

PAT 463/563 (Fall 2025)

Music & AI

Lecture 13: Piano Roll-based Music Generation

Instructor: Hao-Wen Dong

Two Paradigms of Symbolic Music Generation

Text-based

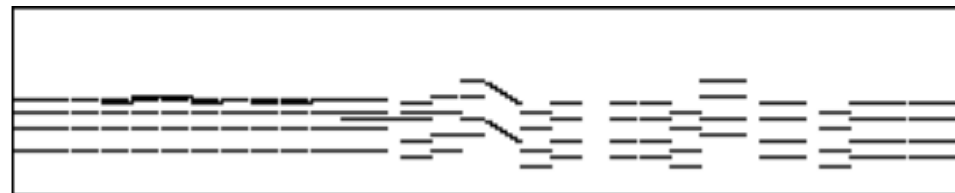
- Treat music like **text**
- Sharing models with **natural language processing (NLP)**
 - RNNs, LSTMs, Transformers, etc.

```
Program_change_0,  
Note_on_60, Time_shift_2, Note_off_60,  
Note_on_60, Time_shift_2, Note_off_60,  
Note_on_76, Time_shift_2, Note_off_67,  
Note_on_67, Time_shift_2, Note_off_67, ...
```

Image-based

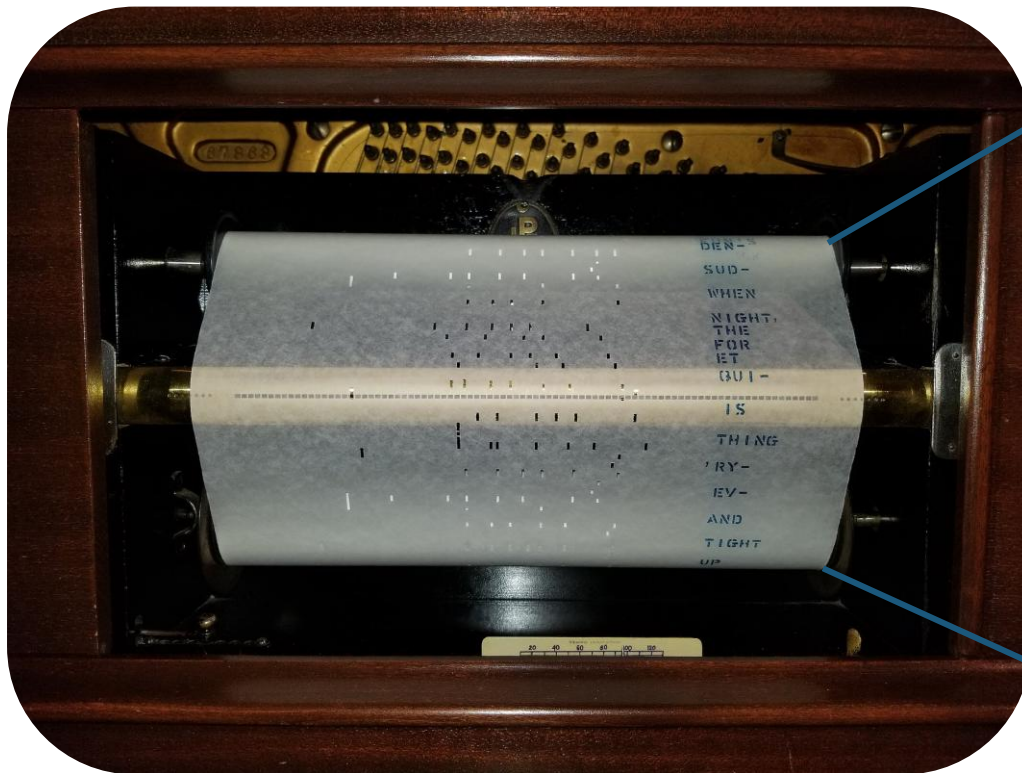
- Treat music like **images**
- Sharing models with **computer vision (CV) & computer graphics (CG)**
 - GANs, VAEs, diffusion models, etc.

Today's topic!

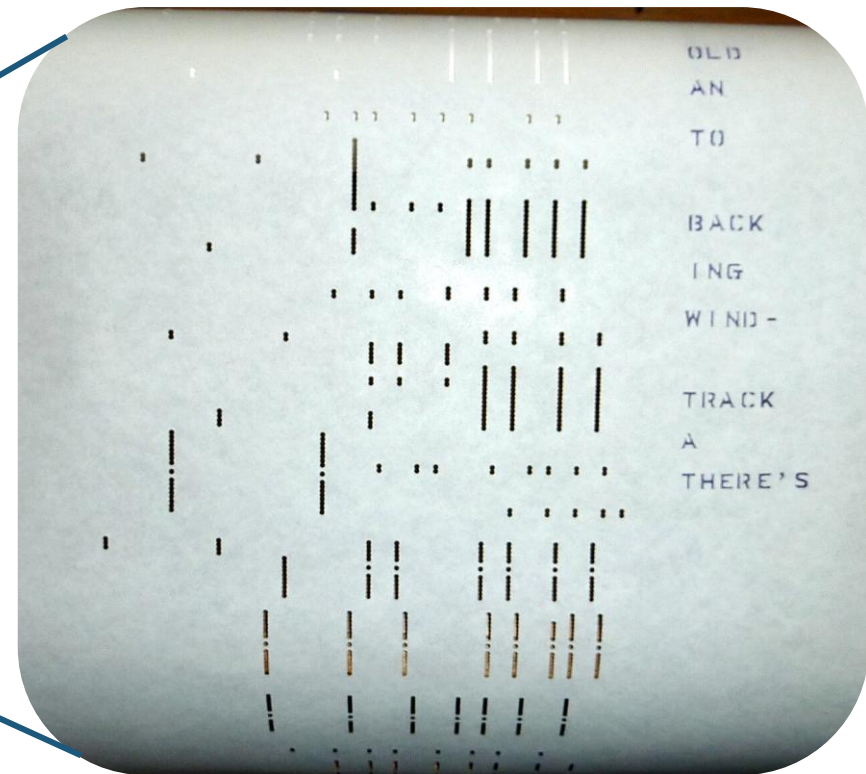


Piano Roll Representation

Piano Rolls



(Source: Draconichiaro)



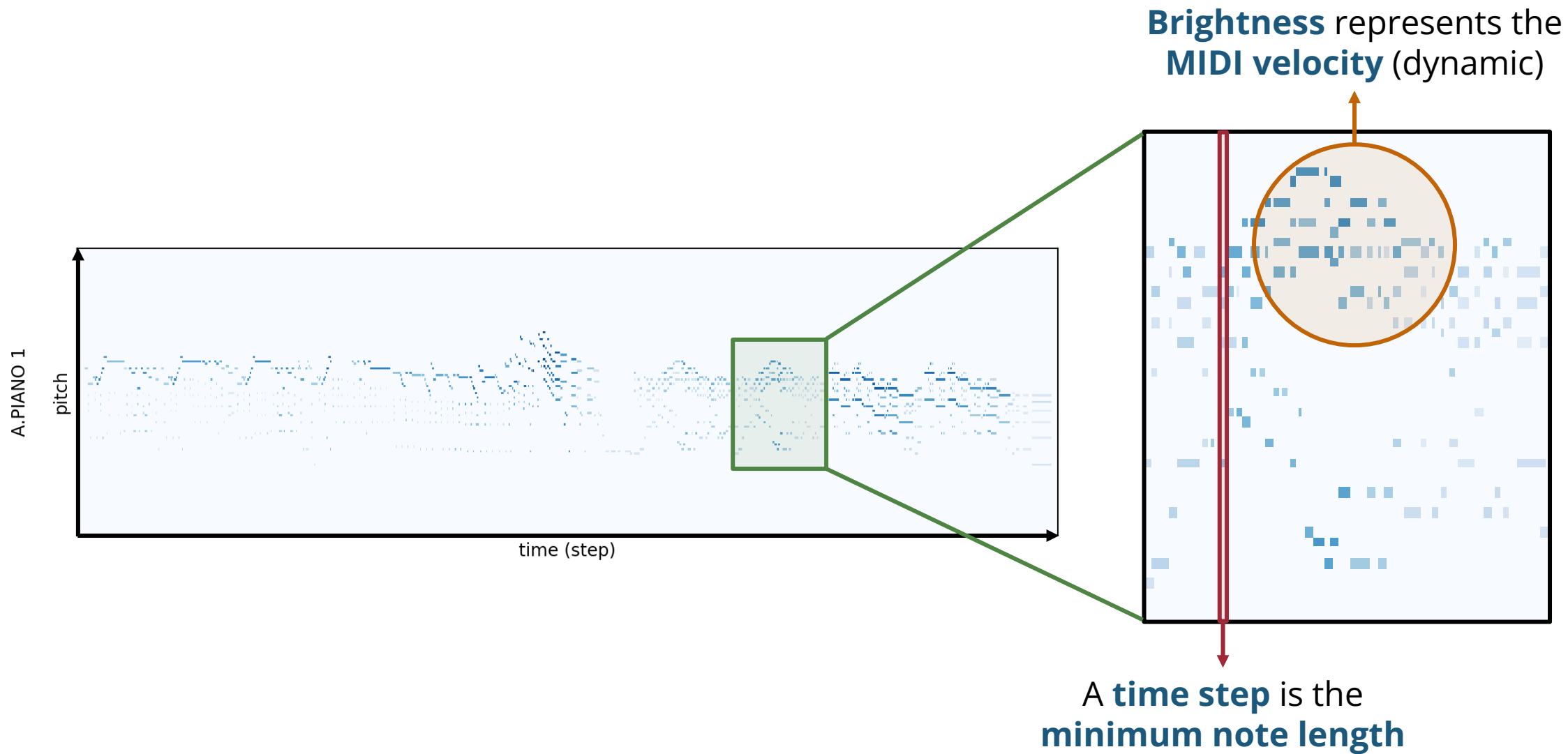
(Source: Tangerineduel)

| Player Pianos



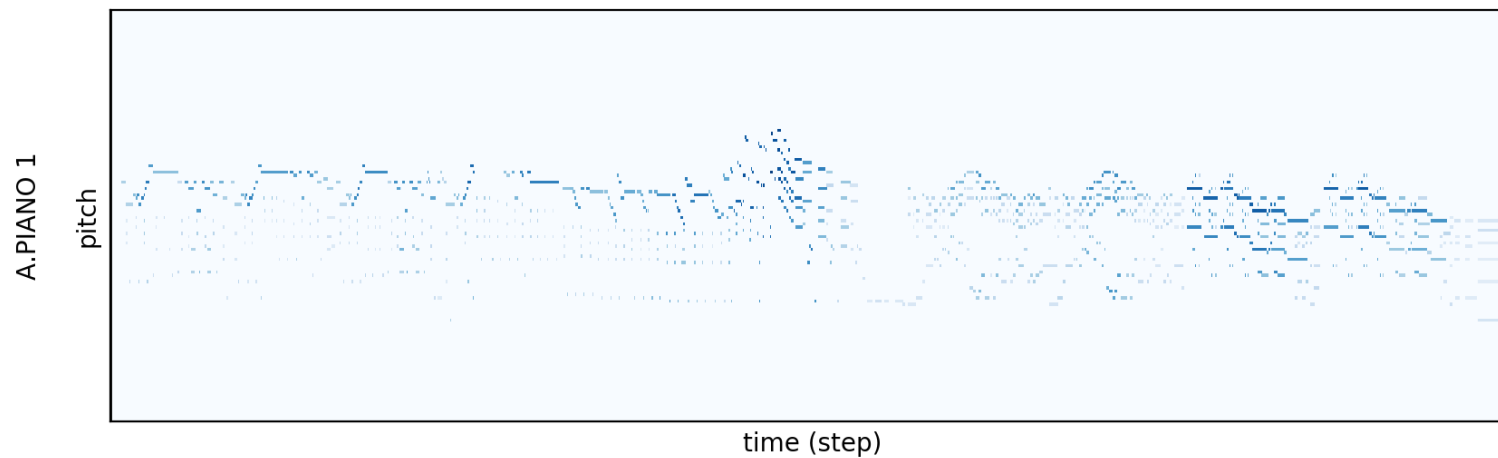
youtu.be/07krQ661fok

Piano Roll Representation

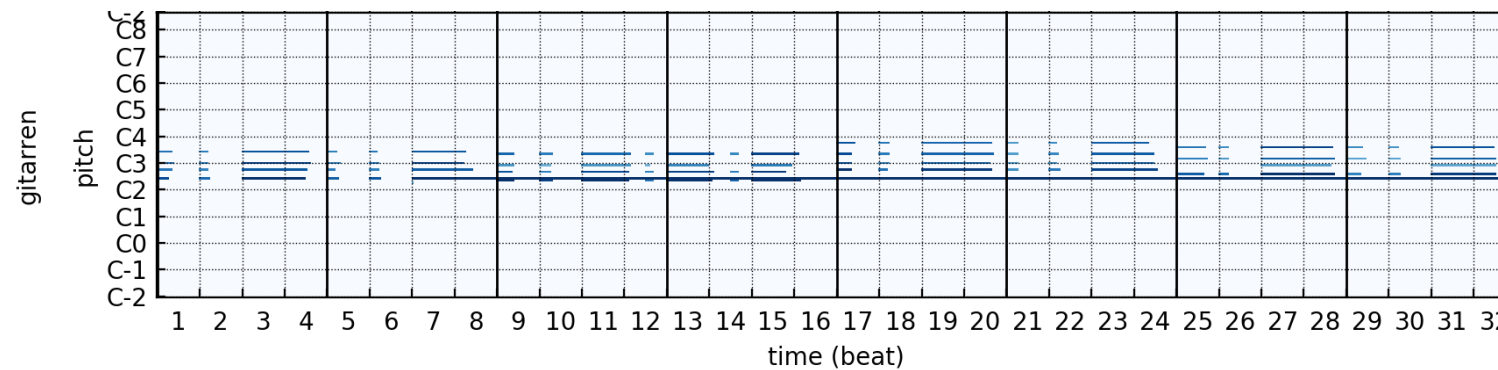


Piano Roll Representation

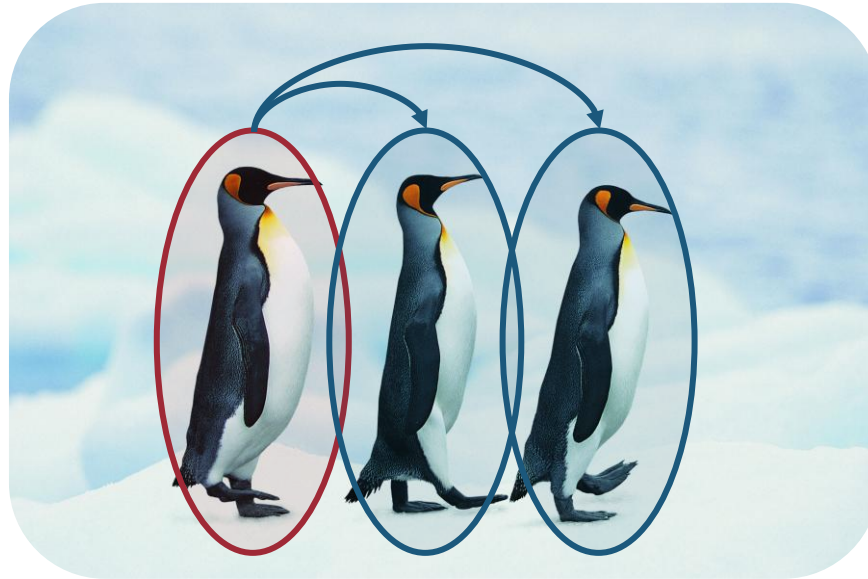
With expressive timing



Without expressive timing



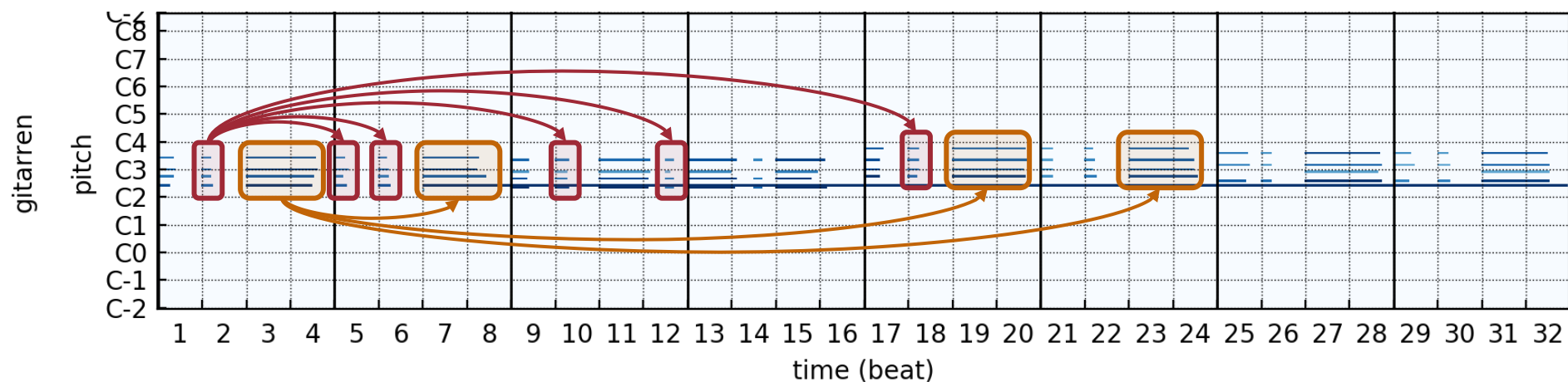
| Reusable Pattern Detectors



Reusable Pattern Detectors



Why Piano Rolls?

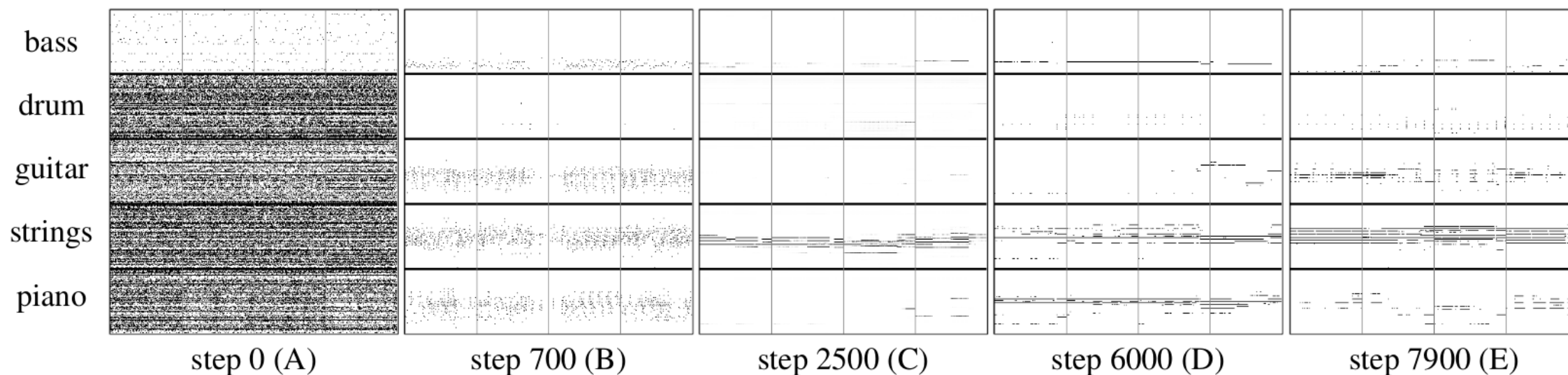


Many musical patterns like melodies, chords, scales and arpeggios
are **translational invariant** in the temporal and pitch axes

MuseGAN: A GAN for Pianorolls (Dong et al., 2018)

The generator improves over time

So does the discriminator!

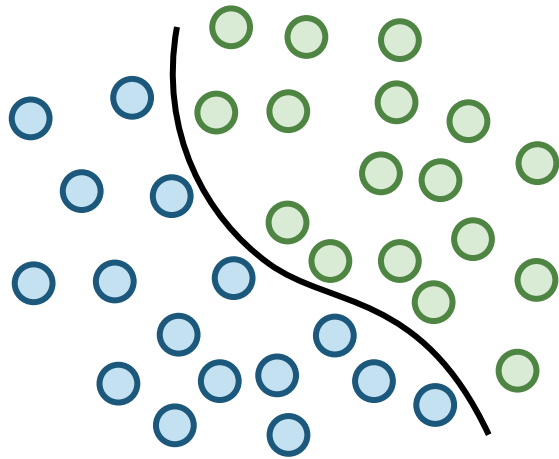


(Source: Dong et al., 2018)

Generative Adversarial Nets (GANs)

Discriminative vs Generative Models

Discriminative



Discriminative models learn the decision boundary

$$P(y|x)$$

Generative

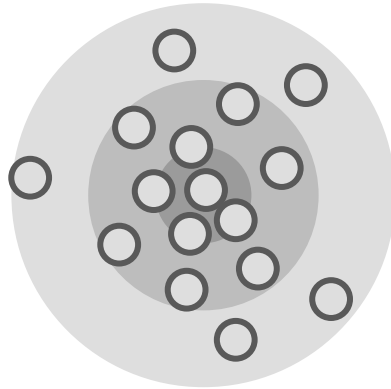


Generative models learn the underlying distribution

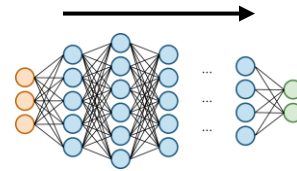
$$P(x) \text{ or } P(x|y)$$

Generating Data from a Random Distribution

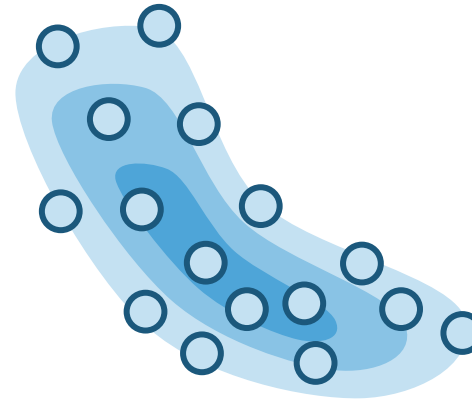
Random distribution



$P(z)$



Data distribution

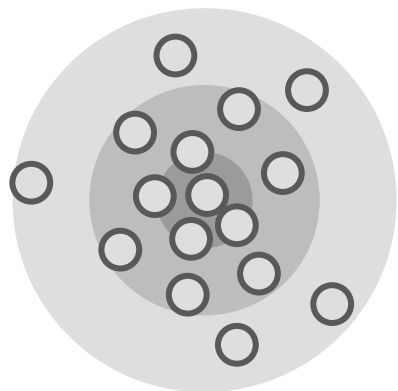


$P(x)$

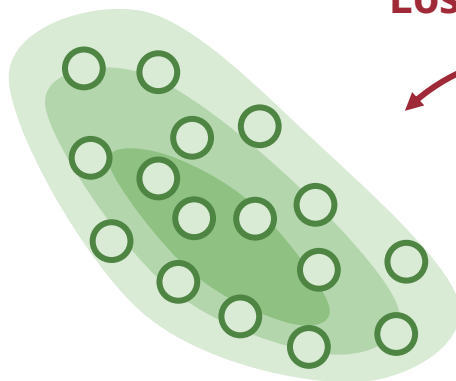
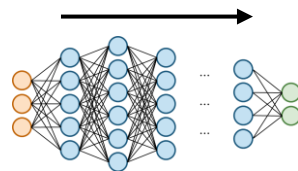
If we can learn this mapping, we can then
generate new samples from the data distribution

A Loss Function for Distributions

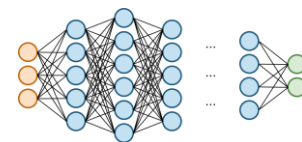
Random distribution



$P(z)$



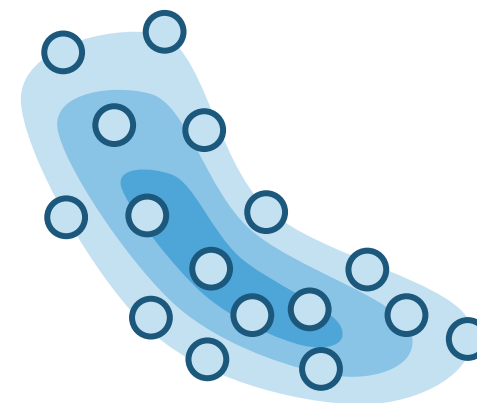
$P(\hat{x})$



Loss function?



