

PAT 463/563 (Fall 2025)

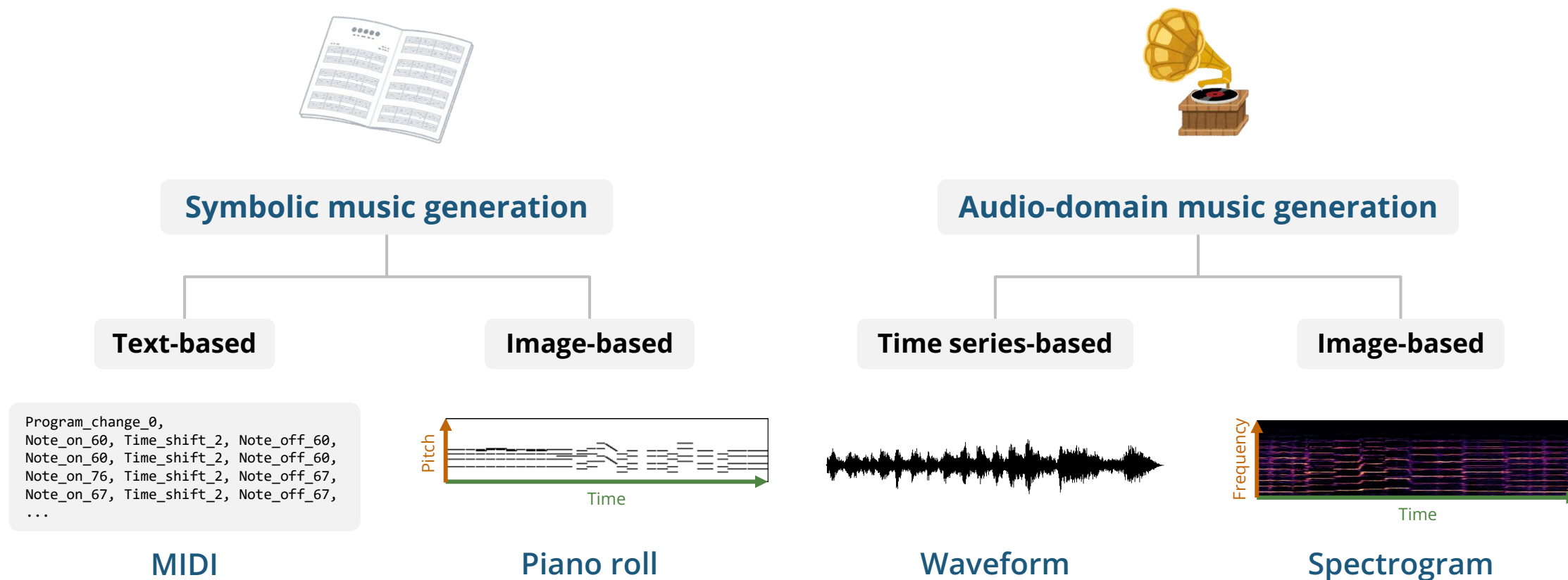
Music & AI

Lecture 14: Audio-domain Music Generation

Instructor: Hao-Wen Dong

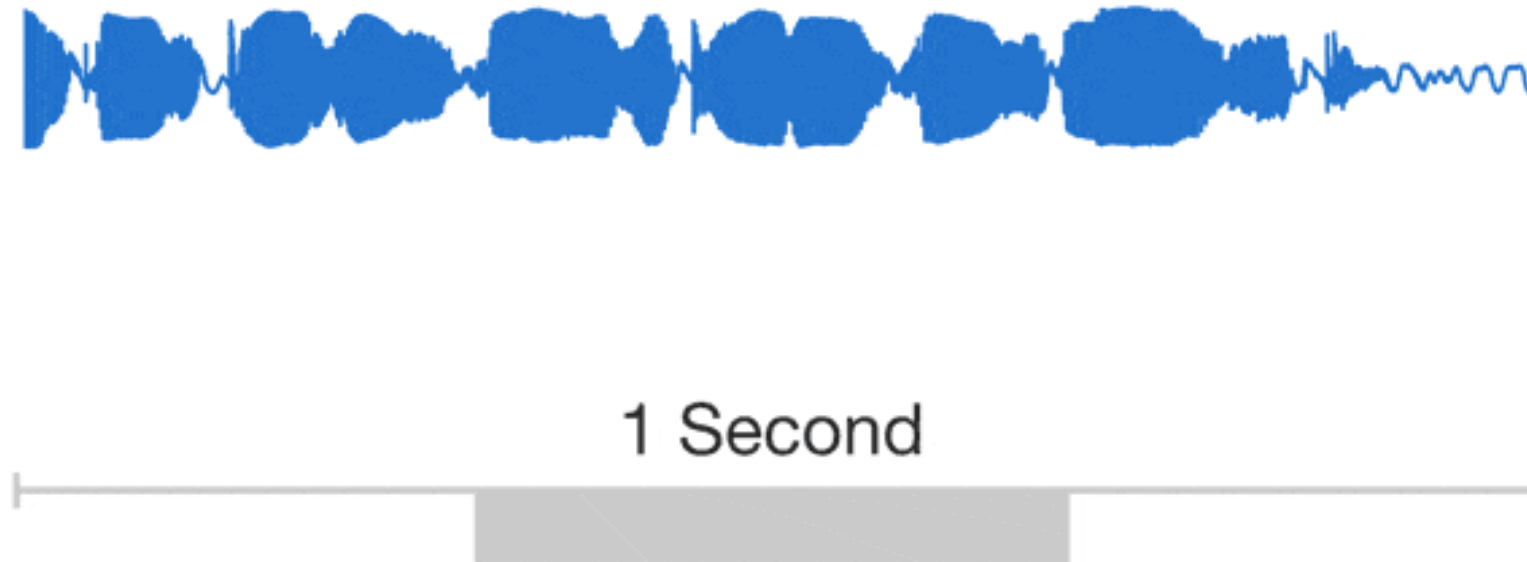
Autoregressive Waveform Synthesis

Four Paradigms of Music Generation



Today, we also have many **latent-space based** systems!

