from 512 to 6000 pixels, using SD's feature extractor (ViT model with 83 million parameters). It takes an average of 1.2 mins on Titan RTX GPU and 7.3 mins on a single Intel i7 CPU to generate a cloak for a single piece of art.

In our initial experiments, we assume *Glaze* generates cloaks using the same image feature extractor as the mimic (e.g. SD's or DALL·E-m's feature extractor). We relax this assumption and evaluate *Glaze*'s performance when artists and mimics use different feature extractors in §6.4.

6.2 Evaluation Metrics

We evaluate our protection performance using both visual assessment and feedback from human artists, and a scalable metric. Here, we describe the setup of our evaluation study and define the exact metrics used for evaluation.

Artist-rated protection success rate (Artist-rated PSR): The user studies ask artists to rate the performance of *Glaze*. We generate a dataset of mimicry attacks on 13 victim artists (the 4 current artists and 9 randomly chosen historical artists) across 23 protection scenarios (including ones in §7). For each participant, we randomly select a set of mimicry attacks out of these 13×23 settings and ask them to evaluate protection success. For each mimicry attempt, we show participants 4 mimicked art pieces and 4 original art pieces from the victim artist. Using original art pieces as an indicator of the human artist's style, we ask participants to consider the mimicked art, and rate the success of *Glaze*'s protection on a 5-level Likert scale (ranging from "not successful at all" to "very successful"). Each mimicry attempt is evaluated by at least 10 participants. We define *artist-rated PSR* as the percent of participants who rated *Glaze*'s protection as "successful" or

²One artist may paint in multiple styles, resulting in multiple candidate target styles from a single artist.

Generic model	Artist dataset	w/o <i>Glaze</i>		w/ <i>Glaze</i> ($p=0.05$)	
		Artist-rated PSR	CLIP-based genre shift	Artist-rated PSR	CLIP-based genre shift
SD	Current	$4.6 \pm 0.3\%$	$2.4 \pm 0.2\%$	$94.3 \pm 0.8\%$	$96.4 \pm 0.5\%$
	Historical	$4.2 \pm 0.2\%$	$1.3 \pm 0.2\%$	$93.3 \pm 0.6\%$	$96.0 \pm 0.3\%$
DALL·E-m	Current	$31.9 \pm 3.5\%$	$6.4 \pm 0.8\%$	$97.4 \pm 0.2\%$	$97.4 \pm 0.3\%$
	Historical	$29.8 \pm 2.4\%$	$5.8 \pm 0.6\%$	$96.8 \pm 0.3\%$	$97.1 \pm 0.2\%$

Table 2. *Glaze* has a high protection success rate, as measured by artists and CLIP, against style mimicry attacks. We compare protection success when artists do not use *Glaze* vs. when they do (with perturbation budget 0.05).

"very successful." Our user studies primarily focus on artists, as they would be most affected by this technology. We found though, that not all current artists despise AI art, and some view it as a new avenue for a different form of artistry.

CLIP-based genre shift: We define a new metric based on CLIP [62], using the intuition that *Glaze* succeeds if the mimicked art has been impacted enough by *Glaze* to be classified into a *different art genre* from the artist's original artwork. We leverage CLIP model's ability to classify art images into art genres. Given a set of mimicked art targeting an artist V , we define *CLIP-based genre shift rate* as the percentage of mimicked art whose top 3 predicted genres do not contain V 's original genre. A higher genre shift rate means more mimicked art belongs to a different genre from the victim artist, and thus means more successful protection.

To calculate the genre shift we use a set of 27 historical genres from WikiArt dataset and 13 digital art genres [33] as the candidate output labels. In Appendix A.3, we show that a pre-trained CLIP model is able to achieve high genre classification performance. We report the average CLIP-based genre shift for all 199 victim artists across all mimicked artworks.

We use CLIP-based genre shift as a supplemental metric to evaluate *Glaze* because it is only able to detect style changes at the granularity of art genres. However, mimicry attacks also fail when *Glaze* causes the mimicked artwork quality to be very low, something that CLIP cannot measure. Measuring the quality of generated image has been a challenging and ongoing research problem in computer vision [5, 36, 41].

6.3 Glaze's Protection Performance

Style mimicry success when *Glaze* is not used. Mimicry attacks are very successful when the mimic has access to a victim's original (unmodified) artwork. Examples of mimicked artwork can be found in Figure 8. The leftmost two columns of Figure 8 show a victim artist's original artwork, while the third column depicts mimicked artwork generated by a style-specific model trained on victim's original artwork when *Glaze* is not used. In our user study, over $> 95\%$ of respondents rated the attack as successful. Table 2, row 1, gives the artist-rated and CLIP-based genre shift for mimicry attacks on unprotected art.

SD models produce stronger mimicry attacks than DALL·E-m models, according to our user study (see Table 2). This is

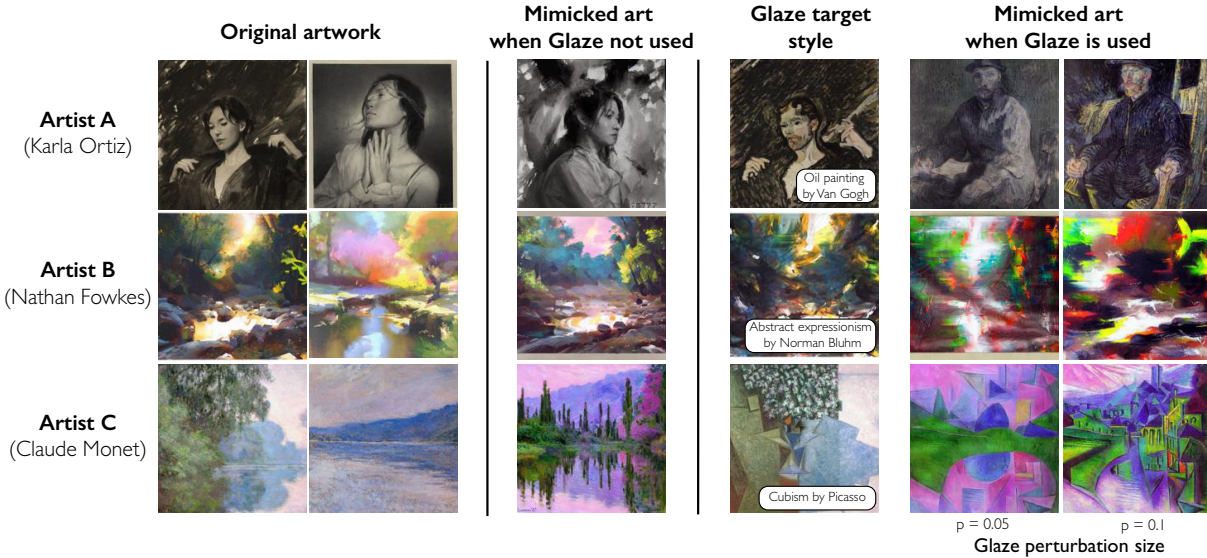


Figure 8. Example *Glaze* protection results for three artists. **Columns 1-2:** artist’s original artwork; **column 3:** mimicked artwork when artist does not use protection; **column 4:** style-transferred artwork (original artwork in column 1 is the source) used for cloak optimization and the name of target style; **column 5-6:** mimicked artwork when artist uses cloaking protection with perturbation budget $p = 0.05$ or $p = 0.1$ respectively. All mimicry examples here use SD-based models.