Data distribution

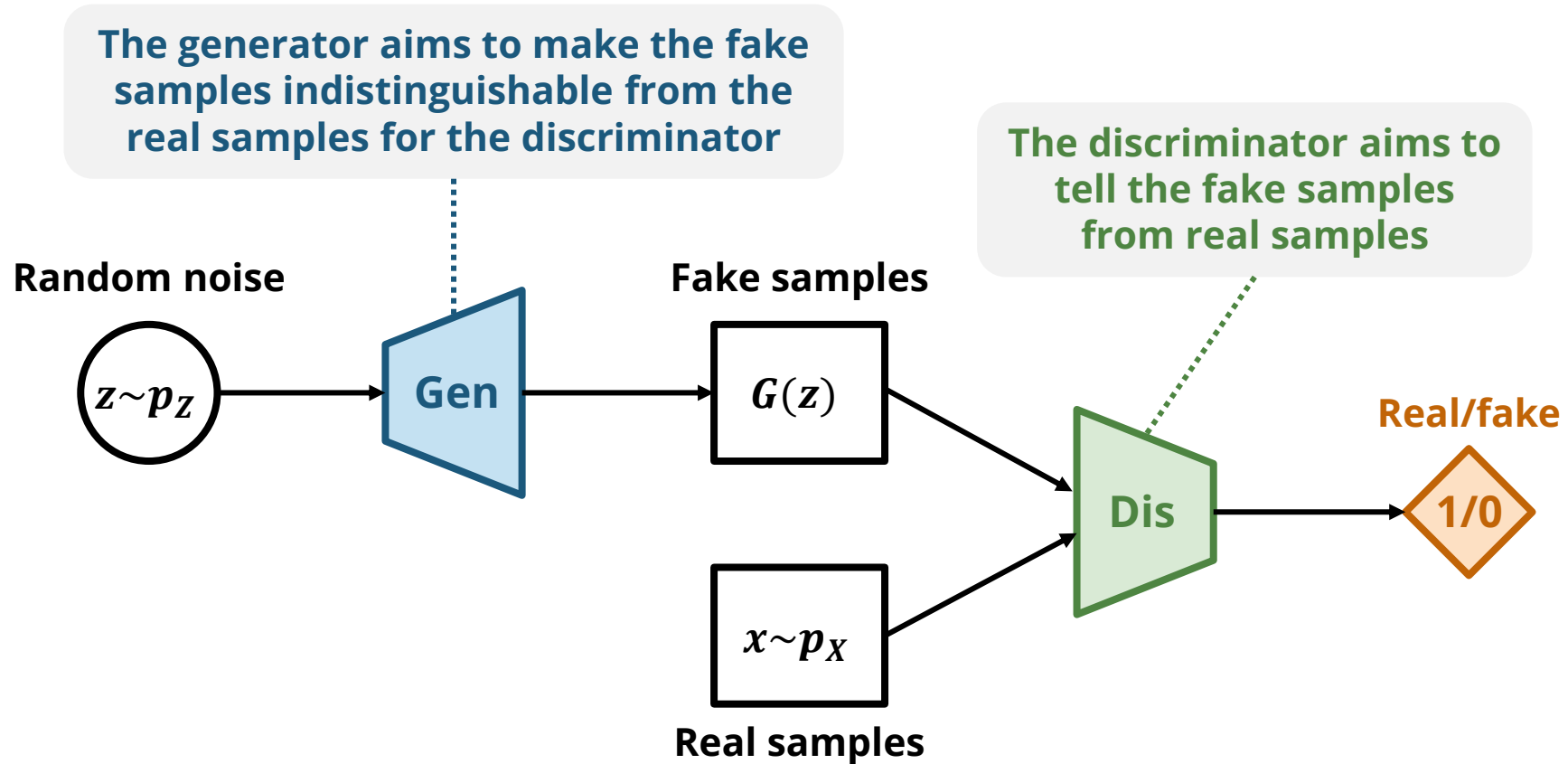


$P(x)$

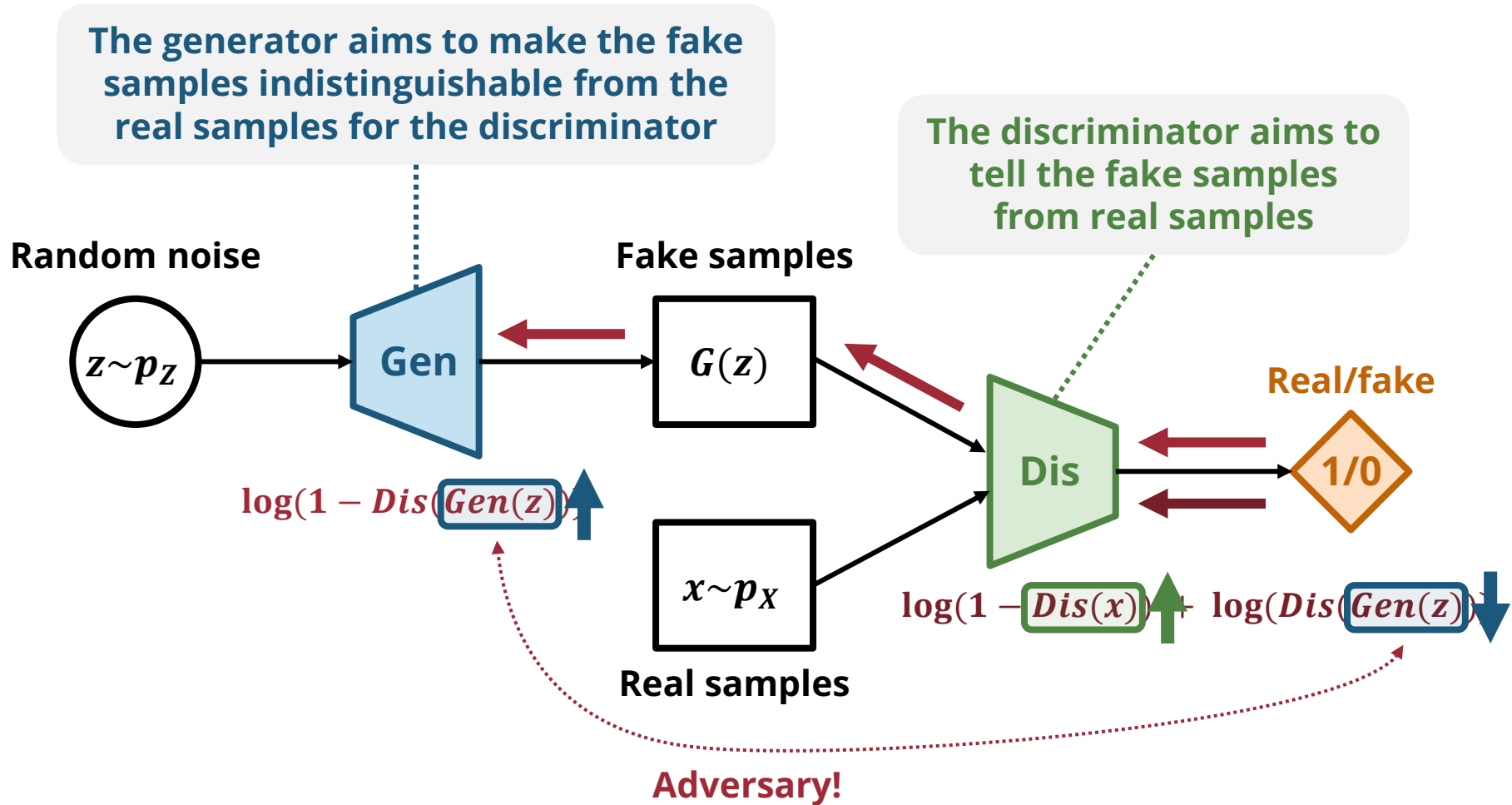
Unfortunately, no easy way to
measure the difference from samples

But what about another neural network!?

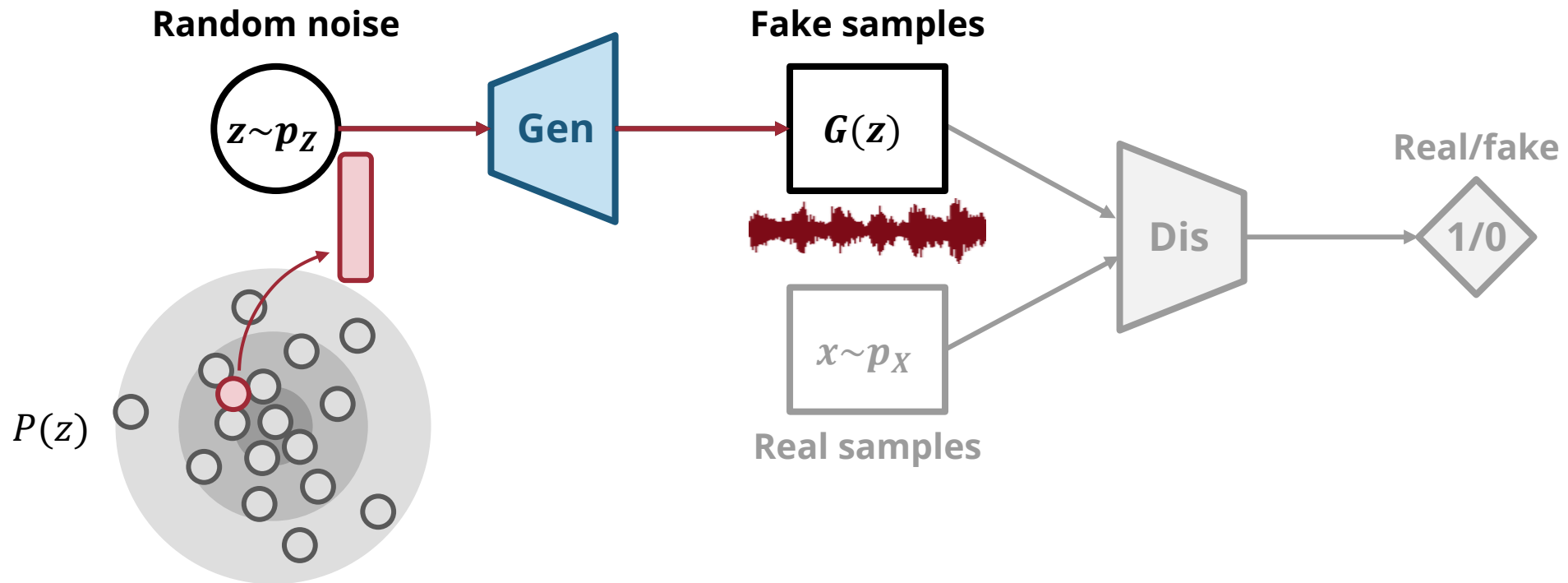
Generative Adversarial Nets (GANs) (Goodfellow et al., 2014)



Generative Adversarial Nets (GANs): Training

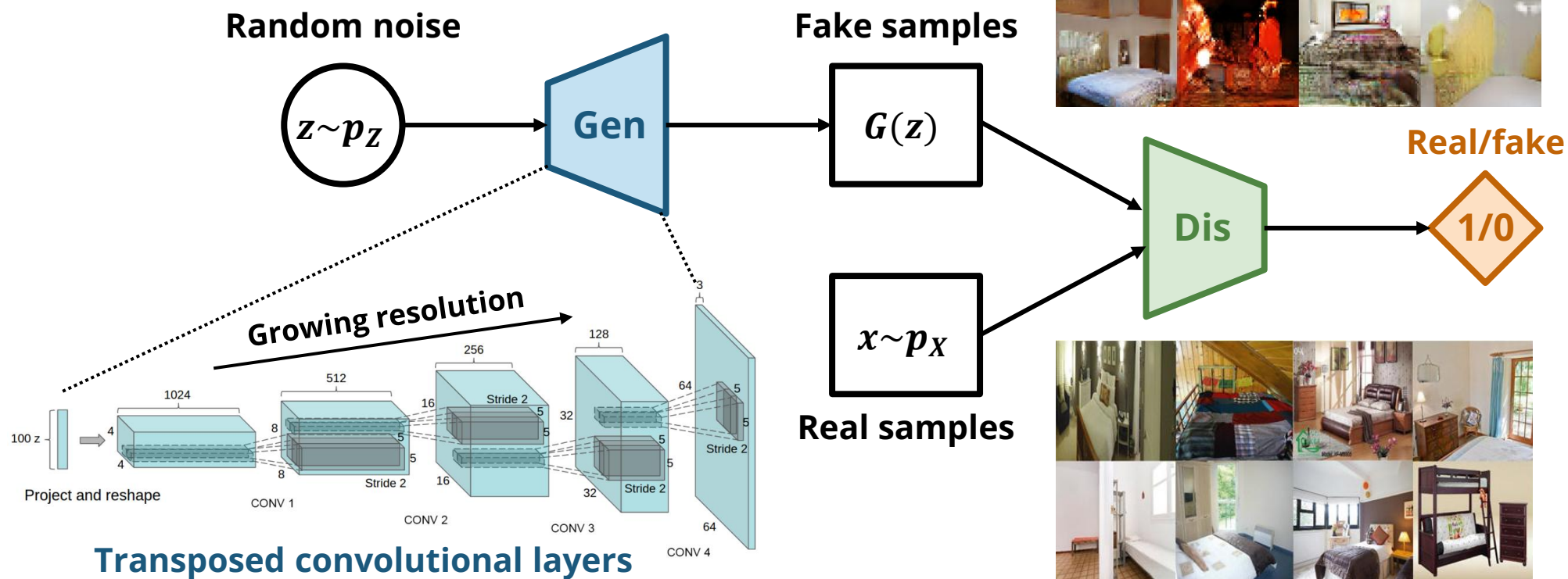


Generative Adversarial Nets (GANs): Generation



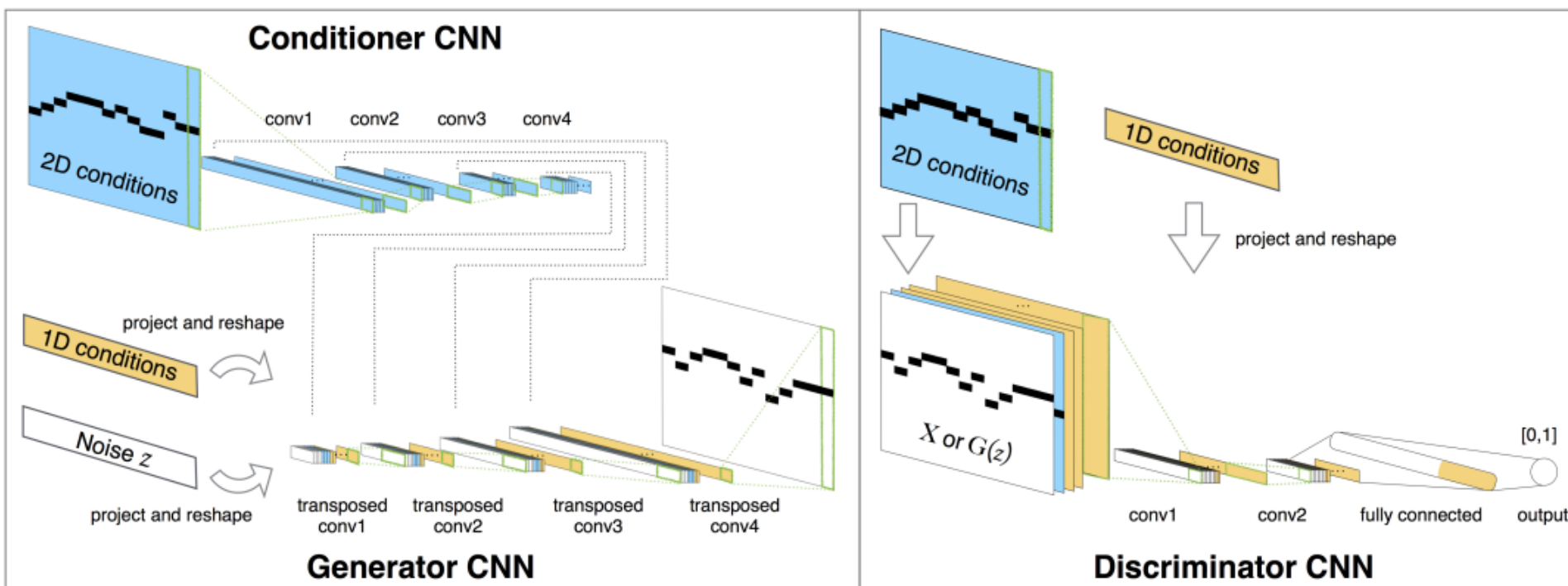
Deep Convolutional GANs (DCGANs) (Radford et al., 2014)

Use CNNs for both the generator and discriminator



Generating Music using GANs

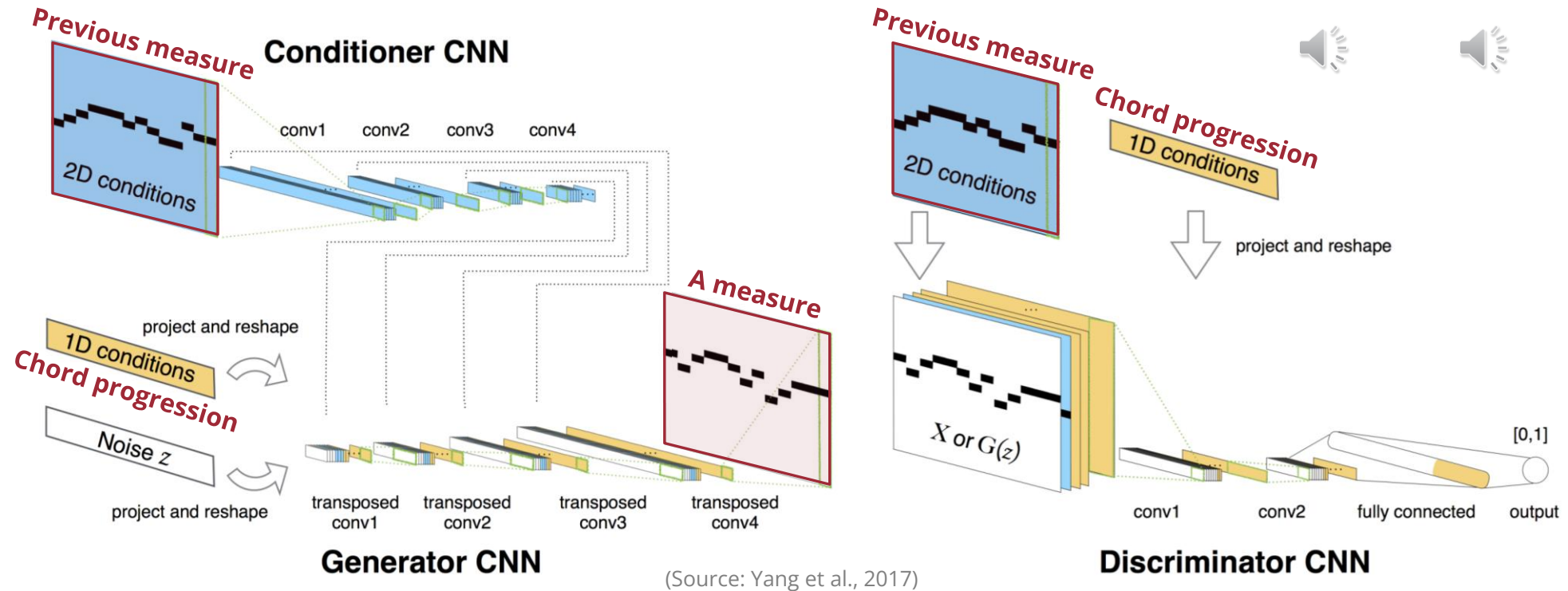
MidiNet (Yang et al., 2017)



(Source: Yang et al., 2017)

