

BRL Seminar

(2024. 03. 05. Tue)

Diffusion Model with Time Series Data 2

통합과정 8학기 이승한

Contents

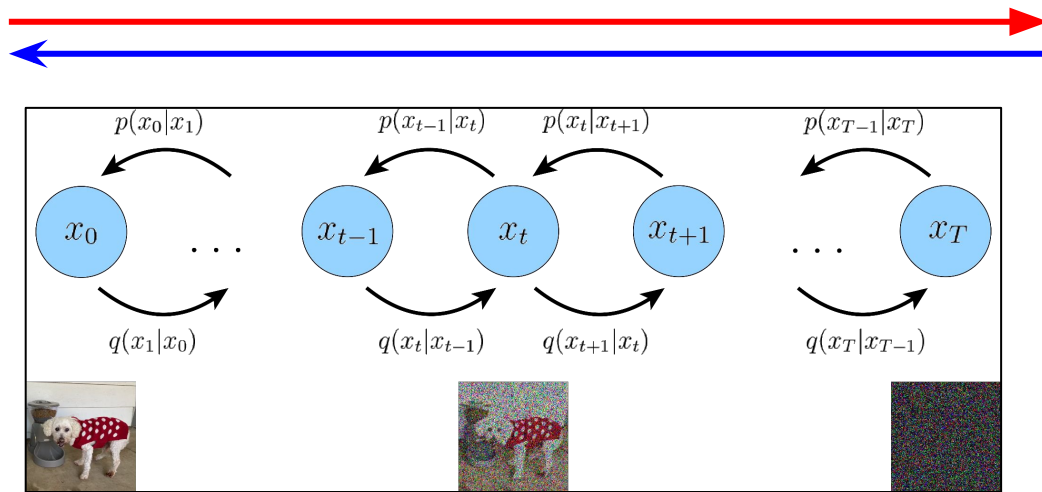
1. Preliminaries: **Time-series Diffusion Model**
2. MG-TSD: **Multi-Granularity** Time Series Diffusion Models with Guided Learning Process (ICLR 2024)
3. **Multi-resolution** Diffusion Model for Time-Series Forecasting (ICLR 2024)

1. Preliminaries: Time-series Diffusion Model

1. Preliminaries: Time-series Diffusion Model

1-1. Diffusion Model

- **Forward** Process: **Add** Noise
- **Backward** Process: **Remove** Noise



1. Preliminaries: Time-series Diffusion Model

1-1. Diffusion Model

- **Forward** Process: **Add** Noise

$$q(x_t | x_{t-1}) = N(x_t; \sqrt{1 - \beta_t}x_{t-1}, \beta_t I)$$

- **Backward** Process: **Remove** Noise

$$p_\theta(x_{t-1} | x_t) = N(x_{t-1}; \mu_\theta(x_t, t), \Sigma_\theta(x_t, t))$$

β_t controls the strength of the noise → **Noise Scheduling**

ex) DDPM: Linear Scheduling $T = 1000, \beta_0 = 0.0001, \beta_{1000} = 0.02$

1. Preliminaries: Time-series Diffusion Model

1-2. Time-series Diffusion Model: TimeGrad (ICML 2021)

- Diffusion model for **TS forecasting**
- Conditional diffusion model
- **“condition = past information”**

$$\prod_{t=t_0}^T p_{\theta}(\mathbf{x}_t^0 \mid \mathbf{h}_{t-1})$$

=

$$\mathbf{h}_t = \text{RNN}_{\theta}(\text{concat}(\mathbf{x}_t^0, \mathbf{c}_t), \mathbf{h}_{t-1})$$

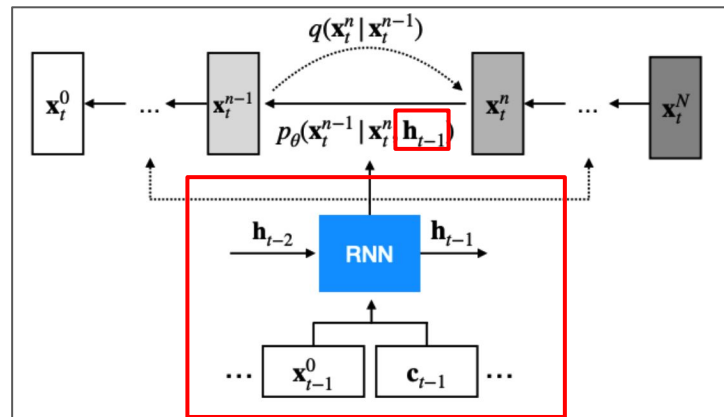
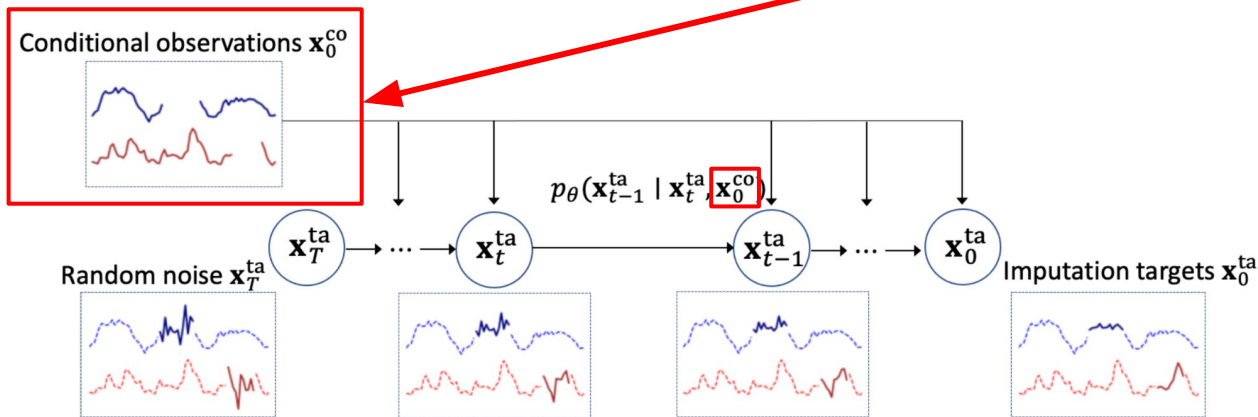


Figure 1. TimeGrad schematic: an RNN conditioned diffusion probabilistic model at some time $t - 1$ depicting the fixed forward process that adds Gaussian noise and the learned reverse processes.

1. Preliminaries: Time-series Diffusion Model

1-2. Time-series Diffusion Model: CSDI (NeurIPS 2021)

- Diffusion model for **TS imputation**
- Conditional diffusion model, where “condition = observed values”



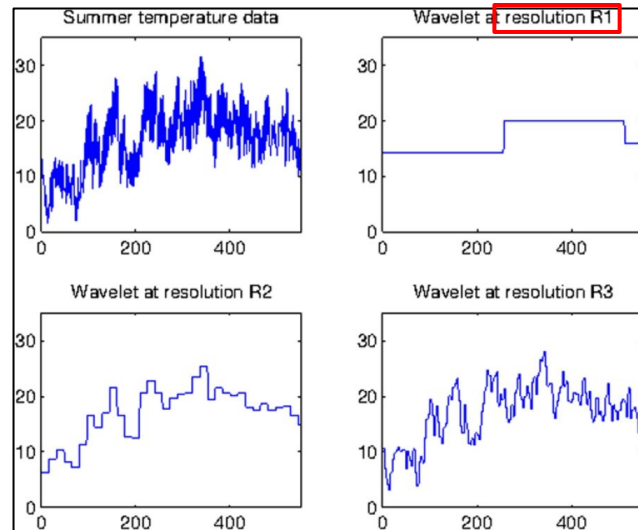
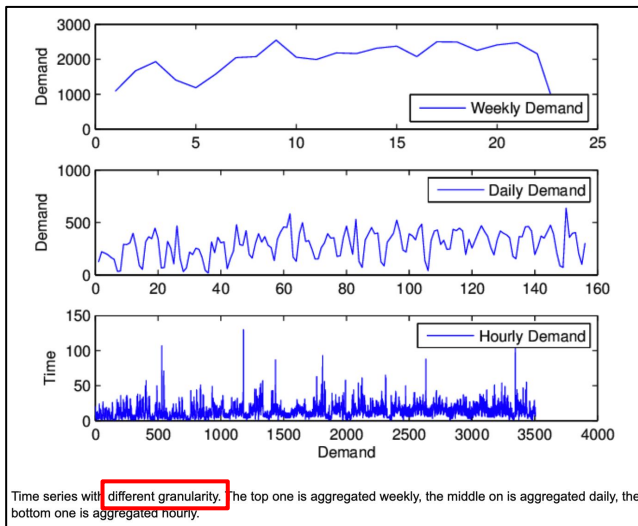
1. Preliminaries: Time-series Diffusion Model

1-3. Multi-granularity & Multi-resolution

TS domain에서 아래의 개념들은 비슷한 맥락

- Granularity: 데이터 분할의 정도
- Resolution: 해상도

- Multi-granularity
- Multi-resolution
- Multi-level
- Hierarchy
- Global & Local
- (Decomposition)



1. Preliminaries: Time-series Diffusion Model

1-3. Multi-granularity & Multi-resolution

TS domain에서 아래의 개념들은 비슷한 맥락

- Multi-granularity
- Multi-resolution
- Multi-level
- Hierarchy
- Global & Local

시계열 패턴을 **coarse & fine** 구분하여 포착

- (Decomposition) ➡ trend를 거시적 & seasonality를 미시적 관점에서 보기도 함!

1. Preliminaries: Time-series Diffusion Model

1-3. Multi-granularity & Multi-resolution

ICLR 2024에 accept 된 4편의 **TS Diffusion Paper**

TS에서의 “Interpretable”
= 99% TS decomposition

★	2024	MG-TSD	MG-TSD: Multi-Granularity Time Series Diffusion Models with Guided Learning Process	ICLR
★	2024	mr-Diff	Multi-resolution Diffusion Models for Time-Series Forecasting	ICLR
	2024	Diffusion-TS	Diffusion-TS: Interpretable Diffusion for General Time Series Generation	ICLR
	2024	TMDM	Transformer-Modulated Diffusion Models for Probabilistic Multivariate Time Series For	ICLR

multi-resolution을 잡아내겠다는 **Goal**은 동일.

이를 포착하는 방법은 알고리즘 별로 상이.

→ 현재 진행 중인 연구에서 제안하는 방법론 또한 **multi-resolution**을 잡아내고자 함.

2. MG-TSD: **Multi-Granularity** Time Series Diffusion

Models with Guided Learning Process (ICLR 2024)

2. MG-TSD

MG-TSD (Multi-Granularity TS Diffusion)

이전까지의 TS Diffusion 연구에서는,
시계열을 분해/계층화 하는 접근이
없었음!

- Model that considers **granularity level** within TS
- Utilize “**coarse-grained data**” across various granularity levels
- MG-TSD
 - (1) MG-TSD architecture
 - (2) Guided diffusion process module

TS Notation: $\mathbf{X}^{(1)} = [\mathbf{x}_1^1, \dots, \mathbf{x}_t^1, \dots, \mathbf{x}_T^1]$, where $\mathbf{x}_t \in \mathbb{R}^D$

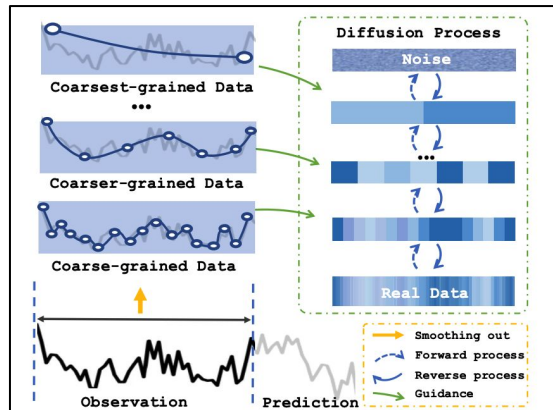


Figure 1: The process of smoothing data from finest-grained to coarsest-grained naturally aligns with the diffusion process.

2. MG-TSD

$$c = f(x_{0:t})$$

$$\hat{x}_{t+1} = g(c, n)$$

- f : Encoder
- g : Diffusion model
- Time step (t) = $1 \sim T$
- Diffusion step (n) = $1 \sim N$

Output = $t+1$ 시점 (Autoregressive)

Input = $1 \sim t$ 시점

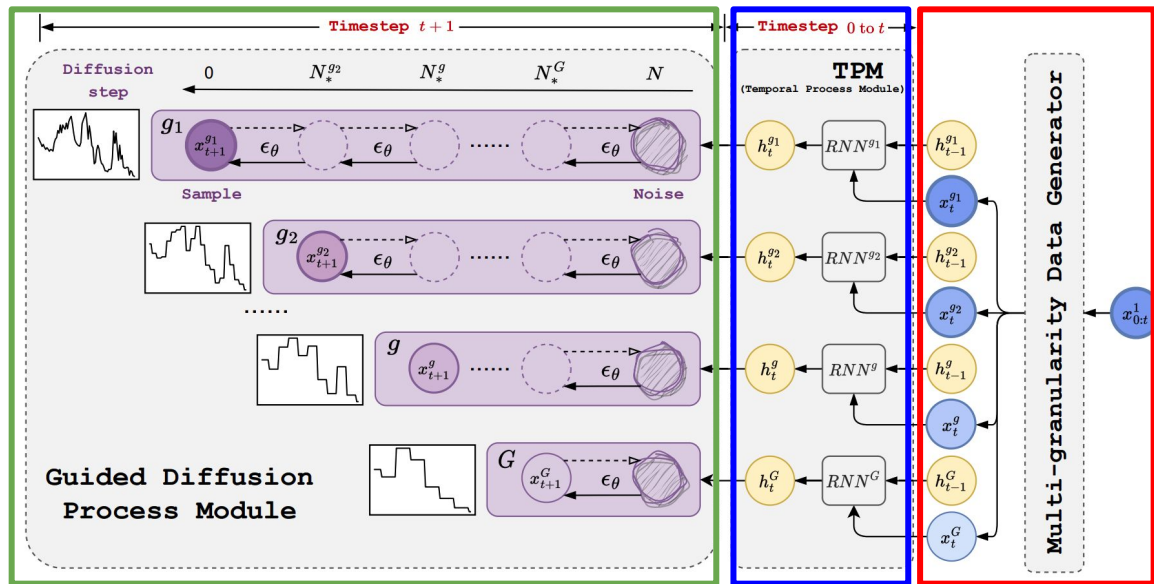


Figure 2: Overview of the Multi-Granularity Time Series Diffusion (MG-TSD) model, consisting of three key modules: **Multi-granularity Data Generator**, **Temporal Process Module (TPM)**, and **Guided Diffusion Process Module** for time series forecasting at a specific granularity level.