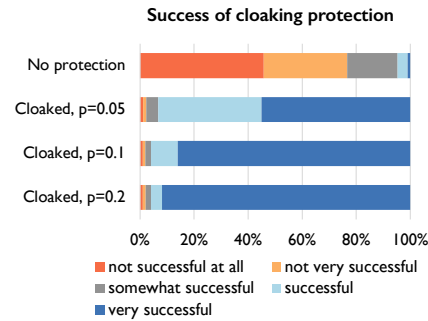


Figure 9. *Glaze*’s cloaking protection success increases as cloak perturbation budget increases. The top row of the figure shows baseline performance with the mimic trains on uncloaked images ($p=0$).

Perturbation budget	Artist-rated PSR	CLIP-based genre shift
0 (no cloak)	$4.6 \pm 1.4\%$	$2.4 \pm 0.8\%$
0.05	$93.3 \pm 0.6\%$	$96.0 \pm 0.3\%$
0.1	$95.9 \pm 0.4\%$	$98.2 \pm 0.1\%$
0.2	$96.1 \pm 0.3\%$	$98.5 \pm 0.1\%$

Table 3. Performance of our system (artist-rated protection success rate and CLIP-based genre shift rate) increases as the perturbation budget increases. (SD model, averaged over all victim artists).

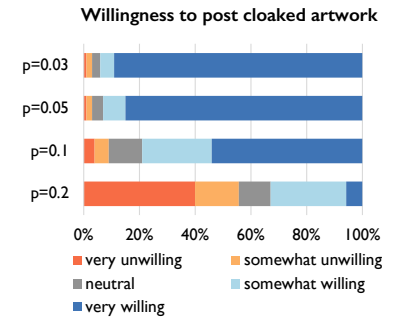


Figure 10. Artists’ willingness to post cloaked artwork in place of the original decreases as perturbation budget of the cloaks increases.

unsurprising, as DALL·E-m models generally produce lower-quality generated images. CLIP-based genre shift does not reflect this phenomenon, as this metric does not assess image quality.

***Glaze*’s success at preventing style mimicry.** *Glaze* makes mimicry attacks markedly less successful, as shown in Figure 8. Columns 5 and 6 (from left) show mimicked artwork when the style-specific models are trained on artwork protected by *Glaze*. For reference, column 4 shows an example style-transferred artwork $\Omega(x, T)$ used to compute *Glaze* cloaks for the protected art pieces. Overall, *Glaze* achieves $> 93.3\%$ artist-rated PSR and $> 96.0\%$ CLIP-based genre shift (see Table 2). *Glaze*’s protection performance is slightly higher for current artists than for historical artists. This is likely because the historical artists’ images are present in the training datasets of our generic models (SD, DALL·E-m), highlighting the additional challenge of protecting well-

known artists whose style was already learned by the generic models.

How large of perturbations will artists tolerate? Increasing the *Glaze* perturbation budget enhances protection performance. We observe that both artist-rated and CLIP-based genre shift increase with perturbation budget (see Figure 9, Table 3, and Figure 20). Given this tradeoff between protection success and *Glaze* protection visibility on original artwork, we evaluate how perturbation size impacts artists’ willingness to use *Glaze*.

We find that artists are willing to add fairly large *Glaze* perturbations to their artwork in exchange for protection against mimicry. To measure this, we show 3 randomly chosen pairs of original/cloaked artwork to each of the 1,156 artists in our first study. For each art pair, we ask the artist whether they would be willing to post the cloaked artwork (instead of the original, unmodified version) on their personal website. More

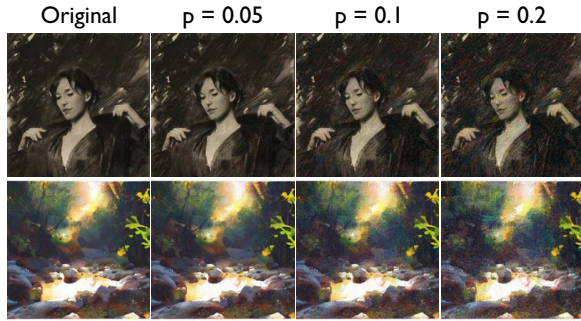


Figure 11. Original artwork and cloaked artwork computed using three different cloak perturbation budgets.

than 92% of artists select “willing” or “very willing” when $p = 0.05$. This number only slightly increases to 94.3% when $p = 0.03$. Figure 10 details artists’ preferences as perturbation budget increases. (see Figure 11 for examples of cloaked artwork with increasing p). Based on these results, we use perturbation budget $p = 0.05$ for all our experiments, since most artists are willing to tolerate this perturbation size.

Surprisingly, over 32.8% artists are willing to use cloaks with $p = 0.2$, which is clearly visible to human eye (see Figure 11). While we are surprised by this high perturbation tolerance, in our follow-up free response artists noted that they would be willing to tolerate large perturbations because of the devastating consequence if their styles are stolen. One participant stated that “I am willing to sacrifice a bit image quality for protection.” Many artists (> 80%) also noted that they have already used traditional, more visually disruptive techniques to protect their artwork online when posting online, *i.e.*, adding watermark or reducing image resolution. One participant stated that “I already use low to medium resolution images only for online posting, thus this would not impact my quality control too much.”

6.4 Glaze’s Protection Robustness

Next, we test *Glaze*’s efficacy in more challenging scenarios. First, we measure performance when the mimic uses a different feature extractor for mimicry than the one used by the artist to generate the cloak. Second, we measure what happens when the mimic has uncloaked artwork samples from the victim. Due to the poor mimicry performance of DALL-E-m, we focus our evaluation using SD as the generic model.

Artist/mimic use different feature extractors. In the real world, it is possible that the mimic will use a different model (and thus a different image feature extractor) for style mimicry than the one used by the victim artist to cloak their artwork. While the feature extractors may still be similar because of the well-known transferability property between large models [18, 58, 79, 84, 101], their differences could reduce the efficacy of cloaking. We test this scenario using three feature extractors— Φ -A, Φ -B, and Φ -C. Φ -A and Φ -B have different model architectures (autoencoder-KL [67] vs. VQ-VAE [65])

Artist dataset	w/o <i>Glaze</i>		w/ <i>Glaze</i> ($p=0.05$)	
	Artist-rated PSR	CLIP-based genre shift	Artist-rated PSR	CLIP-based genre shift
Current	$6.2 \pm 0.5\%$	$3.8 \pm 0.3\%$	$92.5 \pm 0.5\%$	$94.2 \pm 0.3\%$
Historical	$7.2 \pm 0.6\%$	$3.3 \pm 0.4\%$	$92.1 \pm 0.3\%$	$93.9 \pm 0.4\%$

Table 4. Performance of *Glaze* against real-world mimicry service (scenario.gg). Mimicry service achieves high mimicry success when no protection is used. When *Glaze* is used, the mimicry service has low performance.

but are both trained on the ImageNet dataset [19]. Φ -A and Φ -C have different model architectures (autoencoder-KL vs VQ-VAE) and training datasets (ImageNet vs. CelebA [48]).