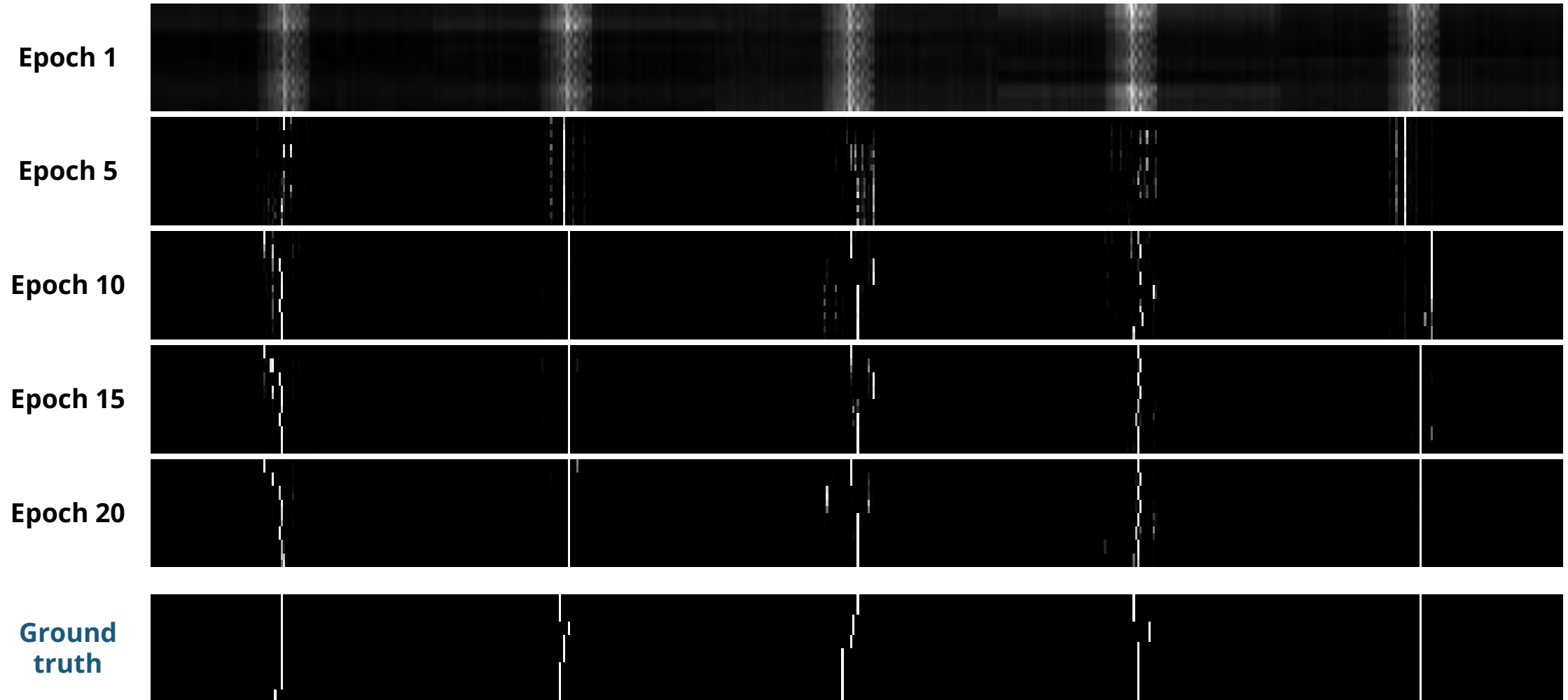
MidiNet (Yang et al., 2017)

Examples of
generated music



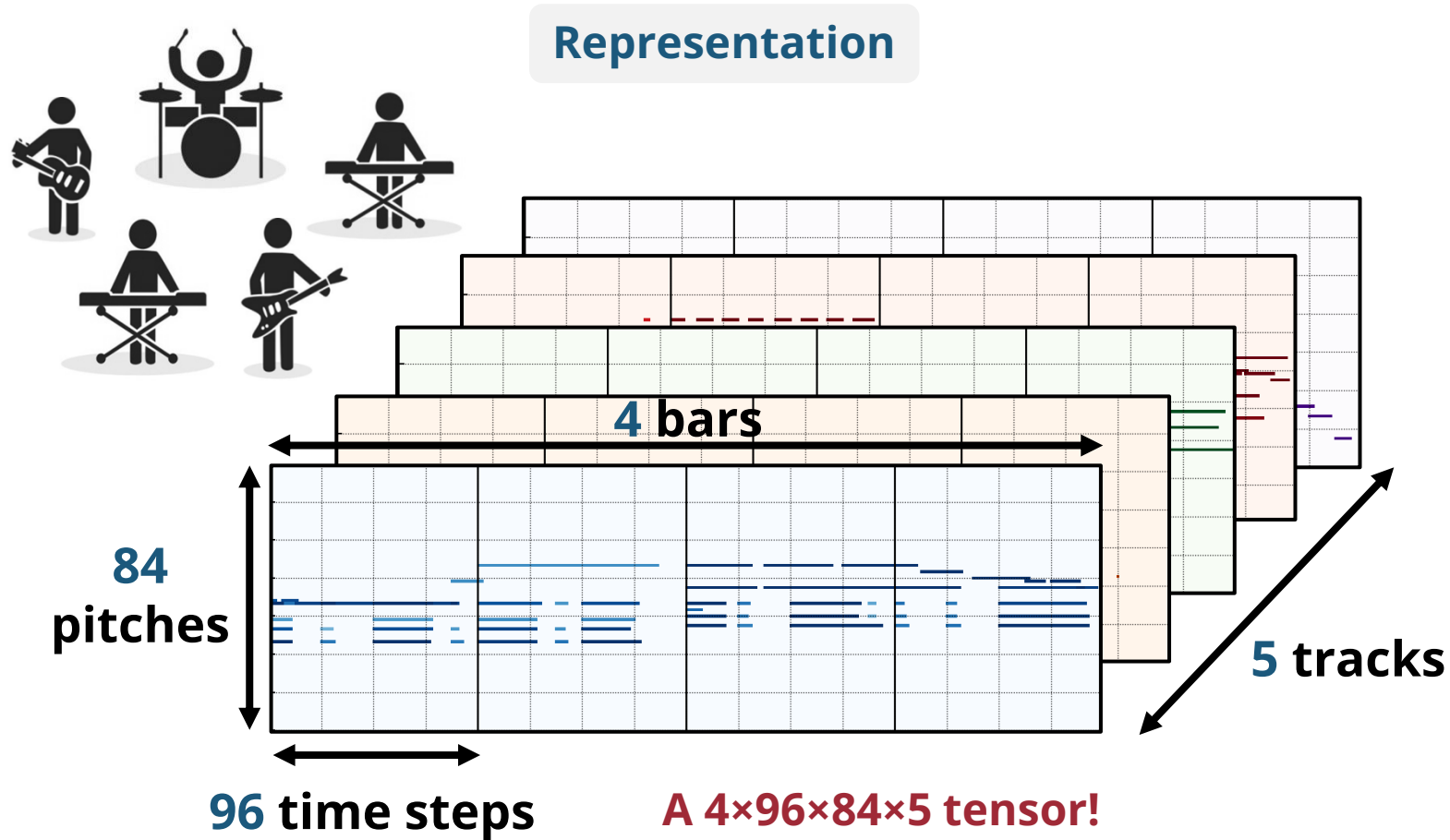
Generate the next measure given the previous one
(measure by measure)

MidiNet (Yang et al., 2017)

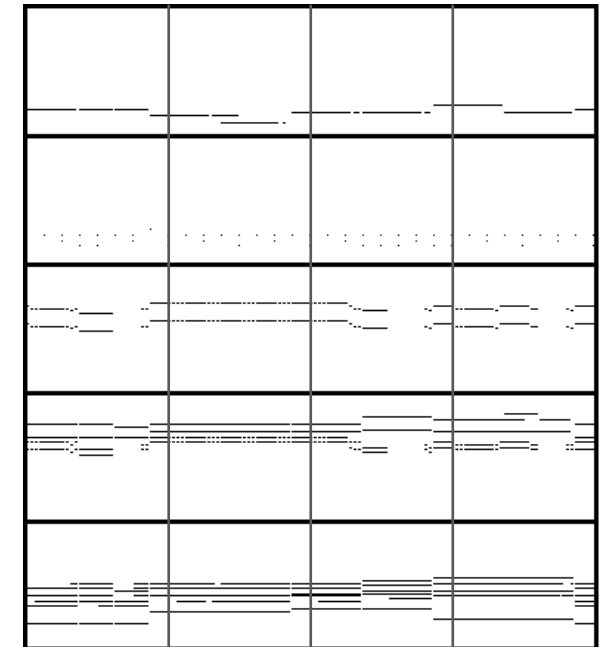


(Source: Yang et al., 2017)

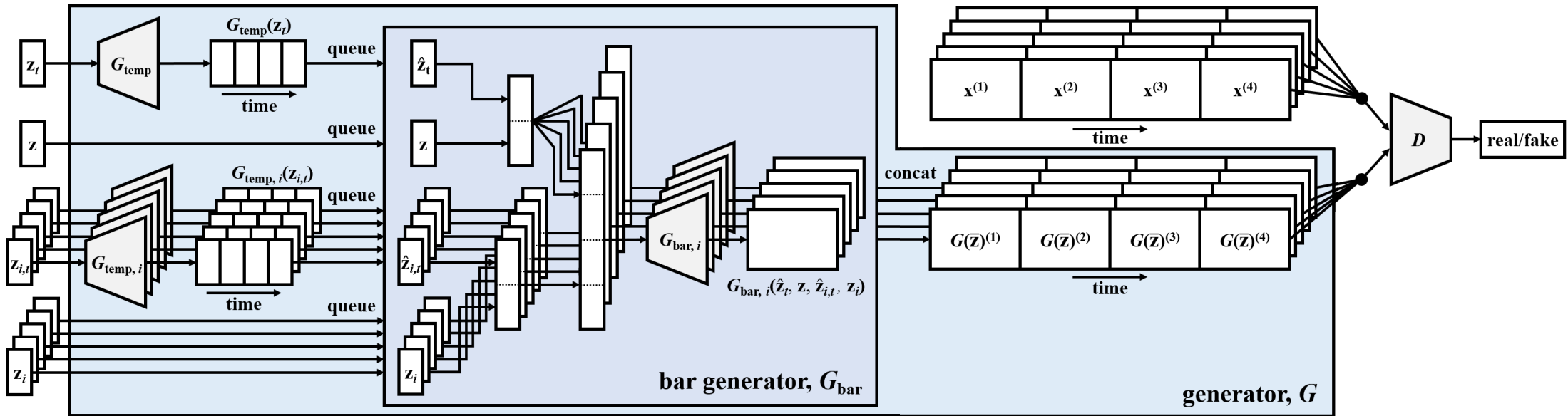
MuseGAN (Dong et al., 2018)



A training sample



MuseGAN (Dong et al., 2018)



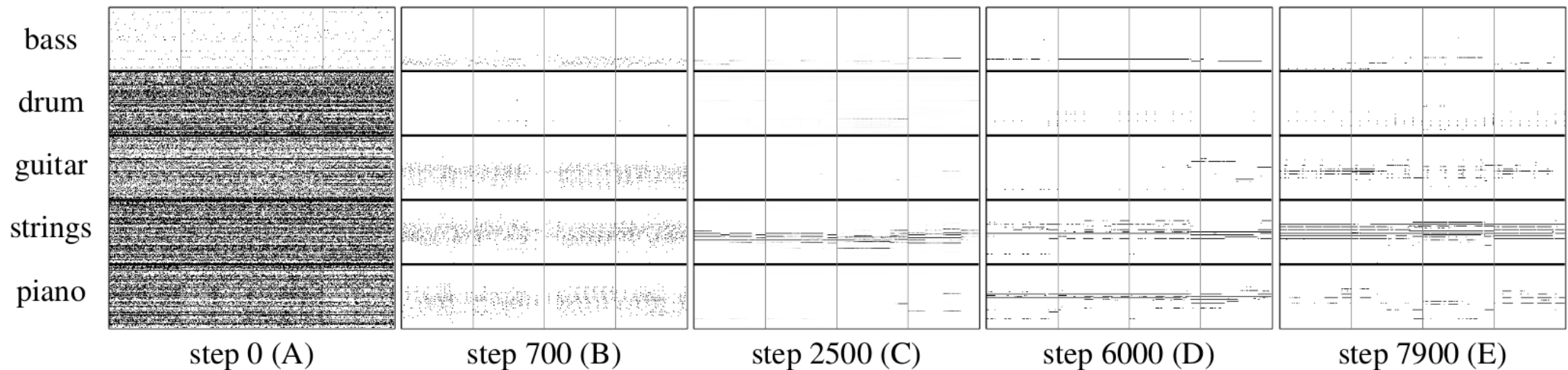
(Source: Dong et al., 2018)

MuseGAN (Dong et al., 2018)

Examples of
generated music



The generator improves over time

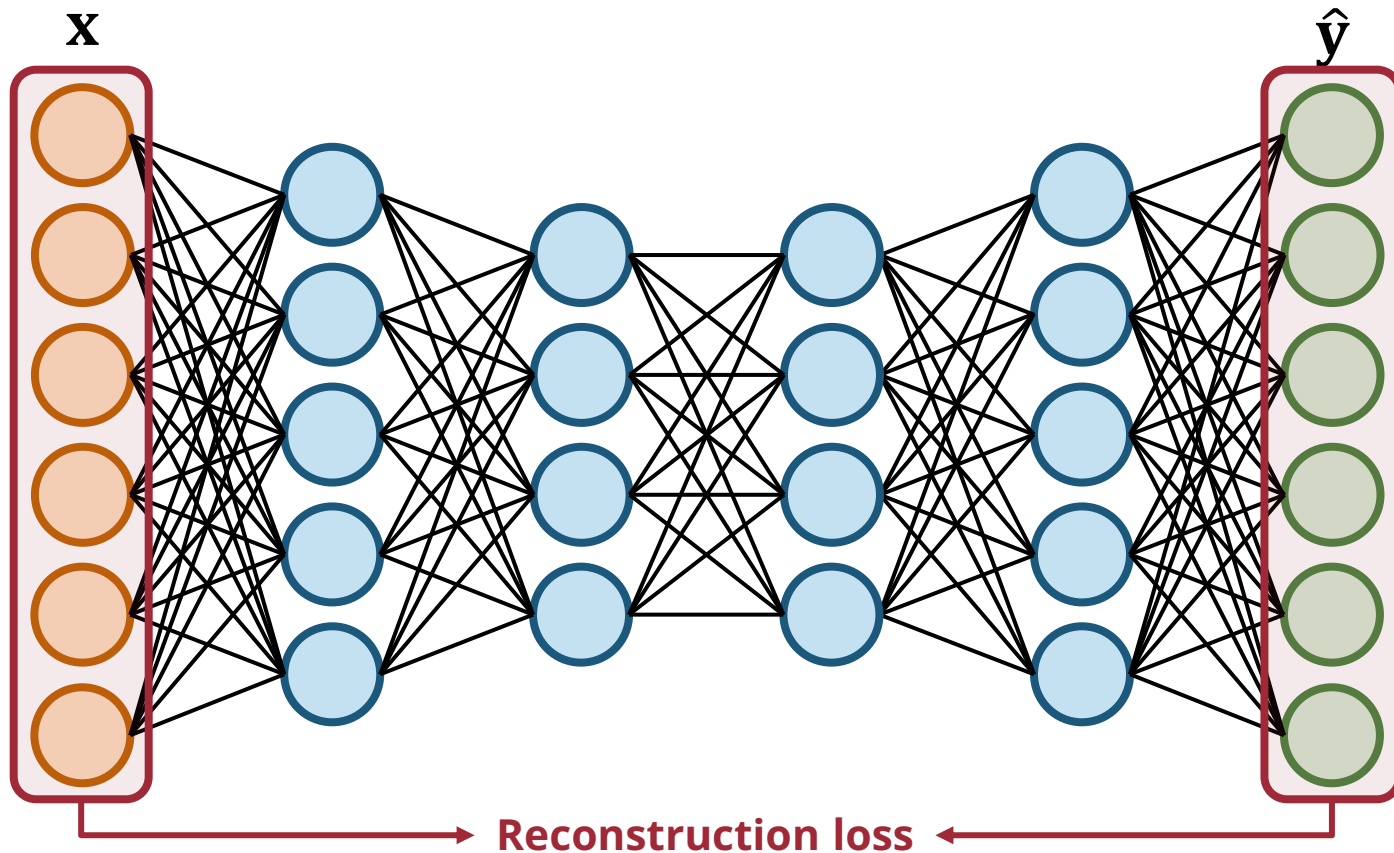


(Source: Dong et al., 2018)

Diffusion Models

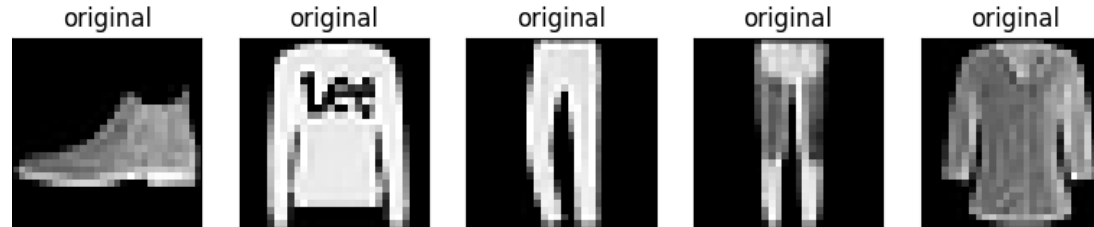
Autoencoders

- A neural network where the **input and output are the same**

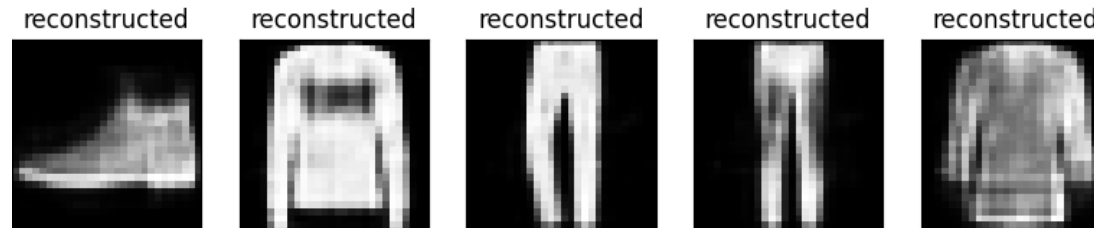


Autoencoders: Reconstruction Examples

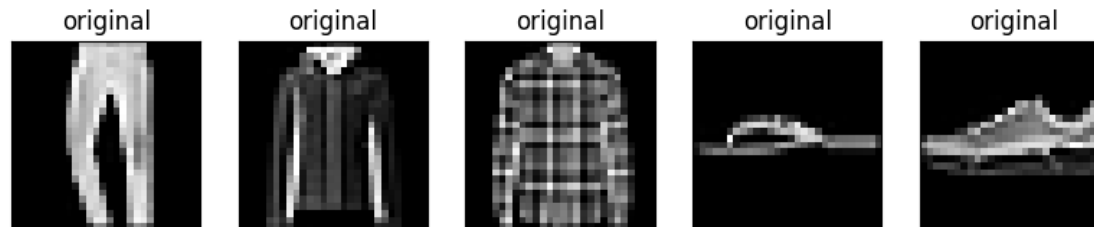
Original



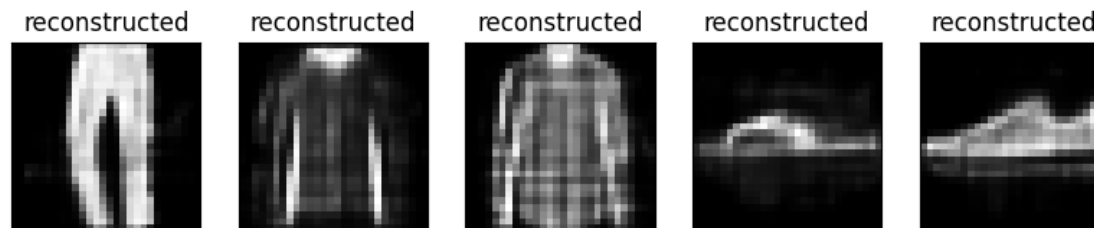
Reconstructed



Original

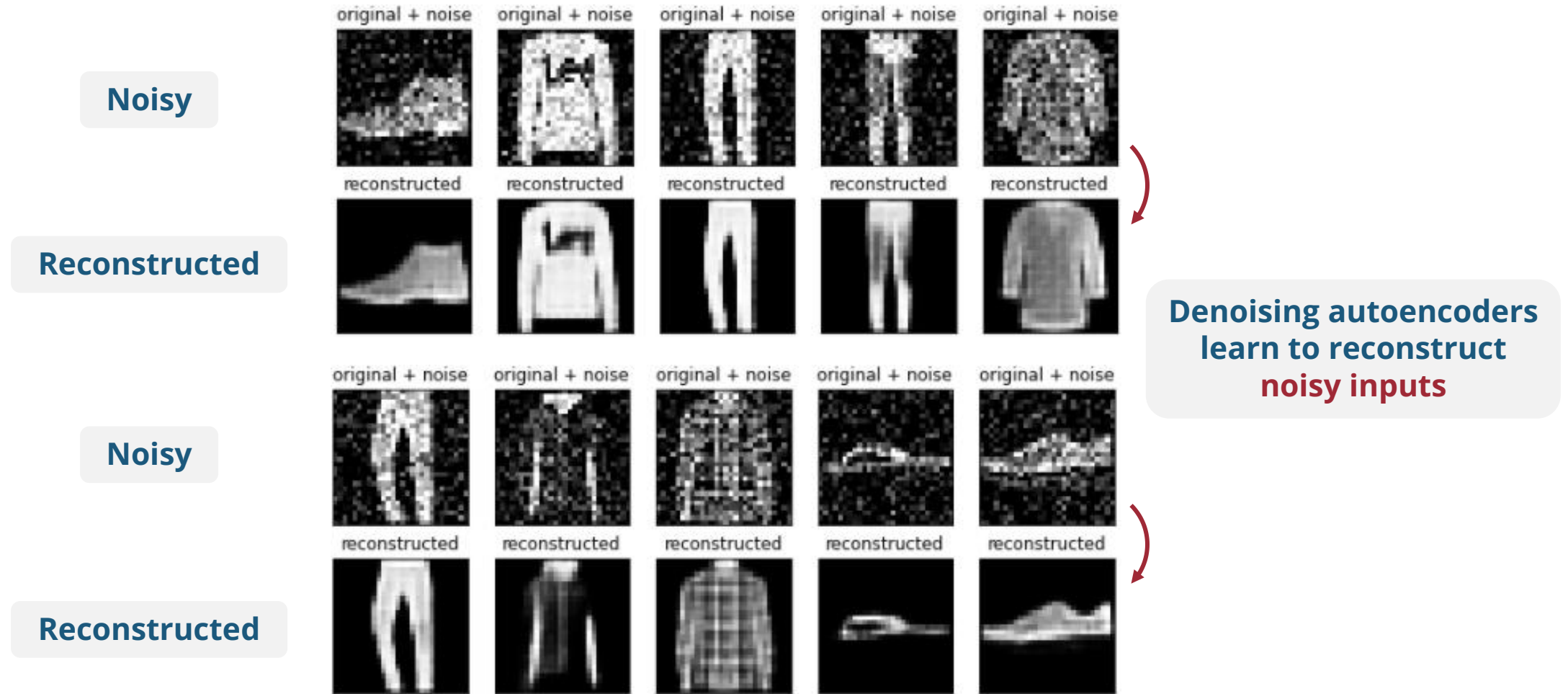


Reconstructed



(Source: tensorflow.org)

Denoising Autoencoders (Pascal et al., 2008)



(Source: tensorflow.org)

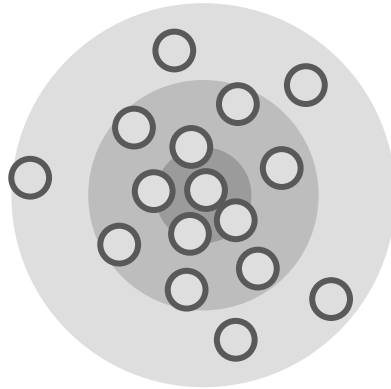
tensorflow.org/tutorials/generative/autoencoder

Pascal Vincent, Hugo Larochelle, Yoshua Bengio, and Pierre-Antoine Manzagol, "Extracting and Composing Robust Features with Denoising Autoencoders," *ICML*, 2008.