잠재 공간 상 Diffusion

<Step 3>

임베딩 (RNN)

<Step 2>

Multi-granularity 데이터

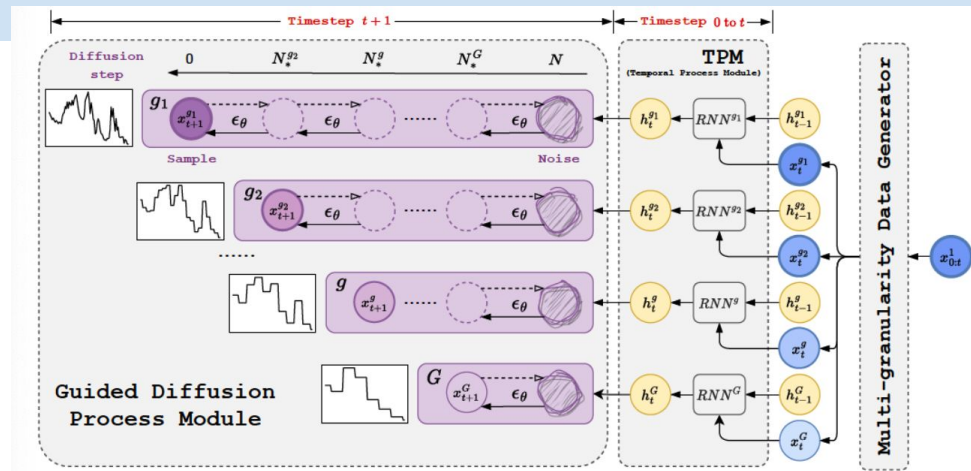
생성

<Step 1>

2. MG-TSD

(1) MG-TSD architecture

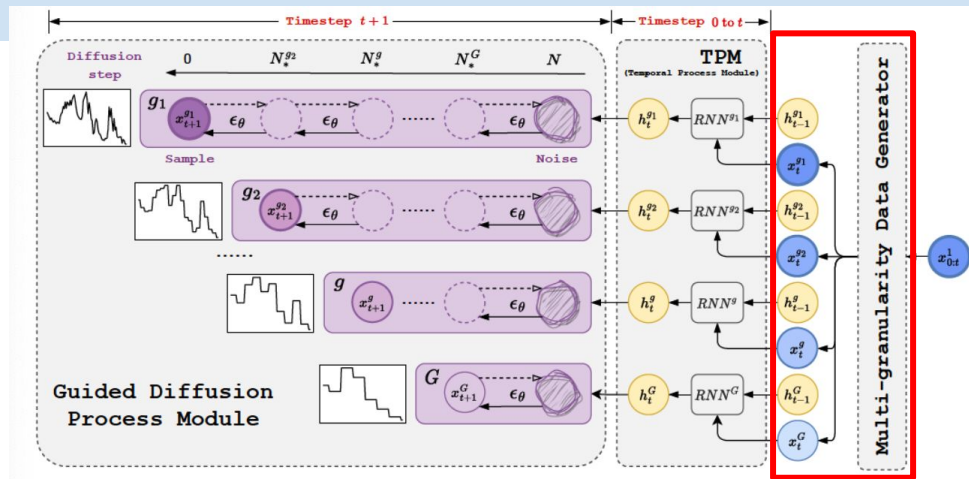
- Composed of three parts
 - a) Multi-granularity data generator
 - b) Temporal process module
 - c) Guided diffusion process module



2. MG-TSD

(1) MG-TSD architecture

- Composed of three parts
 - **a) Multi-granularity data generator**
 - generate multi-granularity data
 - with sliding window + smoothing function (ex. MA)
 - w/o overlapping! -> replicate for same length
 - b) Temporal process module
 - c) Guided diffusion process module

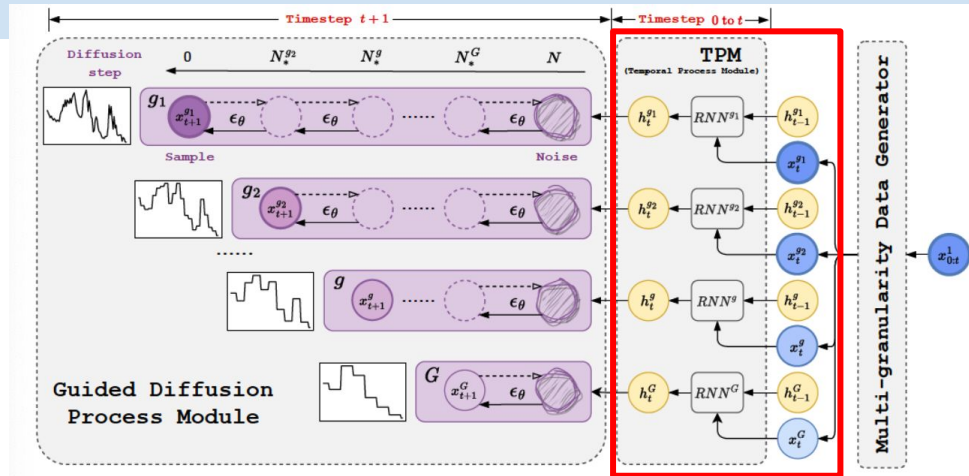


$$\mathbf{X}^{(g)} = f\left(\mathbf{X}^{(1)}, s^g\right)$$

2. MG-TSD

(1) MG-TSD architecture

- Composed of three parts
 - a) Multi-granularity data generator
 - **b) Temporal process module**
 - capture **temporal relationships**
 - use **RNN** for “**each** granularity level”
 - c) Guided diffusion process module

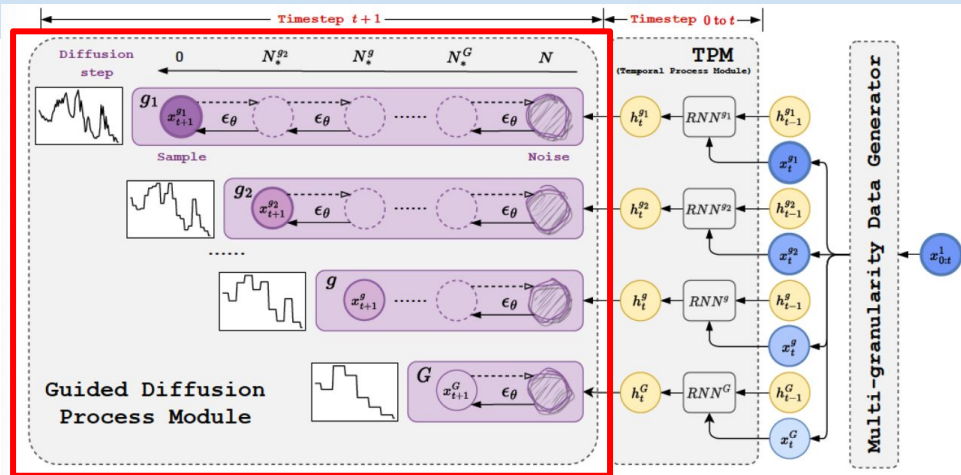


Encoded hidden states : $\mathbf{h}_t^g$

2. MG-TSD

(1) MG-TSD architecture

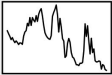
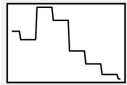
- Composed of three parts
 - a) Multi-granularity data generator
 - b) Temporal process module
 - **c) Guided diffusion process module**
 - Goal: make **TS prediction**
 - For guidance, use **multi-granularity data**



2. MG-TSD

(2) Multi-granularity guided diffusion

Notation

- (**Finest**-grained) $\mathbf{x}_t^{g_1} (g_1 = 1) \dots$ from $\mathbf{X}^{(g_1)}$ 
- (**Coarse**-grained) $\mathbf{x}_t^g \dots$ from $\mathbf{X}^{(g)}$ 
- Variance schedule: $\{\beta_n^1 = 1 - \alpha_n^1 \in (0, 1)\}_{n=1}^N$

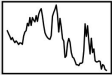
2. MG-TSD

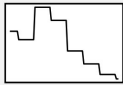
(2) Multi-granularity guided diffusion

Notation

Granularity level (1(fine) ~ G(coarse))

Time step (t 시점)

- (**Finest**-grained) $x_t^{g_1} (g_1 = 1) \dots$ from $\mathbf{X}^{(g_1)}$ 

(**Coarse**-grained) $x_t^g \dots$ from $\mathbf{X}^{(g)}$ 

- Variance schedule: $\{\beta_n^1 = 1 - \alpha_n^1 \in (0, 1)\}_{n=1}^N$

2. MG-TSD

(2) Multi-granularity guided diffusion

Notation

- (**Finest**-grained) $x_t^{g_1} (g_1 = 1) \dots$ from $\mathbf{X}^{(g_1)}$
 - (**Coarse**-grained) $x_t^g \dots \dots \dots$ from $\mathbf{X}^{(g)}$
 - Variance schedule: $\{\beta_n^1 = 1 - \alpha_n^1 \in (0, 1)\}_{n=1}^N$
- Granularity level (1(fine) ~ G(coarse))
- Time step (t 시점)
- 1로 고정 (무시 가능)
- Diffusion step (1~N)

2. MG-TSD

(2) Multi-granularity guided diffusion

Notation

- (**Finest**-grained) $x_t^{g_1} (g_1 = 1) \dots$ from $\mathbf{X}^{(g_1)}$
 - Granularity level (1(fine) ~ G(coarse))
 - Time step (t 시점)
 - 1로 고정 (무시 가능)
- (**Coarse**-grained) $x_t^g \dots \dots \dots$ from $\mathbf{X}^{(g)}$
 - Diffusion step (1~N)
- Variance schedule: $\{\beta_n^1 = 1 - \alpha_n^1 \in (0, 1)\}_{n=1}^N$

Goal: approximate the distn $q(x^{g_1})$

가장 fine 한 (즉, original scale)의 data distn

2. MG-TSD

(2) Multi-granularity guided diffusion

Notation

- (**Finest**-grained) $x_t^{g_1} (g_1 = 1) \dots$ from $\mathbf{X}^{(g_1)}$
- (**Coarse**-grained) $x_t^g \dots \dots \dots$ from $\mathbf{X}^{(g)}$
- Variance schedule: $\{\beta_n^1 = 1 - \alpha_n^1 \in (0, 1)\}_{n=1}^N$

Goal: approximate the distn $q(x^{g_1})$

가장 fine 한 (즉, original scale)의 data distn

Key Idea: Coarse한 정보를, Fine 한 정보에게 주입시키면서 denosing하면 도움이 될 것!

2. MG-TSD

(논문 Notation 실수) 어떨 때는

- TS의 time step
- **Diffusion의 time step**

(2) Multi-granularity guided diffusion

$q(x_{0:N}^{g_1})$: Forward

$x_0^{g_1} \sim q(x_{\boxed{0}}^{g_1})$: Original Data

$p_\theta(x_{0:N}^{g_1})$: Backward

[Goal] Guide the generation of samples by ensuring that the **intermediate latent space retains the underlying time series structure**

[How] By introducing coarse-grained targets x^g at intermediate diffusion step $N_*^g \in [1, N - 1]$

2. MG-TSD

(2) Multi-granularity guided diffusion

Loss Function

$$\mathbb{E}_{\epsilon, \mathbf{x}^g, n} [\|\epsilon - \epsilon_{\theta}(\mathbf{x}_n^g, n)\|^2].$$

$$\circ \boxed{\mathbf{x}_n^g} = \left(\prod_{i=N_*^g}^n \alpha_i^1 \right) \mathbf{x}^g + \sqrt{1 - \prod_{i=N_*^g}^n \alpha_i^1} \epsilon \text{ and } \epsilon \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$$

n번째 diffusion step의 g번째 granularity

2. MG-TSD

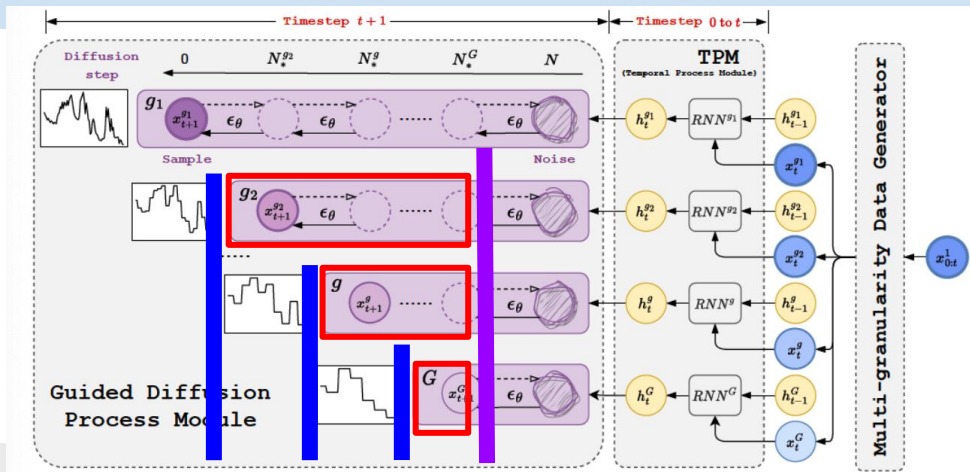
(2) Multi-granularity guided diffusion

Loss Function

$$\mathbb{E}_{\epsilon, \mathbf{x}^g, n} [\| \epsilon - \epsilon_{\theta}(\mathbf{x}_n^g, n) \|^2].$$

$$\circ \mathbf{x}_n^g = \left(\prod_{i=N_*^g}^n \alpha_i^1 \right) \mathbf{x}^g + \sqrt{1 - \prod_{i=N_*^g}^n \alpha_i^1} \epsilon \text{ and } \epsilon \sim \mathcal{N}(\mathbf{0}, I)$$

Granularity level에 depend하는 noise schedule



2. MG-TSD

(2) Multi-granularity guided diffusion

Shared Ratio /

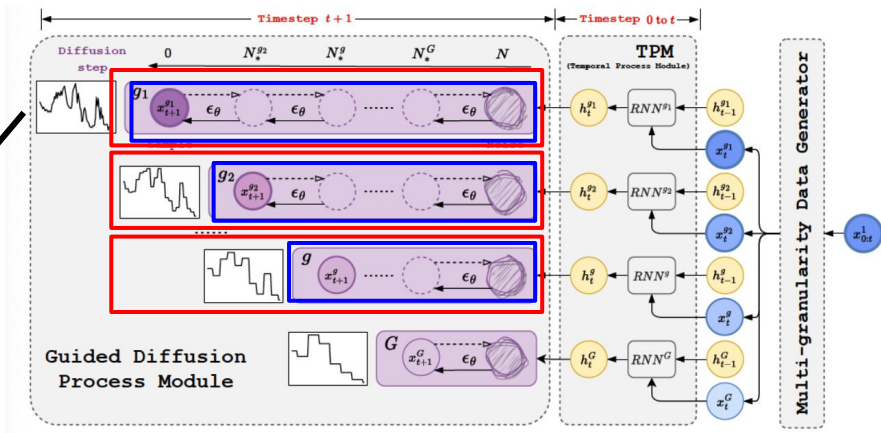
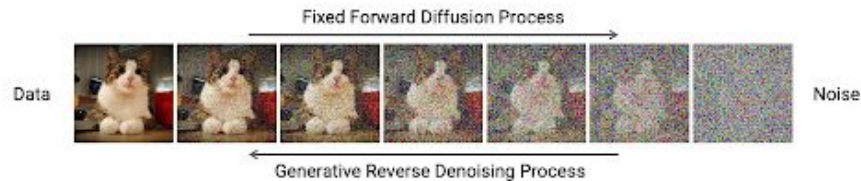
“Noise(Variance) schedule”을 share하는 비율

Share ratio: shared percentage of variance schedule between

- (1) the g th granularity data, where $g \in \{2, \dots, G\}$
- (2) the finest-grained data

→ Define it as $r_g := 1 - (N_*^g - 1) / N$.

- ex) For the finest-grained data, $N_*^1 = 1$ and $r^1 = 1$.