In our experiments, the victim artist uses one feature extractor (either Φ -B or Φ -C) to optimize cloaked artwork, and the mimic trains their style-specific models with SD models using Φ -A. Despite the difference in victim/mimic extractors, *Glaze*’s protection remains highly successful (left half of Figure 12)—the style of mimicked artwork remains distinct from artist’s true style. Artist-rated and CLIP-based genre shift measurements confirm this observation. Artist-rated PSR is > 90.2%, while CLIP-based genre shift is > 94.0%. The PSR is slightly higher when the two feature extractors only differ in architectures (Φ -B to Φ -A) than when they differ in both architecture and training data (Φ -C to Φ -A).

Mimic has access to uncloaked artwork. Another challenging scenario is when the mimic gains access to some *uncloaked* artwork from victim artists. This is a realistic scenario for many prominent artists with a large online presence. As expected, *Glaze*’s protection performance decreases when the mimic has access to more uncloaked artwork (right side of Figure 12). As the ratio of uncloaked/cloaked art in the mimic’s dataset increases, the mimicked artwork becomes more similar to artist’s original style. Yet, *Glaze* is still reasonably effective (87.2% artist-rated PSR) even when artists can only cloak 25% of their artwork. This validates our hypothesis in §5.4 that cloaking will have a noticeable effect as long as the mimic has some cloaked training data.

A mimic with access to a large amount of uncloaked artwork is still an issue for *Glaze*. Fortunately, in our user study, we found that 1) many artists constantly create and share new artwork online, which can be cloaked to offset the percentage of uncloaked artwork, and 2) many artists change their artistic style over time. In our user study, we asked artists to estimate the number of unique art pieces they currently have online (M) and the estimated number of art pieces they anticipate uploading each subsequent year (Y). Among artists with an existing online presence, over 40% have $Y/M > 25\%$, meaning that one year from now, > 20% of their total online artwork would be cloaked (if they start using *Glaze* immediately). More than 81% of artists also stated that their art style has changed over their career, and half of these said that theft of their old, outdated styles is less concerning.

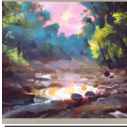
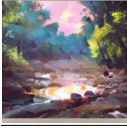
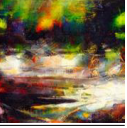
Feature extractors used by artist and mimic					Percentage of artwork cloaked			
	Artist: no cloaking Mimic: Φ -A	Artist: Φ -A Mimic: Φ -A	Artist: Φ -B Mimic: Φ -A	Artist: Φ -C Mimic: Φ -A	0% cloaked	25% cloaked	50% cloaked	75% cloaked
Attempts to mimic artist A								
Attempts to mimic artist B								
Artist-rated PSR	$4.3 \pm 0.2\%$	$93.5 \pm 0.6\%$	$91.3 \pm 0.5\%$	$90.2 \pm 0.8\%$	$4.3 \pm 0.2\%$	$87.2 \pm 1.1\%$	$90.1 \pm 0.8\%$	$91.5 \pm 0.9\%$
CLIP-based genre shift	$1.4 \pm 0.2\%$	$96.0 \pm 0.3\%$	$94.8 \pm 0.4\%$	$94.0 \pm 0.4\%$	$1.4 \pm 0.2\%$	$90.3 \pm 0.8\%$	$93.8 \pm 0.4\%$	$94.7 \pm 0.3\%$

Figure 12. *Glaze* remains successful under two challenging scenarios. Left: when artist and mimic use different feature extractors. Right: when artists can only cloak a portion of their artwork in mimic’s dataset. Bottom of the figure shows artist-rated PSR and CLIP-based genre shift for the corresponding setting.

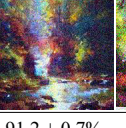
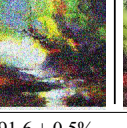
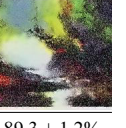
Gaussian noise level				Denoised
$\sigma = 0.05$	$\sigma = 0.1$	$\sigma = 0.15$		
Attempts to mimic artist A				
Attempts to mimic artist B				
Artist-rated PSR	$92.9 \pm 0.5\%$	$91.2 \pm 0.7\%$	$91.6 \pm 0.5\%$	$89.3 \pm 1.2\%$

Figure 13. *Glaze*’s protection performance remains high as mimic adds an increasing amount of Gaussian noise to the cloaked artwork. Even when the mimic adds denoising (last column), *Glaze*’s protection persists.

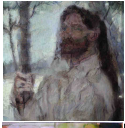
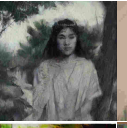
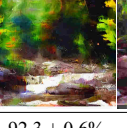
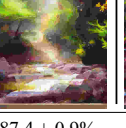
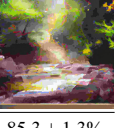
JPEG compression level				Upscaled
20	15	10		
Attempts to mimic artist A				
Attempts to mimic artist B				
Artist-rated PSR	$93.4 \pm 0.8\%$	$92.3 \pm 0.6\%$	$87.4 \pm 0.9\%$	$85.3 \pm 1.3\%$

Figure 14. *Glaze*’s protection performance remains high as mimic adds JPEG compression to the cloaked artwork. Even when the mimic also upscales the mimicked images (last column), *Glaze*’s protection persists.

6.5 Real-World Performance

Next, we test *Glaze* against a real-world style mimicry-as-a-service system, *scenario.gg* [77]. *Scenario.gg* is a web service that allows users to upload a set of images in a specific style. The service then trains a model to mimic the style and returns an API endpoint that allows the user to generate mimicked images in the trained style. The type of model or mimicry method used by the service is unknown.

Glaze remains effective against *scenario.gg*. We ask *scenario.gg* to mimic the style from a set of cloaked or un-cloaked artwork from 4 current artists and 19 historical artists. Table 4 shows that when no protection is used, *scenario.gg* can successfully mimic the victim style ($< 7.2\%$ protection success). The mimicry success of *scenario.gg* is lower than our mimicry technique, likely because *scenario.gg* trains the model for fewer iterations due to computational constraints. When we use *Glaze* to cloak the artwork and upload the cloaked artwork, *scenario.gg* fails to mimic the victim style ($> 92.1\%$ artist-rated PSR and $> 93.9\%$ CLIP-based genre shift rate) as shown in Table 4.

