

STDiff: Spatio-Temporal Diffusion for Continuous Stochastic Video Prediction

Xi Ye*, Guillaume-Alexandre Bilodeau

LITIV, Polytechnique Montréal
xi.ye@polymtl.ca, gabilodeau@polymtl.ca

Abstract

Predicting future frames of a video is challenging because it is difficult to learn the uncertainty of the underlying factors influencing their contents. In this paper, we propose a novel video prediction model, which has infinite-dimensional latent variables over the spatio-temporal domain. Specifically, we first decompose the video motion and content information, then take a neural stochastic differential equation to predict the temporal motion information, and finally, an image diffusion model autoregressively generates the video frame by conditioning on the predicted motion feature and the previous frame. The better expressiveness and stronger stochasticity learning capability of our model lead to state-of-the-art video prediction performances. As well, our model is able to achieve temporal continuous prediction, i.e., predicting in an unsupervised way the future video frames with an arbitrarily high frame rate. Our code is available at <https://github.com/XiYe20/STDiffProject>.

Introduction

Given some observed past frames as input, a video prediction model forecasts plausible future frames, with the aim of mimicking the vision-based precognition ability of humans. Anticipating the future is critical for developing intelligent agents, thus video prediction models have applications in the field of autonomous driving, route planning, and anomaly detection (Lu et al. 2019).

Video frame prediction (VFP) is challenging due to the inherent unpredictability of the future. Theoretically, there are an infinite number of potential future outcomes corresponding to the same past observation. Moreover, the stochasticity increases exponentially as the model predicts towards a more distant future. Early deterministic video prediction models (Finn, Goodfellow, and Levine 2016; Villegas et al. 2017) are incapable of dealing with a multimodal future and thus the prediction tends to be blurry. Subsequently, techniques such as VAE and GANs were introduced for stochastic video prediction (Babaeizadeh et al. 2018; Lee et al. 2018; Denton and Fergus 2018; Castrejon et al. 2019). However, they still fall short in achieving satisfactory results.

Previous works (Castrejon et al. 2019; Wu et al. 2021) empirically observed that the main issue which limits the performance is that the stochastic video prediction model is not expressive enough, specifically, shallow levels of latent variables are incapable of capturing the complex ground-truth latent distribution. Castrejon et al. (2019) introduced hierarchical latent variables into variational RNNs (VRNNs) to mitigate the problem. Recent progress of diffusion-based video prediction models (Voleti et al. 2022; Harvey et al. 2022; Höppe et al. 2022) can also be attributed to the increase of model expressiveness, because diffusion models can be considered as infinitely deep hierarchical VAEs (Huang, Lim, and Courville 2021; Tzen and Raginsky 2019) with a fixed encoder and latent variables across different levels having the same dimensionality as the data.

However, none of the previous models explored extending the stochastic expressiveness along the temporal dimension. Hierarchical VRNNs incorporate randomness at each fixed time step, which is limited by the discrete nature of RNNs. Meanwhile, almost all the previous diffusion-based video prediction models (Voleti et al. 2022; Harvey et al. 2022; Höppe et al. 2022) concatenate a few frames together and learn the distribution of those short video clips as a whole, which ignores an explicit temporal stochasticity estimation between frames. Therefore, there is a need to increase the expressiveness of the stochastic video prediction model over both the spatial and temporal dimensions.

One more issue is that almost all video prediction models can only predict future frames at fixed time steps, ignoring the continuous nature of real-world dynamic scenes. Thus, additional video interpolation models are required to generate videos with a different frame rate. Only a few recent preliminary works have explored this topic (Park et al. 2021; Ye and Bilodeau 2023b), but they are either limited to deterministic predictions (Park et al. 2021) or exhibit low diversity in stochastic predictions (Ye and Bilodeau 2023b). To improve video prediction, a method should be able to do continuous stochastic prediction with a good level of diversity.

In this paper, we propose a novel diffusion-based stochastic video prediction model that addresses the limited expressiveness and the temporal continuous prediction problems for the video prediction task. We propose to increase the expressiveness of the video prediction model by separately learning the temporal and spatial stochasticity. Specifically,

*Corresponding Author.

Copyright © 2024, Association for the Advancement of Artificial Intelligence (www.aaai.org). All rights reserved.

we take the difference images of adjacent past frames as motion information, and those images are fed into a Conv-GRU to extract the initial motion feature of future frames. Given the initial motion feature, a neural stochastic differential equation (SDE) solver predicts the motion feature at an arbitrary future time, which enables continuous temporal prediction. Finally, an image diffusion model conditions on the previous frame and on the motion feature to generate the current frame. Because the diffusion process can also be described by SDE (Song et al. 2021), our model explicitly uses SDEs to describe both the spatial and temporal latent variables, and thus our model is more flexible and expressive than previous ones. Our contributions can be summarized as follows:

- We propose a novel video prediction model with better expressiveness by describing both the spatial and temporal generative process with SDEs;
- Our model attains state-of-the-art (SOTA) FVD and LPIPS score across multiple video prediction datasets;
- Compared to prior approaches, our model significantly enhances the efficiency of generating diverse stochastic predictions;
- To the best of our knowledge, this is the first diffusion-based video prediction model with temporal continuous prediction and with motion content decomposition.

Related Works

Video prediction models can be classified into two types: deterministic or stochastic. Since we aim at addressing uncertainty, we focus on the stochastic models in this section. Most stochastic video prediction models (Villegas et al. 2017; Babaeizadeh et al. 2018; Denton and Fergus 2018; Lee et al. 2018) utilize a variational RNN (VRNN) as the backbone. The performance of VRNN-based models is constrained as there is only one level of latent variables. Hierarchical VAEs are utilized by some works (Castrejon et al. 2019; Chatterjee, Ahuja, and Cherian 2021; Wu et al. 2021) to deal with the stochasticity underfitting problem. FitVid (Babaeizadeh et al. 2021) proposed a carefully designed VRNN architecture to address the underfitting problem. Despite the improvements, the stochasticity of all these models is still characterized by the variance of the spatial latent variables of each frame (Chatterjee, Ahuja, and Cherian 2021).