2. MG-TSD

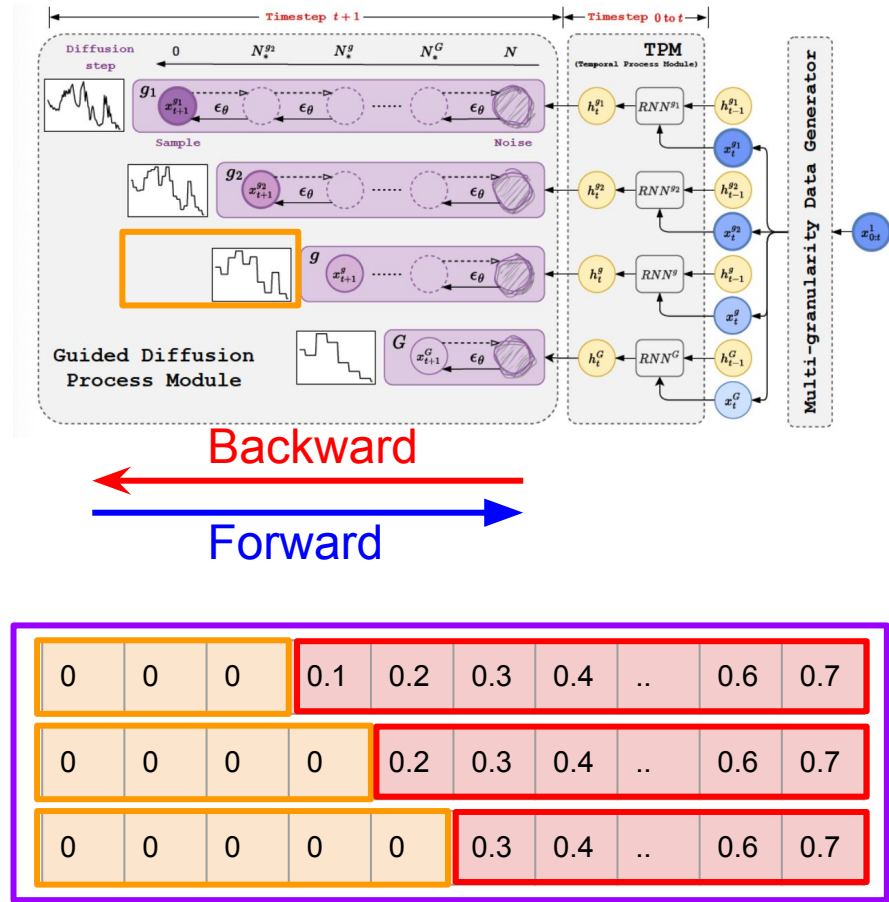
(2) Multi-granularity guided diffusion

Variance Schedule

Variance schedule for granularity g

$$\alpha_n^g(N_*^g) = \begin{cases} 1 & \text{if } n = 1, \dots, N_*^g \\ \alpha_n^1 & \text{if } n = N_*^g + 1, \dots, N \end{cases}$$

$$\text{and } \{\beta_n^g\}_{n=1}^N = \{1 - \alpha_n^g\}_{n=1}^N.$$



2. MG-TSD

(2) Multi-granularity guided diffusion

Guidance Loss Function:

Guidance (X) : $g=1$ (finest)

Guidance (O) : $g=1$ (finest) + **$g=2 \sim G$ (coarse)**

Guidance loss function $L^{(g)}(\theta)$

- for g th-granularity $\mathbf{x}_{n,t}^g$ at timestep t and diffusion step n ,
- $L^{(g)}(\theta) = \mathbb{E}_{\epsilon, \mathbf{x}_{0,t}^g, n} \| (\epsilon - \epsilon_\theta (\sqrt{a_n^g} \mathbf{x}_{0,t}^g + \sqrt{b_n^g} \epsilon, n, \mathbf{h}_{t-1}^g)) \|_2^2$.
- where $\mathbf{h}_t^g = \text{RNN}_\theta(\mathbf{x}_t^g, \mathbf{h}_{t-1}^g)$



Total Guidance loss function

(with $G - 1$ granularity levels of data)

$$\bullet L^{\text{guidance}} = \sum_{g=2}^G \omega^g L^{(g)}(\theta)$$

2. MG-TSD

(2) Multi-granularity guided diffusion

Guidance Loss Function:

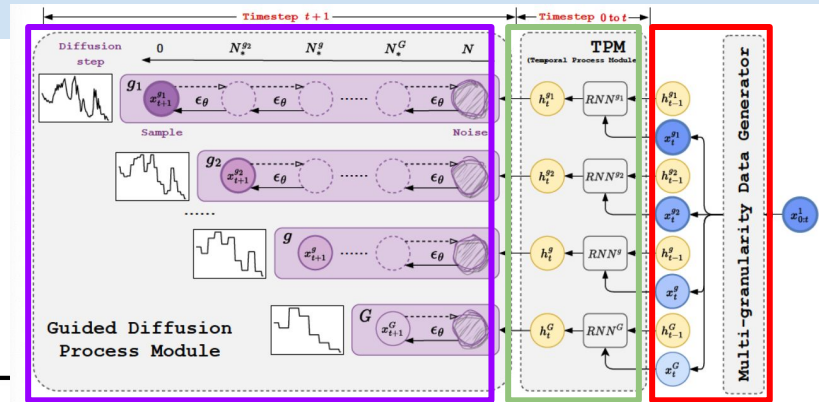
Guidance (X) : $g=1$ (finest)

Guidance (O) : **$g=1$ (finest)** + **$g=2 \sim G$ (coarse)**

$$\begin{aligned} L^{\text{final}} &= \boxed{\omega^1 L^{(1)}(\theta)} + \boxed{L^{\text{guidance}}(\theta)} \\ &= \sum_{g=1}^G \omega^g \mathbb{E}_{\epsilon, \mathbf{x}_{0,t}^g, n} [\|\epsilon - \epsilon_{\theta}(\mathbf{x}_{n,t}^g, n, \mathbf{h}_{t-1}^g)\|^2] \end{aligned}$$

2. MG-TSD

Training & Sampling



Algorithm 1 Training procedure

Input: Context interval $[1, t_0)$; prediction interval $[t_0, T]$; number of diffusion step N ; a set of share ratio for g granularity (or equivalently $\{N_*^g, g \in \{1, \dots, G\}\}$); generated multi-granularity data $[x_1^g, \dots, x_{t_0}^g, \dots, x_T^g], g \in \{1, \dots, G\}$; initial hidden states $\mathbf{h}_0^g, g \in \{1, \dots, G\}$

repeat

- 1: Sample the multi-granularity time series $[x_1^g, \dots, x_T^g], g \in \{1, \dots, G\}$
- 2: Obtain $\mathbf{h}_t^g = \text{RNN}^g(x_t^g, \mathbf{h}_{t-1}^g), g \in \{1, \dots, G\}, t \in [1, \dots, T]$
- 3: **for** $t = t_0$ to T **do**
- 4: Initialize $n \sim \text{Uniform}(1, \dots, N)$ and $\epsilon \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$
- 5: Reset the variance schedule $\{\beta_n^g = 1 - \alpha_n^g(N_*^g)\}_{n=1}^N, g \in \{1, \dots, G\}$.
- 6: Compute loss L^{final} according to Equation 10
- 7: Take the gradient $\nabla_{\theta} L^{\text{final}}$
- 8: **end for**

until converged

x Time step 수 (1~T까지)

2. MG-TSD

Training & Sampling

Reverse process

Algorithm 2 Inference procedure for each timestep $t \in [t_0, T]$

Input: Noise $\mathbf{x}_t^N \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$ and hidden states $\mathbf{h}_{t-1}^g, g \in \{1, \dots, G\}$

1: **for** $n = N$ to 1 **do**

2: **if** $n > 1$ **then**

3: Sample $\mathbf{z} \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$

4: **else**

5: $\mathbf{z} = \mathbf{0}$

6: **end if**

7: **for** $g = 1$ to G **do**

8: $\mathbf{x}_{n-1,t}^g = \frac{1}{\sqrt{\alpha_n^g}} \left(\mathbf{x}_{n,t}^g - \frac{\beta_n^g}{\sqrt{1-\alpha_n^g}} \epsilon_\theta(\mathbf{x}_{n,t}^g, n, \mathbf{h}_{t-1}^g) \right) + \sqrt{\sigma_n^g} \mathbf{z}$, where $\sigma_n^g = \frac{1-\alpha_{n-1}^g}{1-\alpha_n^g} \beta_n^g$.

9: **end for**

10: **end for**

Return $\mathbf{x}_{0,t}^g, g = 1$ (finest-grained data); (Optional: $\mathbf{x}_{0,t}^g, g \in \{2, \dots, G\}$)

Fine ~ Coarse 모두 생성 (순서 상관 X)

최종적으로 얻길 원하는 fine-granularity의 데이터

2. MG-TSD

Experiments

(1) Dataset

Name	Frequency	Number of series	Context length	Prediction length	Multi-granularity dictionary
Solar	1 hour	137	24	24	[1 hour, 4 hour, 12 hour, 24hour, 48 hour]
Electricity	1 hour	370	24	24	[1 hour, 4 hour, 12 hour, 24 hour, 48 hour]
Traffic	1 hour	963	24	24	[1 hour, 4 hour, 12 hour, 24 hour, 48 hour]
Taxi	30 min	1214	24	24	[30 min , 2 hour, 6 hour, 12 hour, 24 hour]
KDD-cup	1 hour	270	48	48	[1 hour, 4 hour, 12 hour, 24hour, 48 hour]
Wikipedia	1 day	2000	30	30	[1 day, 4 day, 7 day, 14 day]

Table 6: Detailed information of the datasets used in our benchmark including data frequency and number of times series (dimension), including the information about context length and prediction length and the multi-granularity dictionary utilized in the multivariate time series forecasting task.

2. MG-TSD

Experiments

(2) MTS forecasting

Baselines. We assess the predictive performance of the proposed **MG-TSD** model in comparison with multivariate time series forecasting models, including Vec-LSTM-ind-scaling (Salinas et al., 2019), GP-scaling (Salinas et al., 2019), GP-Copula (Salinas et al., 2019), Transformer-MAF (Rasul et al., 2020), LSTM-MAF (Rasul et al., 2020), **TimeGrad (Rasul et al., 2021)**, and TACTiS (Drouin et al., 2022). The MG-Input ensemble model serves as the baseline with multi-granularity inputs. It

- **TS Diffusion Model**은 하나 뿐!
- 비교적 오래된 알고리즘과의 비교
(다른 2023,2024 TS Diffusion 논문들도 마찬가지로 경향성 O.
가장 대표적인 **TimeGrad (2021)**, **CSDI (2021)**과만 비교하는 경우가
종종 있음

Table 1: Comparison of $CRPS_{sum}$ (smaller is better) of models on six real-world datasets. The reported mean and standard error are obtained from 10 re-training and evaluation independent runs.

Method	Solar	Electricity	Traffic	KDD-cup	Taxi	Wikipedia
Vec-LSTM ind-scaling	0.4825 $\pm$ 0.0027	0.0949 $\pm$ 0.0175	0.0915 $\pm$ 0.0197	0.3560 $\pm$ 0.1667	0.4794 $\pm$ 0.0343	0.1254 $\pm$ 0.0174
GP-Scaling	0.3802 $\pm$ 0.0052	0.0499 $\pm$ 0.0031	0.0753 $\pm$ 0.0152	0.2983 $\pm$ 0.0448	0.2265 $\pm$ 0.0210	0.1351 $\pm$ 0.0612
GP-Copula	0.3612 $\pm$ 0.0035	0.0287 $\pm$ 0.0005	0.0618 $\pm$ 0.0018	0.3157 $\pm$ 0.0462	0.1894 $\pm$ 0.0087	0.0669 $\pm$ 0.0009
LSTM-MAF	0.3427 $\pm$ 0.0082	0.0312 $\pm$ 0.0046	0.0526 $\pm$ 0.0021	0.2919 $\pm$ 0.1486	0.2295 $\pm$ 0.0082	0.0763 $\pm$ 0.0051
Transformer-MAF	0.3532 $\pm$ 0.0053	0.0272 $\pm$ 0.0017	0.0499 $\pm$ 0.0011	0.2951 $\pm$ 0.0504	0.1531 $\pm$ 0.0038	0.0644 $\pm$ 0.0037
TimeGrad	0.3335 $\pm$ 0.0653	0.0232 $\pm$ 0.0035	0.0414 $\pm$ 0.0112	0.2902 $\pm$ 0.2178	0.1255 $\pm$ 0.0207	0.0555 $\pm$ 0.0088
TACTiS	0.4209 $\pm$ 0.0330	0.0259 $\pm$ 0.0019	0.1093 $\pm$ 0.0076	0.5406 $\pm$ 0.1584	0.2070 $\pm$ 0.0159	—
MG-Input	0.3239 $\pm$ 0.0427	0.0238 $\pm$ 0.0025	0.0658 $\pm$ 0.0065	0.2977 $\pm$ 0.1162	0.1592 $\pm$ 0.0087	0.0567 $\pm$ 0.0001
MG-TSD	0.3081$\pm$0.0099	0.0149$\pm$0.0017	0.0323$\pm$0.0125	0.1837$\pm$0.0865	0.1159$\pm$0.0132	0.0529$\pm$0.0054

2. MG-TSD

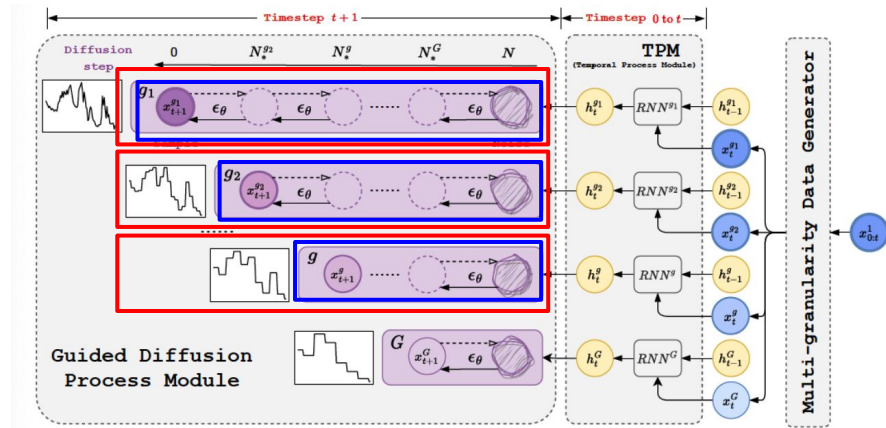
Experiments

(3) Ablation Studies

3-1) Share-ratio

Ratio	4 hour			6 hour		
	CRPS _{sum}	NMAE	NRMSE	CRPS _{sum}	NMAE	NRMSE
20%	0.3489 $\pm$ 0.0190	0.3826 $\pm$ 0.0200	0.7177 $\pm$ 0.0445	0.3378 $\pm$ 0.0305	0.3703 $\pm$ 0.0368	0.6916 $\pm$ 0.0536
40%	0.3405 $\pm$ 0.0415	0.3792 $\pm$ 0.0386	0.6870 $\pm$ 0.0870	0.3275 $\pm$ 0.0250	0.3608 $\pm$ 0.0287	0.6650 $\pm$ 0.0874
60%	0.3268 $\pm$ 0.0473	0.3604 $\pm$ 0.0403	0.6570 $\pm$ 0.0919	0.3166 $\pm$ 0.0376	0.3491 $\pm$ 0.0368	0.6478 $\pm$ 0.0696
80%	0.3172 $\pm$ 0.0249	0.3510 $\pm$ 0.0240	0.6515 $\pm$ 0.051	0.3221 $\pm$ 0.0425	0.3555 $\pm$ 0.0443	0.6542 $\pm$ 0.0747
100%	0.3178 $\pm$ 0.0342	0.3480 $\pm$ 0.0356	0.6391 $\pm$ 0.0503	0.3232 $\pm$ 0.0396	0.3548 $\pm$ 0.0417	0.6550 $\pm$ 0.0660
Ratio	12 hour			24 hour		
	CRPS _{sum}	NMAE	NRMSE	CRPS _{sum}	NMAE	NRMSE
20%	0.3440 $\pm$ 0.0391	0.3767 $\pm$ 0.0450	0.6999 $\pm$ 0.0772	0.3315 $\pm$ 0.0200	0.3603 $\pm$ 0.0298	0.6801 $\pm$ 0.0534
40%	0.3374 $\pm$ 0.0376	0.3713 $\pm$ 0.0340	0.6837 $\pm$ 0.0641	0.3276 $\pm$ 0.0358	0.3612 $\pm$ 0.0361	0.6722 $\pm$ 0.0552
60%	0.3240 $\pm$ 0.0382	0.3597 $\pm$ 0.0388	0.6694 $\pm$ 0.0746	0.3382 $\pm$ 0.0343	0.3737 $\pm$ 0.0365	0.6878 $\pm$ 0.0655
80%	0.3391 $\pm$ 0.0390	0.3719 $\pm$ 0.0403	0.6933 $\pm$ 0.0691	0.3288 $\pm$ 0.0460	0.3639 $\pm$ 0.0476	0.6741 $\pm$ 0.0929
100%	0.3284 $\pm$ 0.0323	0.3538 $\pm$ 0.0450	0.6609 $\pm$ 0.0917	0.3407 $\pm$ 0.0248	0.3692 $\pm$ 0.0244	0.6933 $\pm$ 0.0528

Table 2: Influence of share ratios for different granularities on Solar dataset. The reported mean and standard error are obtained from 10 re-training and evaluation independent runs.