7 Countermeasures

We consider potential countermeasures a mimic could employ to reduce the effectiveness of *Glaze*. We consider the strongest adversarial setting, in which the mimic has white-box access to our protection system, *i.e.*, access to the feature extractor used and protection algorithm. In our experiments, we assume the mimic uses the SD model as the generic model and test the efficacy of each countermeasure on the 13 victim artists from §6.2. Here, we focus on artist-rated PSR metric, because many countermeasures trade off image quality for mimicry efficacy, and CLIP-based metric does not consider image quality.

Image transformation. A popular approach to mitigate the impact of small image perturbations, like those introduced by *Glaze*, is to transform training images before using them for model training [7,25]. In our setting, the mimic could augment the cloaked artwork before fine-tuning their model on them to potentially reduce cloak efficacy. We first test *Glaze*’s resistance to two popular image transformations, adding Gaussian noise and image compression. We also consider a stronger version of this countermeasure that then tries to correct the

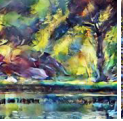
Number of robust training steps				
	1K steps	3K steps	5K steps	10K steps
Attempts to mimic artist A				
Attempts to mimic artist B				
Artist-rated PSR	92.2 ± 0.8%	89.3 ± 1.3%	91.3 ± 0.9%	95.3 ± 0.3%

Figure 15. *Glaze*’s protection performance remains high against robust training countermeasure proposed by Radiya *et al.*. The protection performance first decreases then increases as mimic robustly trains the model with an increasing number of steps.

image quality degradation introduced by the transformations.

Transforming cloaked artwork does not defeat *Glaze*’s protection. Figure 13 shows that as the magnitude of Gaussian noise (σ) increases, the quality of mimicked artwork decreases as fast as or faster than cloak effectiveness. This is because models trained on noisy images learn to generate noisy images. We observe a similar outcome when mimic uses JPEG compression (Figure 14), where image resolution and quality degrade due to heavy compression. Artists-rated PSR decreases slightly but remains above $> 87.4\%$ across both types of data transformations. Artists consider *Glaze*’s protection to be successful when mimicked artwork is of poor quality.

The mimic can take this countermeasure one step further by *reversing* the quality degradation introduced by the noising/compression process. Specifically, a mimic can run image denoising or image upscaling tools on the mimicked artwork (*e.g.*, ones shown in Figure 13 and 14) to increase their quality. We found this approach improves generated image quality but still does not allow for successful mimicry. For denoising, we ran a state-of-the-art CNN-based image denoiser [104] that is specifically trained to remove “additive Gaussian noise” (the same type of noise added to cloaked artwork). The last column of Figure 13 shows the denoised image (using the noisy mimicked image when $\sigma = 0.2$ as the input). While the process removes significant amounts of noise, the denoised artwork still has many artifacts, especially around complex areas of the artwork (*e.g.*, human face). We observe similar results for image upscaling, where we use a diffusion-based image upscaler [83] to improve the quality of compressed images (Figure 14). Overall, our artist-rated protection success rate remains $> 85.3\%$ against this improved countermeasure.

Radiya *et al.* [64] robust training. Radiya *et al.* [64] design a robust training method to defeat cloaking tools like Fawkes [81] and Lowkey [13] in the face recognition setting. At a high level, this method augments the attacker’s training dataset with some cloaked images generated by the cloaking tool and the *correct* output labels. Training on such data makes

the model more robust against cloak perturbations on unseen cloaked images at inference time, and thus, can potentially circumvent the protection.

We test if this robust training approach can defeat *Glaze*. We assume the mimic first robustly trains the feature extractors in their generic models using cloaked artwork generated by *Glaze*, and then trains the generator model to generate images from the robust feature space. Finally, the mimic uses the robust generic model for style mimicry as in §6. We discuss the detailed robust training setup in Appendix A.4.

Glaze performance remains high, even if the mimic robustly trains the generic model for many iterations before using it for style mimicry (see Figure 15). As the model becomes more robust, the mimicked artwork is less impacted by cloaking (less influenced of the target style). However, robust training greatly degrades mimicked image quality, preventing successful mimicry. Overall, the artist-rated PSR remains $> 88.7\%$. To mitigate robust training’s impact on image quality, we explore an alternative robust training method, where we robustly train a new feature extractor designed to remove cloak’s impact while operating in the original feature space (thus no need to change the image generator). We found this robust training approach is also ineffective (details in §A.4).

As discussed in §5.4, *Glaze* remains reasonably effective against Radiya *et al.* because 1) the continuous output space of the generative model, and 2) high quality requirement of art generation. Robust training reduces cloaking’s effectiveness but cannot completely remove its impact. In the classification case (facial recognition), this reduced effectiveness only manifests in small changes in classification confidence (compared to no cloaking) and often does not change the discrete classification outcome. However, in the context of generator models, the continuous output space means that even less-effective cloaks still directly affect the mimicked artwork. Combined with the high quality requirement, the reduced protection effect is enough to disrupt style mimicry, as shown in Figure 15. Additional robust training simply degrades generation quality, rather than reducing cloaking efficacy.

Outlier Detection. Another countermeasure could involve leveraging outlier detection to identify and remove protected images [80, 90, 91]. We test *Glaze*’s robustness to a state-of-the-art outlier detection method that leverages contrastive training [91]. Contrastively trained models project data into a well-separated feature space, which the mimic could leverage.

We assume the mimic has a ground truth set (20) of original artworks from a given artist. The mimic first projects these art pieces into the feature space of a model trained with contrastive loss on ImageNet dataset [91]. The mimic then trains a one-class SVM outlier detector [44] using these ground truth features. Now, given a new artwork from the same artists, the mimic detects whether the artwork is an outlier using the detector. Detection results on 4 current artists (§6) show that outlier detection has limited effectiveness against *Glaze* ($< 65\%$ precision and $< 53\%$ recall at detecting *Glaze* pro-

tected images).

8 Limitations and Releasing Glaze

We conclude with a discussion of the limitations of the current system, then describe our experiences during and after the *Glaze* release.

Limitations. First, protection from *Glaze* relies on artists cloaking a portion of their art in the mimic model’s training dataset. This is challenging for established artists because 1) their styles have matured over the years and are more stable, and 2) many of their art pieces have already been downloaded from art repositories like ArtStation and DeviantArt. These artists’ styles can be mimicked using only older artworks collected before the release of *Glaze*. While artists can prevent mimics from training on newer artwork, they need to rely on opt-out and removal options at art repositories to stop style mimicry.

Second, a system like *Glaze* that protects artists faces an inherent challenge of being *future-proof*. Any technique we use to cloak artworks today might be overcome by a future countermeasure, possibly rendering previously protected art vulnerable. While we are under no illusion that *Glaze* will remain future-proof in the long run, we believe it is an important and necessary first step towards artist-centric protection tools to resist invasive AI mimicry. We hope that *Glaze* and followup projects will provide some protection to artists while longer term (legal, regulatory) efforts take hold.