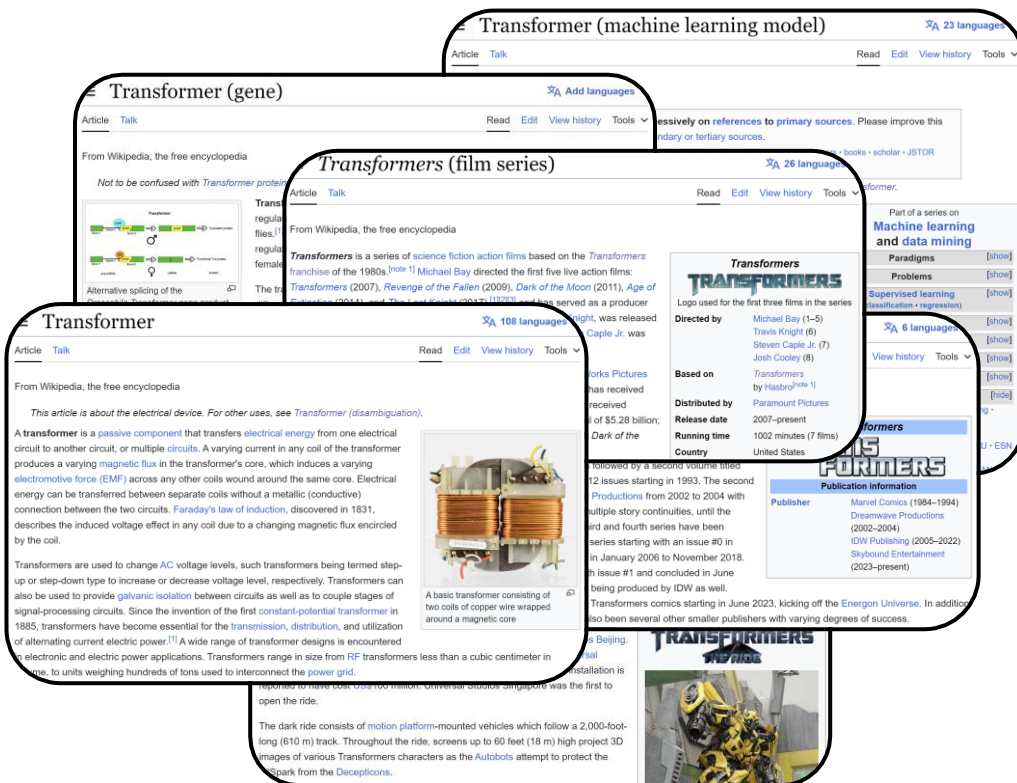
Generating Waveforms using a Neural Network



(Source: van den Oord et al., 2016)

Language Models

- Predicting the next word given the past sequence of words



A transformer is a _____

electrical device ↑

fiction character ↑

deep learning model ↑

family of genes ↑

type of food ↓

musical instrument ↓

| Language Models (Mathematically)

- A class of machine learning models that **learn** the next word probability

$$P(\underset{\substack{\text{Next word}}}{x_i} \mid \underbrace{x_1, x_2, \dots, x_{i-1}}_{\substack{\text{Previous words}}})$$

$P(\text{electrical} \mid \text{A transformer is a})$ ↑

$P(\text{character} \mid \text{A transformer is a})$ ↑

$P(\text{gene} \mid \text{A transformer is a})$ ↑

$P(\text{model} \mid \text{A transformer is a})$ ↑

$P(\text{food} \mid \text{A transformer is a})$ ↓

$P(\text{musical} \mid \text{A transformer is a})$ ↓

Autoregressive Models (Mathematically)

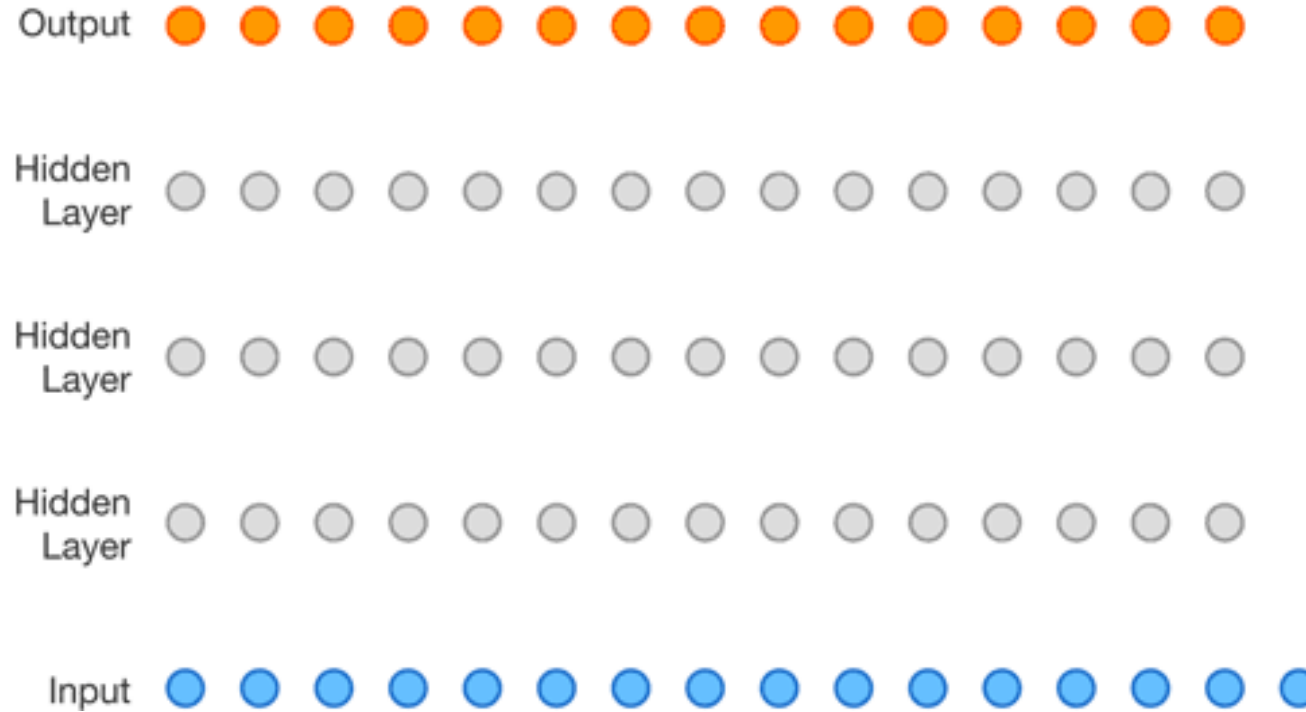
- A class of machine learning models that **learn** the probability of the next value given previous values

$$P(\underbrace{x_i}_{\text{Next number}} \mid \underbrace{x_1, x_2, \dots, x_{i-1}}_{\text{Previous numbers}})$$

$$\begin{array}{l} P(0.1 \mid 0.5, 0.4, 0.3, 0.2) \uparrow \\ P(0.09 \mid 0.5, 0.4, 0.3, 0.2) \uparrow \\ P(0.11 \mid 0.5, 0.4, 0.3, 0.2) \uparrow \\ P(99 \mid 0.5, 0.4, 0.3, 0.2) \downarrow \\ P(-1 \mid 0.5, 0.4, 0.3, 0.2) \downarrow \end{array}$$

The term “autoregressive” has different definitions in machine learning and signal processing.
In signal processing, an autoregressive model needs to be a linear model.