0	0	0	0.1	0.2	0.3	0.4	..	0.6	0.7
0	0	0	0	0.2	0.3	0.4	..	0.6	0.7
0	0	0	0	0	0.3	0.4	..	0.6	0.7

2. MG-TSD

Experiments

(3) Ablation Studies

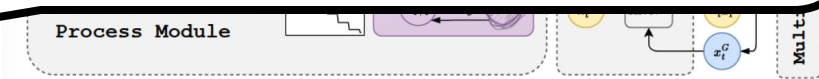
3-1) Share-ratio

For **COARSER** granularities, the model performs better with a **SMALLER** share ratio.

-> Model achieves optimal performance when the **share ratio is chosen at the step where the coarse-grained samples most closely resemble intermediate states.**

Ratio	Fine			6 hour		
	CRPS _{sum}	NMAE	NRMSE	CRPS _{sum}	NMAE	NRMSE
20%	0.3489±0.0190	0.3826±0.0200	0.7177±0.0445	0.3378±0.0305	0.3703±0.0368	0.6916±0.0536
40%	0.3405±0.0415	0.3792±0.0386	0.6870±0.0870	0.3275±0.0350	0.3608±0.0387	0.6650±0.0874
60%	0.3268±0.0473	0.3604±0.0403	0.6579±0.0919	0.3166±0.0376	0.3491±0.0368	0.6478±0.0696
80%	0.3172±0.0249	0.3510±0.0240	0.6515±0.051	0.3221±0.0425	0.3555±0.0443	0.6542±0.0747
100%	0.3178±0.0342	0.3480±0.0356	0.6391±0.0503	0.3232±0.0396	0.3548±0.0417	0.6550±0.0660

Ratio	12 hour			Coarse 24 hour		
	CRPS _{sum}	NMAE	NRMSE	CRPS _{sum}	NMAE	NRMSE
20%	0.3440±0.0391	0.3767±0.0450	0.6999±0.0772	0.3315±0.0200	0.3603±0.0298	0.6801±0.0534
40%	0.3374±0.0376	0.3713±0.0340	0.6837±0.0641	0.3276±0.0358	0.3612±0.0361	0.6722±0.0552
60%	0.3240±0.0382	0.3597±0.0388	0.6694±0.0746	0.3382±0.0343	0.3737±0.0365	0.6878±0.0655
80%	0.3391±0.0390	0.3719±0.0403	0.6933±0.0691	0.3288±0.0460	0.3639±0.0476	0.6741±0.0929
100%	0.3284±0.0323	0.3538±0.0450	0.6609±0.0917	0.3407±0.0248	0.3692±0.0244	0.6933±0.0528

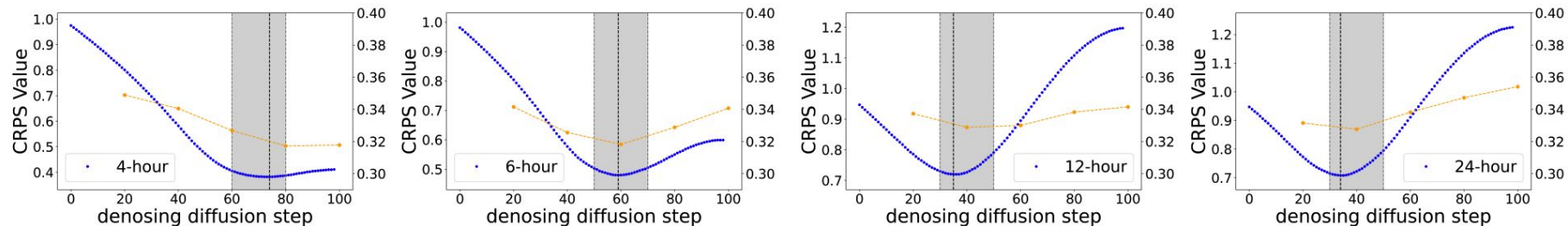


0	0	0	0.1	0.2	0.3	0.4	..	0.6	0.7
0	0	0	0	0.2	0.3	0.4	..	0.6	0.7
0	0	0	0	0	0.3	0.4	..	0.6	0.7

Table 2: Influence of share ratios for different granularities on Solar dataset. The reported mean and standard error are obtained from 10 re-training and evaluation independent runs.

2. MG-TSD

For **coarser** granularities, the model performs better with a **smaller** share ratio.



(a) CRPS_{sum} values between 4-hour targets and samples (b) CRPS_{sum} values between 6-hour targets and samples (c) CRPS_{sum} values between 12-hour targets and samples (d) CRPS_{sum} values between 24-hour targets and samples

Fine

Figure 3: Selection of share ratio for MG-TSD models

Coarse

20%	0.3440 $\pm$ 0.0391	0.3767 $\pm$ 0.0450	0.6999 $\pm$ 0.0772	0.3215 $\pm$ 0.0200	0.3603 $\pm$ 0.0298	0.6801 $\pm$ 0.0554
40%	0.3374 $\pm$ 0.0376	0.3713 $\pm$ 0.0340	0.6837 $\pm$ 0.0641	0.3276 $\pm$ 0.0358	0.3612 $\pm$ 0.0361	0.6722 $\pm$ 0.0552
60%	0.3240 $\pm$ 0.0382	0.3597 $\pm$ 0.0388	0.6694 $\pm$ 0.0746	0.3582 $\pm$ 0.0343	0.3737 $\pm$ 0.0365	0.6878 $\pm$ 0.0655
80%	0.3591 $\pm$ 0.0390	0.3719 $\pm$ 0.0403	0.6935 $\pm$ 0.0691	0.3288 $\pm$ 0.0460	0.3639 $\pm$ 0.0476	0.6741 $\pm$ 0.0929
100%	0.3284 $\pm$ 0.0323	0.3538 $\pm$ 0.0450	0.6609 $\pm$ 0.0917	0.3407 $\pm$ 0.0248	0.3692 $\pm$ 0.0244	0.6933 $\pm$ 0.0528

0	0	0	0	0	0.3	0.4	..	0.6	0.7
---	---	---	---	---	-----	-----	----	-----	-----

Table 2: Influence of share ratios for different granularities on Solar dataset. The reported mean and standard error are obtained from 10 re-training and evaluation independent runs.

2. MG-TSD

Experiments

(3) Ablation Studies

3-1) Number of Granularity

Utilizing **four to five** granularity levels generally suffices

Table 3: Influence of the number of granularities on **MG-TSD** performance for Solar and Electricity Dataset.

Num of gran	Solar			Electricity		
	CRPS _{sum}	NMAE	NRMSE	CRPS _{sum}	NMAE	NRMSE
2	0.3172 $\pm$ 0.0249	0.3510 $\pm$ 0.0240	0.6515 $\pm$ 0.0571	0.0174 $\pm$ 0.0042	0.0226 $\pm$ 0.0071	0.0296 $\pm$ 0.0086
3	0.3110 $\pm$ 0.0329	0.3494 $\pm$ 0.0378	0.6452 $\pm$ 0.0632	0.0160 $\pm$ 0.0020	0.0198 $\pm$ 0.0029	0.0262 $\pm$ 0.0039
4	0.3081$\pm$0.0099	0.3445$\pm$0.0102	0.6245$\pm$0.0268	0.0149$\pm$0.0017	0.0178$\pm$0.0018	0.0241$\pm$0.0030
5	0.3093 $\pm$ 0.0411	0.3430 $\pm$ 0.0451	0.6117 $\pm$ 0.0746	0.0153 $\pm$ 0.0027	0.0181 $\pm$ 0.0043	0.0254 $\pm$ 0.0058
Num of gran	Traffic			KDD-cup		
	CRPS _{sum}	NMAE	NRMSE	CRPS _{sum}	NMAE	NRMSE
2	0.0347 $\pm$ 0.0020	0.0396 $\pm$ 0.0022	0.0593 $\pm$ 0.0043	0.2427 $\pm$ 0.1167	0.3171 $\pm$ 0.1557	0.3745 $\pm$ 0.1652
3	0.0334 $\pm$ 0.0034	0.0382 $\pm$ 0.0035	0.0574 $\pm$ 0.0066	0.2414 $\pm$ 0.1619	0.3030 $\pm$ 0.1789	0.3808 $\pm$ 0.2168
4	0.0326 $\pm$ 0.0041	0.0374 $\pm$ 0.0048	0.0573 $\pm$ 0.0050	0.2198 $\pm$ 0.1162	0.2893 $\pm$ 0.1554	0.3315 $\pm$ 0.1882
5	0.0323$\pm$0.0125	0.0370$\pm$0.0140	0.0563$\pm$0.0230	0.1837$\pm$0.0636	0.2463$\pm$0.0865	0.3001$\pm$0.0997

2. MG-TSD

Experiments

(4) Case Study

Visualization (vs. TimeGrad)

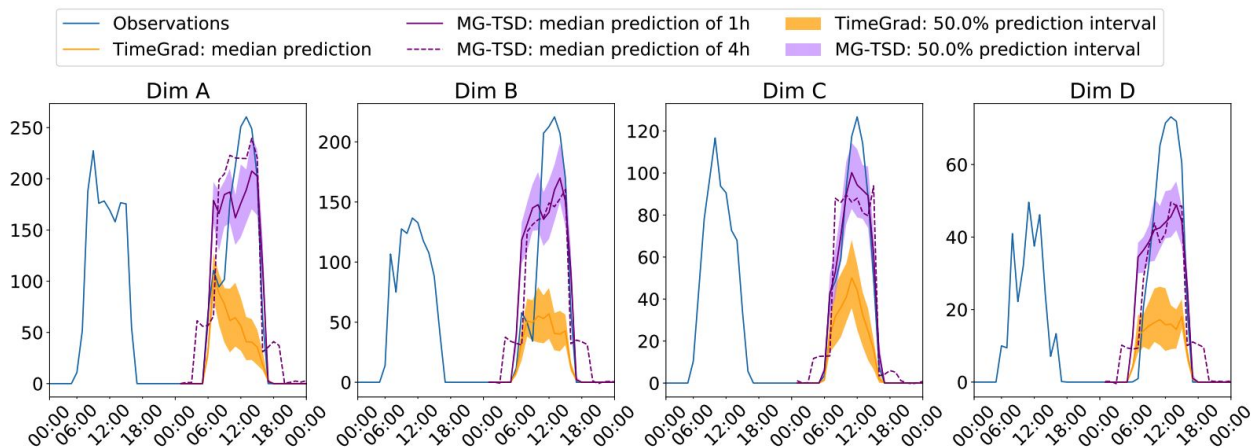
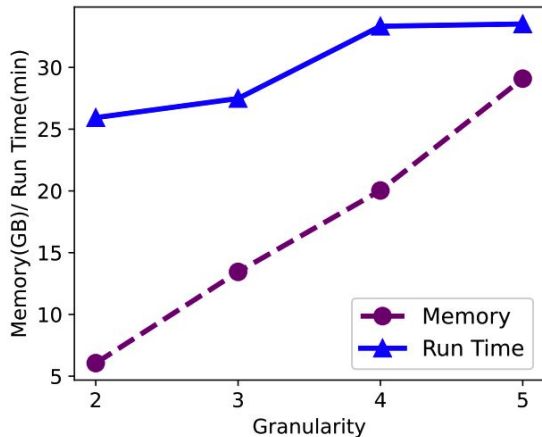


Figure 4: Solar dataset: visualization of the ground-truth, **MG-TSD** predicted mean for 4-hour and 1-hour time series, and TimeGrad predicted mean for the 1-hour time series. Additionally, the 50% prediction intervals for the 1-hour data are also included. These plots represent some illustrative dimensions out of 370 dimensions from the first 24-hour rolling-window.

2. MG-TSD

Appendix



(1) Runtime & Memory Efficiency

Dataset	Num gran	Gran dict	Share ratio	Loss weight
Solar Electricity Traffic KDD-cup	2	[1h,4h]	[1,0.9]	[0.9,0.1]
		[1h,12h]	[1,0.8]	
			[1,0.6]	
	3	[1h,4h,12h]	[1,0.9,0.8]	[0.8, 0.1, 0.1]
		[1h,4h,24h]	[1,0.8,0.8]	[0.9, 0.05, 0.05]
			[1,0.8,0.6]	[0.85, 0.10, 0.05]
	4	[1h,4h,12h,24h]	[1,0.9,0.8,0.8]	[0.8, 0.1, 0.05, 0.05]
		[1h,4h,12h,48h]	[1,0.9,0.8,0.6]	
			[1,0.8,0.6,0.6]	
Taxi	2	[30m,2h]	[1,0.8]	[0.9,0.1]
		[30m,6h]	[1,0.6]	
		[30m,12h]		
	3	[30m,24h]		
	4	[1h,4h,8h,12h,24h]	[1,0.9,0.8,0.6,0.6]	[0.8,0.1,0.05,0.04,0.01]
		[1h,4h,12h,24h,48h]	[1,0.9,0.8,0.6,0.4]	
			[1,0.8,0.6,0.6,0.6]	
Wikipedia	2	[1d,4d]	[1,0.8]	[0.9,0.1]
		[1d,7d]	[1,0.6]	
		[1d,14d]		
	3			
	4			

Table 7: Tested hyper-parameter values for the **MG-TSD** Model. The reported results in the paper are based on a parameter search within these choices.

(2) Hyperparameter Tuning

2. MG-TSD

Summary

- TS를 **multi-granularity**로 나눠서 diffuse한다는 점에서 novel (최초)
- But (상대적으로) **빈약한 실험 및 분석**
 - 적은 수의 실험 및 분석
 - 비교하는 baseline이 단순 (2021 TimeGrad, 2022 CSDI 위주 비교)
(But 대다수의 TS Diffusion 논문들도 마찬가지. 데이터셋/세팅들이 상이해서 (?))
- Openreview Rating: [6,6,6]

3. **Multi-resolution** Diffusion Model for Time-series Forecasting

3. mr-Diff: Multi-Resolution Diffusion Model

Motivation: TS data = different patterns at “**multiple scales**”

=> Why not utilize this “**multi-resolution**” temporal structure?