Releasing *Glaze* and managing expectations. We released *Glaze* as a free application on Mac and Windows in March 2023. We have repeatedly communicated *Glaze*’s limitations to users, both on our website and in communications to artists via our download page, on Twitter, in emails to artists, etc. In these communications, we clearly state that *Glaze* is not a permanent solution against AI mimicry and could potentially be defeated by future attacks.

As of June 2023, *Glaze* has been downloaded > 740K times by artists around the world. Reception on social media and emails to our lab have been extremely enthusiastic and positive. Artists have helped design *Glaze*’s user interface, made how-to videos on YouTube, and managed ad campaigns on Instagram to spur adoption in the community. Based on numerous requests on social media and via emails, we plan to test and deploy a web service in Summer 2023 to expand *Glaze* access to artists who lack compute and GPUs.

One excellent outcome from the *Glaze* release has been the technical discussions it has spurred with a variety of stakeholders. We began and are continuing collaborative efforts to advocate for artists rights, with art-centric social networks, advocate groups in the US (CAA) and the EU (EGAIR), government representatives, and companies who want to protect the IP of their images/characters.

Real-world countermeasures. Finally, we want to describe our experiences deploying *Glaze* in an adversarial setting. In



Glazed art

Plagiarized art by PEZ

Figure 16. Glazed image and generated image from PEZ mimicry method. The original image is *Musa Victoriosa*, a new painting created by Karla Ortiz to be the first artwork to be released publicly under *Glaze* protection.

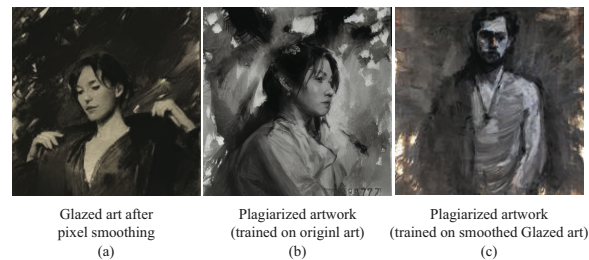


Figure 17. (a) Smoothed artwork by applying pixel smoother on Glazed artwork, (b) plagiarized artwork generated by training on original (unprotected) artwork, and (c) plagiarized artwork generated by training on Glazed artwork that was later pixel-smoothed.

the 3 months since initial release, multiple groups have sought to attack or bypass *Glaze* protection. While several attempts had minimal impact, we describe the two most serious attempts here and evaluate their effectiveness.

The first attack [51] leverages a newer style mimicry method [94], reverse engineering with PEZ. PEZ is able to perform high-quality style mimicry using a *single original image* from the original artist. Initial tests showed *Glaze* is robust against PEZ style mimicry (Figure 16). *Glaze* remains effective likely because *Glaze* directly modifies the feature representation of the art, and is thus effective against stronger mimicry attempts.

A second category of attacks tries to perform pixel-level image smoothing to remove cloaks added by *Glaze* [105]. This applies bilateral filters on Glazed images repeatedly, seeking to remove all added perturbations. We evaluate this attack on Glazed artwork in §6 and fine-tuning a model on the smoothed images. Figure 17 shows *Glaze* remains effective against pixel smoothing. This result is consistent with prior work showing that image smoothing cannot prevent adversarial perturbations [102].

Finally, while we have not yet observed any successful attacks against *Glaze*, we are continuously exploring design

improvements to further enhance robustness against potential future countermeasures.

Acknowledgements

We thank our anonymous reviewers and shepherd for their insightful feedback. We also thank Karla Ortiz, Lyndsey Gallant, Nathan Fowkes, Kim Van Deun, Jon Lam, Eveline Fröhlich, Edit Ballai, Kat Loveland and many other artists, without whom this project would not be possible. This work is supported in part by NSF grants CNS-2241303, CNS-1949650, and the DARPA GARD program. Opinions, findings, and conclusions or recommendations expressed in this material are those of the authors and do not necessarily reflect the views of any funding agencies.