Pascal Vincent, Hugo Larochelle, Isabelle Lajoie, Yoshua Bengio, and Pierre-Antoine Manzagol, "Stacked Denoising Autoencoders: Learning Useful Representations in a Deep Network with a Local Denoising Criterion," *PMLR*, 11(110):3371-2408, 2010.

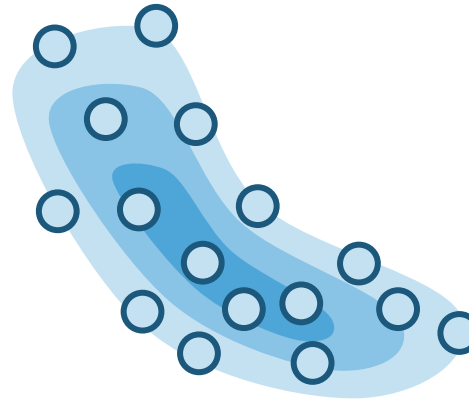
Generating Data from a Random Distribution

Random distribution

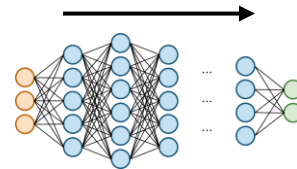


$P(z)$

Data distribution



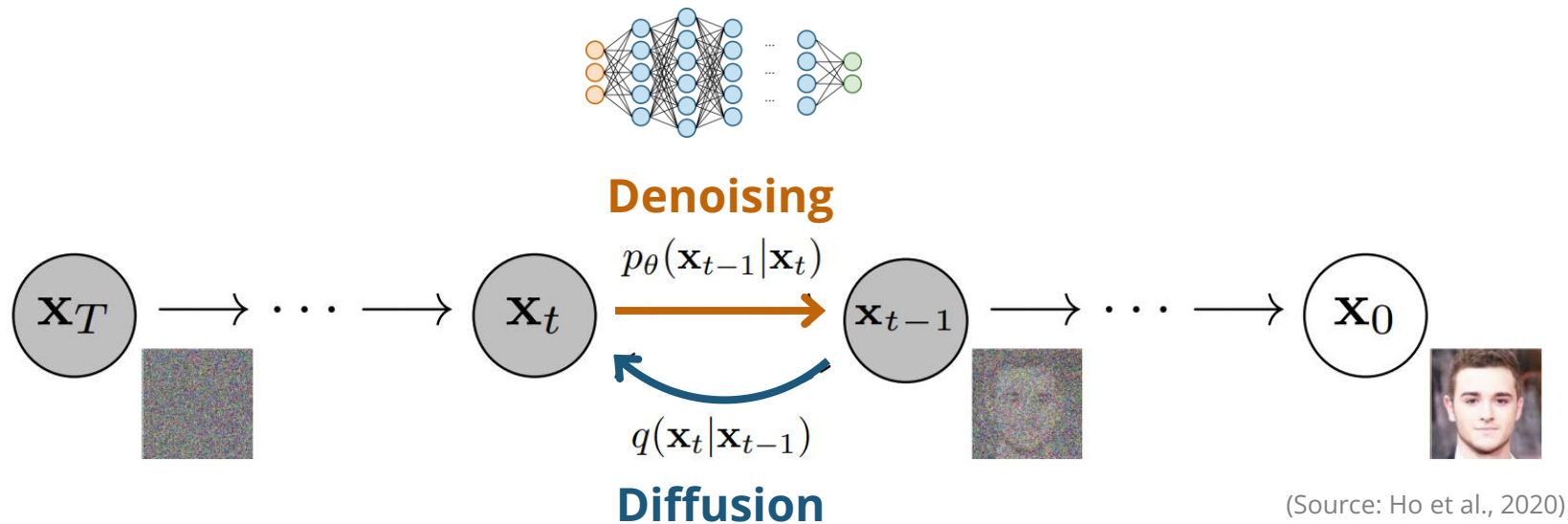
$P(x)$



If we can learn this mapping, we can then
generate new samples from the data distribution

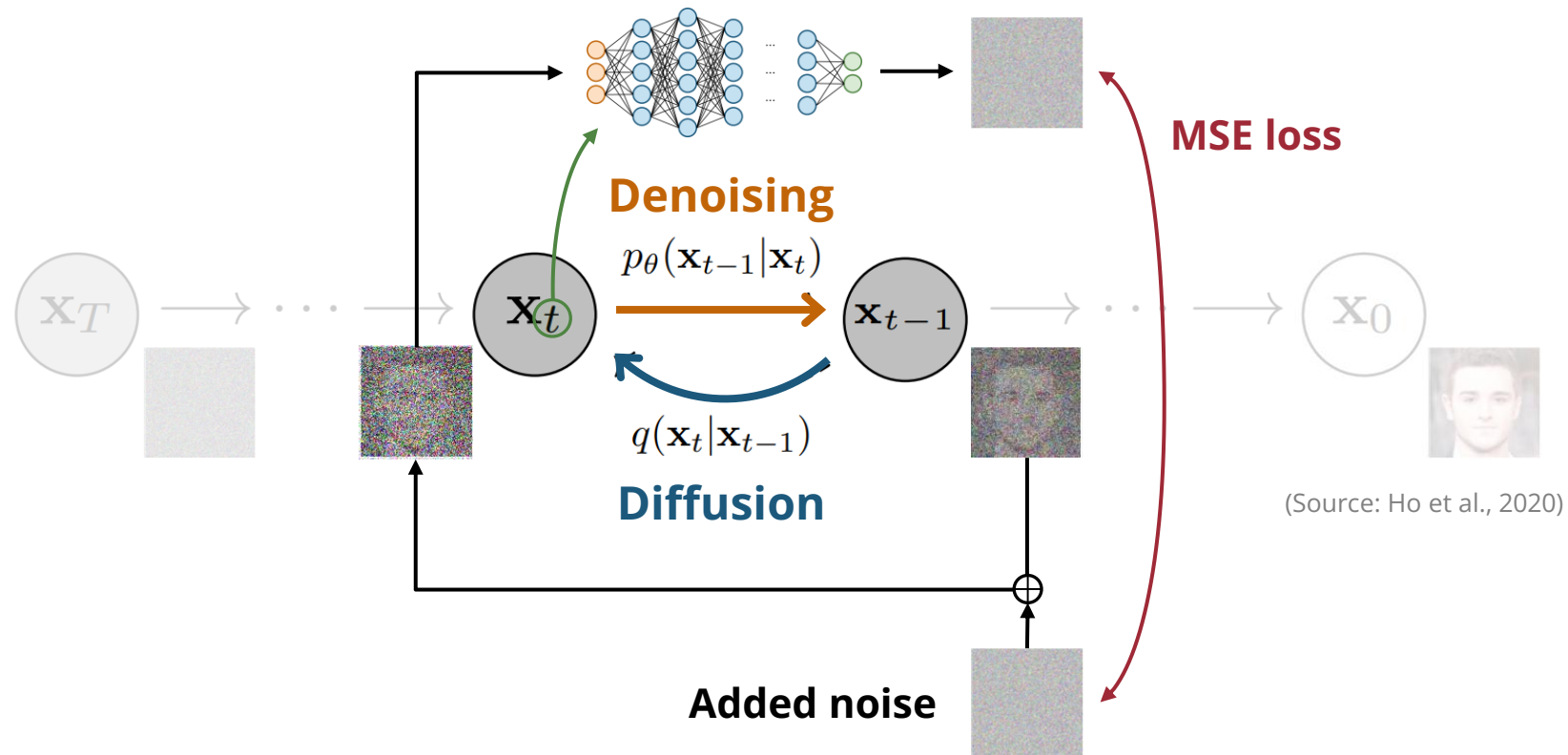
Diffusion Models (Ho et al., 2020)

- **Intuition:** Many denoising autoencoders stacked together



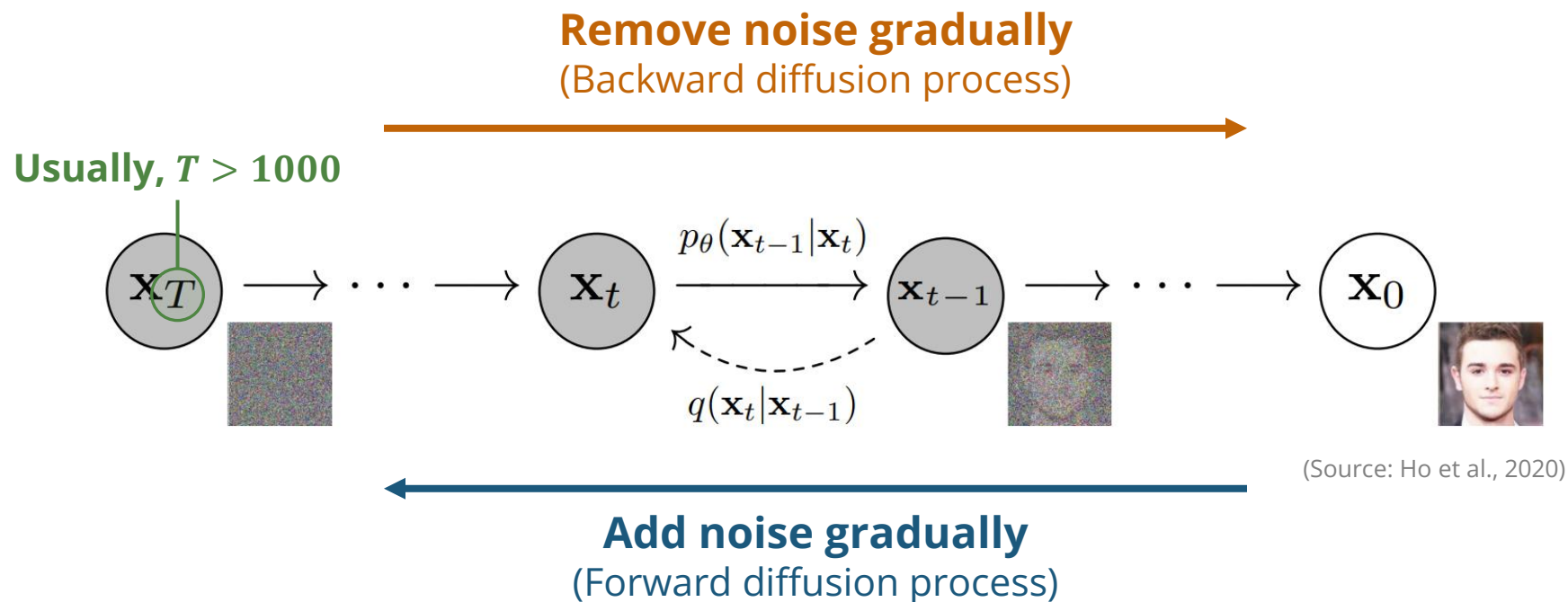
Diffusion Models – Training

- **Intuition:** Many denoising autoencoders stacked together

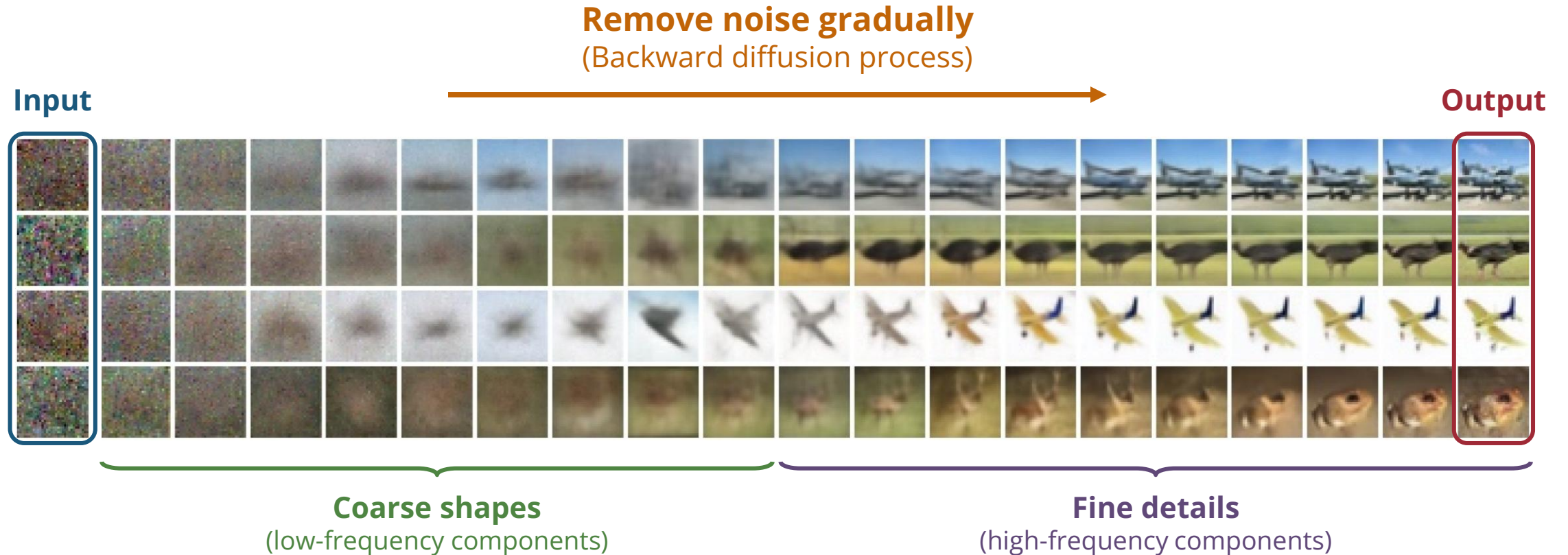


Diffusion Models (Ho et al., 2020)

- **Intuition:** Many denoising autoencoders stacked together



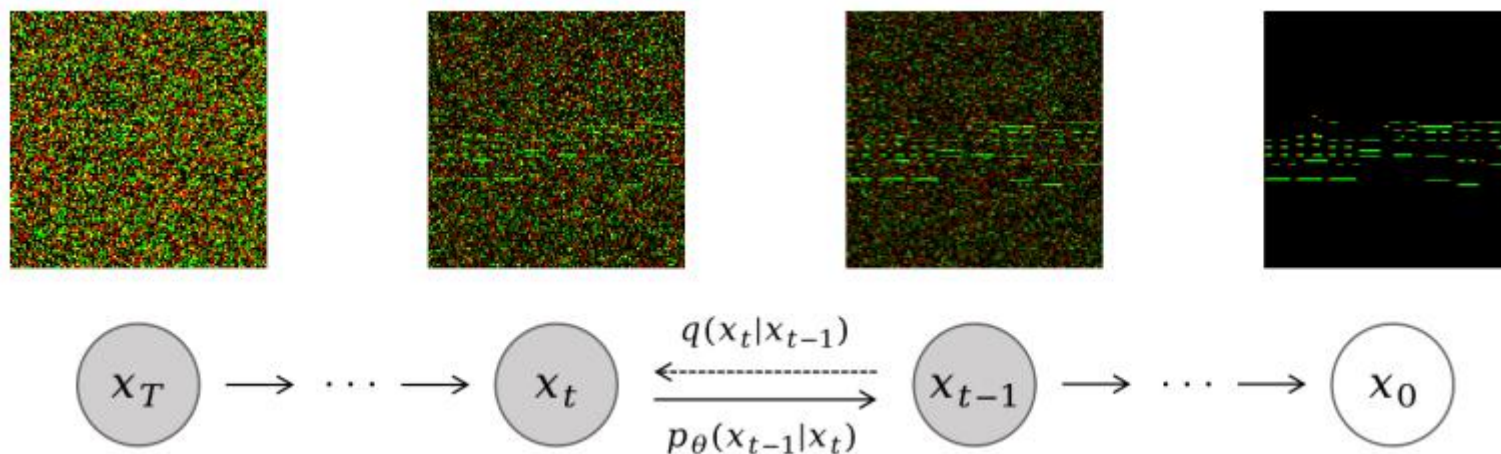
Diffusion Models – Generation



(Source: Ho et al., 2020)

Generating Music using Diffusion Models

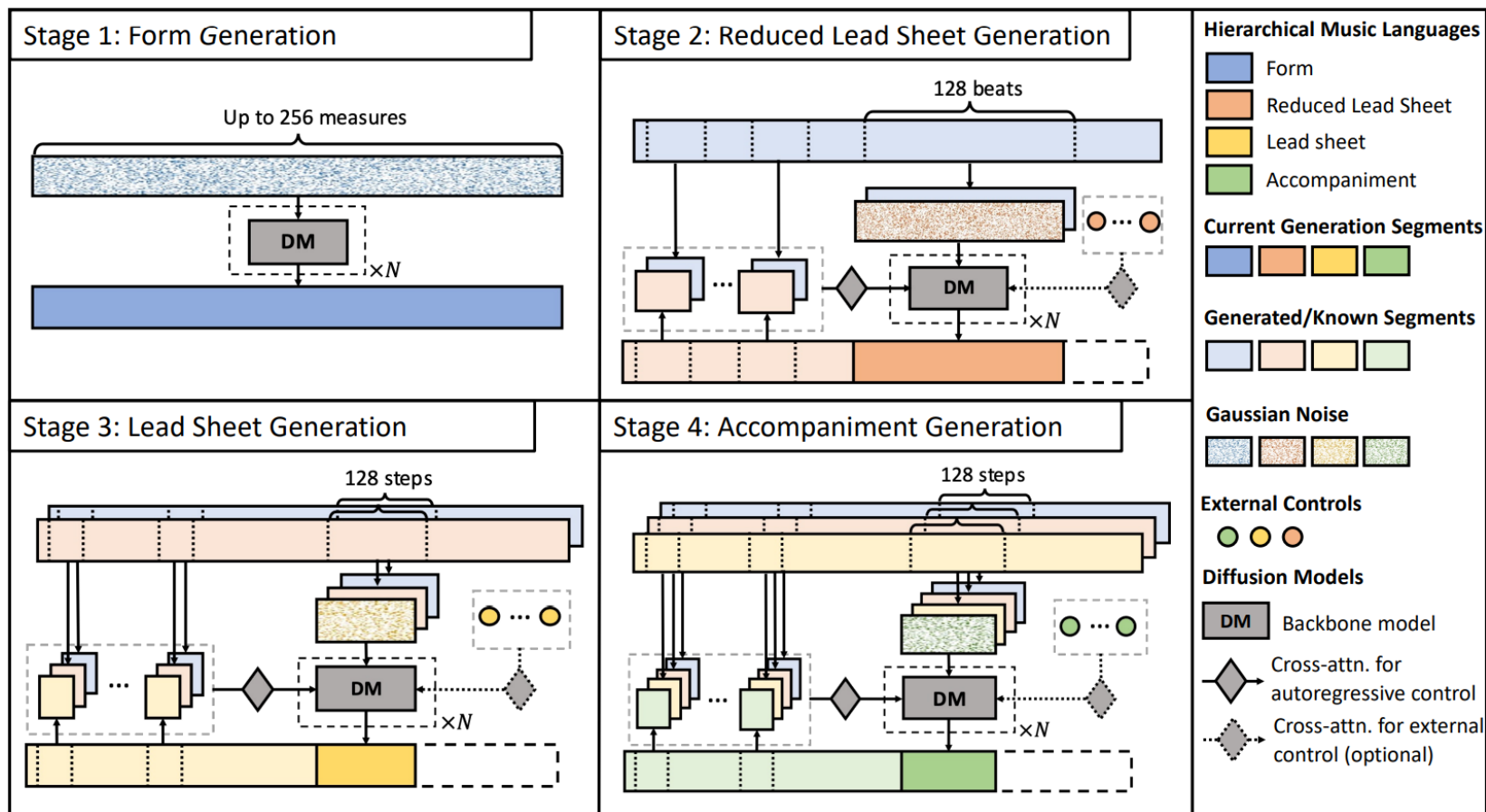
Polyffusion (Min et al., 2023)



(Source: Min et al., 2023)