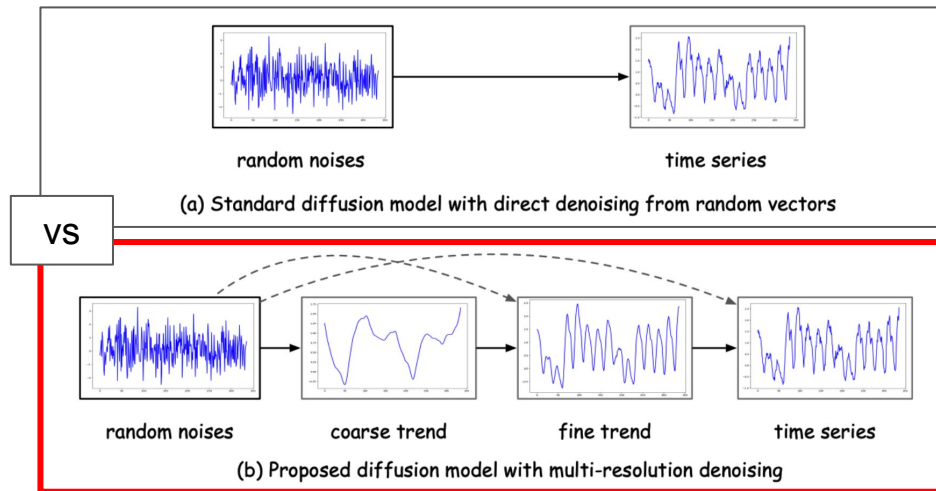
Propose **mr-Diff**

- (1) **Multi-resolution** diffusion model
- (2) **Seasonal-trend** decomposition (Importance: **Trend** >> Seasonality)
- (3) Sequentially extract “**fine-to-coarse**” trends
 - 3-1) Coarsest trend first
 - 3-2) Finer details later! (with **coarse trends as condition**)
- (4) **Non-autoregressive**

3. mr-Diff: Multi-Resolution Diffusion Model

Details & Contributions of mr-Diff

- 1) Decompose the denoising objective into **several sub-objectives**
- 2) First work to integrate “**seasonal-trend decomposition**” based **multi-resolution** analysis
- 3) **Progressive** denoising
(coarse -> fine)



3. mr-Diff: Multi-Resolution

In each stage $s + 1 \dots$

a conditional diffusion model is learned to reconstruct the "trend $\mathbf{Y}_s$ extracted from the forecast window"

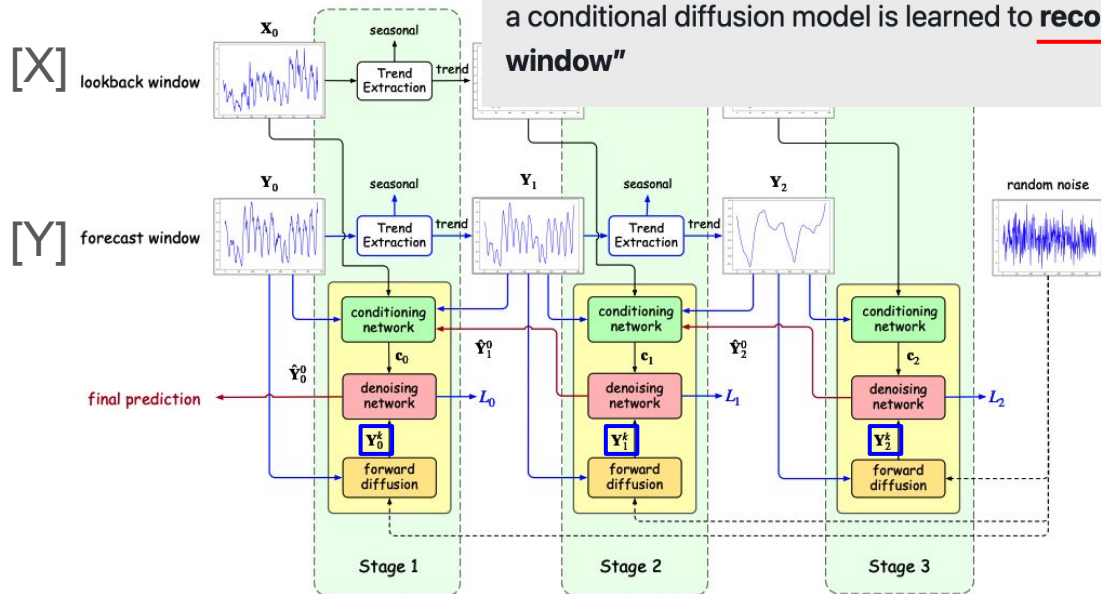


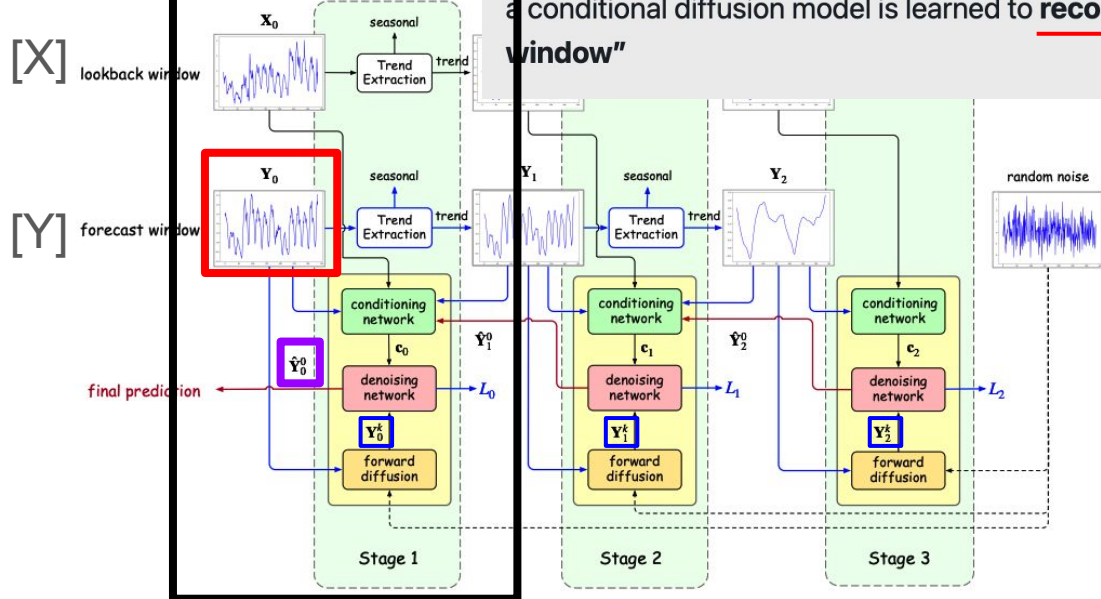
Figure 2: The proposed multi-resolution diffusion model mr-Diff. For simplicity of illustration, we use $S = 3$ stages. At stage $s + 1$, $\mathbf{Y}_s$ denotes the corresponding trend component extracted from the segment in the forecast window, $\mathbf{Y}_s^k$ is the diffusion sample at diffusion step k , and $\hat{\mathbf{Y}}_s^0$ is the denoised output.

3. mr-Diff

Stage 1

In each stage $s + 1 \dots$

a conditional diffusion model is learned to reconstruct the "trend Y_s extracted from the forecast window"



Ground Truth
Prediction

Figure 2: The proposed multi-resolution diffusion model mr-Diff. For simplicity of illustration, we use $S = 3$ stages. At stage $s + 1$, Y_s denotes the corresponding trend component extracted from the segment in the forecast window, Y_s^k is the diffusion sample at diffusion step k , and $\hat{Y}_s^0$ is the denoised output.

3. mr-Diff: Multi-Resolution Diffusion

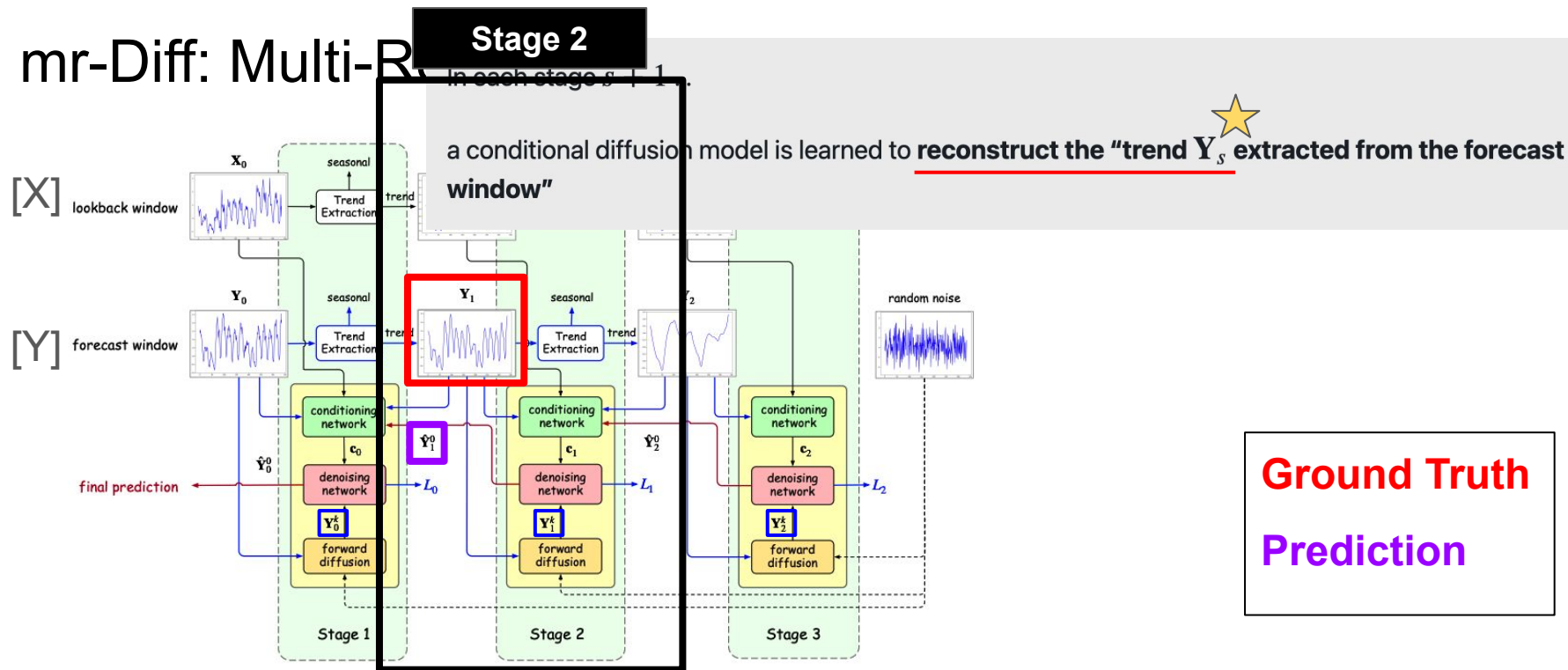


Figure 2: The proposed multi-resolution diffusion model mr-Diff. For simplicity of illustration, we use $S = 3$ stages. At stage $s + 1$, Y_s denotes the corresponding trend component extracted from the segment in the forecast window, Y_s^k is the diffusion sample at diffusion step k , and $\hat{Y}_s^0$ is the denoised output.

3. mr-Diff: Multi-Resolution

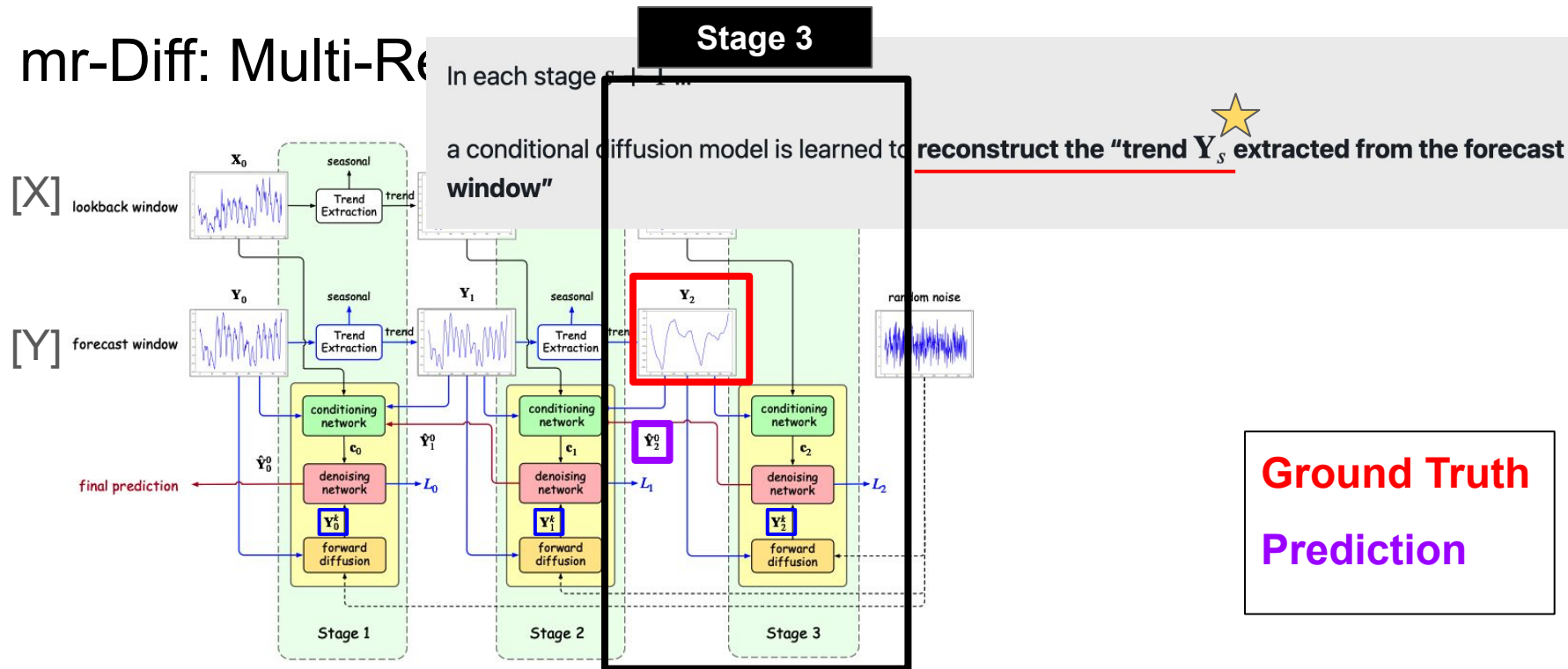


Figure 2: The proposed multi-resolution diffusion model mr-Diff. For simplicity of illustration, we use $S = 3$ stages. At stage $s + 1$, Y_s denotes the corresponding trend component extracted from the segment in the forecast window, Y_s^k is the diffusion sample at diffusion step k , and $\hat{Y}_s^0$ is the denoised output.