WaveNet (van den Oord et al., 2016)

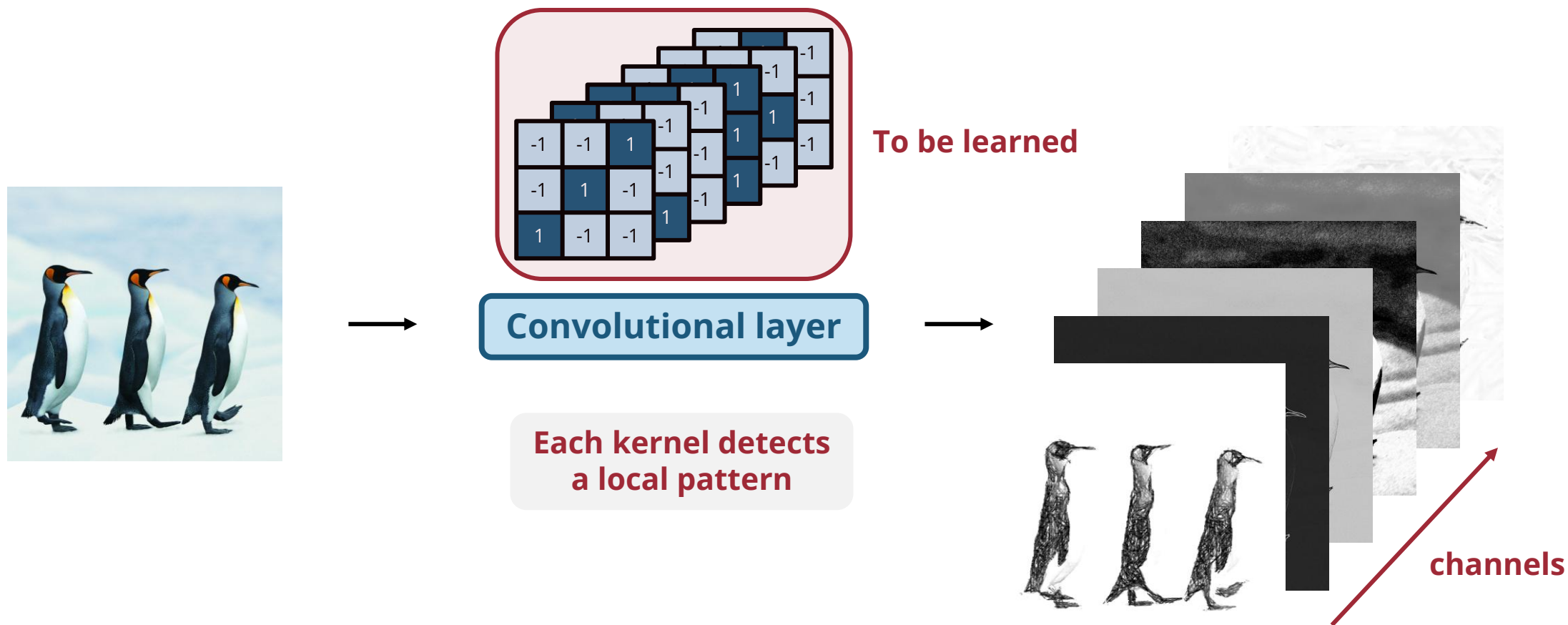


(Source: van den Oord et al., 2016)

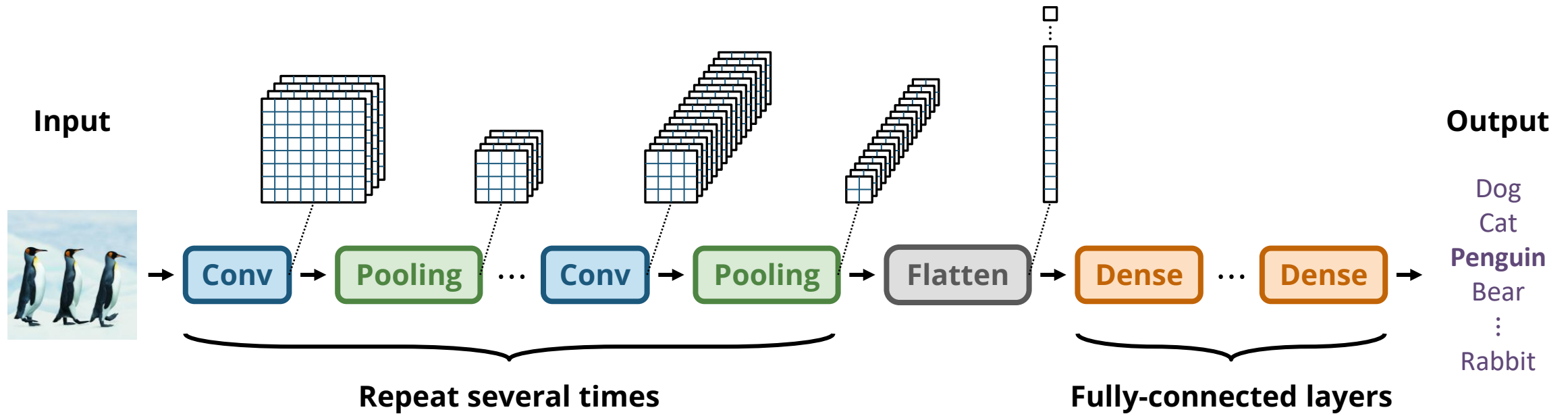
A convolutional neural network for raw waveform generation

Convolutional Layer

- A convolutional layer consists of many **learnable kernels** (channels)