References

- [1] AI RENDER. AI Render - Stable Diffusion in Blender. , 2022.
- [2] ANDERSEN, S. The Alt-Right Manipulated My Comic. Then A.I. Claimed It. The New York Times, 2022.
- [3] BAIO, A. Invasive Diffusion: How one unwilling illustrator found herself turned into an AI model, 2022.
- [4] BALAJI, Y., ET AL. ediffi: Text-to-image diffusion models with an ensemble of expert denoisers. *arXiv:2211.01324* (2022).
- [5] BLAU, Y., AND MICHAELI, T. The perception-distortion tradeoff. In *Proc. of CVPR* (2018).
- [6] BRITAIN, B. AI-created images lose U.S. copyrights in test for new technology. Reuters, February 2023.
- [7] CARLINI, N., AND WAGNER, D. Adversarial examples are not easily detected: Bypassing ten detection methods. In *Proc. of AISec* (2017).
- [8] CARLINI, N., AND WAGNER, D. Towards evaluating the robustness of neural networks. In *Proc. of IEEE S&P* (2017).
- [9] CHANDRASEKARAN, V., ET AL. Face-off: Adversarial face obfuscation. *PETS 2* (2020), 369–390.
- [10] CHANGPINYO, S., SHARMA, P., DING, N., AND SORICUT, R. Conceptual 12m: Pushing web-scale image-text pre-training to recognize long-tail visual concepts. In *Proc. of CVPR* (2021).
- [11] CHEN, W., HU, H., SAHARIA, C., AND COHEN, W. W. Re-imagen: Retrieval-augmented text-to-image generator. *arXiv:2209.14491* (2022).
- [12] CHENEY, A. Hundreds of animation guild members join WGA writers on picket line for first time. ABC7 News, May 2023.
- [13] CHEREPANOVA, V., ET AL. Lowkey: Leveraging adversarial attacks to protect social media users from facial recognition. In *ICLR* (2021).
- [14] CIVITAI. <https://civitai.com>, 2022.
- [15] CLARKE, L. When AI can make art – what does it mean for creativity? The Guardian, Nov 2022.
- [16] DAYMA, B. DALLE Mega - Training Journal, 2022.
- [17] DAYMA, B., ET AL. Dalle mini, 2021.
- [18] DEMONTIS, A., ET AL. Why do adversarial attacks transfer? explaining transferability of evasion and poisoning attacks. In *Proc. of USENIX Security* (2019).
- [19] DENG, J., DONG, W., SOCHER, R., LI, L.-J., LI, K., AND FEI-FEI, L. Imagenet: A large-scale hierarchical image database. In *Proc. of CVPR* (2009).
- [20] DIXIT, P. Meet the three artists behind a landmark lawsuit against AI art generators. BuzzFeedNews, January 2023.
- [21] DRYHURST, M. AI art and the problem of consent. ArtReview, January 2023.
- [22] EDWARDS, B. Artists stage mass protest against AI-generated artwork on artstation. Ars Technica, December 2022.
- [23] ELIAÇIK, E. Does ArtStation become PromptStation? Data Economy, Jan. 2023.
- [24] EVTIMOV, I., STURMFELS, P., AND KOHNO, T. Foggysight: A scheme for facial lookup privacy. *PETS 3* (2021).
- [25] FEINMAN, R., CURTIN, R. R., SHINTRE, S., AND GARDNER, A. B. Detecting adversarial samples from artifacts. *arXiv:1703.00410* (2017).
- [26] FLUX, E. What does the rise of AI mean for the future of art? Sydney Morning Herald, Dec. 2022.
- [27] FOX, V. AI art & the ethical concerns of artists. Beautiful Bizarre, March 2023.
- [28] GAL, R., ET AL. An image is worth one word: Personalizing text-to-image generation using textual inversion. *arXiv:2208.01618* (2022).
- [29] GONDIM-RIBEIRO, G., TABACOF, P., AND VALLE, E. Adversarial attacks on variational autoencoders. *arXiv:1806.04646* (2018).
- [30] GROWCOOT, M. Midjourney founder admits to using a ‘hundred million’ images without consent. PetaPixel, Dec 2022.
- [31] HEIDORN, C. Mind-Boggling Midjourney Statistics in 2022. Tokenized, 2022.
- [32] HEIKKILA, M. This artist is dominating AI-generated art. and he’s not happy about it. MIT Technology Review, Sept 2022.
- [33] HOARE, A. Digital Illustration Styles. Illustrators, 2021.
- [34] IVANENKO, N. Midjourney v4: an incredible new version of the AI image generator. Mezha, 2022.
- [35] JOSEPH SAVERI LAW FIRM LLP. Class Action Filed Against Stability AI, Midjourney, and DeviantArt for DMCA Violations, Right of Publicity Violations, Unlawful Competition, Breach of TOS, 2023.
- [36] KARRAS, T., ET AL. Training generative adversarial networks with limited data. In *Proc. of NeurIPS* (2020).
- [37] KAWAR, B., ET AL. Imagic: Text-based real image editing with diffusion models. *Proc. of CVPR* (2023).
- [38] KE, Y., ET AL. Efficient near-duplicate detection and sub-image retrieval. In *Proc. of MM* (2004).
- [39] KOS, J., FISCHER, I., AND SONG, D. Adversarial examples for generative models. In *Proc. of SPW* (2018).
- [40] KURAKIN, A., GOODFELLOW, I., AND BENGIO, S. Adversarial examples in the physical world. *arXiv:1607.02533* (2016).
- [41] KYNKÄÄNNIEMI, T., KARRAS, T., AITTALA, M., AILA, T., AND LEHTINEN, J. The role of imagenet classes in fréchet inception distance. *arXiv:2203.06026* (2022).
- [42] LAIDLAW, C., SINGLA, S., AND FEIZI, S. Perceptual adversarial robustness: Defense against unseen threat models. *arXiv:2006.12655* (2020).
- [43] LEVINE, G. A New Stable Diffusion Plug-In For GIMP & Krita. 80.lv, Sept. 2022.
- [44] LI, K.-L., HUANG, H.-K., TIAN, S.-F., AND XU, W. Improving one-class SVM for anomaly detection. In *Proc. of IEEE CAT* (2003).
- [45] LI, W., ET AL. Object-driven text-to-image synthesis via adversarial training. In *Proc. of CVPR* (2019).
- [46] LI, Y., FAN, H., HU, R., FEICHTENHOFER, C., AND HE, K. Scaling language-image pre-training via masking. *arXiv:2212.00794* (2022).