3. mr-Diff: Multi-Resolution

Stage 3

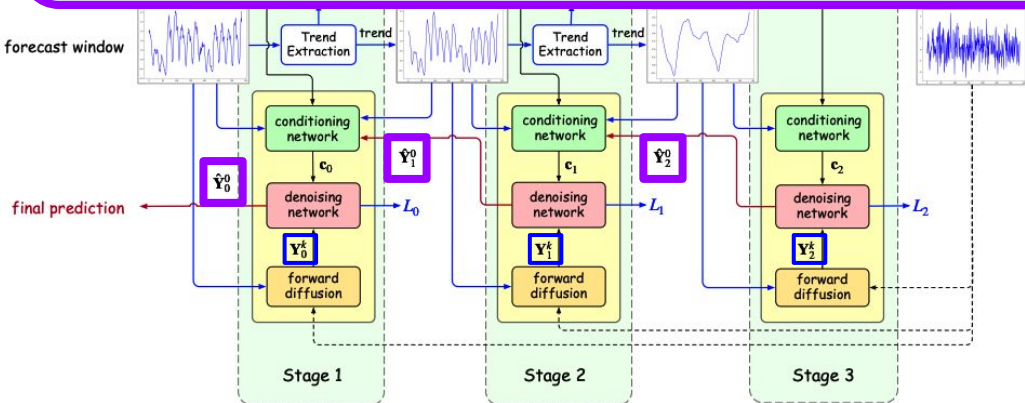
In each stage $s + 1 \dots$

extracted from the forecast

$[X]_{loc}$

How to predict the trend of the forecast window?

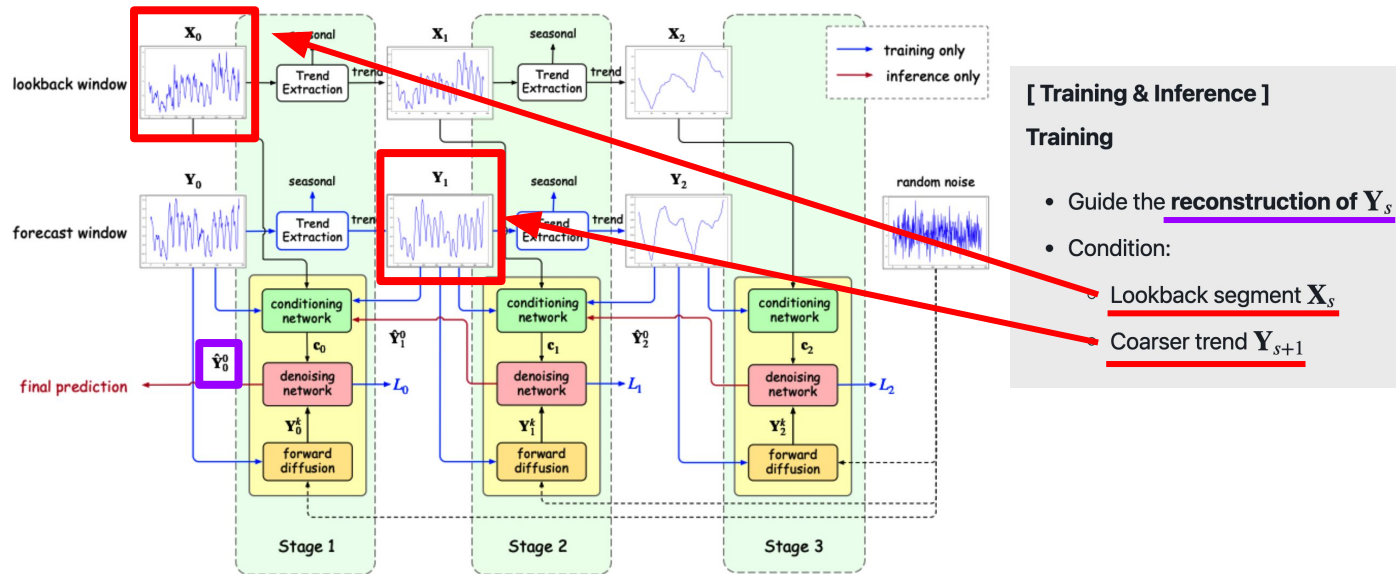
$[Y]$ forecast window



Ground Truth
Prediction

Figure 2: The proposed multi-resolution diffusion model mr-Diff. For simplicity of illustration, we use $S = 3$ stages. At stage $s + 1$, Y_s denotes the corresponding trend component extracted from the segment in the forecast window, Y_s^k is the diffusion sample at diffusion step k , and $\hat{Y}_s^0$ is the denoised output.

3. mr-Diff: Multi-Resolution Diffusion Model



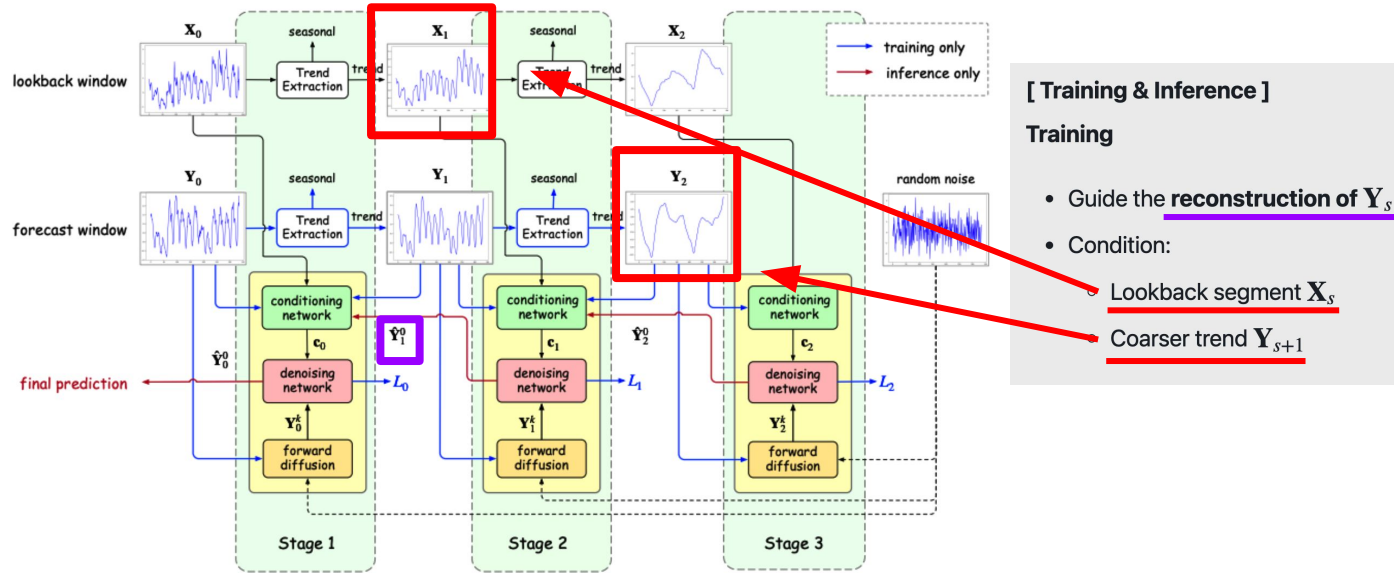
[Training & Inference]

Training

- Guide the reconstruction of Y_s
- Condition:
 - Lookback segment X_s
 - Coarser trend Y_{s+1}

[Training]

3. mr-Diff: Multi-Resolution Diffusion Model



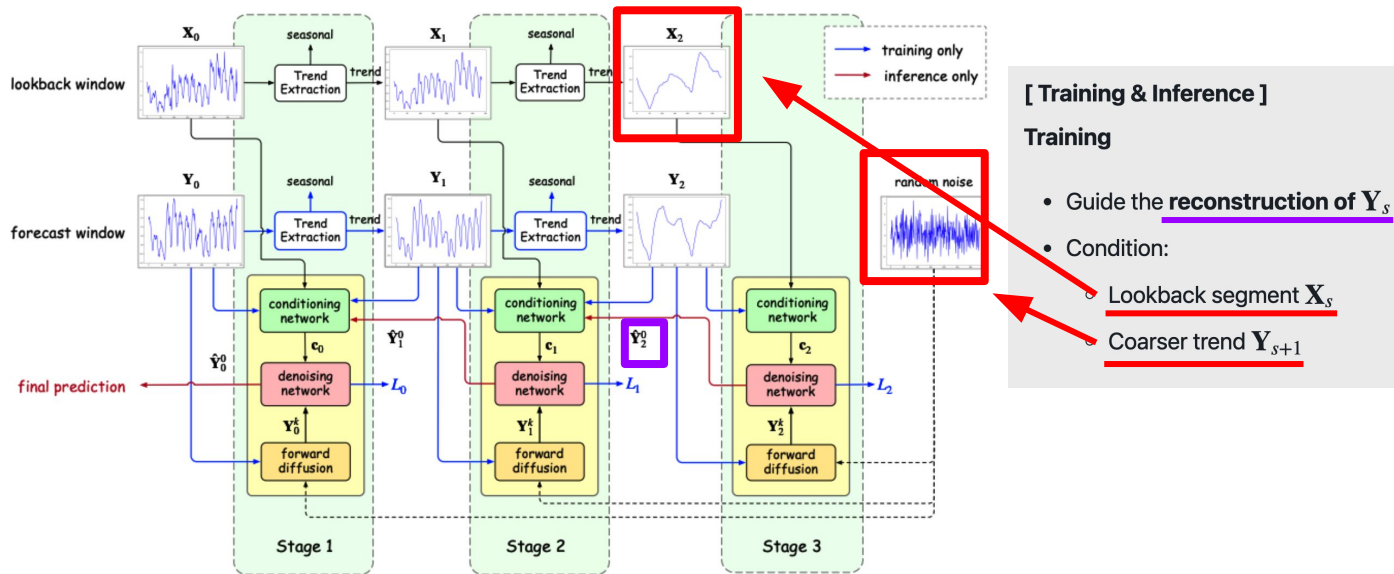
[Training & Inference]

Training

- Guide the reconstruction of Y_s
- Condition:
 - Lookback segment X_s
 - Coarser trend Y_{s+1}

[Training]

3. mr-Diff: Multi-Resolution Diffusion Model



[Training]

3. r

Y의 (s) 단계의 trend를 예측하기 위해

- (1) **X의 (s) 단계**의 trend
- (2) **Y의 (s+1) 단계**의 trend

시점의 차이: $X < Y$

Granularity의 차이: $(s) < (s+1)$

[Training]

3. r

Y의 (s) 단계의 trend를 예측하기 위해

- (1) X의 (s) 단계의 trend
- (2) Y의 (s+1) 단계의 trend

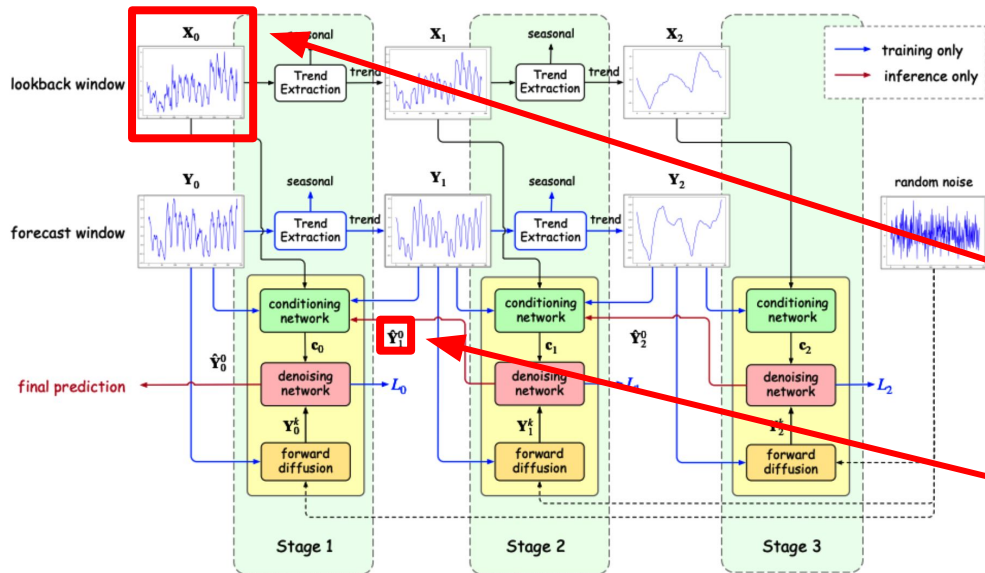
Q) Inference 시에는,
ground truth를 모르지
않나?

시점의 차이: $X < Y$

Granularity의 차이: $(s) < (s+1)$

[Training]

3. mr-Diff: Multi-Resolution Diffusion Model



[Training & Inference]

Training

- Guide the **reconstruction of $\mathbf{Y}_s$**
- Condition:
 - Lookback segment $\mathbf{X}_s$
 - Coarser trend $\mathbf{Y}_{s+1}$

Inference

- Ground-truth $\mathbf{Y}_{s+1}$ is not available
- Replaced by its estimate $\hat{\mathbf{Y}}_{s+1}^0$ produced by the denoising process at stage $s + 1$.

Use the “estimated” value!

[Inference]

3. mr-Diff: Multi-Resolution Diffusion Model

1. Extending Fine-to-Coarse Trends

- Seasonal & Trend Decomposition

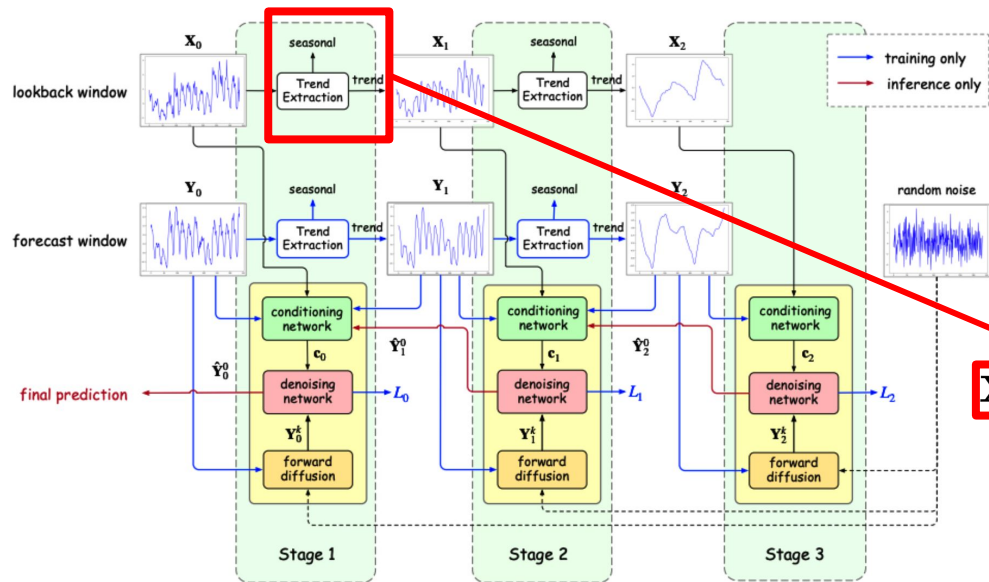
2. Temporal Multi-resolution Reconstruction

- Sinusoidal Temporal Embedding
- Forward & Backward Process
- Conditioning Network

3. mr-Diff: Multi-Resolution Diffusion Model

1. Extending Fine-to-Coarse Trends

- Seasonal & Trend Decomposition



- (Intuition) Easier to predict a finer trend **from a coarser trend**
- Finer seasonal component from a coarser seasonal component may be difficult

****Focus only on the TREND**

$$\mathbf{X}_s = \text{AvgPool}(\text{Padding}(\mathbf{X}_{s-1}), \tau_s), s = 1, \dots, S-1$$

3. mr-Diff: Multi-Resolution Diffusion Model

2. Temporal Multi-resolution Reconstruction

- **Sinusoidal Temporal Embedding**
- Forward & Backward Process
- Conditioning Network