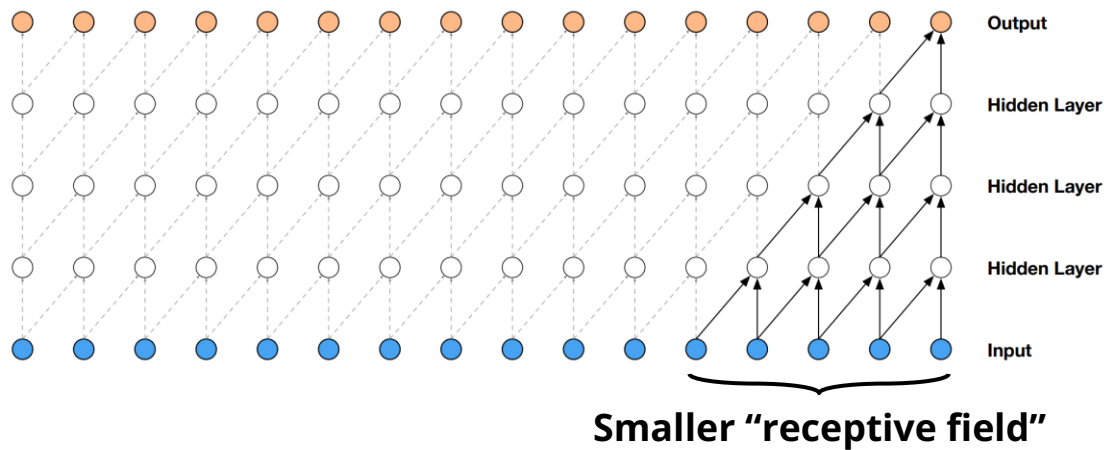


Convolutional Neural Network (CNNs)

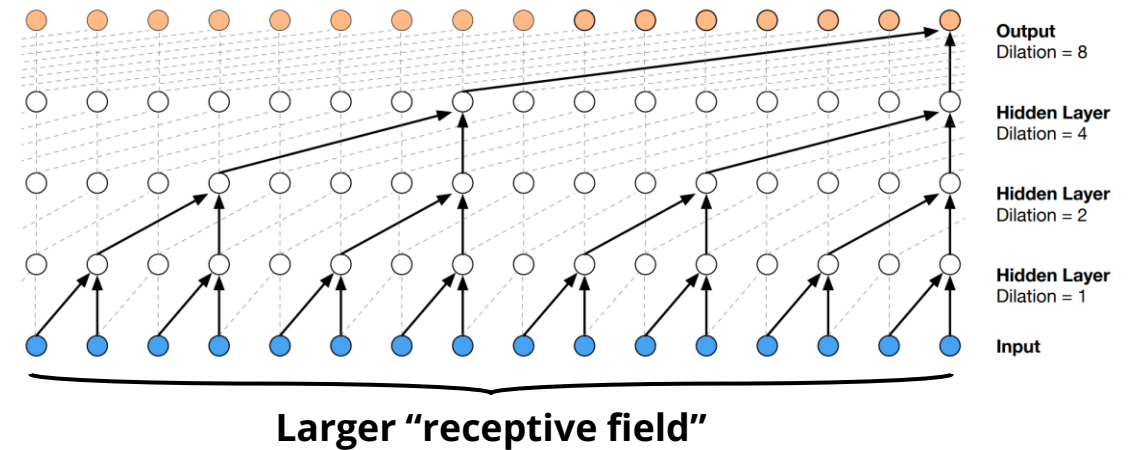


WaveNet (van den Oord et al., 2016)

Standard convolution



Dilated convolution



[deepmind.google/discover/
blog/wavenet-a-generative-
model-for-raw-audio](https://deepmind.google/discover/blog/wavenet-a-generative-model-for-raw-audio)