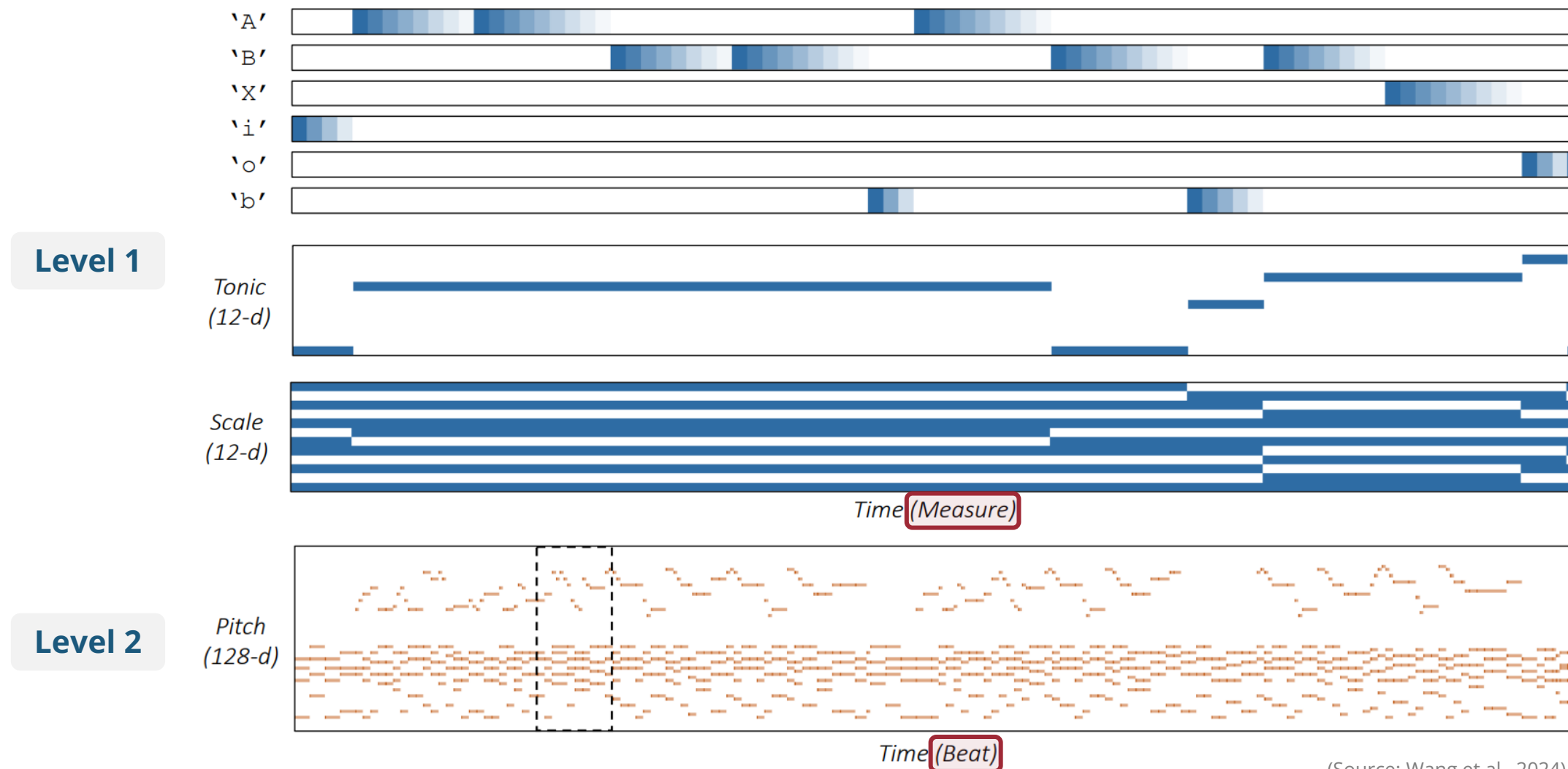
polyffusion.github.io

Cascaded Diffusion Models (Wang et al., 2024)



(Source: Wang et al., 2024)

Cascaded Diffusion Models (Wang et al., 2024)

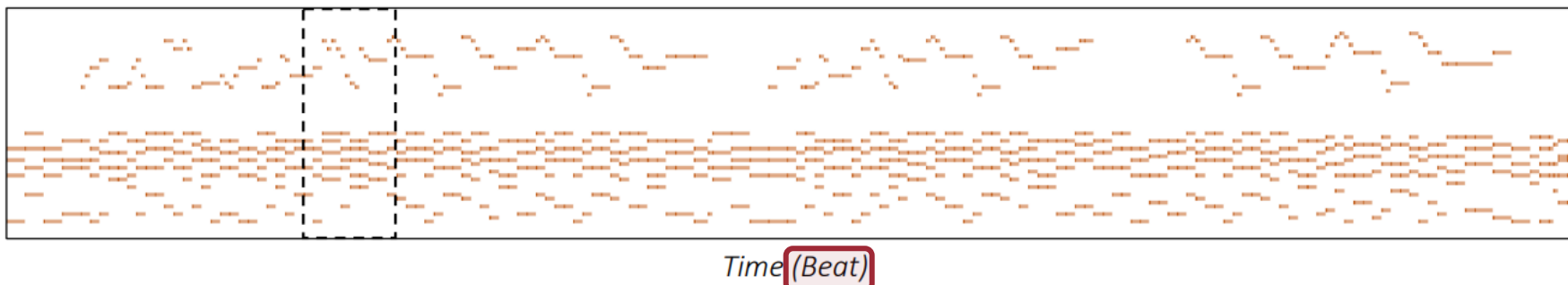


(Source: Wang et al., 2024)

Cascaded Diffusion Models (Wang et al., 2024)

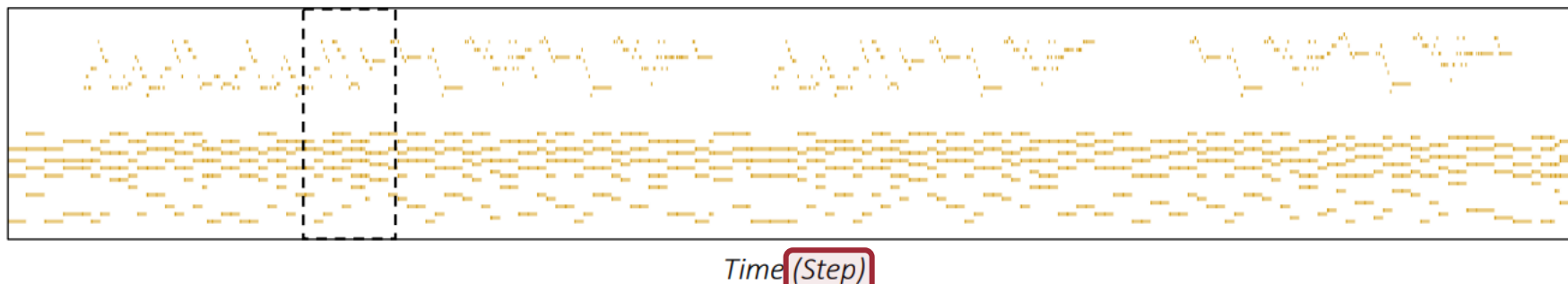
Level 2

Pitch
(128-d)



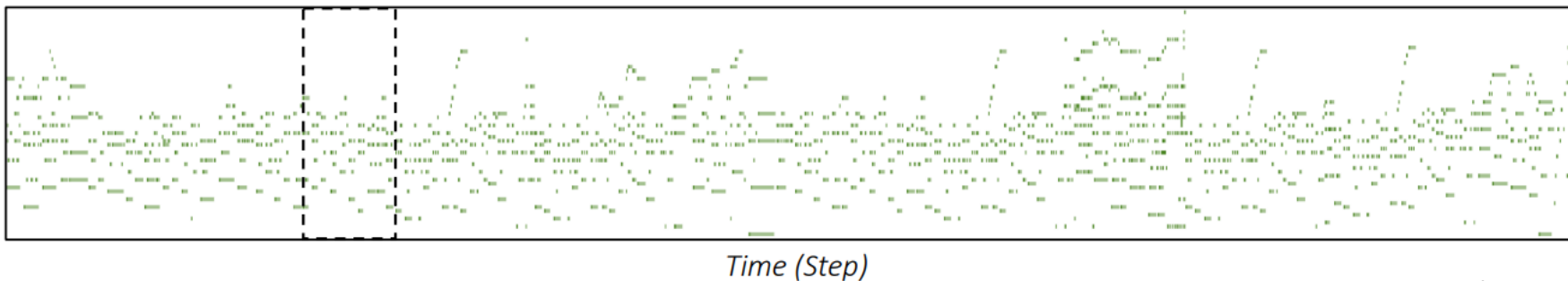
Level 3

Pitch
(128-d)



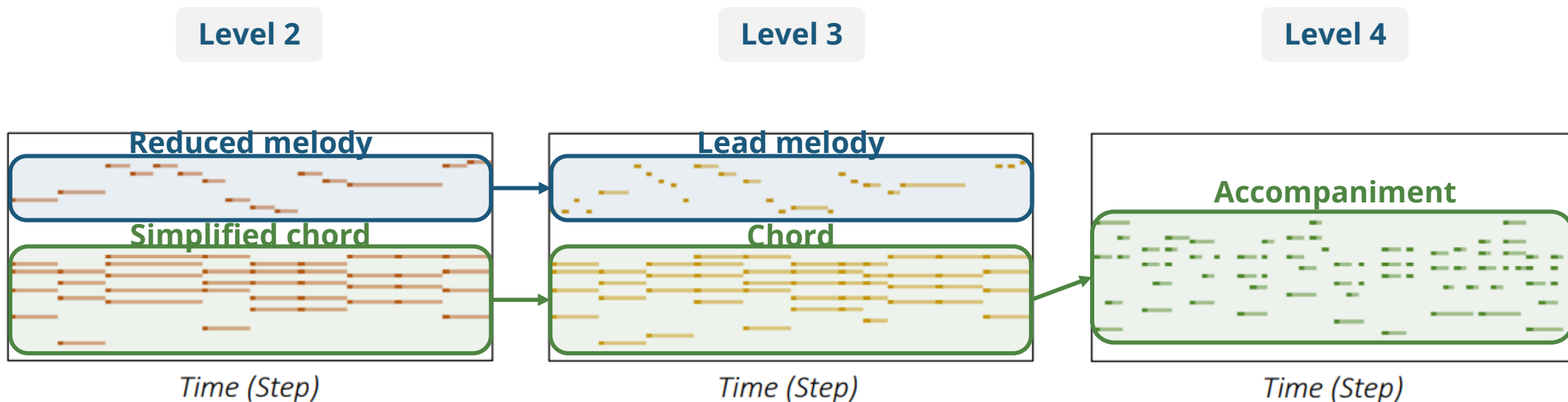
Level 4

Pitch
(128-d)



(Source: Wang et al., 2024)

Cascaded Diffusion Models (Wang et al., 2024)

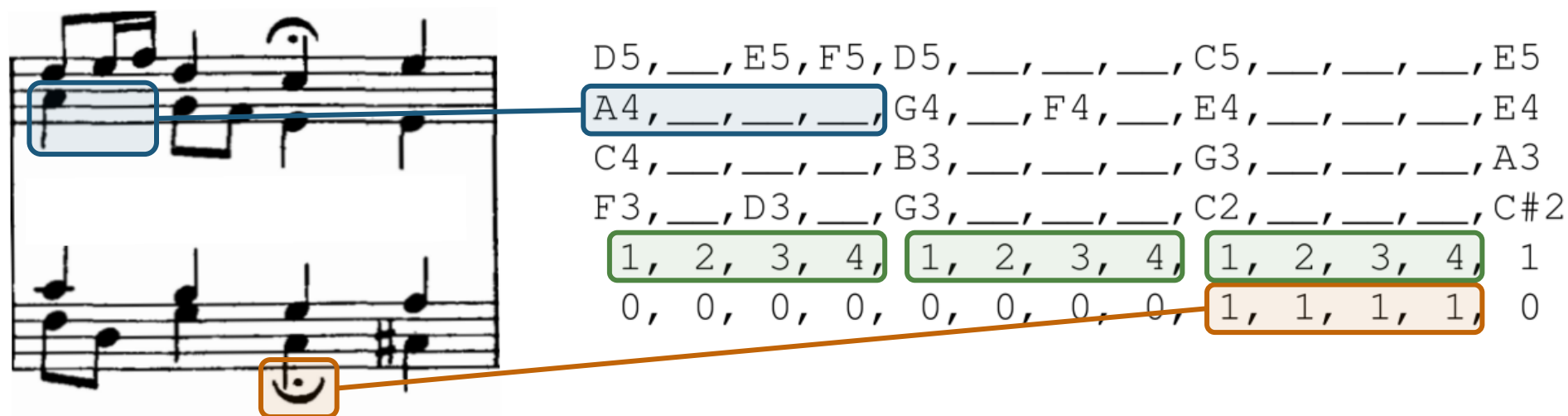


(Source: Wang et al., 2024)

wholesonggen.github.io

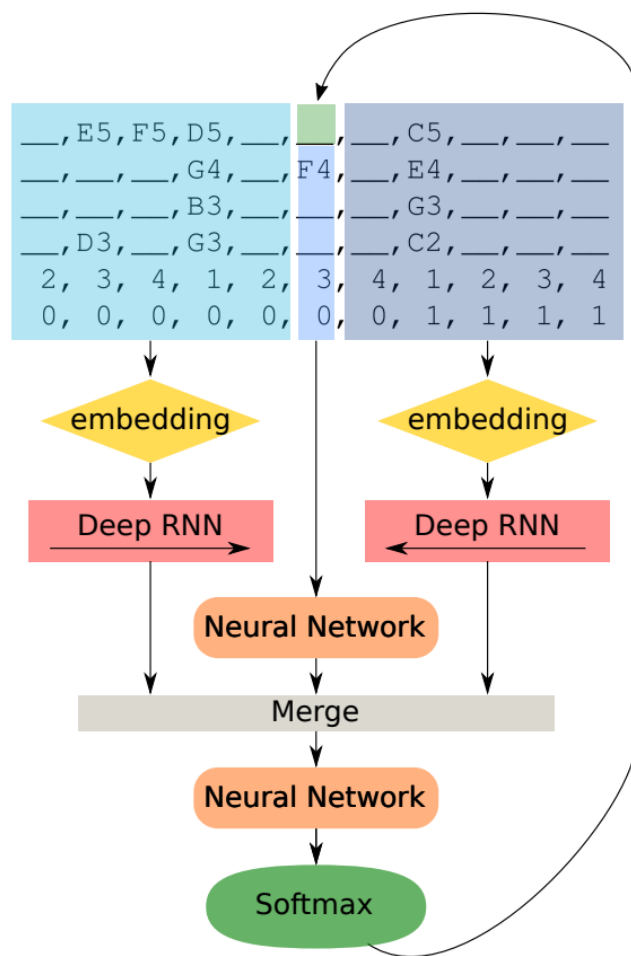
Music Infilling Models

DeepBach (Hadjeres et al., 2017)



(Source: Hadjeres et al., 2017)

DeepBach (Hadjeres et al., 2017)



(Source: Hadjeres et al., 2017)

Algorithm 1 Pseudo-Gibbs sampling

- 1: **Input:** Chorale length L , metadata $\mathcal{M}$ containing lists of length L , probability distributions (p_1, p_2, p_3, p_4) , maximum number of iterations M
- 2: Create four lists $\mathcal{V} = (\mathcal{V}_1, \mathcal{V}_2, \mathcal{V}_3, \mathcal{V}_4)$ of length L
- 3: {The lists are initialized with random notes drawn from the ranges of the corresponding voices (sampled uniformly or from the marginal distributions of the notes)}
- 4: **for** m from 1 to M **do**
- 5: Choose voice i uniformly between 1 and 4
- 6: Choose time t uniformly between 1 and L
- 7: Re-sample $\mathcal{V}_i^t$ from $p_i(\mathcal{V}_i^t | \mathcal{V}_{\setminus i, t}, \mathcal{M}, \theta_i)$
- 8: **end for**
- 9: **Output:** $\mathcal{V} = (\mathcal{V}_1, \mathcal{V}_2, \mathcal{V}_3, \mathcal{V}_4)$

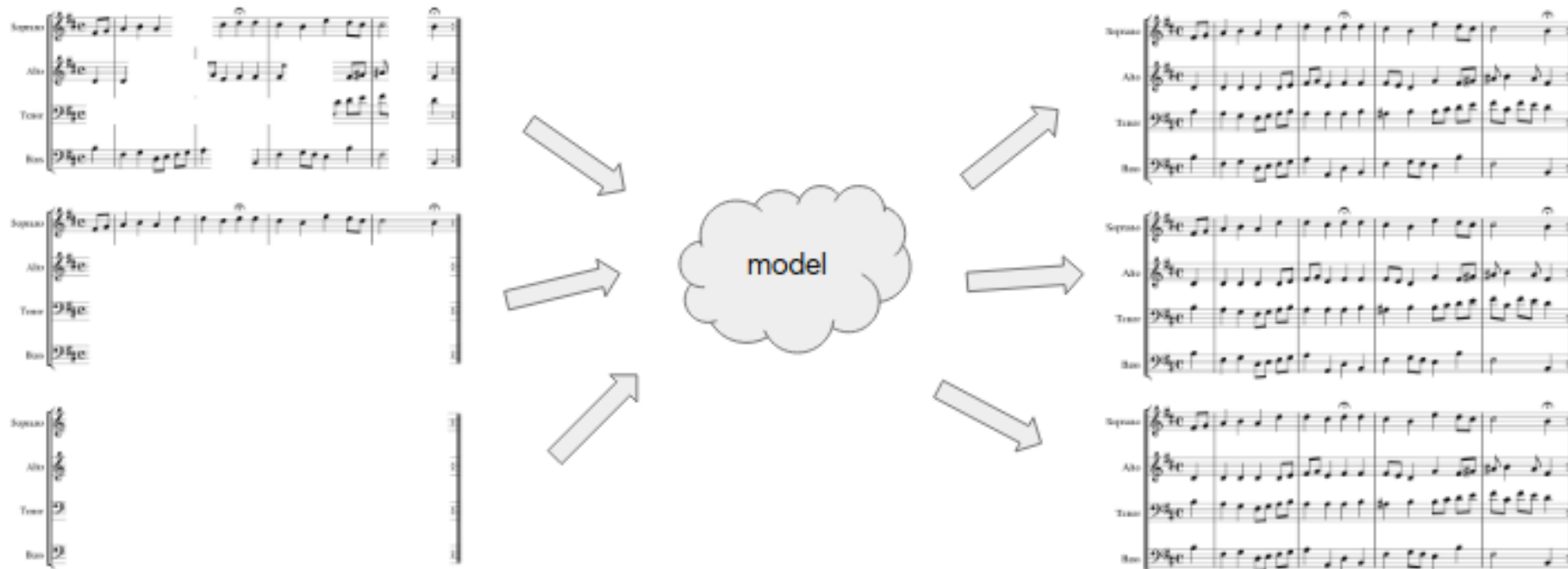
DeepBach: ReharmORIZATION Examples (Hadjeres et al., 2017)



youtu.be/QiBM7-5hA6o

Coconet (Huang et al., 2017)

- Based on Orderless NADE (Uria et al, 2014)



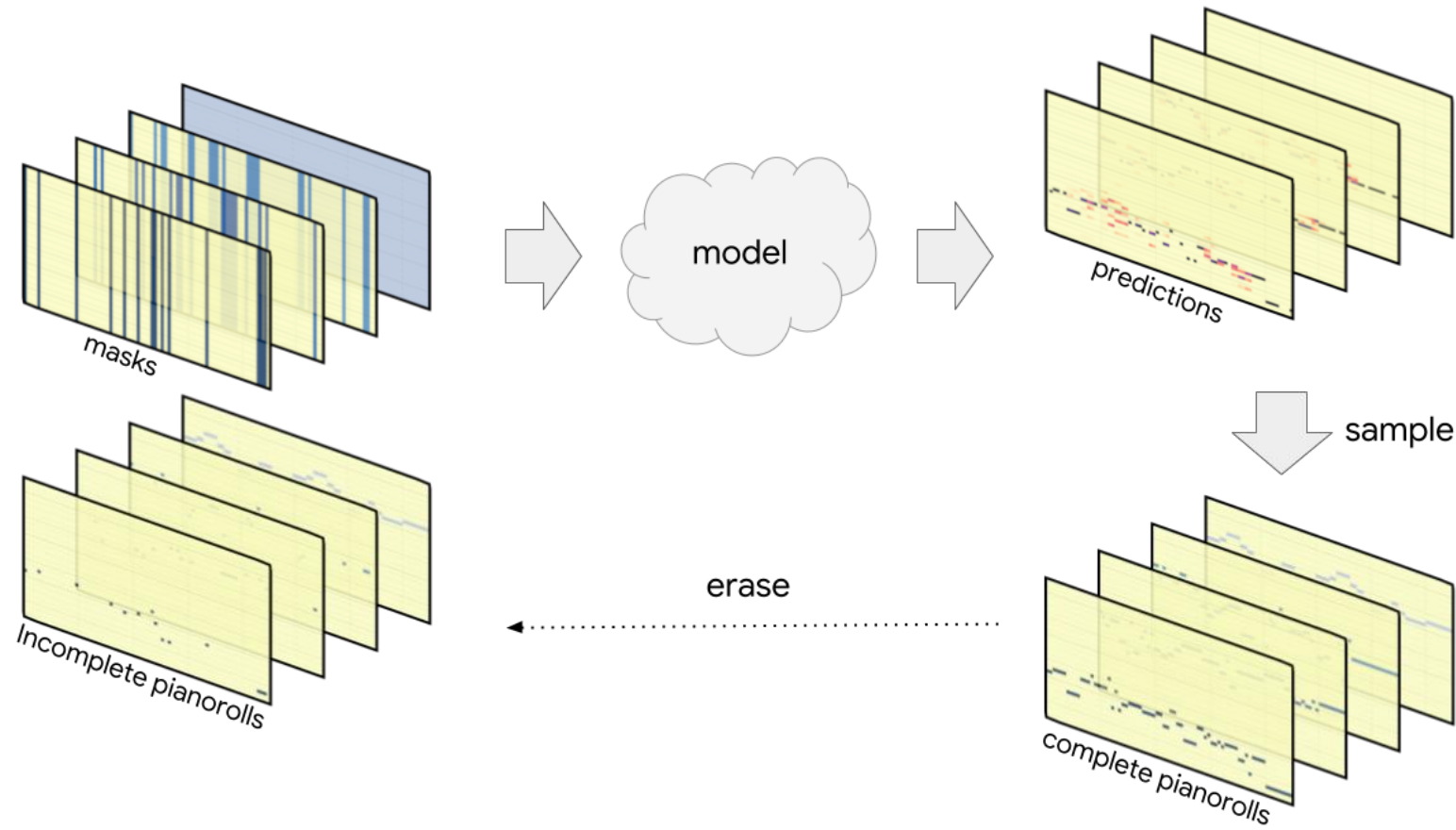
(Source: Huang et al., 2019)

Benigno Uria, Iain Murray, and Hugo Larochelle, "A Deep and Tractable Density Estimator," *ICML*, 2014.

Cheng-Zhi Anna Huang, Tim Cooijmans, Adam Roberts, Aaron Courville, and Douglas Eck, "Counterpoint by Convolution," *ISMIR*, 2017.

Cheng-Zhi Anna Huang, Tim Cooijmans, Monica Dinculescu, Adam Roberts, and Curtis Hawthorne, "Coconet: the ML model behind today's Bach Doodle," *Magenta Blog*, 2019.

Coconet (Huang et al., 2017)



(Source: Huang et al., 2019)

Cheng-Zhi Anna Huang, Tim Cooijmans, Adam Roberts, Aaron Courville, and Douglas Eck, "Counterpoint by Convolution," *ISMIR*, 2017.