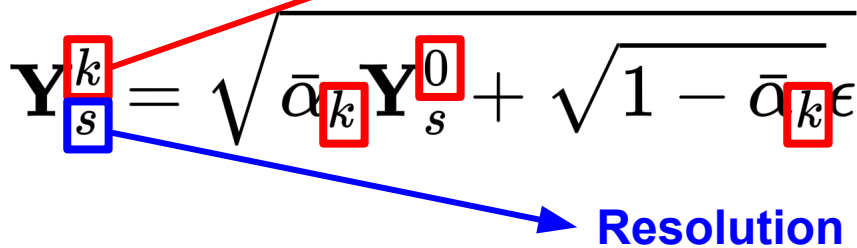
$$k_{\text{embedding}} = \left[\sin\left(10^{\frac{0 \times 4}{w-1}} t\right), \dots, \sin\left(10^{\frac{w \times 4}{w-1}} t\right), \cos\left(10^{\frac{0 \times 4}{w-1}} t\right), \dots, \cos\left(10^{\frac{w \times 4}{w-1}} t\right) \right]$$

$$\mathbf{p}^k = \text{SiLU}(\text{FC}(\text{SiLU}(\text{FC}(k_{\text{embedding}}))))$$

3. mr-Diff: Multi-Resolution Diffusion Model

2. Temporal Multi-resolution Reconstruction

- Sinusoidal Temporal Embedding
- **Forward** & Backward Process
- Conditioning Network

$$\mathbf{Y}_{s,k} = \sqrt{\bar{\alpha}_k} \mathbf{Y}_s^0 + \sqrt{1 - \bar{\alpha}_k} \epsilon$$


3. mr-Diff: Multi-Resolution Diffusion Model

2. Temporal Multi-resolution Reconstruction

- Sinusoidal Temporal Embedding
- Forward & **Backward** Process
- **Conditioning Network**

Decompose the denoising objective into S sub-objectives

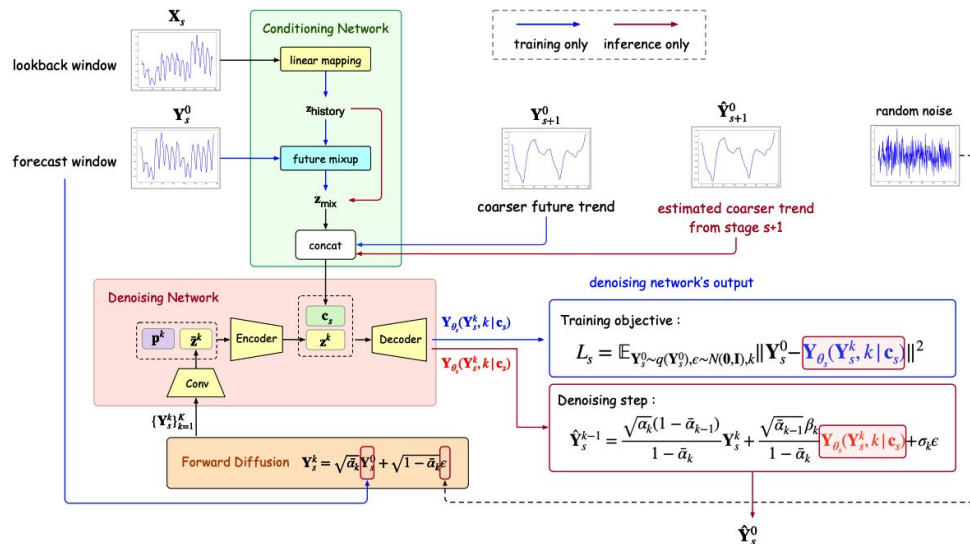
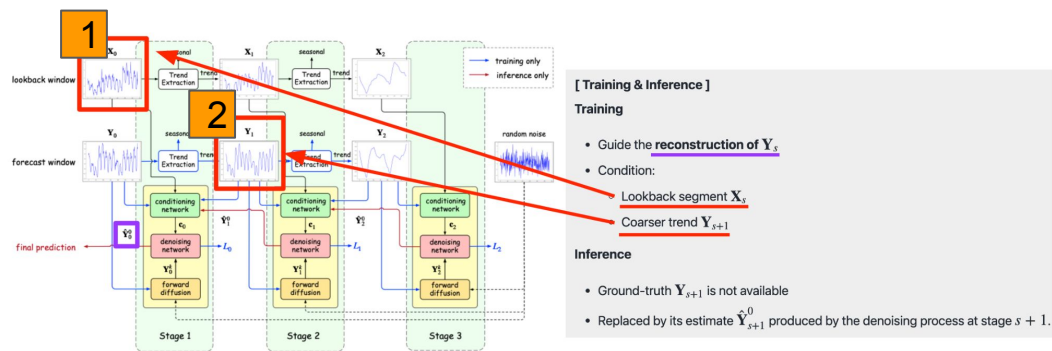


Figure 3: The conditioning network and denoising network.

3. mr-Diff: Multi-Resolution Diffusion Model

2. Temporal Multi-resolution Reconstruction

3. mr-Diff: Multi-Resolution Diffusion Model



[Training]

How to incorporate

1 and 2 as condition?

3. mr-Diff: Multi-Resolution Diffusion Model

2. Temporal Multi-resolution Reconstruction

Existing works

- Use the **original TS lookback segment $\mathbf{X}_0$** as condition $\mathbf{c}$

mr-Diff

- Use the **lookback segment $\mathbf{X}_s$** at the same decomposition **stage s** .

→ Allows better and easier reconstruction

($\because \mathbf{X}_s$ has the same resolution as $\mathbf{Y}_s$ to be reconstructed)

↔ When $\mathbf{X}_0$ is used as in existing TS diffusion models, the denoising network may overfit temporal details at the finer level.

VS

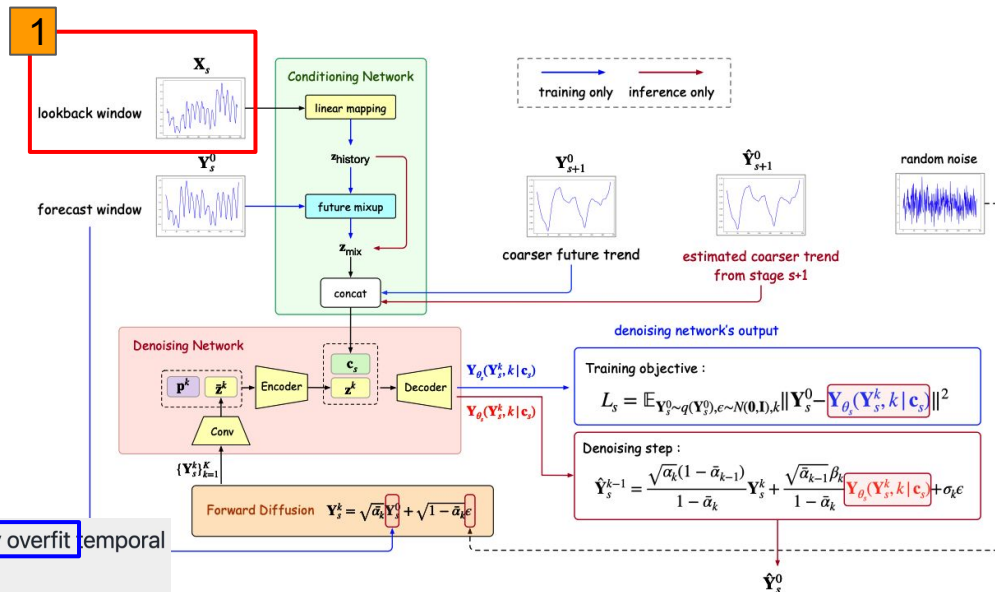


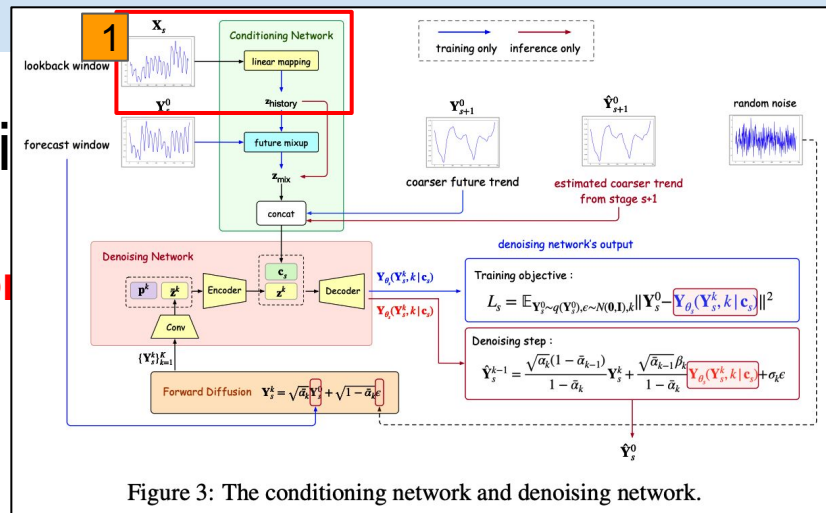
Figure 3: The conditioning network and denoising network.

3. mr-Diff: Multi-Resolution Diffusion

2. Temporal Multi-resolution Reconstruction

Procedures

- Step 1) **Linear mapping** is applied on $\mathbf{X}_s$ to produce a $\mathbf{z}_{\text{history}} \in \mathbb{R}^{d \times H}$.
- Step 2) **Future-mixup**: to enhance $\mathbf{z}_{\text{history}}$.
 - $\mathbf{z}_{\text{mix}} = \mathbf{m} \odot \mathbf{z}_{\text{history}} + (1 - \mathbf{m}) \odot \mathbf{Y}_s^0$.
 - Similar to teacher forcing, which mixes the ground truth with previous prediction output
- Step 3) **Coarser trend** $\mathbf{Y}_{s+1}$ ($= \mathbf{Y}_{s+1}^0$) can also be useful for conditioning
 - $\rightarrow \mathbf{z}_{\text{mix}}$ is concatenated with $\mathbf{Y}_{s+1}^0$ to produce the condition $\mathbf{c}_s$ (a $2d \times H$ tensor).

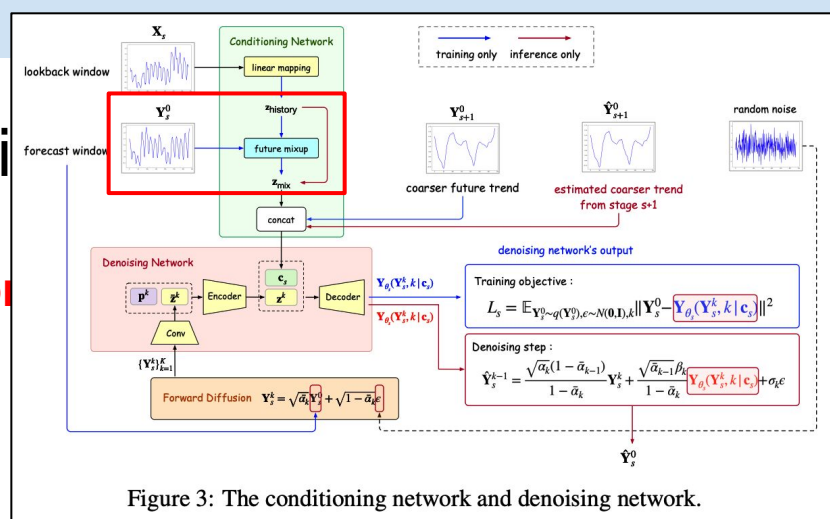


3. mr-Diff: Multi-Resolution Diffusion

2. Temporal Multi-resolution Reconstruction

Procedures

- Step 1) **Linear mapping** is applied on $\mathbf{X}_s$ to produce a $\mathbf{z}_{\text{history}} \in \mathbb{R}^{d \times H}$.
- Step 2) **Future-mixup** to enhance $\mathbf{z}_{\text{history}}$. (feat. TimeDiff, ICML 2023)
 - $\mathbf{z}_{\text{mix}} = \mathbf{m} \odot \mathbf{z}_{\text{history}} + (1 - \mathbf{m}) \odot \mathbf{Y}_s^0$
 - Similar to teacher forcing, which mixes the ground truth with previous prediction output
- Step 3) **Coarser trend** $\mathbf{Y}_{s+1}$ ($= \mathbf{Y}_{s+1}^0$) can also be useful for conditioning
 - $\mathbf{z}_{\text{mix}}$ is concatenated with $\mathbf{Y}_{s+1}^0$ to produce the condition $\mathbf{c}_s$ (a $2d \times H$ tensor).

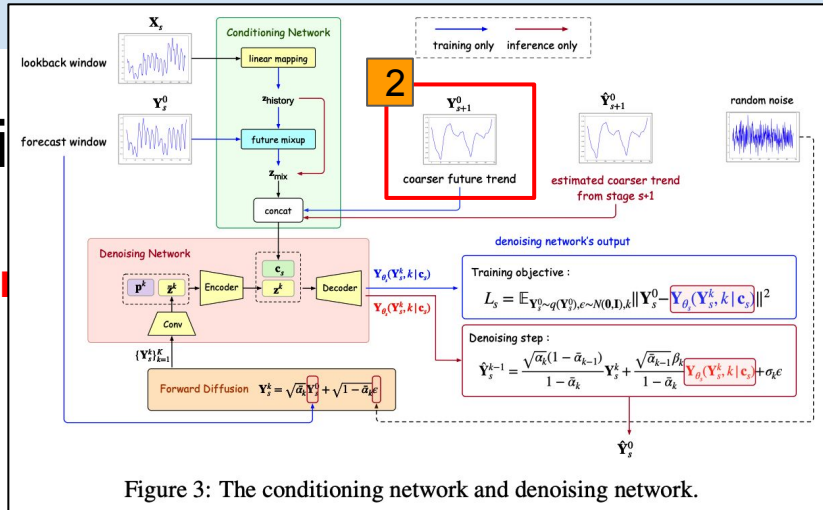


3. mr-Diff: Multi-Resolution Diffusion

2. Temporal Multi-resolution Reconstruction

Procedures

- Step 1) **Linear mapping** is applied on $\mathbf{X}_s$ to produce a $\mathbf{z}_{\text{history}} \in \mathbb{R}^{d \times H}$.
- Step 2) **Future-mixup**: to enhance $\mathbf{z}_{\text{history}}$.
 - $\mathbf{z}_{\text{mix}} = \mathbf{m} \odot \mathbf{z}_{\text{history}} + (1 - \mathbf{m}) \odot \mathbf{Y}_s^0$.
 - Similar to teacher forcing, which mixes the ground truth with previous prediction output
- Step 3) **Coarser trend** $\mathbf{Y}_{s+1}$ ($= \mathbf{Y}_{s+1}^0$) can also be useful for conditioning
 → $\mathbf{z}_{\text{mix}}$ is concatenated with $\mathbf{Y}_{s+1}^0$ to produce the condition $\mathbf{c}_s$ (a $2d \times H$ tensor).



3. mr-Diff: Multi-Resolution Diffusion

2. Temporal Multi-resolution Reconstruction

At test phase...

- Ground-truth $\mathbf{Y}_s^0$ is no longer available
→ No future-mixup ... simply set $\mathbf{z}_{\text{mix}} = \mathbf{z}_{\text{history}}$.
- Coarser trend $\mathbf{Y}_{s+1}$ is also not available
→ Concatenate $\mathbf{z}_{\text{mix}}$ with the estimate $\hat{\mathbf{Y}}_{s+1}^0$ generated from stage $s + 2$.