Some methods, such as MCNet (Villegas et al. 2017), decompose the motion and content of videos with the assumption that different video frames share similar content but with different motion, thus the decomposition of motion and content facilitates the learning. It was applied only spatially. Given its benefit, this strategy inspired us to learn the stochasticity over the temporal and spatial dimension separately. NUQ (Chatterjee, Ahuja, and Cherian 2021) is the only VRNN-based model that considers the temporal predictive uncertainty. It enforces an uncertainty regularization to the vanilla loss function for a better convergence without modification of the VRNN architecture. In contrast, our proposed model employ an SDE to explicitly account for the randomness of the temporal motion.

Given the success of image diffusion models, some works have adapted them for video prediction (Höppe et al. 2022; Voleti et al. 2022; Nikankin, Haim, and Irani 2022; Harvey et al. 2022; Yang, Srivastava, and Mandt 2022). Almost all diffusion-based models are non-autoregressive models that learn to estimate the distribution of a short future video clip. Therefore, they ignore the explicit learning of motion stochasticity and generate multiple future frames with a fixed frame rate all at once. Yang, Srivastava, and Mandt (2022) combined a deterministic autoregressive video prediction model with a diffusion model, where the latter is used to generate stochastic residual change between frames to correct the deterministic prediction. TimeGrad (Rasul et al. 2021) combines an autoregressive model with a diffusion model, but it is only for low dimensional time series data prediction. In contrast, our method combines an autoregressive model with a diffusion model for videos and can make a continuous temporal prediction.

Besides, few recent works investigated continuous video prediction, including Vid-ODE (Park et al. 2021) and NPVP (Ye and Bilodeau 2023b). Vid-ODE (Park et al. 2021) combines neural ODE and Conv-GRU to predict the features of future frames and a CNN decoder is taken to decode the frames in a compositional manner. It achieves temporal continuous predictions, but it is deterministic. NPVP (Ye and Bilodeau 2023b) is a non-autoregressive model that tackles the continuous problem by formulating the video prediction as an attentive neural process. However, the diversity of the stochastic predictions is low because NPVP only takes a single global latent variable for the whole sequence to account for the stochasticity. Our use of SDEs allows us to obtain better diversity.

Background

Neural SDEs

Stochastic differential equations are widely used to describe the dynamics in engineering. Compared with an ordinary differential equation (ODE), a SDE takes into account stochasticity by incorporating randomness into the differential equation. Given an observation X_t at time t , an Itô SDE is formulated as

$$dX_t = \mu(X_t, t)dt + \sigma(X_t, t)dW_t, \quad (1)$$

where W denotes a Wiener process (Brownian motion). The first term and second term of the right hand side are the drift term and diffusion term respectively. We can integrate Eq. 1 by Itô's calculus. In this paper, we use a simpler version of SDE, where μ and σ are independent of t .

Given observations X_t , we can fit an SDE by parameterizing μ and σ with neural networks. This is referred to as neural SDEs (Li et al. 2020; Kidger et al. 2021). Li et al. (2020) generalized the adjoint sensitivity method for ODEs into SDEs for an efficient gradient computation to enable the learning of neural SDEs.

Diffusion Models

There are two types of equivalent diffusion models. The first one is denoising diffusion probabilistic models (DDPM)

(Ho, Jain, and Abbeel 2020; Sohl-Dickstein et al. 2015). DDPM consists of a fixed forward diffusion and a learned reversed denoising process. The forward diffusion process gradually adds noise to the data until it converges to a standard Gaussian. During the reverse process (generative process), a random Gaussian noise is sampled, and a trained neural network is repeatedly applied to denoise and finally converts the random Gaussian noise to a clean data sample. The second type is the score matching-based method (Song et al. 2021), which describes both the forward and reversed diffusion processes as SDEs. Song et al. (2021) studied the connection between the two methods and proved that a DDPM is equivalent to a discretization of a score matching model based on continuous variance preserving SDEs (VP-SDEs). As we also model the motion feature of videos by a SDE in this paper, we present the image diffusion model under the framework of SDEs, specifically, the VP-SDE score matching model.

The forward and reversed processes of the VP-SDE score matching model are given in Eq. 2 and Eq. 3 respectively:

$$dx_t = -\frac{1}{2}\beta(t)x_t dt + \sqrt{\beta(t)}dW_t \quad (2)$$

$$dx_t = [-\frac{1}{2}\beta(t)x_t - \beta(t)\nabla_{x_t}\log q_t(x_t)]dt + \sqrt{\beta(t)}d\bar{W}_t, \quad (3)$$

where $t \sim U[0, T]$ denotes a random diffusion timestep with a maximum value of T , $x_T \sim \mathcal{N}(0, I)$, $x_t \sim q_t(x_t)$ is the perturbed image at diffusion step t , $\beta(t)$ denotes the forward noise schedule, and $\bar{W}_t$ denotes a backward Wiener process. The VP-SDE score matching model is trained by minimizing the following score matching loss:

$$E_t E_{x_0} E_{x_t \sim q_t(x_t|x_0)} \|\mathbf{s}_\theta(x_t, t) - \nabla_{x_t} \log q_t(x_t|x_0)\|^2, \quad (4)$$

where $\mathbf{s}_\theta$ denotes the score estimation neural network, and $q_t(x_t|x_0) = \mathcal{N}(x_t; \gamma_t x_0, \sigma_t^2 I)$ denotes the distribution of perturbed image x_t with $\gamma_t = e^{-\frac{1}{2} \int_0^t \beta(s) ds}$ and $\sigma_t^2 = 1 - e^{-\int_0^t \beta(s) ds}$. Therefore, the forward diffusion process described by Eq. 2 can be achieved by direct re-parameterized sampling, i.e., $x_t = \gamma_t x_0 + \sigma_t \epsilon$, where $\epsilon \sim \mathcal{N}(0, I)$. In this case, we can parameterize $\mathbf{s}_\theta(x_t, t)$ by $-\frac{\epsilon_\theta(x_t, t)}{\sigma_t}$, where ϵ_θ denotes a noise predictor neural network. Then the score matching loss in Eq. 4 can be reduced to an equivalent noise prediction loss as DDPM:

$$E_t E_{x_0} E_{\epsilon \sim \mathcal{N}(0, I)} \frac{1}{\sigma_t^2} \|\epsilon - \epsilon_\theta(x_t, t)\|^2, \quad (5)$$

except that here, the diffusion timestep t is continuous instead of discrete. Finally, we can apply loss weighting $\lambda(t) = \sigma_t^2$ to Eq. 5 for a good perceptual quality, which also further simplifies the loss function to be:

$$E_t E_{x_0} E_{\epsilon \sim \mathcal{N}(0, 1)} \|\epsilon - \epsilon_\theta(x_t, t)\|^2. \quad (6)$$

Under the variational framework, Huang, Lim, and Courville (2021) proved that continuous-time diffusion models, like the VP-SDE score matching model, can be viewed as “the continuous limit of hierarchical VAEs” as purposed by Tzen and Raginsky (2019). In other words, a continuous diffusion model is equivalent to an infinitely deep hierarchical VAE.

Methodology

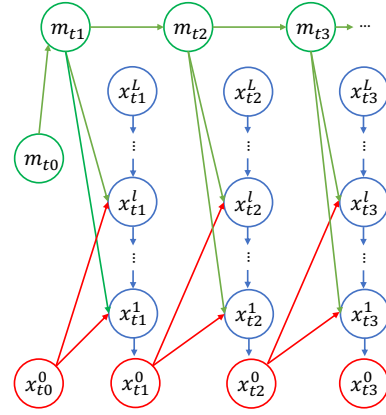


Figure 1: Graphical model for the generation process of STDiff. Green arrows denote the temporal motion connections, and blue arrows denote the connections between latent variables of each frame at timestep t , i.e., the reverse image diffusion process. Red arrows denote the recurrent connection from the previous frame to each level of latent variable in the next time step. m_{t0} denotes the initial motion feature extracted from observed frames. x_{t0}^0 denotes the most recent observed frame.

Note that in this section, t is used to denote video temporal coordinates. We address the video prediction problem that predicts P future frames $\mathbf{x} = \{x_{t_1}^0, x_{t_2}^0, \dots, x_{t_P}^0\}$ given N past frames $\mathbf{c} = \{x_1^0, x_2^0, \dots, x_N^0\}$. During training, t_i is a discrete future frame integer index. However, t_i can be an arbitrary positive real number timestep during inference. The training objective is to maximize $p(\mathbf{x}|\mathbf{c})$. In order to simplify the learning problem, we decompose the motion and content features and we also factorize the joint distribution along the temporal dimension, i.e., autoregressive prediction.

Considering the close relationship between diffusion models and hierarchical VAEs (Huang, Lim, and Courville 2021), we formulate our generative probabilistic model under the framework of VRNNs (Figure 1) as an extended hierarchical recurrent RNN, but with infinite-dimensional latent variables over the spatio-temporal domain. To avoid confusion, we use l instead of t to denote the spatial diffusion timesteps. $l = 0$ denotes the pixel space. Given the initial motion feature m_{t_0} extracted from $\mathbf{c}$ and the most recent observed frame $x_{t_0}^0$, i.e., x_N^0 , we can first factorize the probability as follows with the assumption that the temporal process is Markovian:

$$p(\mathbf{x}|\mathbf{c}) = p(\mathbf{x}|x_{t_0}^0, m_{t_0}) = \prod_{i=1}^P p(x_{t_i}^0|x_{t_{i-1}}^0, m_{t_{i-1}}), \quad (7)$$

where the motion feature m_{t_i} follows a transitional distribution $p(m_{t_i}|m_{t_{i-1}})$ defined by the temporal neural SDE in Eq.14. Following the same technique of conditional DDPM (Ho, Jain, and Abbeel 2020), we can derive the following learning objective:

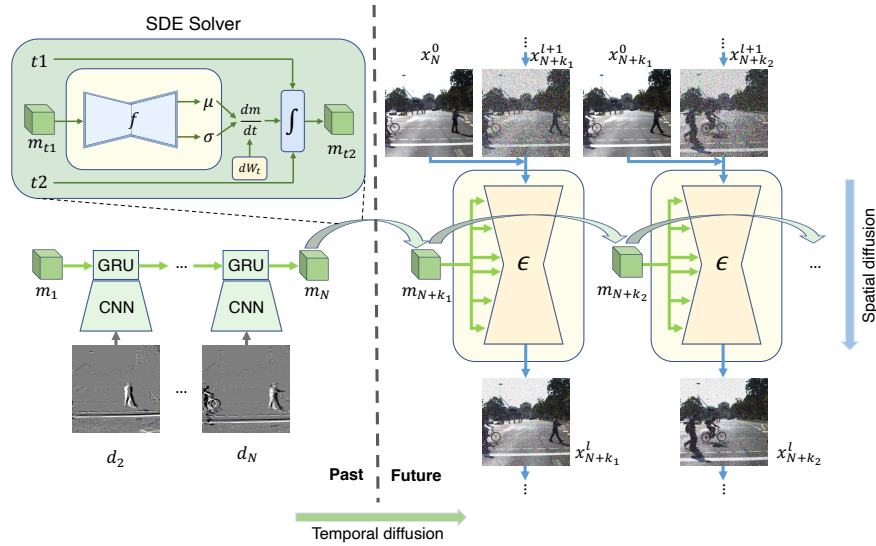


Figure 2: Neural network architecture of STDiff. Difference images of past frames are encoded as motion feature m_N by a ConvGRU for future motion prediction. The curved arrow denotes one step of random future motion prediction (SDE integration), i.e., the temporal motion diffusion process. Detail computation flow of the SDE solver is shown in the top left green box. The conditional image diffusion model recurrently predicts each future frame given motion feature and the previous frame.