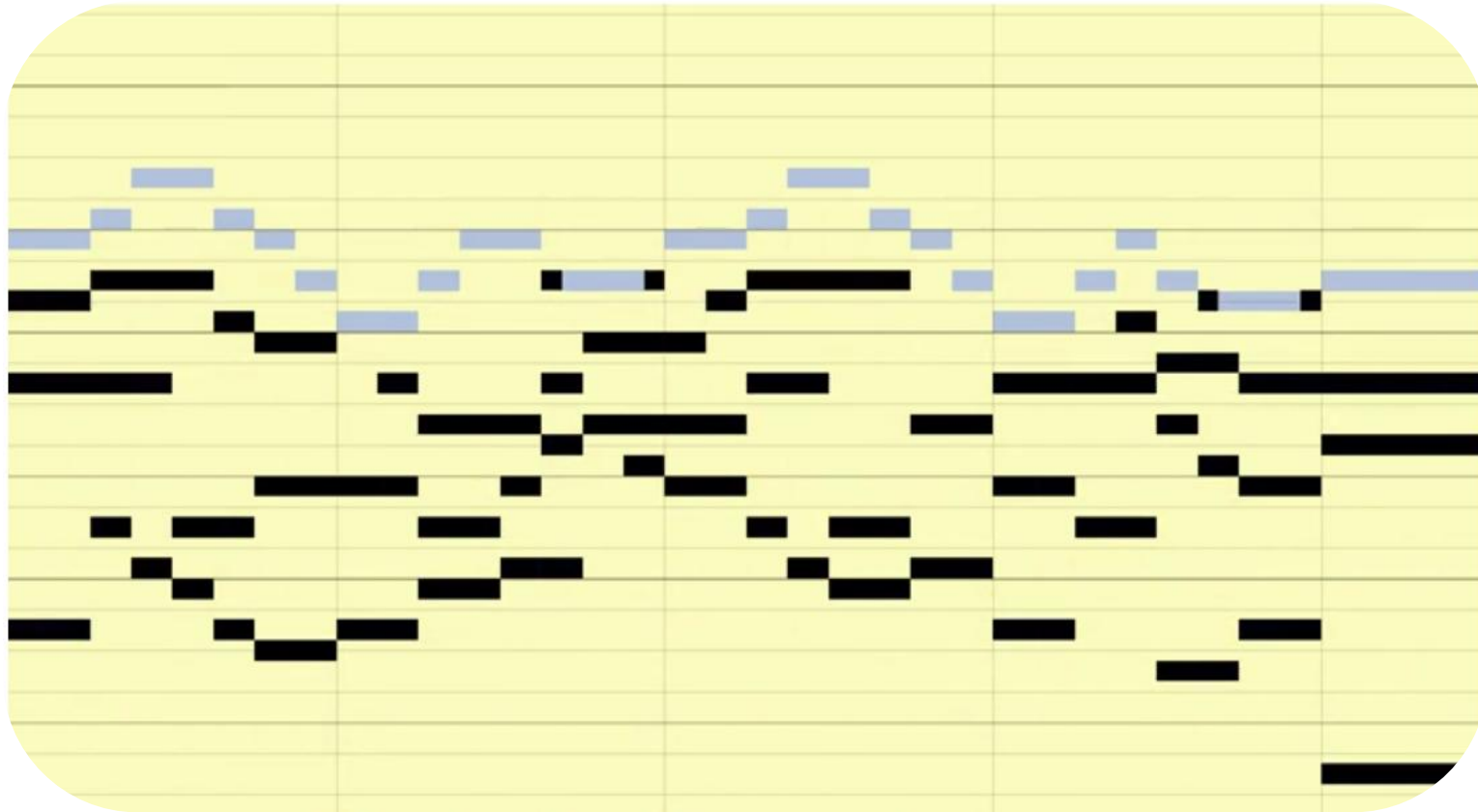
Cheng-Zhi Anna Huang, Tim Cooijmans, Monica Dinulescu, Adam Roberts, and Curtis Hawthorne, "Coconet: the ML model behind today's Bach Doodle," *Magenta Blog*, 2019.

Coconet (Huang et al., 2017)



(Source: Huang et al., 2017)

Coconet (Huang et al., 2017)



(Source: Huang et al., 2017)

| JS Bach Doodle with Coconet (2019)



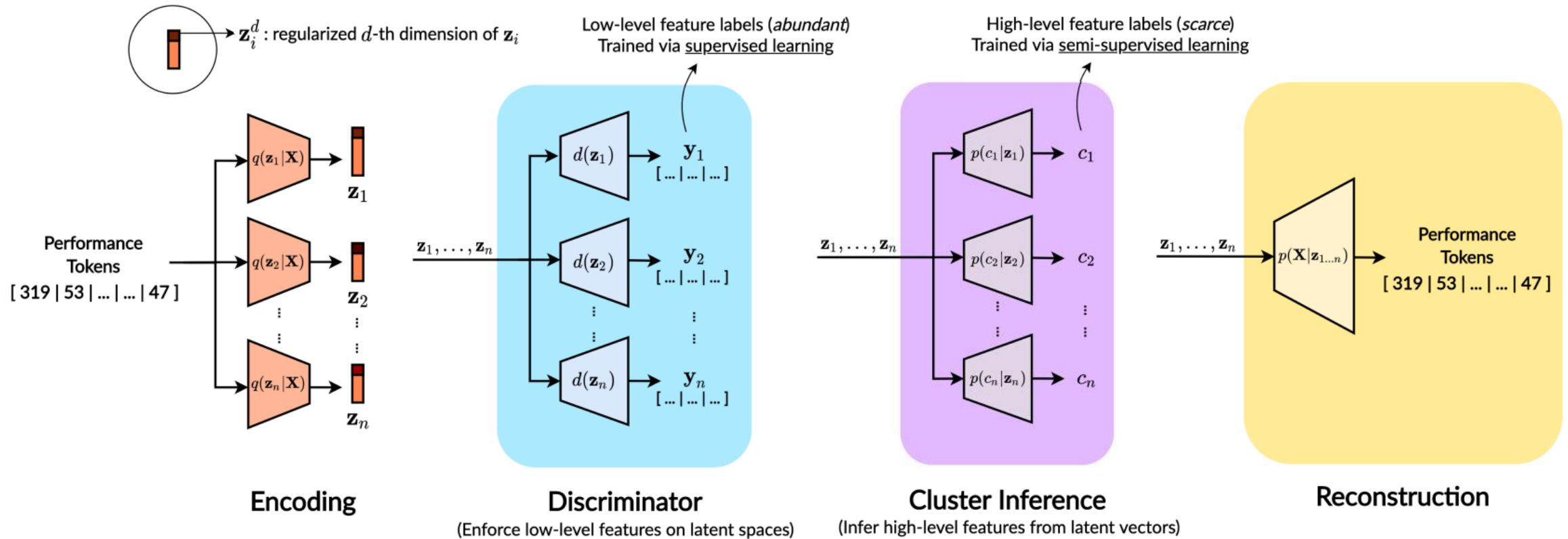
doodles.google/doodle/celebrating-johann-sebastian-bach/



youtu.be/XBfYPp6KF2g & magenta.tensorflow.org/coconet

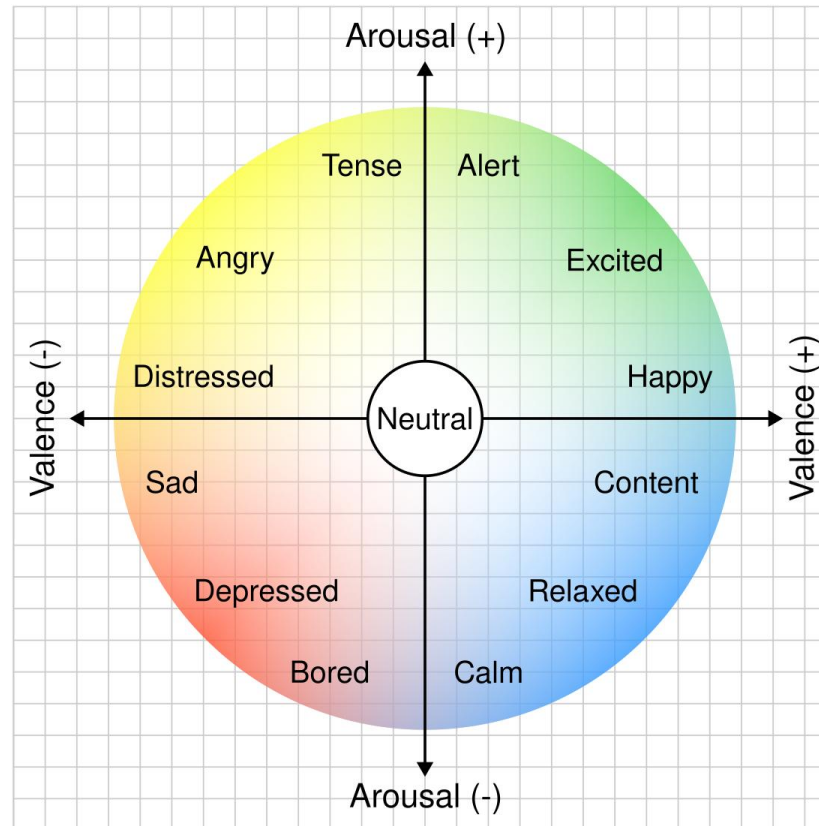
Controllable Music Generation

Music FaderNet (Tan & Herremans, 2020)



(Source: Tan & Herremans, 2020)

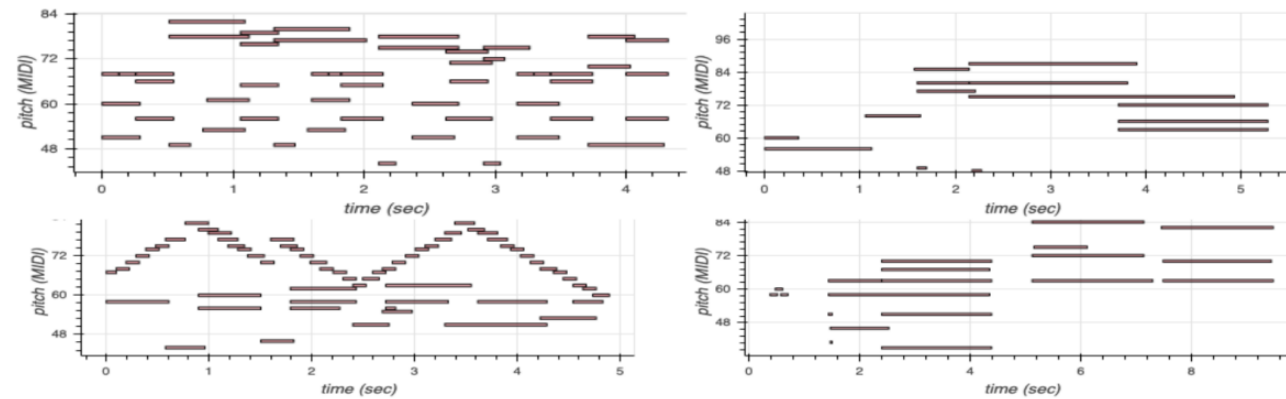
Valence-Arousal Model for Emotion



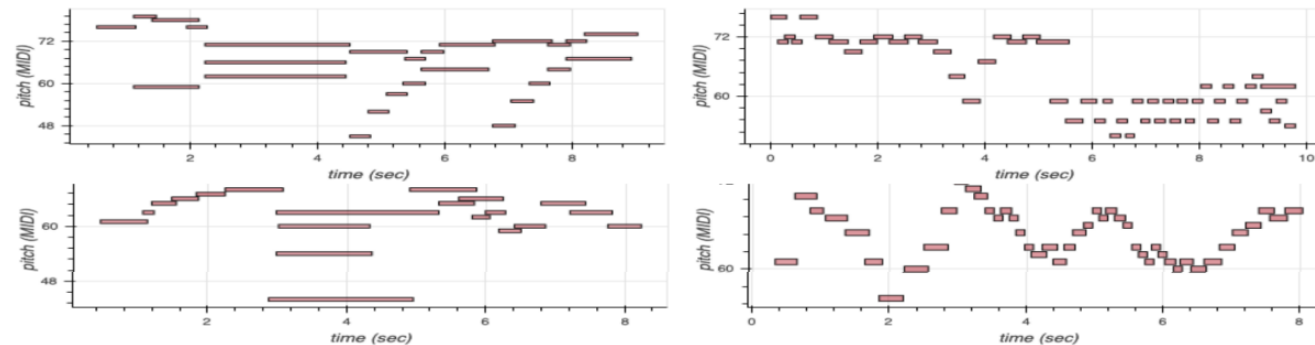
(Source: mrAnmol)

Music FaderNet (Tan & Herremans, 2020)

High Arousal → Low Arousal



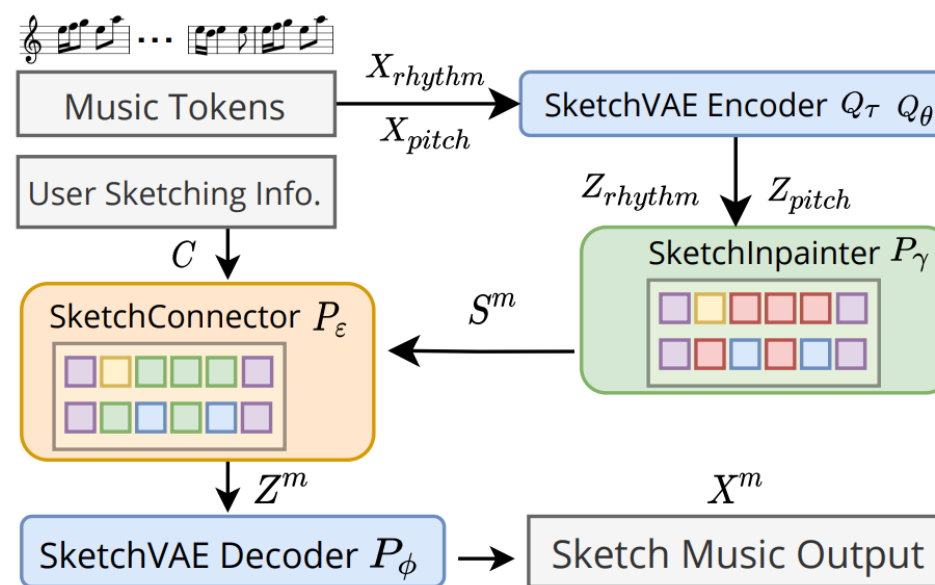
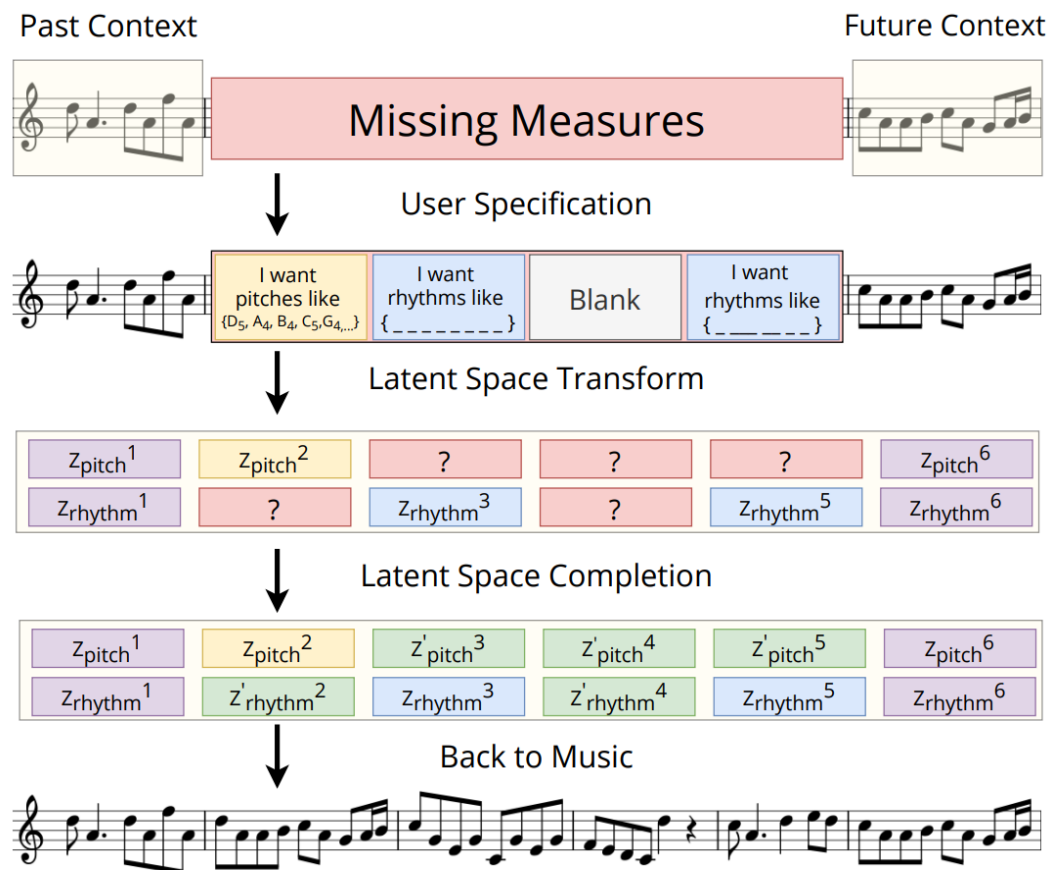
Low Arousal → High Arousal



(Source: Tan & Herremans, 2020)

music-fadernets.github.io

Music SketchNet (Chen et al., 2020)



(Source: Chen et al., 2020)

Music SketchNet (Chen et al., 2020)

The diagram illustrates the Music SketchNet architecture across four tracks: Original, Control Pitch, Control Rhythm, and Control Both. The timeline is divided into three sections: Past Context, Generation, and Future Context.

- Original:** A continuous musical melody in 4/4 time.
- Control Pitch:** Shows pitch control points (blue notes) and sets of pitch classes (e.g., {Ab5, Db6, Eb6, Gb6}, {C6, Eb6, Db6, F6, Db6}, {F6, Gb6, Ab6, Ab6, F6}, {Db6, F6, Ab6, Bb6, Db6}).
- Control Rhythm:** Shows rhythm control points (pink rectangles) and a "No Sketch" region (grey rectangle).
- Control Both:** Combines pitch and rhythm control points.

Labels at the bottom indicate the sections: Past Context, Generation, and Future Context.

(Source: Chen et al., 2020)

| Topics of Symbolic Music Generation

Unconditional

Symbolic music generation

- $\emptyset \rightarrow$ melody
- $\emptyset \rightarrow$ lead sheet
- $\emptyset \rightarrow$ sheet music

Melody
& chords

Conditional

Automatic arrangement

- Melody $\rightarrow$ lead sheet
- Melody $\rightarrow$ multitrack
- Lead sheet $\rightarrow$ multitrack
- Solo $\rightarrow$ multitrack
- Multitrack $\rightarrow$ simple version

Performance rendering

- Sheet music $\rightarrow$ performance

Improvisation systems

- Performance $\rightarrow$ performance

Multimodal

X-to-music generation

- Text $\rightarrow$ music
- Video $\rightarrow$ music
- Gestures $\rightarrow$ music
- Motions $\rightarrow$ music
- Brain waves $\rightarrow$ music
- X $\rightarrow$ music

Recap

Two Paradigms of Symbolic Music Generation

Text-based

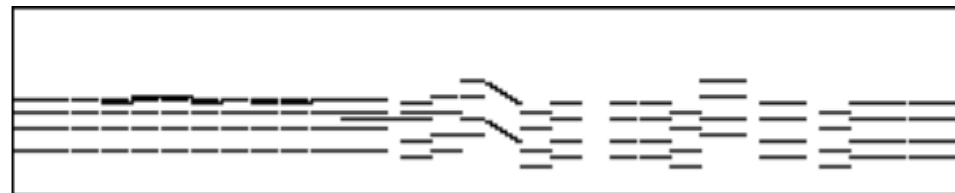
- Treat music like **text**
- Sharing models with **natural language processing (NLP)**
 - RNNs, LSTMs, Transformers, etc.

```
Program_change_0,  
Note_on_60, Time_shift_2, Note_off_60,  
Note_on_60, Time_shift_2, Note_off_60,  
Note_on_76, Time_shift_2, Note_off_67,  
Note_on_67, Time_shift_2, Note_off_67, ...
```

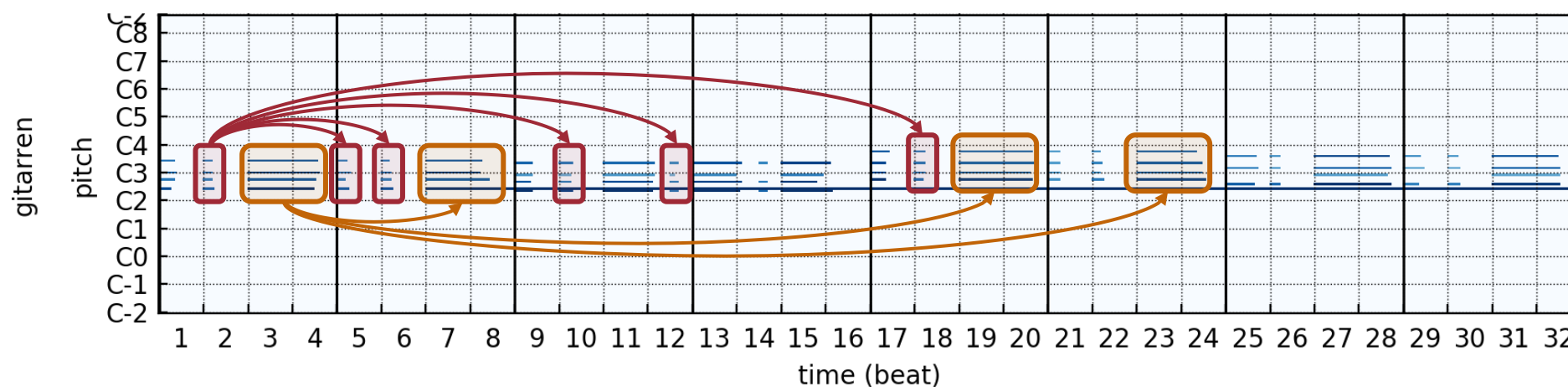
Image-based

- Treat music like **images**
- Sharing models with **computer vision (CV) & computer graphics (CG)**
 - GANs, VAEs, diffusion models, etc.

Today's topic!