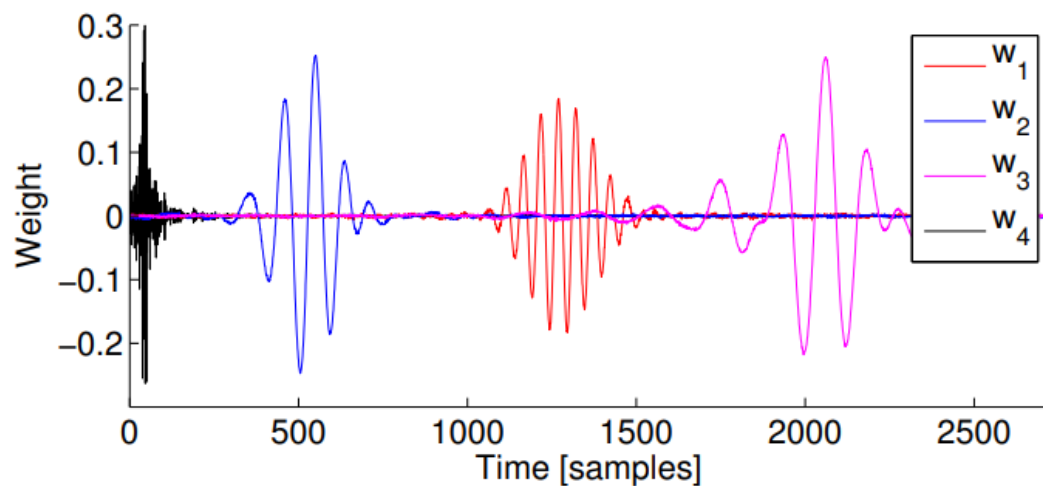
Example of generated music



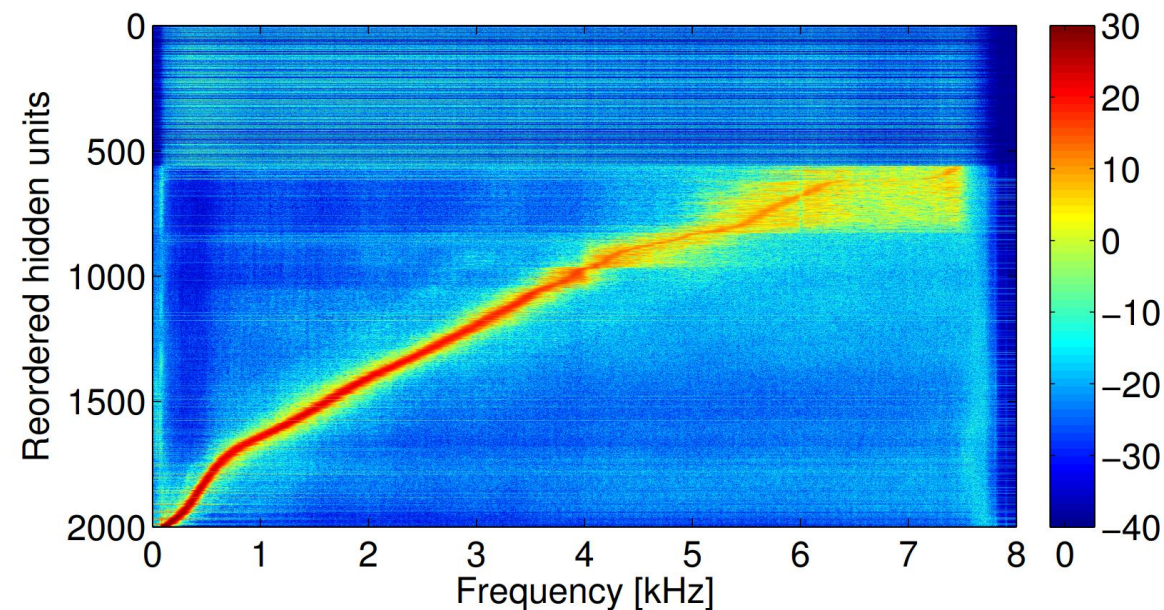
(Source: van den Oord et al., 2016)