$$\mathcal{L}_\theta = \sum_{i=1}^P \mathbb{E}_q[\mathcal{L}_L + \sum_{l>1} \mathcal{L}_{l-1} + \mathcal{L}_0], \quad (8)$$

where each term is defined as follows:

$$\mathcal{L}_L = D_{KL}(q(x_{t_i}^L | x_{t_i}^0) || p_\theta(x_{t_i}^L)), \quad (9)$$

$$\mathcal{L}_{l-1} = D_{KL}(q(x_{t_i}^{l-1} | x_{t_i}^l, x_{t_i}^0) || p_\theta(x_{t_i}^{l-1} | x_{t_i}^l, x_{t_{i-1}}^0, m_{t_i})), \quad (10)$$

$$\mathcal{L}_0 = -\log p_\theta(x_{t_i}^0 | x_{t_i}^1, x_{t_{i-1}}^0, m_{t_i}) \quad (11)$$

By applying the techniques described in (Ho, Jain, and Abbeel 2020), we can simplify the objective in Eq. 8 to be

$$\mathcal{L}_\theta = \sum_{i=1}^P E_l E_{x_{t_i}^0} E_{\epsilon \sim \mathcal{N}(0,1)} \left\| \epsilon - \epsilon_\theta(x_{t_i}^l, x_{t_{i-1}}^0, m_{t_i}, l) \right\|^2. \quad (12)$$

Please refer to the appendix in the arXiv version of this paper (Ye and Bilodeau 2023a) for a detailed derivation of the loss function.

Spatio-Temporal Diffusion (STDiff) Architecture

We call our proposed method STDiff (Spatio-Temporal Diffusion) and its architecture is shown in Figure 2. It consists of two parts, a motion predictor and a recurrent diffusion frame predictor. The motion predictor encodes all the past motion features and predicts future motion features. Given the predicted motion feature at different future time steps and the previous frame, the recurrent diffusion frame predictor generates the desired future frame. The detailed architecture of the two modules is described as follows.

Motion Predictor. We decompose the motion and content feature because the SDEs are naturally designed for dynamic

information modeling, and also it eases the learning burden of the neural SDE. The motion predictor is divided into two parts: 1) a Conv-GRU for past motion encoding and 2) a neural SDE for future motion prediction. Assuming a regular temporal step in the observed past frames, we utilize a Conv-GRU for past motion encoding due to its flexibility and efficiency. In order to achieve the motion and appearance decomposition, we calculate the difference images d_i of adjacent past frames as the input of the Conv-GRU. Given the zero-initialized motion hidden state m_1 and $N-1$ difference images, the Conv-GRU outputs the motion feature m_N for a future prediction.

Given m_N as the initial value of future motion features, we fit the future motion dynamic by a neural SDE, which is equivalent to a learned diffusion process. Specifically, given the motion feature $m_{t_{i-1}}$ at time step t_{i-1} , a small neural network f_θ is taken to parameterize the drift coefficient and diffusion coefficient respectively. Then, the motion feature at t_i is integrated with

$$\mu, \sigma = f_\theta(m_{t_{i-1}}) \quad (13)$$

$$m_{t_i} = m_{t_{i-1}} + \int_{t_{i-1}}^{t_i} \mu dt + \int_{t_{i-1}}^{t_i} \sigma dW_t. \quad (14)$$

In Figure 2, each curved arrow denotes one future motion prediction step, i.e., one integration step for the neural SDE. For better learning of the temporal dynamics, we randomly sample t_i from the future time steps during training (see Algorithm 1 for details about the training).

Recurrent diffusion predictor. At a future time step t_i , given a noisy frame $x_{t_i}^l$ derived from the forward diffusion process, a UNet is trained to predict the noise condition on the previous clean frame $x_{t_{i-1}}^0$ and the motion feature m_{t_i} . In detail, $x_{t_{i-1}}^0$ is concatenated with $x_{t_i}^l$ as the input of the

Models	SMMNIST, $5 \rightarrow 10$		
	FVD↓	SSIM↑	LPIPS↓
SVG-LP (Denton and Fergus 2018)	90.81	0.688	153.0
Hier-VRNN (Castrejon et al. 2019)	57.17	0.760	103.0
MCVD-concat (Voleti et al. 2022)	25.63	<u>0.786</u>	-
MCVD-spatin (Voleti et al. 2022)	<u>23.86</u>	0.780	-
NPVP (Ye and Bilodeau 2023b)	95.69	0.817	188.7
<i>STDiff</i> (ours)	14.93	0.748	<u>146.2</u>

Models	BAIR, $2 \rightarrow 28$		
	FVD↓	SSIM↑	LPIPS↓
SAVP (Lee et al. 2018)	116.4	0.789	63.4
Hier-VRNN (Castrejon et al. 2019)	143.4	0.829	55.0
STMFANet (Jin et al. 2020)	159.6	0.844	93.6
VPTR-NAR (Ye and Bilodeau 2022)	-	0.813	70.0
NPVP (Ye and Bilodeau 2023b)	923.62	<u>0.842</u>	<u>57.43</u>
MCVD-concat (Voleti et al. 2022)	120.6	0.785	70.74
MCVD-spatin (Voleti et al. 2022)	132.1	0.779	75.27
FitVid (Babaeizadeh et al. 2021)	<u>93.6</u>	-	-
<i>STDiff</i> (ours)	88.1	0.818	69.40

Table 1: VFP results on SMMNIST (left) and BAIR (right). Boldface: best results. Underlined: second best results.