- [47] LIU, G. The world's smartest artificial intelligence just made its first magazine cover. *Cosmopolitan*, June 2022.
- [48] LIU, Z., LUO, P., WANG, X., AND TANG, X. Deep learning face attributes in the wild. In *Proc. of ICCV* (Dec. 2015).
- [49] LUO, Z., XI, Y., ZHANG, R., AND MA, J. VC-GPT: Visual conditioned GPT for end-to-end generative vision-and-language pre-training. *arXiv:2201.12723* (2022).
- [50] MANSIMOV, E., PARISOTTO, E., BA, J. L., AND SALAKHUTDINOV, R. Generating images from captions with attention. *arXiv:1511.02793* (2015).
- [51] MARX, D. <https://twitter.com/DigThatData/status/1636386617392009224>.
- [52] MENG, C., ET AL. On distillation of guided diffusion models. *arXiv:2210.03142* (2022).
- [53] MIDJOURNEY. Community Showcase, 2022.
- [54] MURPHY, B. P. Is Lensa AI Stealing From Human Art? An Expert Explains The Controversy. *ScienceAlert*, 2022.
- [55] NGUYEN, K. AI is causing student artists to rethink their creative career plans. *KQED Arts*, April 2023.
- [56] NOCEDAL, J., AND WRIGHT, S. *Numerical optimization, series in operations research and financial engineering*. 2006.
- [57] NOVELAI. NovelAI changelog, 2022.
- [58] PAPERNOT, N., MCDANIEL, P., AND GOODFELLOW, I. Transferability in machine learning: From phenomena to black-box attacks using adversarial samples. *arXiv:1605.07277* (2016).
- [59] PEREZ, M. AI Art Generator Cupixel Rakes in 5M From Craft Store JOANN. Built in Boston, August 2022.
- [60] PLUNKETT, L. AI creating 'art' is an ethical and copyright nightmare. *Kotaku*, August 2022.
- [61] POPLI, N. He Used AI to Publish a Children's Book in a Weekend. Artists Are Not Happy About It. *Time*, Dec. 2022.
- [62] RADFORD, A., ET AL. Learning transferable visual models from natural language supervision. In *Proc. of ICML* (2021).
- [63] RADFORD, A., METZ, L., AND CHINTALA, S. Unsupervised representation learning with deep convolutional generative adversarial networks. *arXiv:1511.06434* (2015).
- [64] RADIYA-DIXIT, E., ET AL. Data poisoning won't save you from facial recognition. *arXiv:2106.14851* (2021).
- [65] RAMESH, A., ET AL. Zero-shot text-to-image generation. In *Proc. of ICML* (2021).
- [66] RAMESH, A., ET AL. Hierarchical text-conditional image generation with clip latents. *arXiv:2204.06125* (2022).
- [67] ROMBACH, R., ET AL. High-resolution image synthesis with latent diffusion models. In *Proc. of CVPR* (2022).
- [68] RONY, J., GRANGER, E., PEDERSOLI, M., AND BEN AYED, I. Augmented lagrangian adversarial attacks. In *Proc. of ICCV* (2021).
- [69] ROOSE, K. An A.I.-Generated Picture Won an Art Prize. Artists Aren't Happy. *The New York Times*, Sept. 2022.
- [70] RUIZ, N., ET AL. Dreambooth: Fine tuning text-to-image diffusion models for subject-driven generation. *arxiv:2208.12242* (2022).
- [71] SABOUR, S., CAO, Y., FAGHRI, F., AND FLEET, D. J. Adversarial manipulation of deep representations. *arXiv:1511.05122* (2015).
- [72] SAHARIA, C., ET AL. Palette: Image-to-image diffusion models. In *Proc. of SIGGRAPH* (2022).
- [73] SALEH, B., AND ELGAMMAL, A. Large-scale classification of fine-art paintings: Learning the right metric on the right feature. *arXiv:1505.00855* (2015).
- [74] SALEHI, M., ET AL. Arae: Adversarially robust training of autoencoders improves novelty detection. *Neural Networks* 144 (2021), 726–736.
- [75] SALKOWITZ, R. AI is coming for commercial art jobs. can it be stopped? *Forbes*, Sept 2022.
- [76] SALMAN, H., KHADDAJ, A., LECLERC, G., ILYAS, A., AND MADRY, A. Raising the cost of malicious AI-powered image editing. *arXiv:2302.06588* (2022).
- [77] SCENARIO.GG. AI-generated game assets, 2022.
- [78] SCHUHMANN, C., ET AL. Laion-5b: An open large-scale dataset for training next generation image-text models. *arXiv:2210.08402* (2022).
- [79] SHAN, S., DING, W., WENGER, E., ZHENG, H., AND ZHAO, B. Y. Post-breach recovery: Protection against white-box adversarial examples for leaked dnn models. In *Proc. of ACM CCS* (2022).
- [80] SHAN, S., WENGER, E., WANG, B., LI, B., ZHENG, H., AND ZHAO, B. Y. Gotta catch 'em all: Using honeypots to catch adversarial attacks on neural networks. In *Proc. of ACM CCS* (2020).
- [81] SHAN, S., WENGER, E., ZHANG, J., LI, H., ZHENG, H., AND ZHAO, B. Y. Fawkes: Protecting privacy against unauthorized deep learning models. In *Proc. of USENIX Security* (2020).
- [82] SHARMA, P., DING, N., GOODMAN, S., AND SORICUT, R. Conceptual captions: A cleaned, hypernymed, image alt-text dataset for automatic image captioning. In *Proc. of ACL* (2018).
- [83] STABILITY AI. Stable Diffusion v2.1 and DreamStudio Updates, Dec. 2022.
- [84] SUCIU, O., ET AL. When does machine learning fail? generalized transferability for evasion and poisoning attacks. In *Proc. of USENIX Security* (2018).
- [85] TABACOF, P., TAVARES, J., AND VALLE, E. Adversarial images for variational autoencoders. *arXiv:1612.00155* (2016).
- [86] TAO, M., ET AL. DF-GAN: A simple and effective baseline for text-to-image synthesis. In *Proc. of CVPR* (2022).
- [87] THOMEE, B., ET AL. Yfcc100m: The new data in multimedia research. *CACM* 59, 2 (2016), 64–73.
- [88] TRAN, T. H. Image Apps Like Lensa AI Are Sweeping the Internet, and Stealing From Artists. *Yahoo News*, 2022.
- [89] UPCHURCH, P., ET AL. Deep feature interpolation for image content changes. In *Proc. of CVPR* (2017).
- [90] WANG, B., YAO, Y., SHAN, S., LI, H., VISWANATH, B., ZHENG, H., AND ZHAO, B. Y. Neural cleanse: Identifying and mitigating backdoor attacks in neural networks. In *Proc. of IEEE S&P* (2019).
- [91] WANG, F., AND LIU, H. Understanding the behaviour of contrastive loss. In *Proc. of CVPR* (2021).
- [92] WEATHERBED, J. ArtStation is hiding images protesting AI art on the platform. *The Verge*, 2022.
- [93] WEEKMAN, K. People Have Raised Serious Concerns About The AI Art App That's All Over Your Instagram Feed. *BuzzFeed News*, 2022.
- [94] WEN, Y., ET AL. Hard prompts made easy: Gradient-based discrete optimization for prompt tuning and discovery. *arXiv:2302.03668* (2023).
- [95] WENGER, E., SHAN, S., ZHENG, H., AND ZHAO, B. Y. SoK: Anti-facial recognition technology. In *Proc. of IEEE S&P* (2023).
- [96] WHITE, T. Sampling generative networks. *arXiv:1609.04468* (2016).
- [97] WIGGERS, K. Scenario lands \$6M for its AI platform that generates game art assets. *TechCrunch*, 2022.
- [98] WIGGERS, K. Stability AI, the startup behind Stable Diffusion, raises 101M. *TechCrunch*, 2022.

- [99] XU, T., ET AL. AttnGAN: Fine-grained text to image generation with attentional generative adversarial networks. In *Proc. of CVPR* (2018).
- [100] YANG, S. Why Artists are Fed Up with AI Art. *Fayden Art*, Dec. 2022.
- [101] YOSINSKI, J., CLUNE, J., BENGIO, Y., AND LIPSON, H. How transferable are features in deep neural networks? In *Proc. of NeurIPS* (2014).
- [102] ZHANG, H., AVRITHIS, Y., FURON, T., AND AMSALEG, L. Smooth adversarial examples. *EURASIP Journal on Information Security* (2020), 1–12.
- [103] ZHANG, H., ET AL. StackGAN: Text to photo-realistic image synthesis with stacked generative adversarial networks. In *Proc. of ICCV* (2017).
- [104] ZHANG, K., ET AL. Beyond a Gaussian denoiser: Residual learning of deep CNN for image denoising. *IEEE Trans. on Image Processing* (2017), 3142–3155.
- [105] ZHANG, L. Adversecleaner. <https://github.com/llyasviel/AdverseCleaner>.
- [106] ZHANG, R., ET AL. The unreasonable effectiveness of deep features as a perceptual metric. In *Proc. of CVPR* (2018).
- [107] ZHU, M., PAN, P., CHEN, W., AND YANG, Y. Dm-gan: Dynamic memory generative adversarial networks for text-to-image synthesis. In *Proc. of CVPR* (2019).

A Appendix

A.1 Adapting Existing Cloaking Systems

Here, we consider whether prior image cloaking systems can be adapted to provide protection against art style mimicry. Our results show adapting existing cloaking systems has limited effectiveness for our goals.

Adapting existing cloaking systems. Fawkes [81] generates a cloak on user face images by optimizing the feature space difference between the cloaked image and a target image. The target image is simply a face image of a different person. We adapt Fawkes to anti-mimicry protection by switching the feature extractor from facial recognition to the same one we use for *Glaze*. For the target image used, we assume Fawkes randomly picks an artwork from a different artist. Fawkes uses DSSIM to bound the input perturbation. For a fair comparison, we change Fawkes perturbation from DSSIM to LPIPS, ones used by *Glaze*.