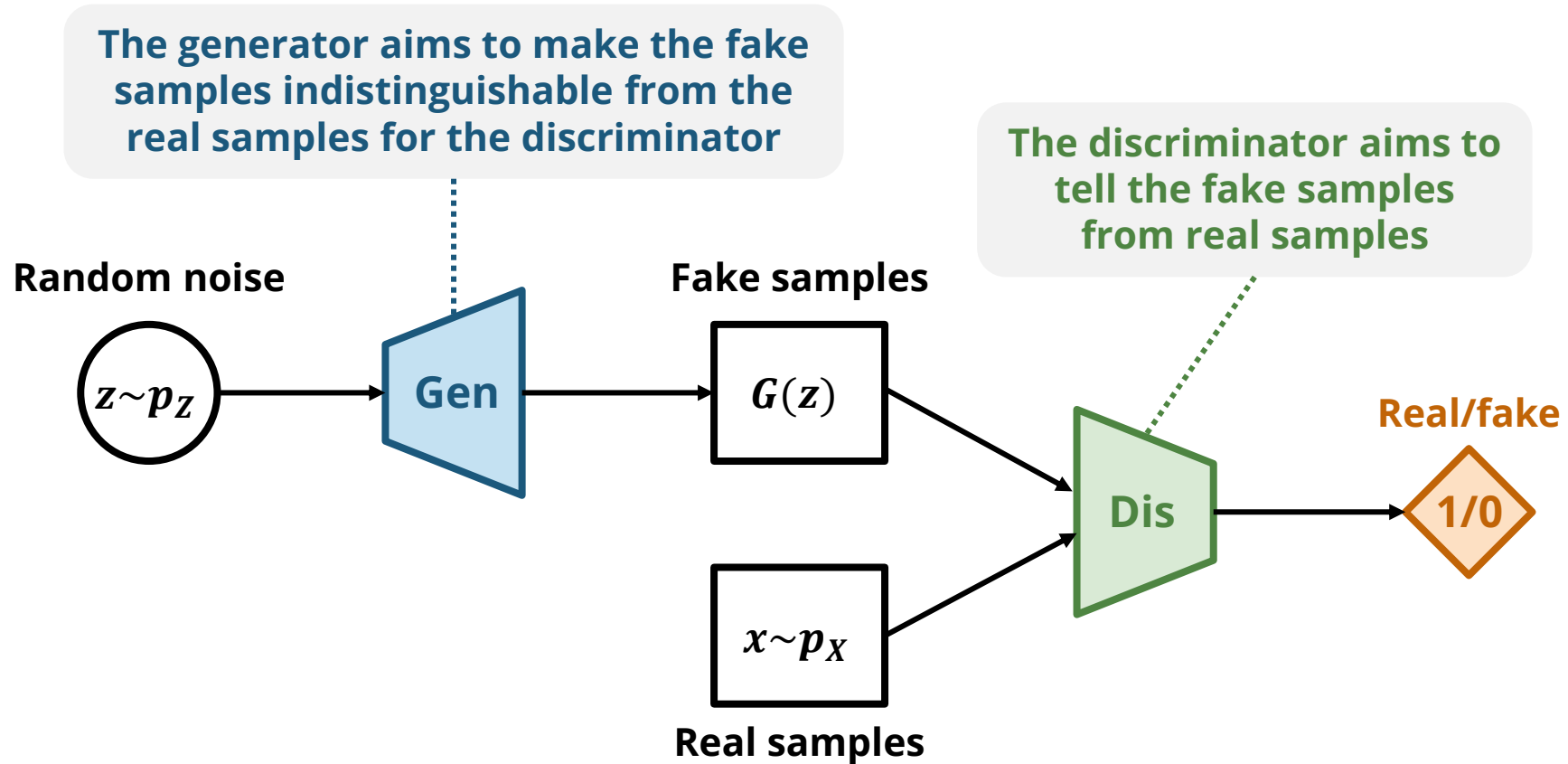


Why Piano Rolls?



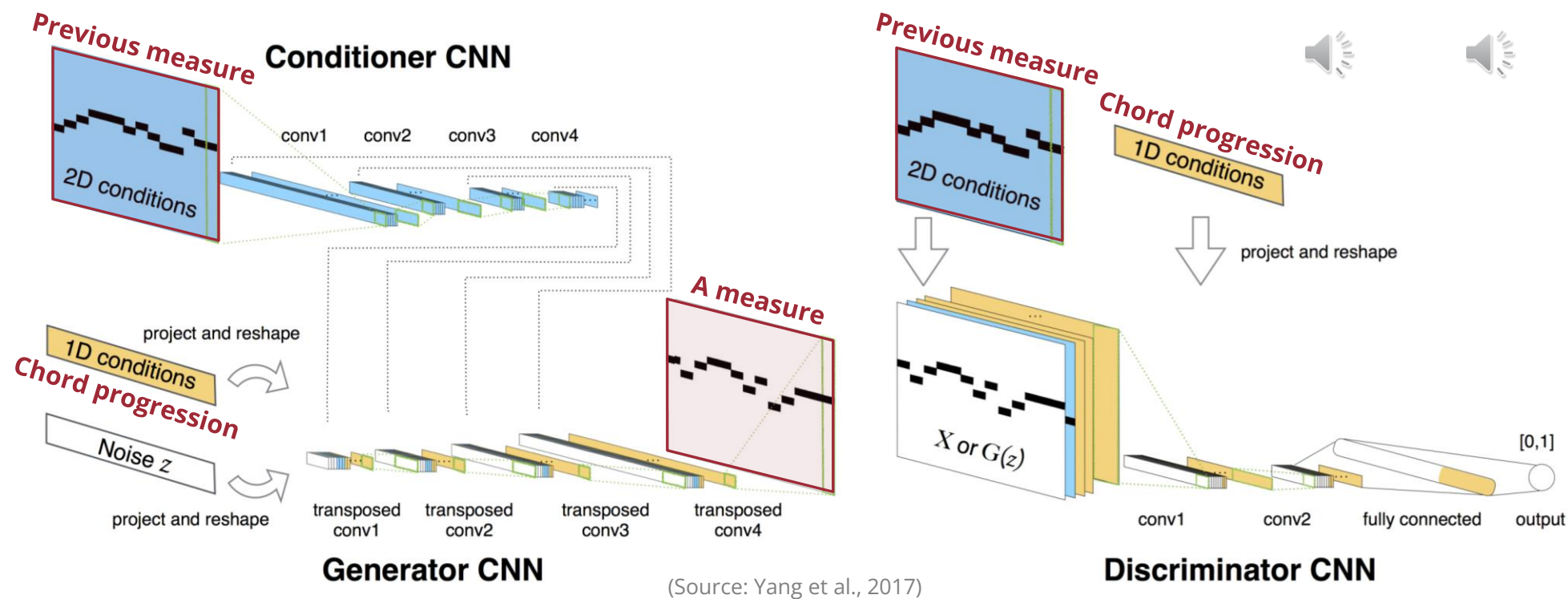
Many musical patterns like melodies, chords, scales and arpeggios
are **translational invariant** in the temporal and pitch axes

Generative Adversarial Nets (GANs) (Goodfellow et al., 2014)



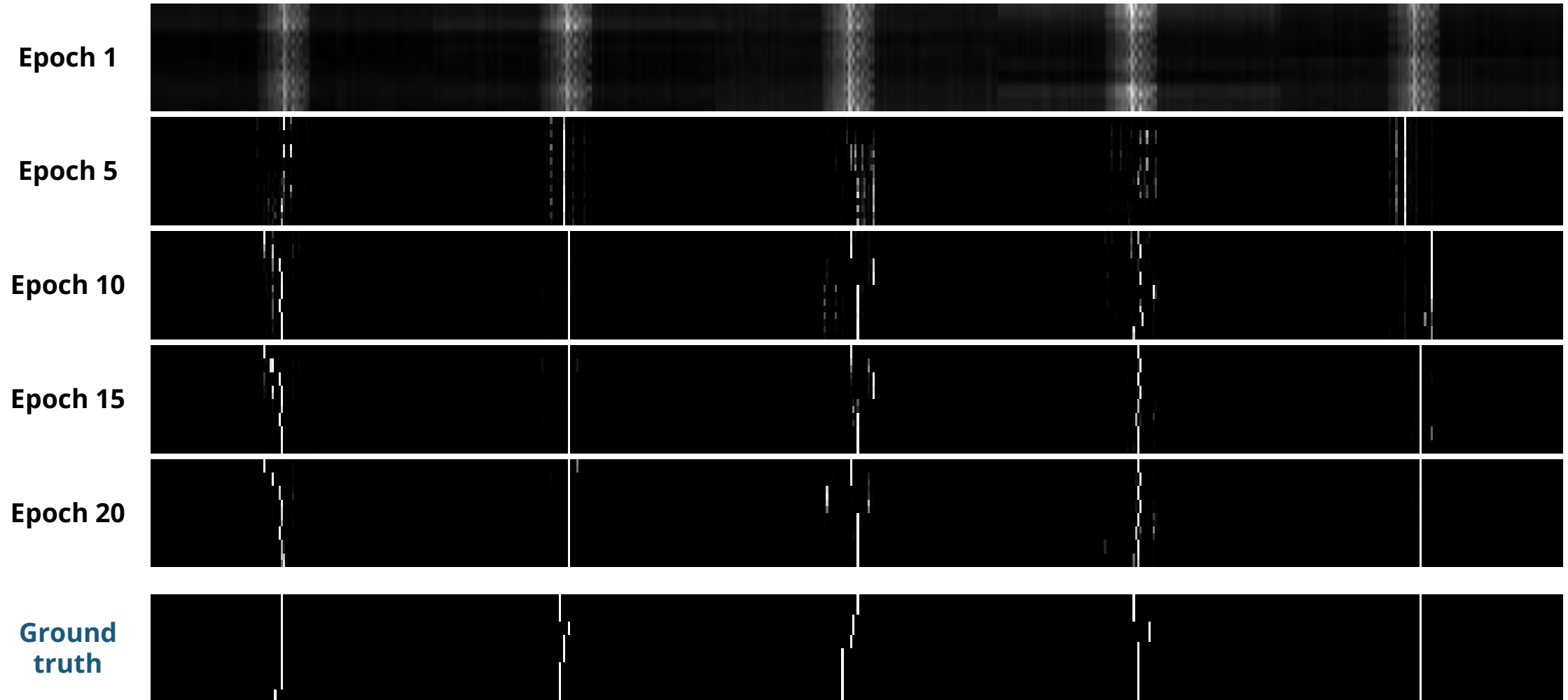
MidiNet (Yang et al., 2017)

Examples of
generated music



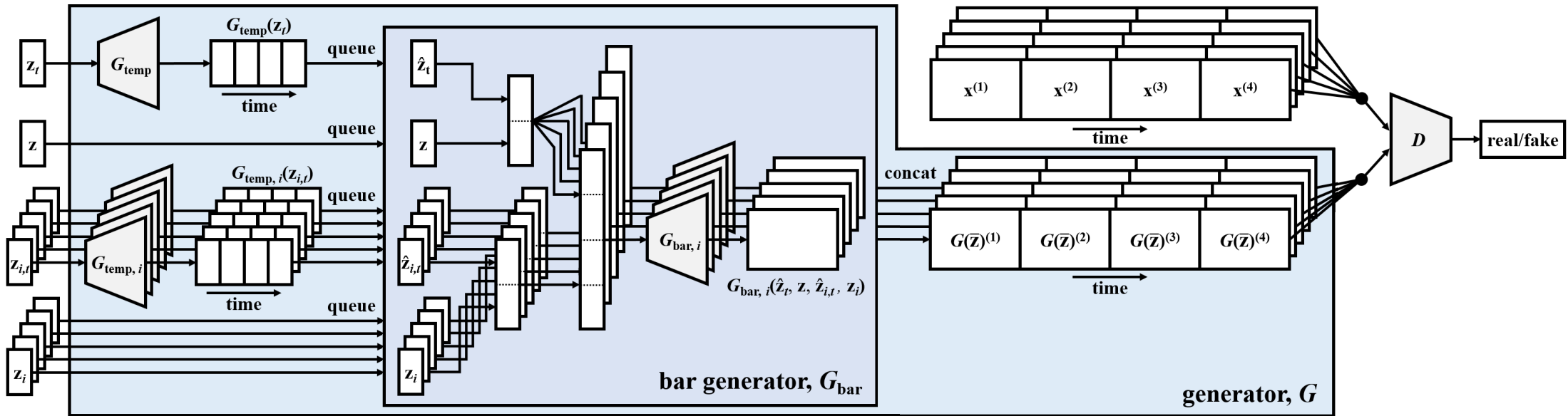
**Generate the next measure given the previous one
(measure by measure)**

MidiNet (Yang et al., 2017)



(Source: Yang et al., 2017)

MuseGAN (Dong et al., 2018)



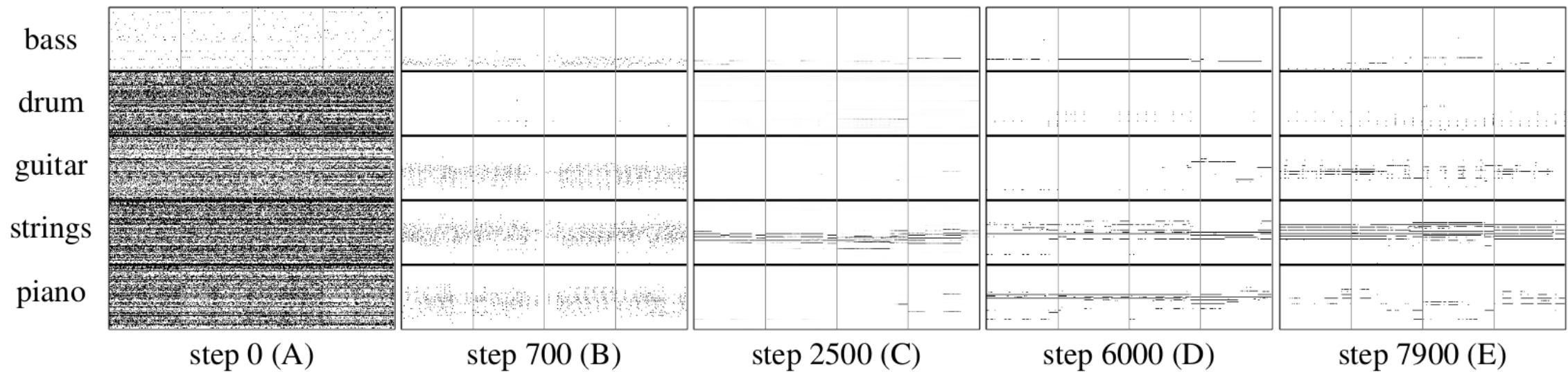
(Source: Dong et al., 2018)

MuseGAN (Dong et al., 2018)

Examples of
generated music



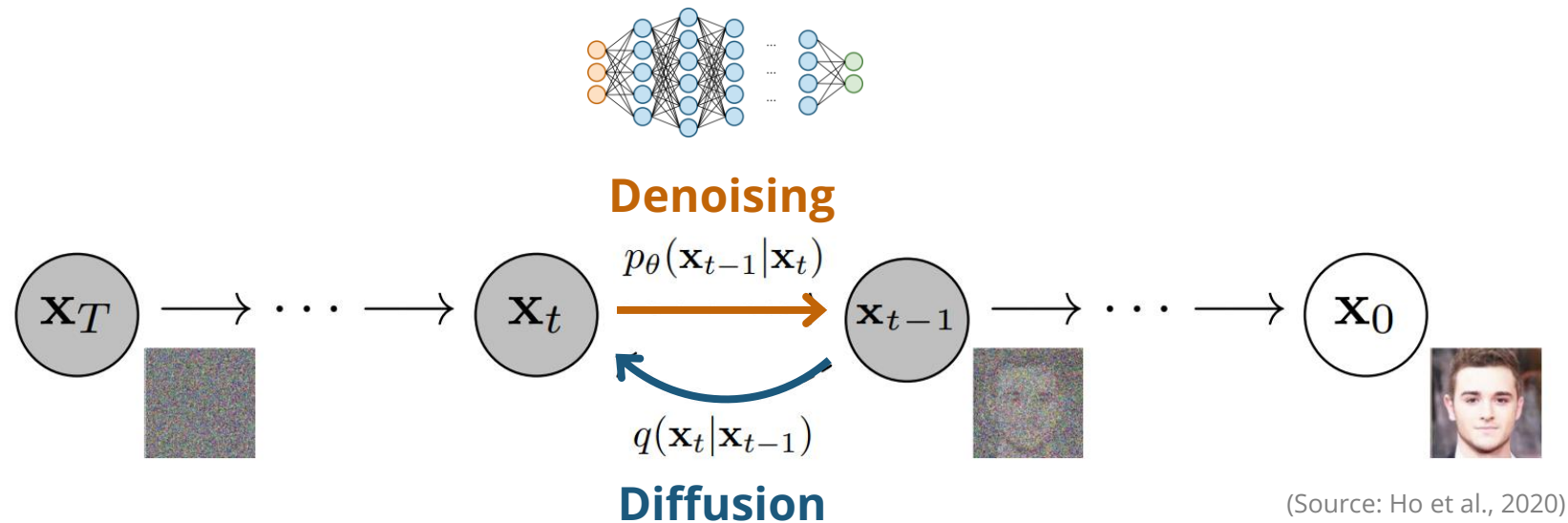
The generator improves over time



(Source: Dong et al., 2018)

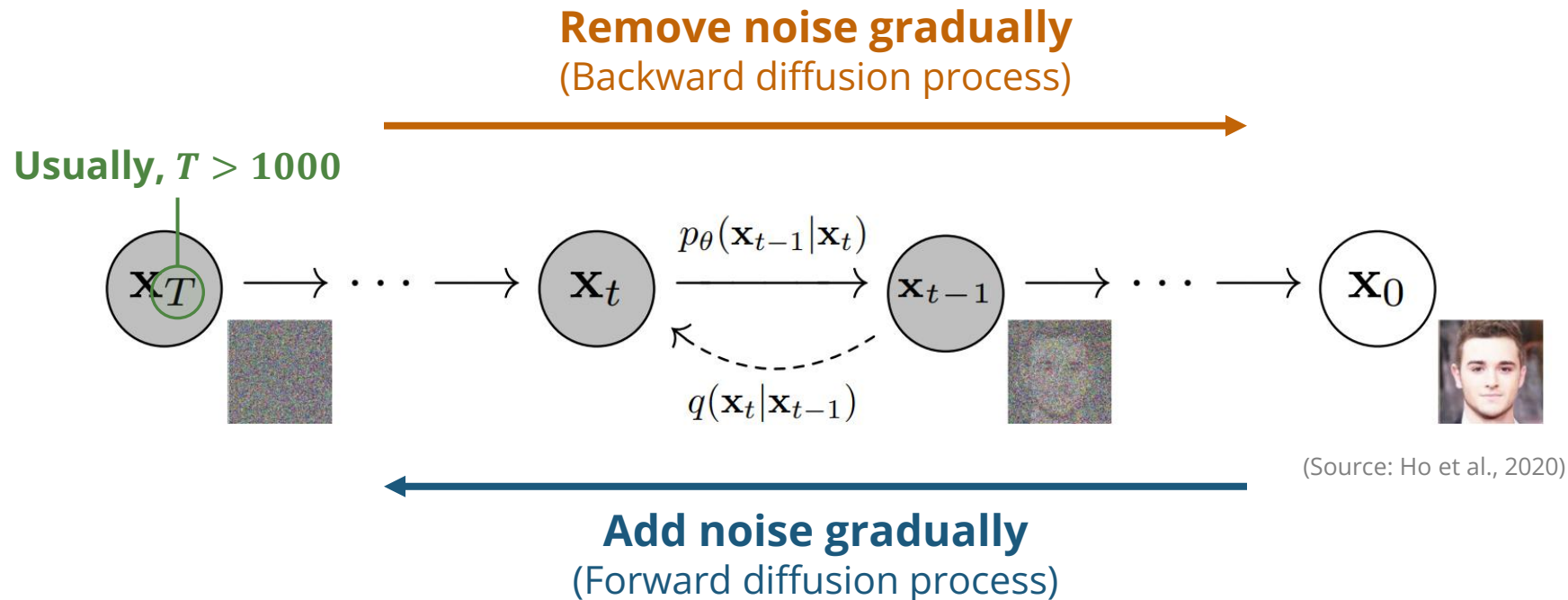
Diffusion Models (Ho et al., 2020)

- **Intuition:** Many denoising autoencoders stacked together

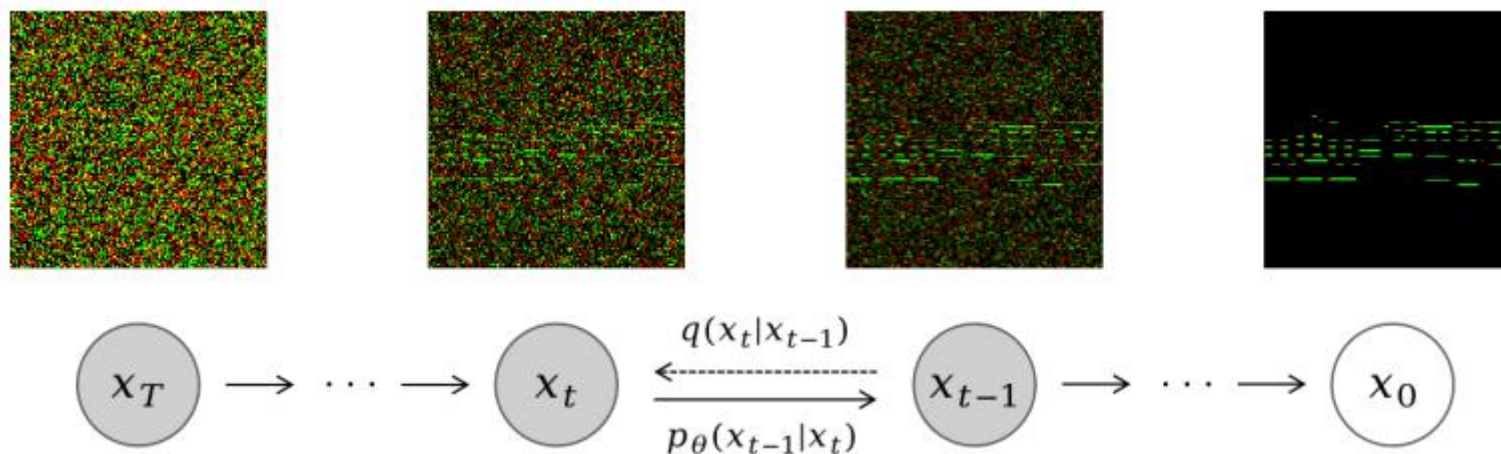


Diffusion Models (Ho et al., 2020)