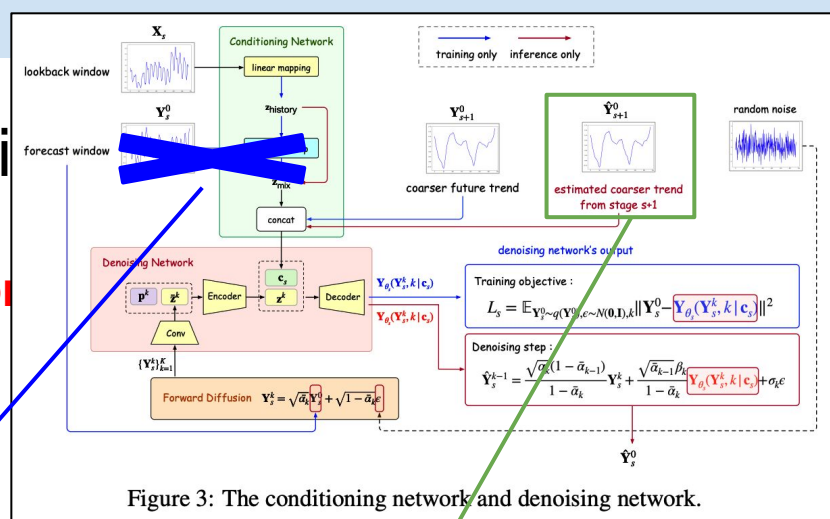
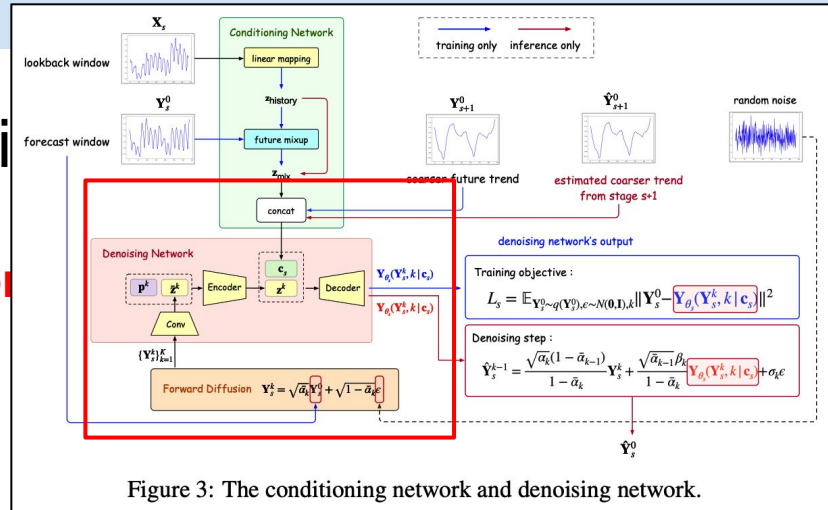


Figure 3: The conditioning network and denoising network.

3. mr-Diff: Multi-Resolution Diffusion

2. Temporal Multi-resolution Reconstruction



Denoising Network

Outputs a $\mathbf{Y}_{\theta_s}(\mathbf{Y}_s^k, k | \mathbf{c}_s)$ with guidance from the condition $\mathbf{c}_s$

Denoising process at step k of stage $s + 1$:

- $p_{\theta_s}(\mathbf{Y}_s^{k-1} | \mathbf{Y}_s^k, \mathbf{c}_s) = \mathcal{N}(\mathbf{Y}_s^{k-1}; \mu_{\theta_s}(\mathbf{Y}_s^k, k | \mathbf{c}_s, \sigma_k^2 \mathbf{I})), k = K, \dots, 1.$
- $\mathbf{Y}_{\theta_s}(\mathbf{Y}_s^k, k | \mathbf{c}_s)$ is an estimate of $\mathbf{Y}_s^0$.

3. mr-Diff: Multi-Resolution Diffusion Model

Experiments

(1) Dataset

Table 4: Summary of dataset statistics, including the dimension, total number of observations, sampling frequency, and prediction length H used in the experiments.

	dimension	#observations	frequency	steps (H)
<i>NorPool</i>	18	70,128	1 hour	720 (1 month)
<i>Caiso</i>	10	74,472	1 hour	720 (1 month)
<i>Traffic</i>	862	17,544	1 hour	168 (1 week)
<i>Electricity</i>	321	26,304	1 hour	168 (1 week)
<i>Weather</i>	21	52,696	10 mins	672 (1 week)
<i>Exchange</i>	8	7,588	1 day	14 (2 weeks)
<i>ETTh1</i>	7	17,420	1 hour	168 (1 week)
<i>ETTm1</i>	7	69,680	15 mins	192 (2 days)
<i>Wind</i>	7	48,673	15 mins	192 (2 days)

3. mr-Diff: Multi-Resolution Diffusion Model

Experiments

(2) TS forecasting

(2021,22) TS Diffusion models

(2023) TimeDiff

TimeDiff

Table 1: Univariate prediction MAEs on the real-world time series datasets (subscript is the rank). Results of all baselines (except NLinear, SCINet, and N-Hits) are from (Shen & Kwok, 2023).

	NorPool	Caiso	Traffic	Electricity	Weather	Exchange	ETH1	ETTm1	Wind	avg rank
mr-Diff	0.609 ₍₂₎	0.212 ₍₄₎	0.197 ₍₂₎	0.332 ₍₁₎	0.032 ₍₁₎	0.094 ₍₁₎	0.196 ₍₁₎	0.149 ₍₁₎	1.168 ₍₂₎	1.7
TimeDiff	0.613 ₍₃₎	0.209 ₍₃₎	0.207 ₍₃₎	0.341 ₍₃₎	0.035 ₍₄₎	0.107 ₍₁₃₎	0.202 ₍₂₎	0.154 ₍₆₎	1.209 ₍₅₎	4.0
TimeGrad	0.841 ₍₂₃₎	0.386 ₍₂₂₎	0.894 ₍₂₃₎	0.898 ₍₂₃₎	0.036 ₍₆₎	0.107 ₍₁₃₎	0.212 ₍₈₎	0.167 ₍₁₂₎	1.239 ₍₁₁₎	16.6
CSDI	0.763 ₍₂₀₎	0.282 ₍₁₄₎	0.468 ₍₂₀₎	0.540 ₍₁₉₎	0.037 ₍₇₎	0.107 ₍₁₃₎	0.221 ₍₁₂₎	0.170 ₍₁₄₎	1.218 ₍₇₎	15.1
SSSD	0.770 ₍₂₁₎	0.263 ₍₁₂₎	0.226 ₍₆₎	0.403 ₍₈₎	0.041 ₍₈₎	0.118 ₍₁₆₎	0.250 ₍₂₀₎	0.169 ₍₁₃₎	1.356 ₍₂₂₎	14.3
D ³ VAE	0.774 ₍₂₂₎	0.613 ₍₂₃₎	0.237 ₍₉₎	0.539 ₍₁₈₎	0.041 ₍₈₎	0.107 ₍₁₃₎	0.221 ₍₁₂₎	0.160 ₍₉₎	1.321 ₍₁₉₎	14.9
CPF	0.710 ₍₁₅₎	0.338 ₍₁₈₎	0.385 ₍₁₉₎	0.592 ₍₂₁₎	0.039 ₍₄₎	0.094 ₍₁₎	0.221 ₍₁₂₎	0.153 ₍₅₎	1.256 ₍₁₂₎	11.9
PSA-GAN	0.623 ₍₅₎	0.250 ₍₈₎	0.355 ₍₁₇₎	0.373 ₍₆₎	0.039 ₍₂₂₎	0.109 ₍₁₄₎	0.225 ₍₁₆₎	0.174 ₍₁₆₎	1.287 ₍₁₆₎	13.3
N-Hits	0.646 ₍₇₎	0.276 ₍₁₃₎	0.232 ₍₇₎	0.412 ₍₁₂₎	0.033 ₍₂₎	0.100 ₍₅₎	0.228 ₍₁₇₎	0.157 ₍₈₎	1.256 ₍₁₂₎	8.9
FiLM	0.654 ₍₉₎	0.290 ₍₁₅₎	0.315 ₍₁₄₎	0.412 ₍₁₂₎	0.069 ₍₁₄₎	0.104 ₍₁₀₎	0.210 ₍₆₎	0.149 ₍₁₎	1.189 ₍₃₎	8.6
SCINet	0.616 ₍₄₎	0.285 ₍₁₄₎	0.241 ₍₁₀₎	0.412 ₍₁₂₎	0.102 ₍₁₉₎	0.106 ₍₁₂₎	0.202 ₍₂₎	0.165 ₍₁₀₎	1.472 ₍₂₃₎	10.3
SCINet	0.616 ₍₄₎	0.285 ₍₁₄₎	0.225 ₍₅₎	0.439 ₍₁₃₎	0.130 ₍₂₁₎	0.096 ₍₃₎	0.242 ₍₁₈₎	0.165 ₍₁₀₎	1.236 ₍₉₎	10.4
SCINet	0.616 ₍₄₎	0.285 ₍₁₄₎	0.375 ₍₁₈₎	0.430 ₍₁₀₎	0.083 ₍₁₇₎	0.148 ₍₁₉₎	0.302 ₍₂₂₎	0.210 ₍₂₂₎	1.348 ₍₂₁₎	17.4
SCINet	0.616 ₍₄₎	0.285 ₍₁₄₎	0.269 ₍₁₁₎	0.478 ₍₁₇₎	0.098 ₍₁₈₎	0.111 ₍₁₅₎	0.260 ₍₂₁₎	0.174 ₍₁₆₎	1.338 ₍₂₀₎	14.4
SCINet	0.616 ₍₄₎	0.285 ₍₁₄₎	0.278 ₍₁₂₎	0.453 ₍₁₄₎	0.057 ₍₁₃₎	0.168 ₍₂₂₎	0.212 ₍₈₎	0.195 ₍₂₀₎	1.271 ₍₁₄₎	14.3
SCINet	0.616 ₍₄₎	0.285 ₍₁₄₎	0.495 ₍₂₁₎	0.623 ₍₂₂₎	0.040 ₍₁₀₎	0.152 ₍₂₀₎	0.220 ₍₁₁₎	0.174 ₍₁₆₎	1.319 ₍₁₈₎	17.3
SCINet	0.616 ₍₄₎	0.285 ₍₁₄₎	0.215 ₍₄₎	0.455 ₍₁₅₎	0.107 ₍₂₀₎	0.104 ₍₁₀₎	0.211 ₍₇₎	0.179 ₍₁₉₎	1.284 ₍₁₅₎	13.1
SCINet	0.616 ₍₄₎	0.285 ₍₁₄₎	0.308 ₍₁₃₎	0.433 ₍₁₁₎	0.069 ₍₁₄₎	0.118 ₍₁₆₎	0.212 ₍₈₎	0.172 ₍₁₅₎	1.236 ₍₉₎	13.2
SCINet	0.616 ₍₄₎	0.285 ₍₁₄₎	0.336 ₍₁₆₎	0.469 ₍₁₆₎	0.071 ₍₁₆₎	0.103 ₍₉₎	0.247 ₍₁₉₎	0.196 ₍₂₁₎	1.212 ₍₆₎	15.4
SCINet	0.616 ₍₄₎	0.285 ₍₁₄₎	0.322 ₍₁₅₎	0.377 ₍₇₎	0.037 ₍₇₎	0.101 ₍₆₎	0.205 ₍₅₎	0.150 ₍₄₎	1.167 ₍₁₎	6.7
NLinear	0.637 ₍₆₎	0.238 ₍₆₎	0.192 ₍₁₎	0.334 ₍₂₎	0.033 ₍₂₎	0.097 ₍₄₎	0.203 ₍₄₎	0.149 ₍₁₎	1.197 ₍₄₎	3.3
DLinear	0.671 ₍₁₀₎	0.206 ₍₂₎	0.236 ₍₈₎	0.348 ₍₄₎	0.310 ₍₂₃₎	0.102 ₍₇₎	0.222 ₍₁₅₎	0.155 ₍₇₎	1.221 ₍₈₎	9.3
LSTMa	0.707 ₍₁₄₎	0.333 ₍₁₇₎	0.757 ₍₂₂₎	0.557 ₍₂₀₎	0.053 ₍₁₂₎	0.136 ₍₁₈₎	0.332 ₍₂₃₎	0.239 ₍₂₃₎	1.298 ₍₁₇₎	18.4

TimeDiff

- No code
- Different result (use MSE, not MAE)

3. mr-Diff: Multi-Resolution Diffusion Model

Experiments

(3) Analysis

Visualization (TimeDiff, mr-Diff: no code implementation)

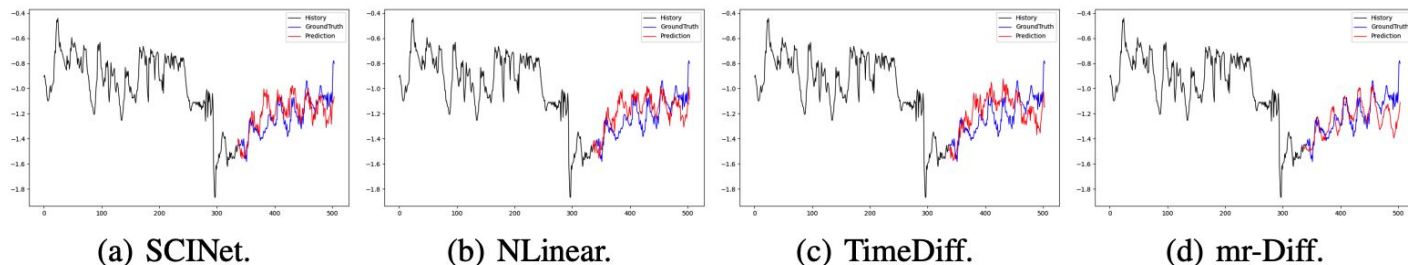


Figure 4: Visualizations on *ETTh1* by (a) SCINet (Liu et al., 2022), (b) NLinear (Zeng et al., 2023), (c) TimeDiff (Shen & Kwok, 2023) and (d) the proposed mr-Diff.

3. mr-Diff: Multi-Resolution Diffusion Model

Experiments

(3) Analysis

Efficiency analysis

Table 3: Inference time (in ms) of various time series diffusion models with different prediction horizons (H) on the univariate *ETTh1*.

	$H=96$	$H=168$	$H=192$	$H=336$	$H=720$
mr-Diff ($S=2$)	8.3	9.5	9.8	11.9	21.6
mr-Diff ($S=3$)	12.5	14.3	14.9	16.8	27.5
mr-Diff ($S=4$)	16.7	19.1	19.7	28.5	36.4
mr-Diff ($S=5$)	30.0	30.2	30.2	35.0	43.6
TimeDiff	16.2	17.3	17.6	26.5	34.6
TimeGrad	870.2	1620.9	1854.5	3119.7	6724.1
CSDI	90.4	128.3	142.8	398.9	513.1
SSSD	418.6	590.2	645.4	1054.2	2516.9

3. mr-Diff: Multi-Resolution Diffusion Model

Summary

- TS를 multi-granularity로 나눠서 **diffuse**한다는 점에서 novel
- 마찬가지로, 빈약한 실험 및 분석
- Openreview Rating: [8,6,6,6]

4. Conclusion

- Time-Series Diffusion Model의 **연구 초창기**
- 다른 도메인의 모델을 Time Series에 접목할 때, **1차원적으로 쉽게 적용할 수 있는 접근:**
 - **Time Series Decomposition**
 - **Hierarchical Approach**
 - **Seasonality Analysis (feat. Fourier Transform)**
 - ...
- 이러한 연구들이 올해 서서히 나오고 있는 추세.
- 향후 연구 방향) 새로운 아이디어로 위의 **TS domain-specific** 특징을 포착하는 알고리즘 고안