The general design of Lowkey [13] is similar to Fawkes, except Lowkey does not optimize cloak images towards a target in feature space but simply optimizes cloaked images to be different from the original one. We directly apply LowKey for anti-mimicry protection: Lowkey maximizes the cloaked artwork to have a different feature representation from the original artwork.

Photoguard [76] works by minimizing the norm of the image feature vector. It is equivalent to Fawkes when Fawkes selects the zero feature vector as the target for optimization. For anti-mimicry, we adapted Photoguard to minimize the norm of feature representation of the cloaked artwork.

Performance comparison. Figure 22 show Fawkes, Lowkey, and Photoguard have limited effectiveness at protection against mimicry. Out of the three existing systems, Fawkes achieves the best performance with 41.0% artist-rated protection success rate. While we can see small artifacts introduced by Fawkes and Lowkey, they are not sufficient to prevent mimicry. In our tests, we use the same LPIPS perturbation level and the same feature extractor for optimization for all cloaking systems.

A.2 Additional Information on Style Mimicry

Impact of fine-tuning on mimicry success. Figure 21 compares the mimicry performance when a mimic attack fine-tunes on the victim artist’s artwork or directly using a generic model. For artists who are not household names (e.g. iconic artists like Van Gogh), fine-tuning significantly improves mimicry performance. We generate images using text captions containing the artist’s name, e.g., “a river by Nathan Fowkes.”

Details on training parameters. For stable diffusion, we follow the same training parameters as the original paper [67]. We use $5 \cdot 10^{-6}$ learning rate and batch size of 32. For a generation, we follow the default setting using the PNDM sampler and 50 sampling steps. For DALL·E-m, we also follow the same training setup as [17] with a learning rate of $2 \cdot 10^{-5}$ and batch size 32. To generate images, we use the default setting with a condition scale equal to 10.

Impact of selecting random seed. For diffusion-based models (e.g., SD), artwork generation is controlled by a random seed (random noise input at the beginning). Different random seeds lead to very different images, and thus it is common practice for mimics to generate a set of artwork using different seed and select the best artwork. A relevant question is, can a mimic use sheer randomness to generate a plagiarized artwork that succeeds despite *Glaze* protection.

We investigate the impact of random seed selection on mimicry success in the presence of *Glaze*. Given a style-specific model and a given text prompt, the mimic generates 100 plagiarized artworks using different random seeds. Similar to how we calculate CLIP-based genre shift, we then use the CLIP model to identify any artwork that belongs to the same genre as the target artist’s style. The results show that 4.3% of the time, the mimic is able to find at least 1 out of the 100 plagiarized artwork that passed CLIP filtering. While the filtered artwork does belong to the same genre as the artist, we found they tend to have lower image quality. We verify this observation in our user study, and $> 94.1\%$ human artists rated the protection remains successful (i.e. these art pieces failed to mimick the art style). We believe the reason that some plagiarized artwork still shares the same genre as victim style after protection, is that text-to-image models today are still imperfect and often output poor-quality images in rare cases with some random seed.

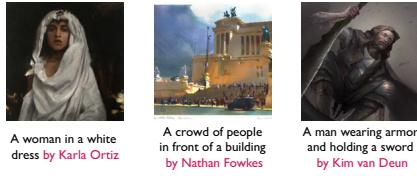


Figure 18. Example data used for fine-tuning, including artwork from different artists and their text captions.

	1K steps	2K steps	3K steps	4K steps
Attempts to mimic artist A				
Attempts to mimic artist B				
Artist-rated PSR	89.4 ± 1.5%	32.9 ± 4.3%	4.3 ± 0.2%	6.4 ± 0.8%

Figure 19. The success of style mimicry when the mimic fine-tunes the model for an increasing number of iterations.

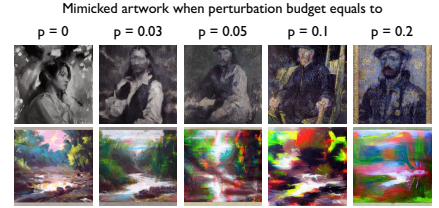


Figure 20. Mimicked artwork when artist uses an increasingly high perturbation budget to protect their original art.



Figure 21. Comparing performance of art mimicry directly using a generic model to that of mimicry on a model that has been fine-tuned on the victim's art pieces. **Column 1-2:** artists' original artwork; **column 3-4:** plagiarized artwork generated from a style-specific model fine-tuned on artist's art; **column 5-6:** plagiarized artwork generated from the generic SD model using the artist's name as prompt.

	No protection	Fawkes	LowKey	PhotoGuard	Ours
Attempts to mimic artist A					
Attempts to mimic artist B					
Artist-rated PSR	4.3 ± 0.2%	41.0 ± 6.4%	37.6 ± 5.5%	11.5 ± 3.2%	93.5 ± 0.6%

Figure 22. Comparing protection levels provided by different cloaking systems, including adapted versions of Fawkes, Lowkey, and Photoguard for style protection. *Glaze* significantly outperforms these adapted alternatives.

A.3 CLIP-based metric

We test CLIP's performance in classifying artwork into the correct art genre. We take 27 historical genres from WikiArt and 13 digital art genres [33] as the candidate labels. We collect a test dataset consisting of 1000 artwork from WikiArt dataset, each containing the ground truth labels from the Wikiart dataset. Then we collect 100 artwork for each of the 13 digital art genres by searching the name of the genre on ArtStation, one of the largest digital art-sharing platforms. We evaluate CLIP performance using top-3 accuracy as many art genres are similar to each other (*e.g.*, impressionism vs fauvism). CLIP achieves 96.4% top-3 accuracy on artwork from WikiArt and 94.2% for artwork from ArtStation.

A.4 Additional Countermeasures

Details on robust training. Here, we give details on the robust training method we used. We follow prior work [74] on robust training of autoencoder models. The mimic first uses *Glaze* to generate a large number of cloaked artwork using

artwork from WikiArt dataset. Given the feature extractor Φ used by mimic's text-to-image model, the mimic trains Φ to minimize the following loss function:

$$\min_{\Phi} \|\Phi(x_{\text{cloaked}}) - \Phi(x_{\text{org}})\|_2^2 \quad (3)$$

where x_{cloaked} and x_{org} is a pair of cloaked and original artworks. This optimization effectively forces Φ to extract the same feature representation for cloaked and original artwork. To prevent the extractor from collapsing (*e.g.*, output zero vectors for all inputs), we regularized the training with the standard VAE reconstruction loss and train the decoder D at the same time. Given the high discrepancy between features of cloaked and original artwork, this training process significantly modifies the internals of Φ as well as the feature space. Thus, the mimic needs to fine-tune the decoder D and generator G on the new robust feature space. We assume the mimic trains Φ for K steps on K different pairs of cloaked/original artwork, and then fine-tunes D and G until convergence.