Algorithm 1: Training

Input Observed frames: $V_o = \{x_1, x_2, \dots, x_N\}$
 Frames to predict: $V_p = \{x_{N+1}, x_{N+2}, \dots, x_{N+M}\}$
Initialize $\text{ConvGRU}_\theta, f_\theta, \epsilon_\theta$,

- 1: **repeat**
- 2: Initialize motion states: $m_1 = \mathbf{0}$
- 3: **for** $i = 2, \dots, N$ **do**
- 4: $d_i = x_i - x_{i-1}$
- 5: $m_i = \text{ConvGRU}_\theta(d_i, m_{i-1})$
- 6: **end for**
- 7: $\{k_1, \dots, k_P\} = \text{Sort}(\text{RandomChoice}(\{1, \dots, M\}))$
- 8: $\{t_0, t_1, \dots, t_P\} = \{N, N + k_1, \dots, N + k_P\}$
- 9: **for** $i = 1, \dots, P$ **do**
- 10: $m_{t_i} = \text{SDESolver}(f_\theta, m_{t_{i-1}}, (t_{i-1}, t_i))$
- 11: $l \sim \text{Uniform}(0, L), \epsilon \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$
- 12: $x_{t_i}^l = \gamma_l x_{t_i}^0 + \sigma_l \epsilon$
- 13: Take gradient descent step on
- 14: $\nabla_\theta \|\epsilon - \epsilon_\theta(x_{t_i}^l, x_{t_{i-1}}^0, m_{t_i}, l)\|^2$
- 15: **end for**
- 16: **until** convergence

UNet, and m_{t_i} is fed into each ResNet block of the UNet by the manner of spatially-adaptive denormalization (Park et al. 2019). Please see Algorithm 1 for training details.

The inference process of the recurrent diffusion predictor is almost the same as the prediction process, with the distinction that the model generates the previous clean frame $x_{t_{i-1}}^0$ during the previous time step, and subsequently feeds it back into the UNet input for forecasting the next frame.

Experiments

We evaluated the performance of the proposed *STDiff* model on KITTI (Geiger et al. 2013), Cityscapes (Cordts et al. 2016), KTH (Schuldt, Laptev, and Caputo 2004), BAIR (Ebert et al. 2017) and stochastic moving MNIST (SMMNIST) (Denton and Fergus 2018) datasets. KITTI and Cityscapes are high-resolution traffic video datasets used to evaluate our model real-world application capabilities. The KTH dataset consists of different human motion videos, BAIR includes videos of a randomly moving robot arm that pushes different objects. BAIR and SMMNIST pose greater challenges due to their higher levels of stochasticity com-

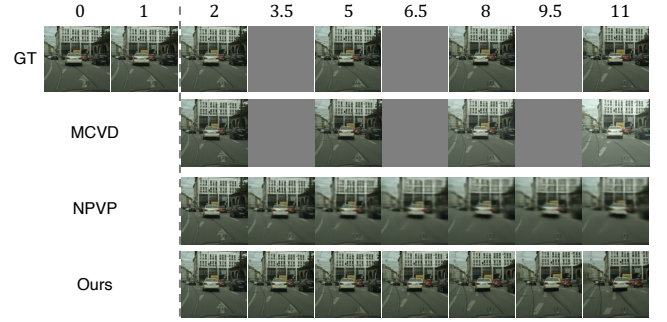


Figure 3: Prediction examples on the Cityscapes dataset for MCVD (Voleti et al. 2022), NPVP (Ye and Bilodeau 2023b), and our model. Gray frames indicate non-existent or unpredictable frames. MCVD exhibits issues with brightness changes and lacks the ability of continuous predictions. NPVP predictions are noticeably more blurry than the other two models.

pared to the three others. The number of past frames and future frames to predict is determined according to the experimental protocols of previous works (see appendix in the arXiv version of this paper for more details (Ye and Bilodeau 2023a)). All models are trained with 4 NVIDIA V100 Volta GPU (32G memory).

Frame Prediction Performance

We firstly tested *STDiff* given the same protocol as previous work, i.e., report the best SSIM, LPIPS out of multiple different random predictions, together with the FVD score. Note that we only sample 10 different random predictions for each test example as for MCVD, instead of 100 different predictions for all other previous stochastic models. The results are presented in Table 1 and Table 2.

In Table 1, we observe that *STDiff* achieves the new SOTA for the FVD score on SMMNIST, and also the second best LPIPS score. For BAIR, *STDiff* also outperforms all previous work in terms of FVD score. The results show that the random predictions of *STDiff* have a better temporal coherence and match the ground-truth distribution much better, i.e., the predictions are more natural and plausible.

The datasets presented in Table 2 contains less stochas-

Models	KITTI, 4 → 5			Models	Cityscapes, 2 → 28		
	FVD ↓	SSIM ↑	LPIPS ↓		FVD ↓	SSIM ↑	LPIPS ↓
PredNet (Lotter, Kreiman, and Cox 2017)	-	0.48	629.5	SVG-LP (Denton and Fergus 2018)	1300.26	0.574	549.0
Voxel Flow (Liu et al. 2017)	-	0.43	415.9	VRNN 1L (Castrejon et al. 2019)	682.08	0.609	304.0
SADM (Bei, Yang, and Soatto 2021)	-	0.65	311.6	Hier-VRNN (Castrejon et al. 2019)	567.51	0.628	264.0
Wu et al. (Wu, Wen, and Chen 2022)	-	0.61	263.5	GHVAEs (Wu et al. 2021)	418.00	<u>0.740</u>	194.0
DMVFN (Hu et al. 2023)	-	0.71	<u>260.5</u>	MCVD-concat (Voleti et al. 2022)	<u>141.31</u>	0.690	112.0
NPVP (Ye and Bilodeau 2023b)	<u>134.69</u>	0.66	279.0	NPVP (Ye and Bilodeau 2023b)	768.04	0.744	183.2
STDiff (ours)	51.39	0.54	114.6	STDiff (ours)	107.31	0.658	<u>136.26</u>

Table 2: VFP results on KITTI (left) and Cityscapes (right). Boldface: best results. Underlined: second best results.

ticity compared with SMMNIST and BAIR. For the KITTI dataset, STDiff outperforms previous work with a significant margin in terms of LPIPS and FVD. STDiff also achieves the best FVD score and second-best LPIPS score on the Cityscapes dataset. Some prediction examples of Cityscapes dataset are shown in Figure 3. Our model avoids the brightness changing issue seen in MCVD, and our predictions are sharper and more realistic compared to the NPVP model.