1D CNNs & Fourier Transform

Convolution kernels learned

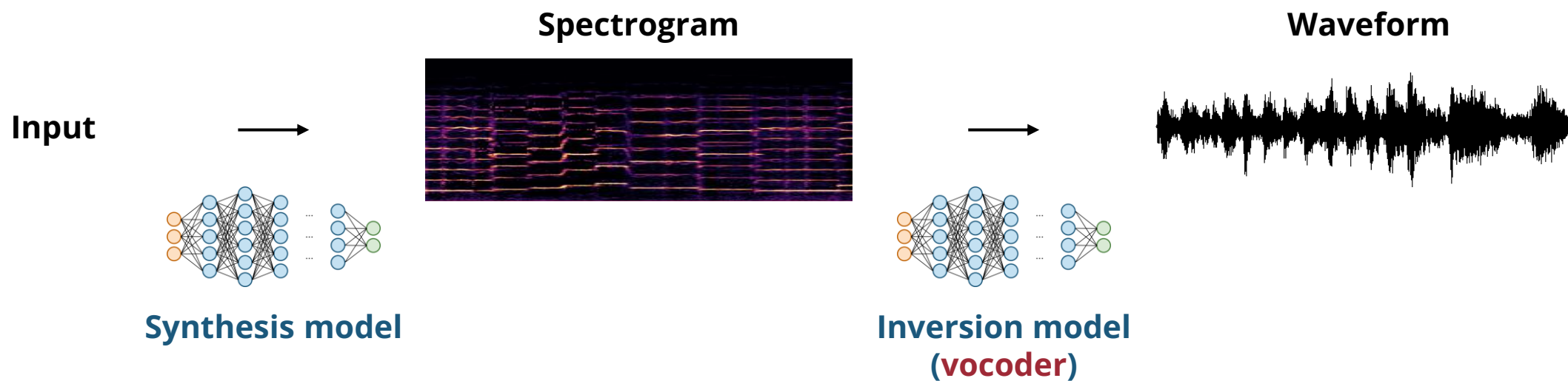


Peak frequency detected by the learned kernels

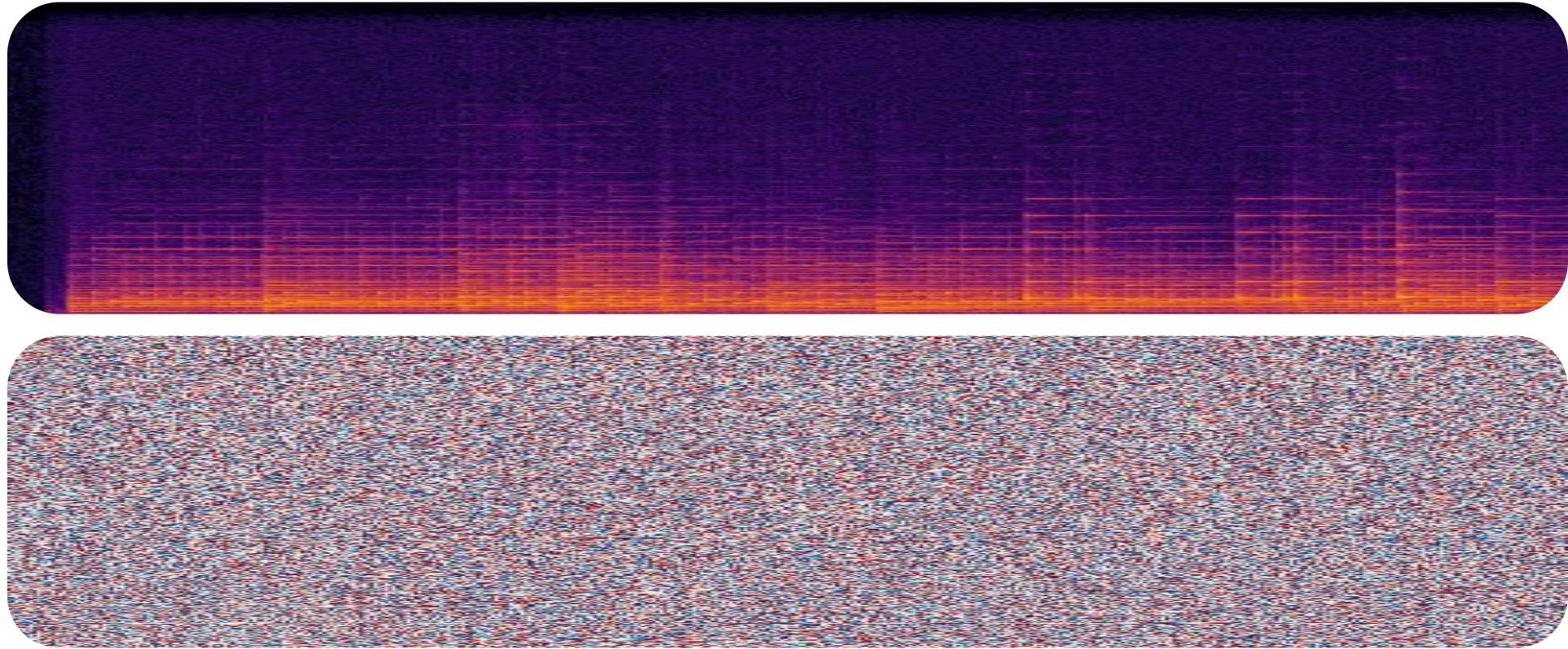


Frequency-domain Audio Synthesis

Frequency-domain Audio Synthesis



| Importance of the **Phase** Information

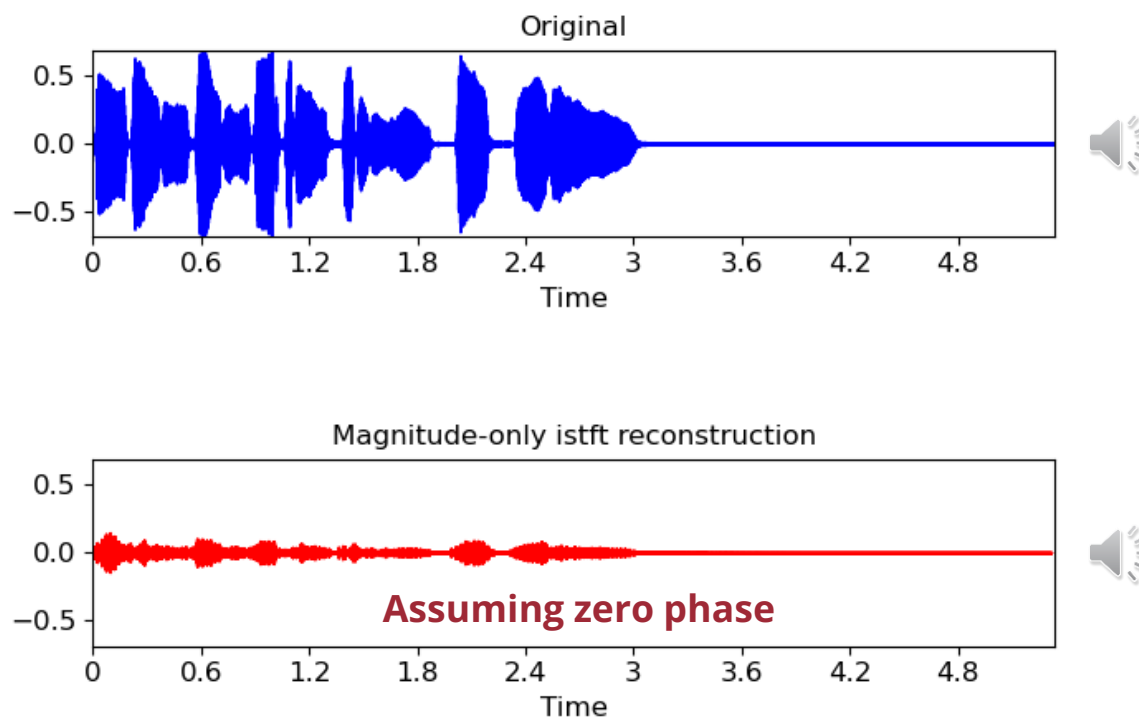


(Source: Dieleman et al., 2020)

Real phase 

Random phase 

Inverse STFT without Phase Information



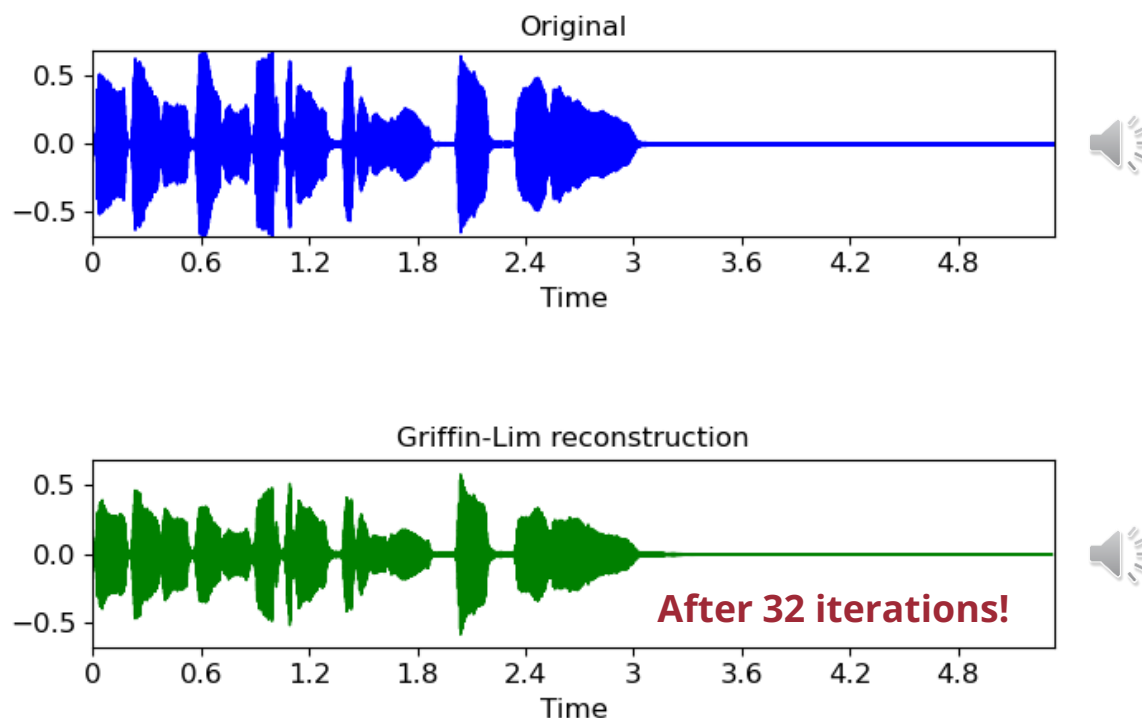
(Source: librosa documentation)

Complex-valued
STFT matrix

$$\text{ISTFT}(\boxed{M}) = \arg \min_y (M - \text{STFT}(y))^2$$

Find the signal y that minimize the
MSE between the input and $\text{STFT}(y)$

Griffin-Lim Algorithm (Griffin & Lim, 1984)



(Source: librosa documentation)

Given a magnitude-only STFT matrix



Randomly initialize the phase



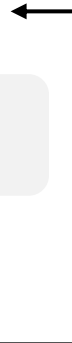
$$y' = \arg \min_y (M - \text{STFT}(y))^2$$