- **Intuition:** Many denoising autoencoders stacked together



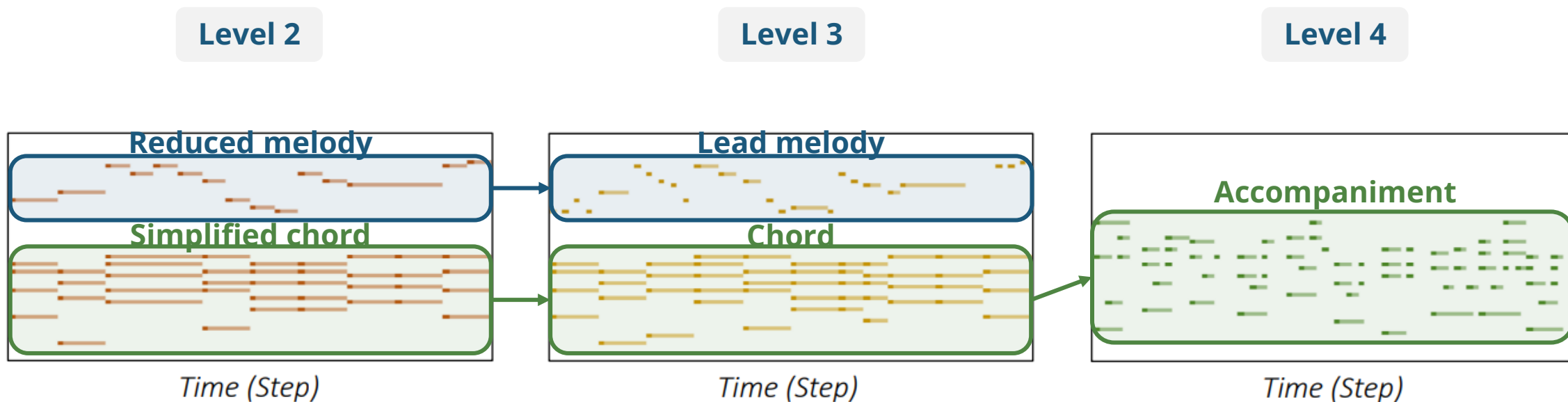
Polyffusion (Min et al., 2023)



(Source: Min et al., 2023)

polyffusion.github.io

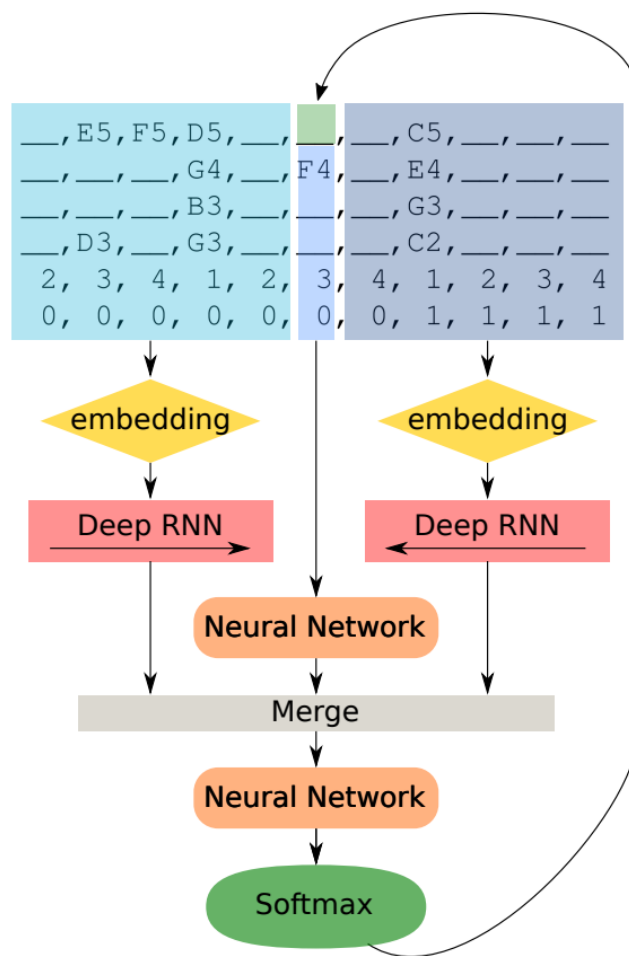
Cascaded Diffusion Models (Wang et al., 2024)



(Source: Wang et al., 2024)

wholesonggen.github.io

DeepBach (Hadjeres et al., 2017)

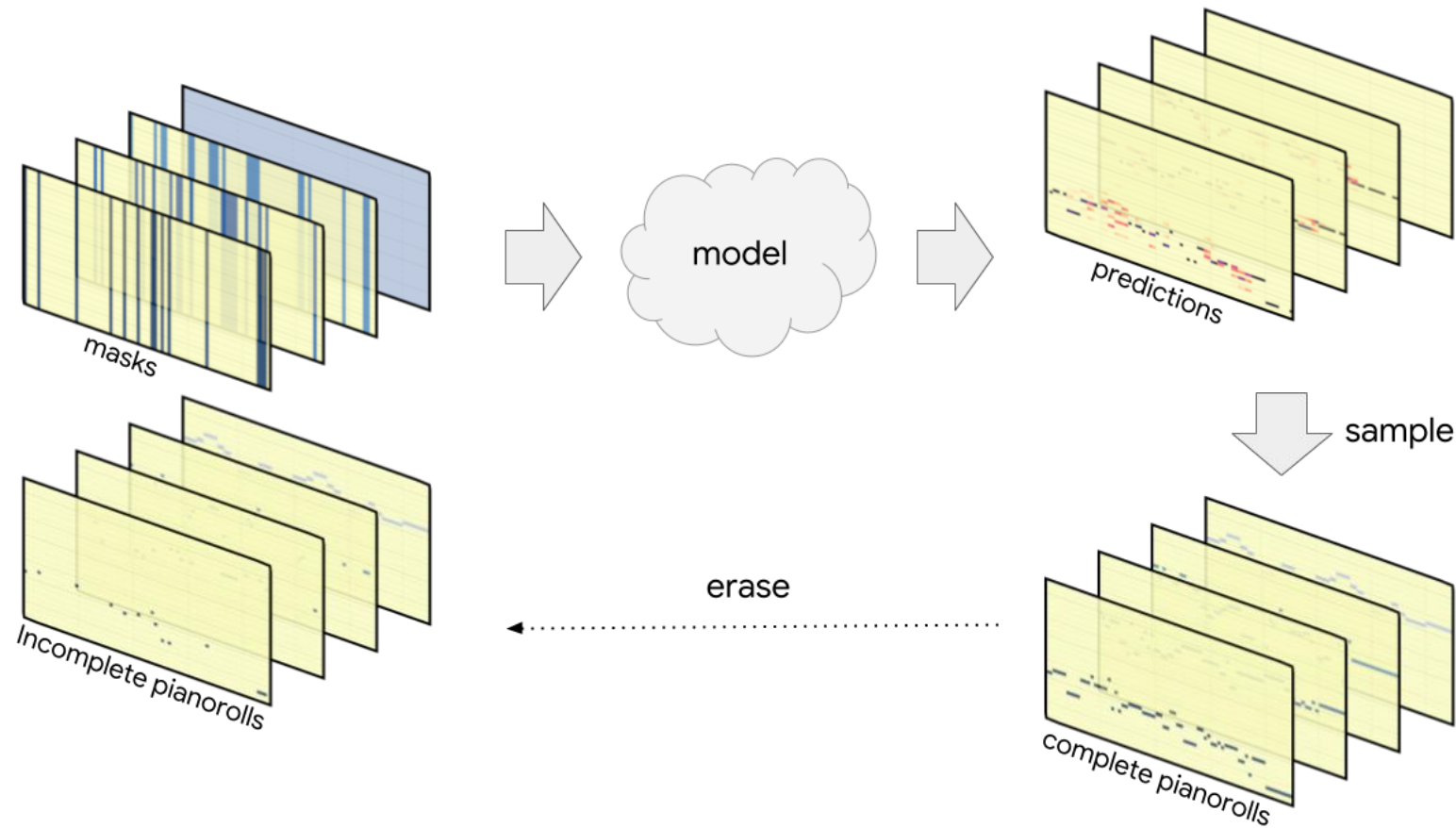


(Source: Hadjeres et al., 2017)

Algorithm 1 Pseudo-Gibbs sampling

- 1: **Input:** Chorale length L , metadata $\mathcal{M}$ containing lists of length L , probability distributions (p_1, p_2, p_3, p_4) , maximum number of iterations M
 - 2: Create four lists $\mathcal{V} = (\mathcal{V}_1, \mathcal{V}_2, \mathcal{V}_3, \mathcal{V}_4)$ of length L
 - 3: {The lists are initialized with random notes drawn from the ranges of the corresponding voices (sampled uniformly or from the marginal distributions of the notes)}
 - 4: **for** m from 1 to M **do**
 - 5: Choose voice i uniformly between 1 and 4
 - 6: Choose time t uniformly between 1 and L
 - 7: Re-sample $\mathcal{V}_i^t$ from $p_i(\mathcal{V}_i^t | \mathcal{V}_{\setminus i, t}, \mathcal{M}, \theta_i)$
 - 8: **end for**
 - 9: **Output:** $\mathcal{V} = (\mathcal{V}_1, \mathcal{V}_2, \mathcal{V}_3, \mathcal{V}_4)$
-

Coconet (Huang et al., 2017)



(Source: Huang et al., 2019)

Cheng-Zhi Anna Huang, Tim Cooijmans, Adam Roberts, Aaron Courville, and Douglas Eck, "Counterpoint by Convolution," *ISMIR*, 2017.

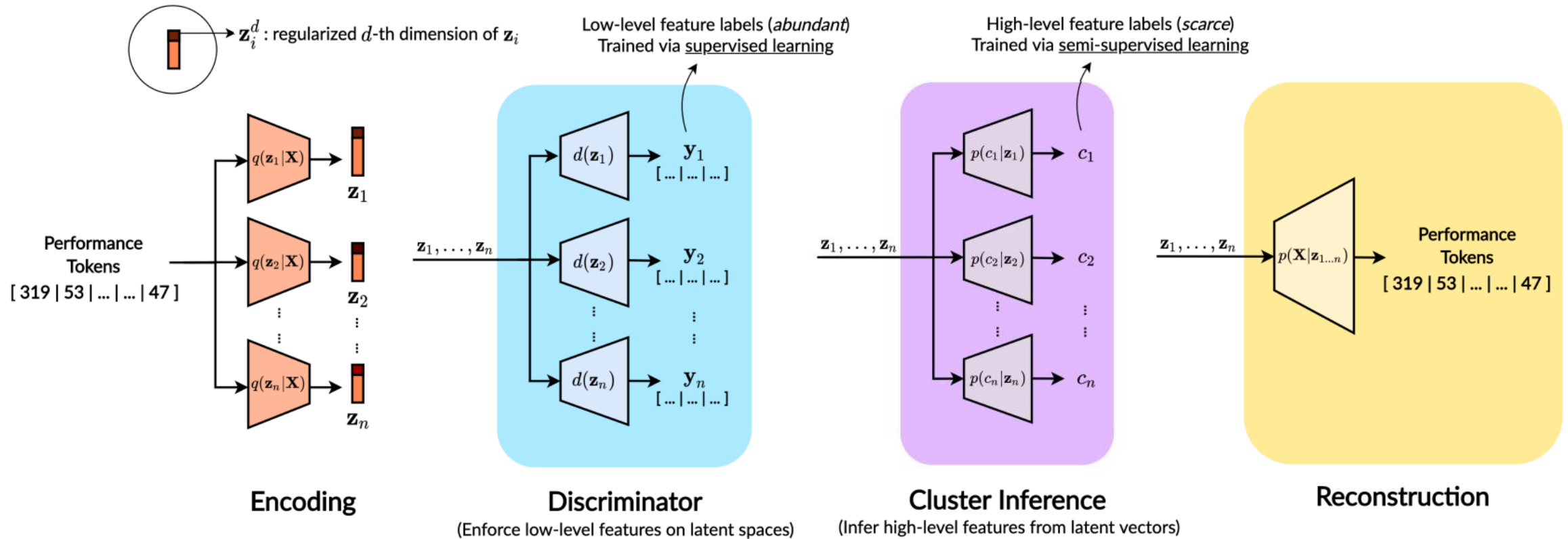
Cheng-Zhi Anna Huang, Tim Cooijmans, Monica Dinculescu, Adam Roberts, and Curtis Hawthorne, "Coconet: the ML model behind today's Bach Doodle," *Magenta Blog*, 2019.

Coconet (Huang et al., 2017)



(Source: Huang et al., 2017)

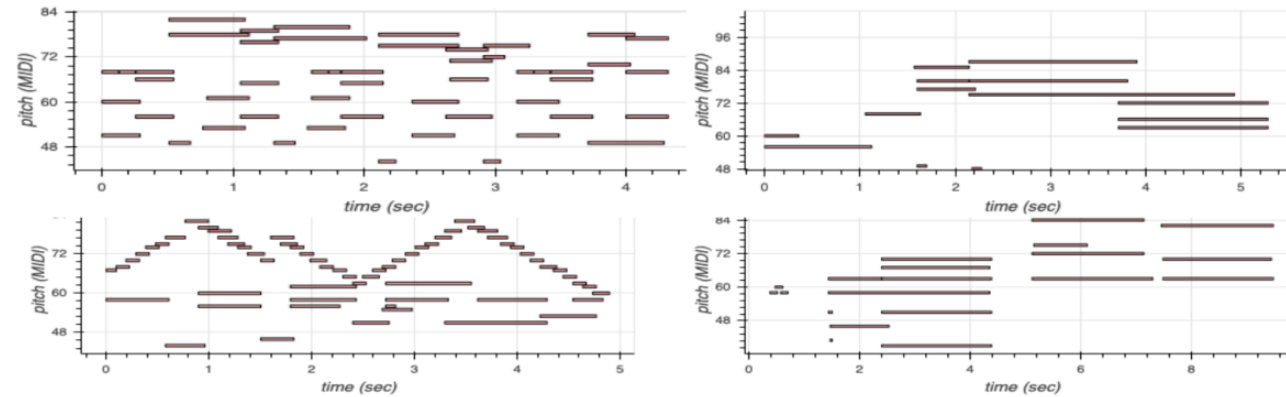
Music FaderNet (Tan & Herremans, 2020)



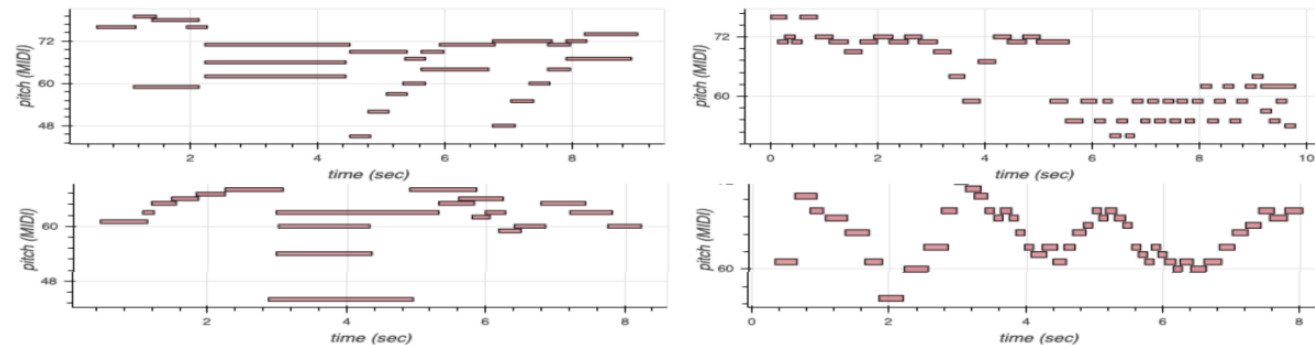
(Source: Tan & Herremans, 2020)

Music FaderNet (Tan & Herremans, 2020)

High Arousal → Low Arousal



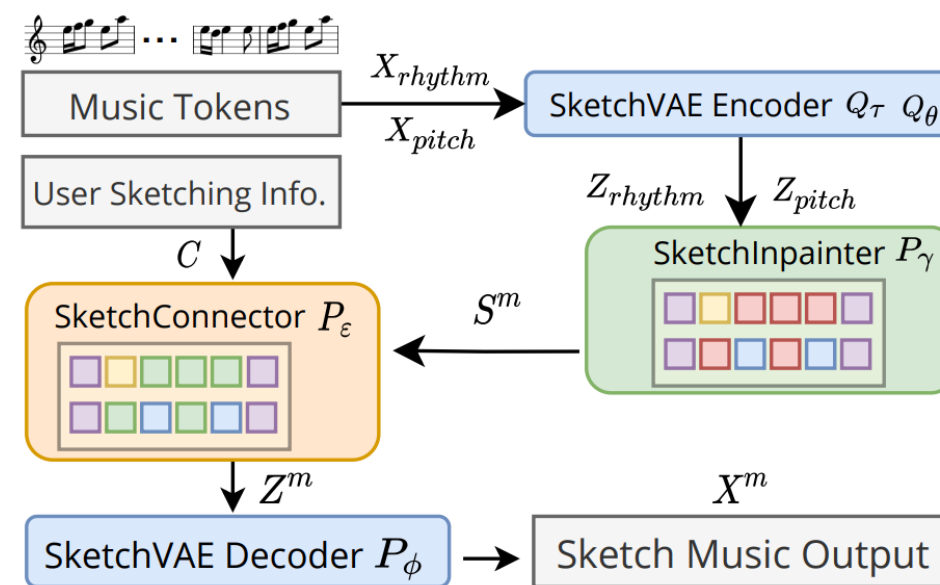
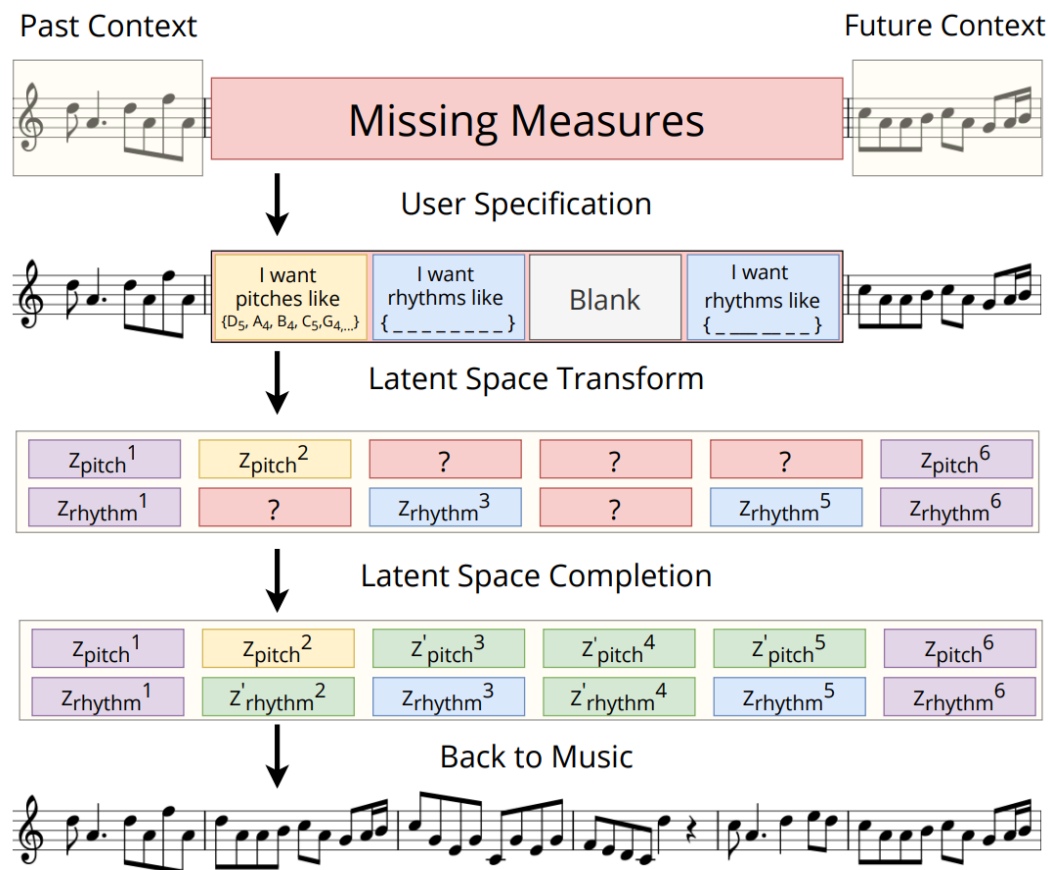
Low Arousal → High Arousal



(Source: Tan & Herremans, 2020)

music-fadernets.github.io

Music SketchNet (Chen et al., 2020)



(Source: Chen et al., 2020)

Music SketchNet (Chen et al., 2020)

Original

Control Pitch

Control Rhythm

Control Both

Past Context

Generation

Future Context

{Ab5, Db6, Eb6, Gb6}

{C6, Eb6, Db6, F6, Db6}

{F6, Gb6, Ab6, Ab6, F6}

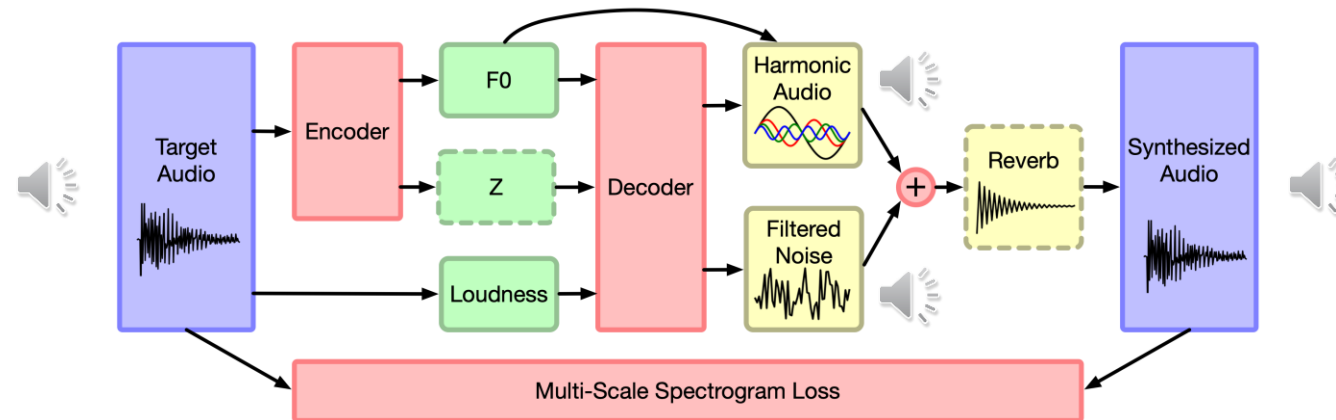
{Db6, F6, Ab6, Bb6, Db6}

No Sketch

(Source: Chen et al., 2020)

Next Lecture

Audio-domain Music Generation



(Source: Engel et al., 2020)