In general, the evaluation results show that STDiff has a better performance than previous deterministic and stochastic models in terms of either FVD or LPIPS. It is desirable because it is well known that the quality assessments produced by FVD and LPIPS are more aligned with human assessment than SSIM or PSNR (Unterthiner et al. 2019; Zhang et al. 2018).

Stochastic Video Prediction Diversity

In Wang et al. (2022), generated images diversity is evaluated as LPIPS distance between different generated images. A greater LPIPS value between two generated images indicates increased dissimilarity in terms of content and structure. Likewise, we quantify the video diversity as average frame-by-frame LPIPS between different predicted video clips given the same past frames. To distinguish it from the standard LPIPS score between predictions and ground-truth, this diversity measure is termed as inter-LPIPS (iLPIPS). The iLPIPS results are listed in Table 3. All the iLPIPS values are reported at a 10^{-3} scale. Since iLPIPS does not incorporate the ground-truth, we use the Frechet Video Distance (FVD) as a supplementary metric. FVD ensures that randomly generated predictions exhibit satisfactory visual quality and that their distribution closely approximates the ground-truth distribution, i.e., the generated random predictions are plausible.

For all methods, we sampled 10 different random predictions for each test example to calculate the iLPIPS score. For KTH, we evaluated on the whole test set. For SMMNIST, we evaluated the results on 256 different test examples, similarly to MCVD (Voleti et al. 2022). Besides, we take the same number of reverse diffusion sampling steps as MCVD, which is 100. Increasing the diffusion sampling steps would further improve the performance. NPVP (Ye and Bilodeau 2023b) and Hier-VRNN (Castrejon et al. 2019) are trained based on their official code and MCVD (Voleti et al. 2022) is tested with their officially released trained models.

Comparing with the neural process-based NPVP method,

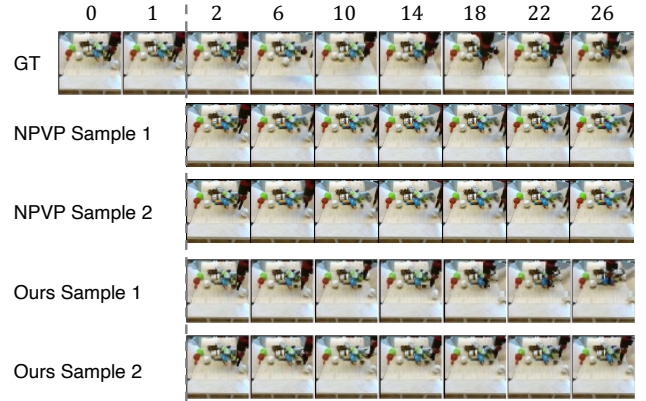


Figure 4: Random predictions for BAIR dataset by our model and NPVP (Ye and Bilodeau 2023b).

The iLPIPS score of our STDiff model is more than 50 times bigger on KTH and twice bigger on SMMNIST. NPVP has a SOTA frame-by-frame SSIM or LPIPS performance, but the generated future frames lack diversity as assessed by iLPIPS. A visual examination of the results of NPVP validates this. We believe this is explained by the fact that NPVP only uses a single layer of latent variable for the VAE and uses a single global latent variable to account for the randomness of the whole sequence. Thus, NPVP has a very limited expressiveness for stochastic modeling. Figure 4 presents several uncured predicted examples, demonstrating the better plausibility and greater diversity of our predictions compared to NPVP. Notably, NPVP tends to move the robot arm outside the field of view to minimize MSE.

We also compare our model with Hier-VRNN that has 10 layers of latent variables. STDiff outperforms Hier-VRNN by a large margin on both datasets, despite Hier-VRNN having a larger model size. Compared with the SOTA diffusion-based model MCVD, STDiff outperforms it on SMMNIST and obtains a comparable iLPIPS for KTH. STDiff accomplishes this with a smaller model that has 100M less parameters. We believe that MCVD is not as efficient as STDiff for stochastic prediction because it learns the randomness of a whole video clip at once, similarly to NPVP, which requires a bigger model and also ignores an explicit temporal stochasticity learning. These results indicate that the good performance of our proposed method is not only attributed

Models	#Parameters	KTH		SMMNIST	
		$10 \rightarrow 20$		$5 \rightarrow 10$	
		iLPIPS $\uparrow$	FVD $\downarrow$	iLPIPS $\uparrow$	FVD $\downarrow$
NPVP (Ye and Bilodeau 2023b)	122.68M	0.46	93.49	67.83	51.66
Hier-VRNN (Castrejon et al. 2019)	260.68M	1.22	278.83	6.72	22.65
MCVD(Voleti et al. 2022)	328.60M	26.04	93.38	<u>123.93</u>	<u>21.53</u>
STDiff-ODE	201.91M	15.60	<u>90.68</u>	91.07	46.62
<i>STDiff</i> (ours)	204.28M	<u>25.08</u>	89.67	136.27	14.93

Table 3: Stochastic video prediction diversity on KTH and SMMNIST datasets. A bigger iLPIPS denotes larger diversity, any deterministic prediction model has a iLPIPS of 0. Smaller FVD indicated better visual quality and more plausible predictions. Boldface: best results. Underlined: second best results.

Models	BAIR				KITTI			
	$2 \rightarrow 28, 2 \times fps$				$4 \rightarrow 5, 2 \times fps$			
	FVD $\downarrow$	SSIM $\uparrow$	LPIPS $\downarrow$	iLPIPS $\uparrow$	FVD $\downarrow$	SSIM $\uparrow$	LPIPS $\downarrow$	iLPIPS $\uparrow$
Vid-ODE (Park et al. 2021)	2948.82	0.310	322.58	0	615.98	0.23	591.54	0
NPVP (Ye and Bilodeau 2023b)	1159.14	0.795	58.71	<u>1.06</u>	<u>248.82</u>	0.56	<u>313.94</u>	<u>6.08</u>
<i>STDiff</i> (ours)	122.85	0.808	<u>75.37</u>	109.33	78.39	<u>0.47</u>	160.92	136.94

Table 4: Continuous VFP results on BAIR and KITTI. Boldface: best results. Underlined: second best results.

to the powerful image diffusion model, but also from our novel SDE-based recurrent motion predictor.