Find the signal y that minimize the MSE between the input and $\text{STFT}(y)$

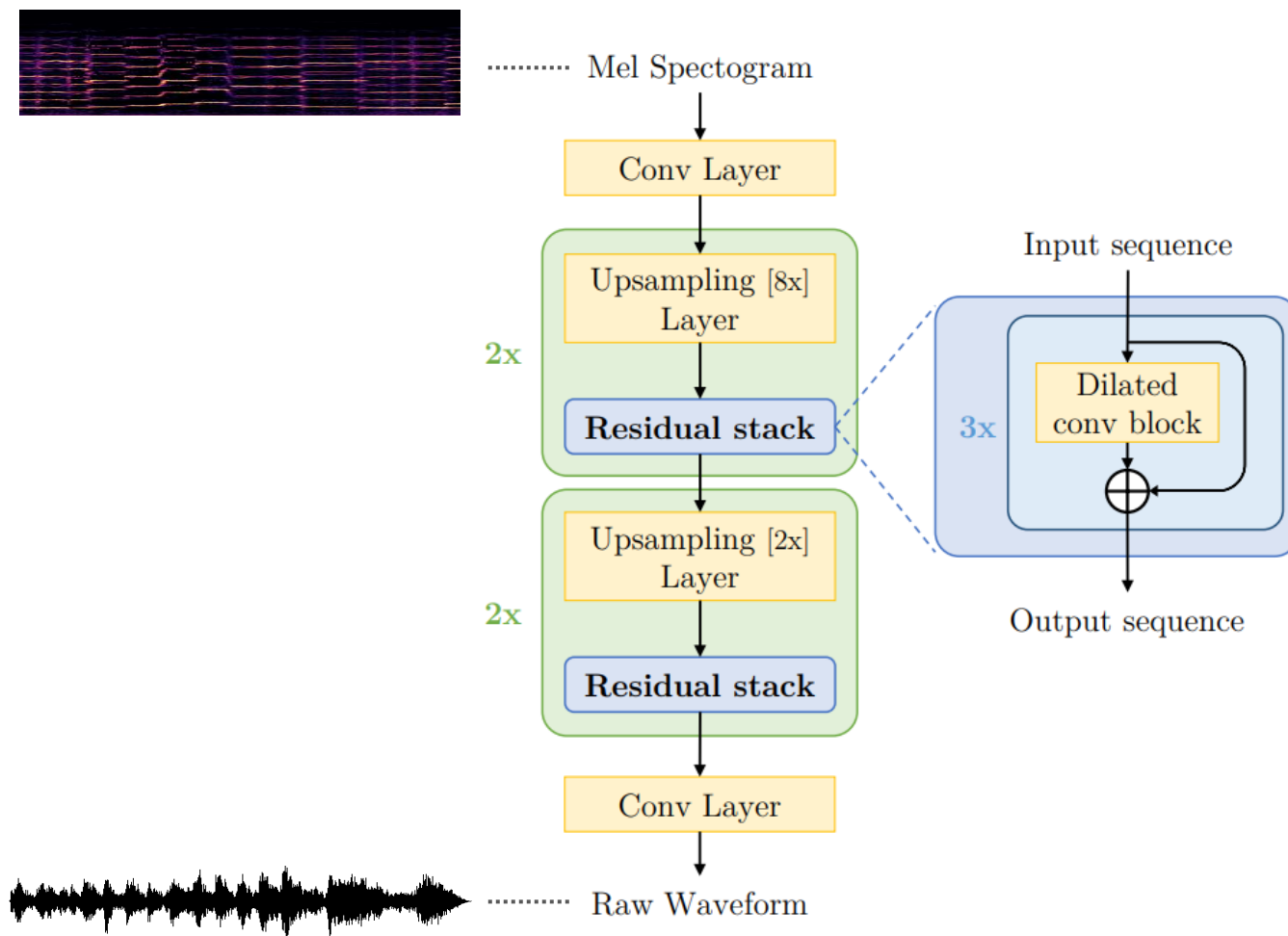


$$M' = \text{STFT}(y')$$

Find the STFT of the signal y



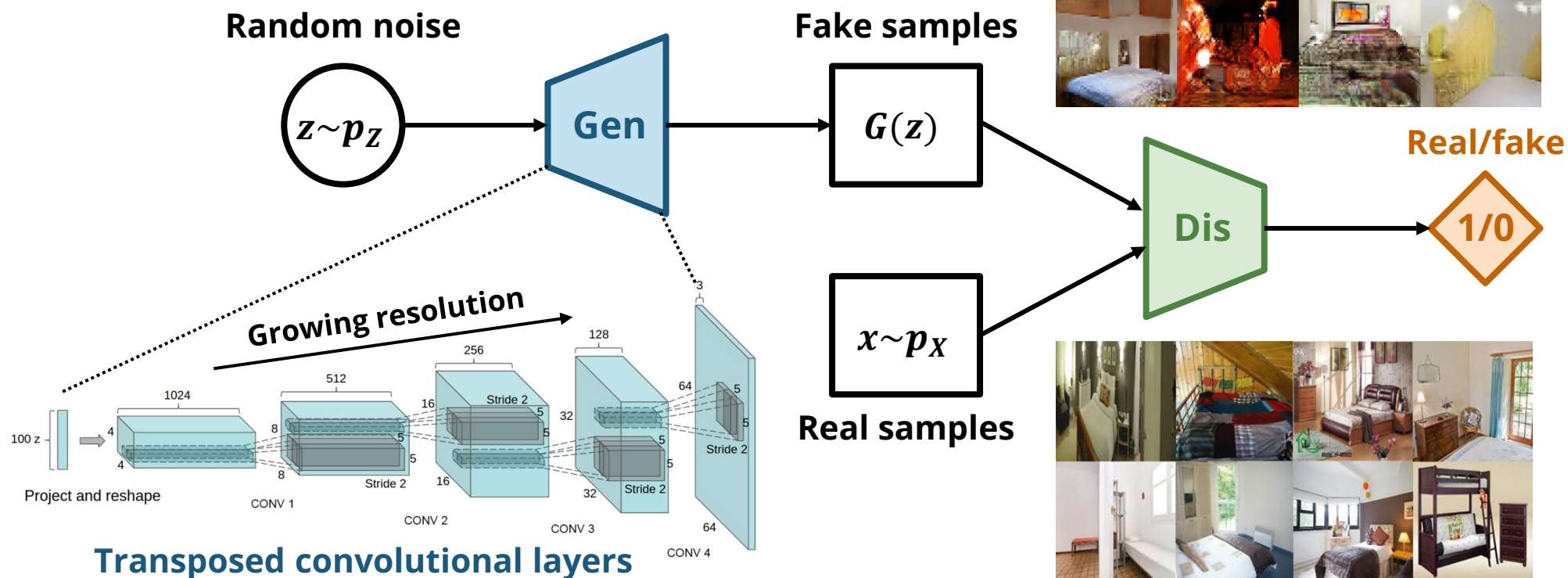
MelGAN (Kumar et al., 2019)