In order to further investigate the effectiveness of our stochastic motion predictor, we implemented an ODE version of STDiff, which has the same architecture as STDiff, except that the motion predictor is ODE-based instead of SDE-based. We observe that STDiff has almost 1.5 times more diversity than STDiff-ODE on both datasets in terms of iLPIPS. Notably, our STDiff also achieves the best FVD on both datasets, highlighting that the prediction of STDiff, thanks to the use of an SDE, has large diversity and good visual quality simultaneously.

We can draw the conclusion that our STDiff has a better stochasticity modeling performance than previous stochastic models. The comparison with NPVP and Hier-VRNN validates the motivation that increasing the layers of latent variable for Hierarchical VAE is beneficial. And the comparison with STDiff-ODE and MCVD experimentally proves our claim that explicitly temporal stochasticity learning is also critical for a better diversity in future predictions.

Continuous Prediction

We summarize the continuous prediction performance of three continuous prediction models in Table 4. MCVD is not included because it cannot conduct temporal continuous prediction. For this evaluation, we downsampled two datasets to 0.5 frame rate for training, then make the models predict with $2\times$ frame rate during test. This way, we get access to ground-truth high-frame rate test videos for metric calculations. The stochastic NPVP and STDiff predict 10 different random trajectories for each test example. In addition to iLPIPS and FVD, we also use the traditional evaluation protocol to report the best SSIM and LPIPS scores out of 10 different random predictions.

STDiff outperforms Vid-ODE by a large margin in terms of all metrics. In addition to be deterministic and not decom-

posing the motion and content as we do, the capacity of Vid-ODE is too small. Indeed, performance of Vid-ODE on the original KTH dataset is not bad (Park et al. 2021), but down-sampling largely increases the difficulty as the model needs to predict motion with a larger temporal gap. For the big, realistic and high resolution KITTI dataset, Vid-ODE fails to achieve reasonable performance mainly due to its limited model size. Examination of visual examples on BAIR shows that Vid-ODE cannot predict the stochastic motion of the robot arm and the image quality quickly degrades because of a large accumulated error.

As shown in Table 4, for both datasets, STDiff achieves comparable or better performance in terms of SSIM and LPIPS compared to NPVP. However, we observe a significant performance gap in terms of FVD and iLPIPS with STDiff performing much better, especially for the dataset with more stochasticity, i.e., BAIR. This indicates that the capacity of randomness modeling also influences the performance of continuous prediction. Continuous prediction example on Cityscapes are shown in Figure 3, STDiff is able to predict frames at non-existent training dataset coordinates (e.g., 3.5, 6.5, and 9.5). Remarkably, STDiff holds the theoretical potential to predict videos at arbitrary frame rates.

Conclusion

In this paper, we propose a novel temporal continuous stochastic video prediction model. Specifically, we model both the spatial and temporal generative process as SDEs by integrating a SDE-based temporal motion predictor with a recurrent diffusion predictor, which greatly increases the stochastic expressiveness and also enables temporal continuous prediction. In this way, our model is able to predict future frames with an arbitrary frame rate and greater diversity. Our model reaches the SOTA in terms of FVD, LPIPS, and iLPIPS on multiple datasets.

Acknowledgements

We thank FRQ-NT and REPARTI for the support of this research via the strategic cluster program.

References

- Babaeizadeh, M.; Finn, C.; Erhan, D.; Campbell, R.; and Levine, S. 2018. Stochastic variational video prediction. In *ICLR*.
- Babaeizadeh, M.; Saffar, M. T.; Nair, S.; Levine, S.; Finn, C.; and Erhan, D. 2021. FitVid: Overfitting in Pixel-Level Video Prediction. ArXiv:2106.13195 [cs].
- Bei, X.; Yang, Y.; and Soatto, S. 2021. Learning Semantic-Aware Dynamics for Video Prediction. In *CVPR*.
- Castrejon, L.; Ballas, N.; Courville, A.; and A. 2019. Improved Conditional VRNNs for Video Prediction. In *ICCV*.
- Chatterjee, M.; Ahuja, N.; and Cherian, A. 2021. A Hierarchical Variational Neural Uncertainty Model for Stochastic Video Prediction. 9751–9761.
- Cordts, M.; Omran, M.; Ramos, S.; Rehfeld, T.; Enzweiler, M.; Benenson, R.; Franke, U.; Roth, S.; and Schiele, B. 2016. The Cityscapes Dataset for Semantic Urban Scene Understanding. 3213–3223.
- Denton, E.; and Fergus, R. 2018. Stochastic Video Generation with a Learned Prior. In *ICML*.
- Ebert, F.; Finn, C.; Lee, A. X.; and Levine, S. 2017. Self-supervised visual planning with temporal skip connections. In *CoRL*.
- Finn, C.; Goodfellow, I.; and Levine, S. 2016. Unsupervised learning for physical interaction through video prediction. In *NIPS*, 64–72.
- Geiger, A.; Lenz, P.; Stiller, C.; and Urtasun, R. 2013. Vision meets robotics: The KITTI dataset. *The International Journal of Robotics Research*, 32(11): 1231–1237.
- Harvey, W.; Naderiparizi, S.; Masrani, V.; Weilbach, C.; and Wood, F. 2022. Flexible Diffusion Modeling of Long Videos.
- Ho, J.; Jain, A.; and Abbeel, P. 2020. Denoising Diffusion Probabilistic Models. In *NeurIPS*, volume 33, 6840–6851.
- Hu, X.; Huang, Z.; Huang, A.; Xu, J.; and Zhou, S. 2023. A Dynamic Multi-Scale Voxel Flow Network for Video Prediction. 6121–6131.
- Huang, C.-W.; Lim, J. H.; and Courville, A. C. 2021. A Variational Perspective on Diffusion-Based Generative Models and Score Matching. In *Advances in Neural Information Processing Systems*, volume 34, 22863–22876. Curran Associates, Inc.
- Höppe, T.; Mehrjou, A.; Bauer, S.; Nielsen, D.; and Dittadi, A. 2022. Diffusion Models for Video Prediction and Infilling.
- Jin, B.; Hu, Y.; Tang, Q.; Niu, J.; Shi, Z.; Han, Y.; and Li, X. 2020. Exploring Spatial-Temporal Multi-Frequency Analysis for High-Fidelity and Temporal-Consistency Video Prediction. In *CVPR*.
- Kidger, P.; Foster, J.; Li, X.; and Lyons, T. J. 2021. Neural SDEs as Infinite-Dimensional GANs. In *Proceedings of the 38th International Conference on Machine Learning*, 5453–5463. PMLR. ISSN: 2640-3498.
- Lee, A. X.; Zhang, R.; Ebert, F.; Abbeel, P.; Finn, C.; and Levine, S. 2018. Stochastic Adversarial Video Prediction.
- Li, X.; Wong, T.-K. L.; Chen, R. T. Q.; and Duvenaud, D. 2020. Scalable Gradients for Stochastic Differential Equations. In *Proceedings of the Twenty Third International Conference on Artificial Intelligence and Statistics*, 3870–3882. PMLR. ISSN: 2640-3498.
- Liu, Z.; Yeh, R. A.; Tang, X.; Liu, Y.; and Agarwala, A. 2017. Video Frame Synthesis Using Deep Voxel Flow. In *International Conference on Computer Vision (ICCV)*, 4473–4481.
- Lotter, W.; Kreiman, G.; and Cox, D. 2017. Deep predictive coding networks for video prediction and unsupervised learning. In *ICLR*, 1–18.
- Lu, Y.; Kumar, K. M.; Nabavi, S. s.; and Wang, Y. 2019. Future Frame Prediction Using Convolutional VRNN for Anomaly Detection. In *AVSS*.
- Nikankin, Y.; Haim, N.; and Irani, M. 2022. SinFusion: Training Diffusion Models on a Single Image or Video. ArXiv:2211.11743 [cs].
- Park, S.; Kim, K.; Lee, J.; Choo, J.; Lee, J.; Kim, S.; and Choi, E. 2021. Vid-ODE: Continuous-Time Video Generation with Neural Ordinary Differential Equation. In *AAAI*.
- Park, T.; Liu, M.-Y.; Wang, T.-C.; and Zhu, J.-Y. 2019. Semantic Image Synthesis With Spatially-Adaptive Normalization. 2337–2346.
- Rasul, K.; Seward, C.; Schuster, I.; and Vollgraf, R. 2021. Autoregressive Denoising Diffusion Models for Multivariate Probabilistic Time Series Forecasting. In *Proceedings of the 38th International Conference on Machine Learning*, 8857–8868. PMLR. ISSN: 2640-3498.
- Schuldt, C.; Laptev, I.; and Caputo, B. 2004. Recognizing human actions: a local SVM approach. In *ICPR*.
- Sohl-Dickstein, J.; Weiss, E.; Maheswaranathan, N.; and Ganguli, S. 2015. Deep Unsupervised Learning using Nonequilibrium Thermodynamics. In *Proceedings of the 32nd International Conference on Machine Learning*, 2256–2265. PMLR. ISSN: 1938-7228.
- Song, Y.; Sohl-Dickstein, J.; Kingma, D. P.; Kumar, A.; Ermon, S.; and Poole, B. 2021. Score-Based Generative Modeling through Stochastic Differential Equations.
- Tzen, B.; and Raginsky, M. 2019. Neural Stochastic Differential Equations: Deep Latent Gaussian Models in the Diffusion Limit. ArXiv:1905.09883 [cs, stat].
- Unterthiner, T.; Steenkiste, S. v.; Kurach, K.; Marinier, R.; Michalski, M.; and Gelly, S. 2019. FVD: A new Metric for Video Generation. In *ICLR Workshop*.
- Villegas, R.; Yang, J.; Hong, S.; Lin, X.; and Lee, H. 2017. Decomposing motion and content for natural video sequence prediction. In *ICLR*.

- Voleti, V.; Jolicoeur-Martineau, A.; Pal, and Christopher. 2022. Masked Conditional Video Diffusion for Prediction, Generation, and Interpolation. In *Advances in Neural Information Processing Systems*.
- Wang, W.; Bao, J.; Zhou, W.; Chen, D.; Chen, D.; Yuan, L.; and Li, H. 2022. SinDiffusion: Learning a Diffusion Model from a Single Natural Image. ArXiv:2211.12445 [cs].
- Wu, B.; Nair, S.; Martin-Martin, R.; Fei-Fei, L.; and Finn, C. 2021. Greedy Hierarchical Variational Autoencoders for Large-Scale Video Prediction. In *CVPR*, 2318–2328.
- Wu, Y.; Wen, Q.; and Chen, Q. 2022. Optimizing Video Prediction via Video Frame Interpolation. 17814–17823.
- Yang, R.; Srivastava, P.; and Mandt, S. 2022. Diffusion Probabilistic Modeling for Video Generation. ArXiv:2203.09481 [cs, stat].
- Ye, X.; and Bilodeau, G.-A. 2022. VPTR: Efficient Transformers for Video Prediction. In *ICPR*.
- Ye, X.; and Bilodeau, G.-A. 2023a. STDiff: Spatio-temporal Diffusion for Continuous Stochastic Video Prediction. arXiv:2312.06486.
- Ye, X.; and Bilodeau, G.-A. 2023b. A Unified Model for Continuous Conditional Video Prediction. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR) Workshops*, 3603–3612.
- Zhang, R.; Isola, P.; Efros, A. A.; Shechtman, E.; and Wang, O. 2018. The unreasonable effectiveness of deep features as a perceptual metric. In *CVPR*, 586–595.