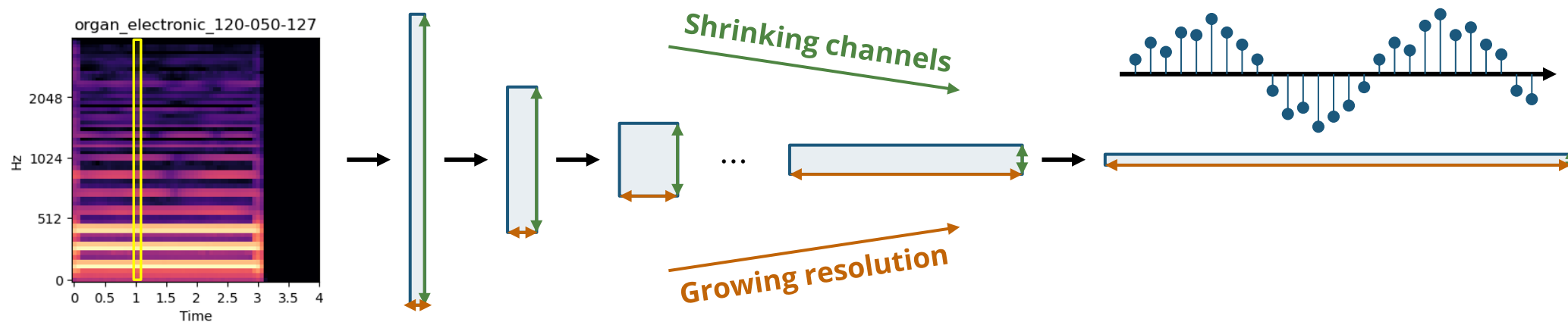
(Source: Kumar et al., 2019)

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

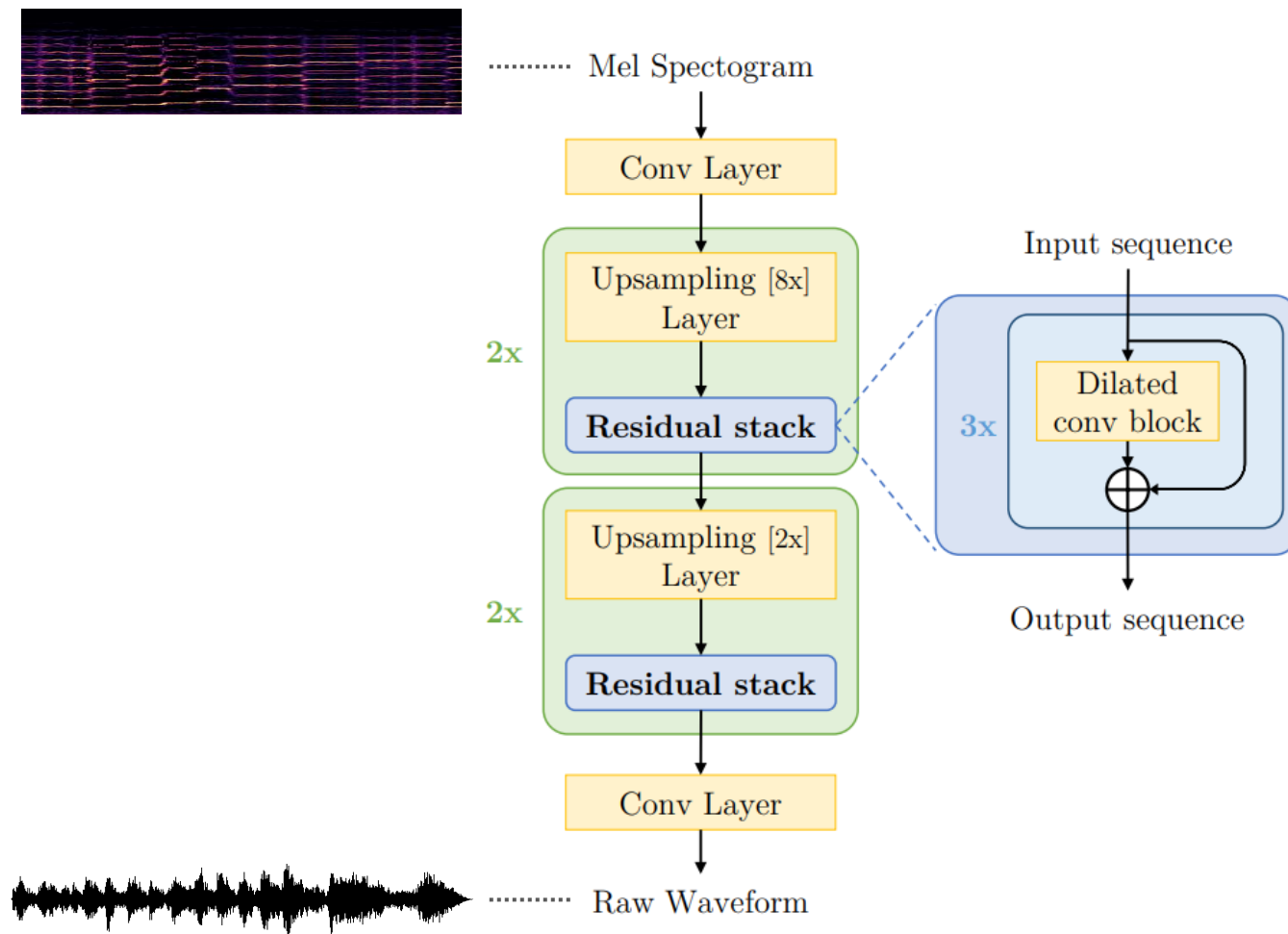
Use CNNs for both the generator and discriminator



Upsampling for Vocoders



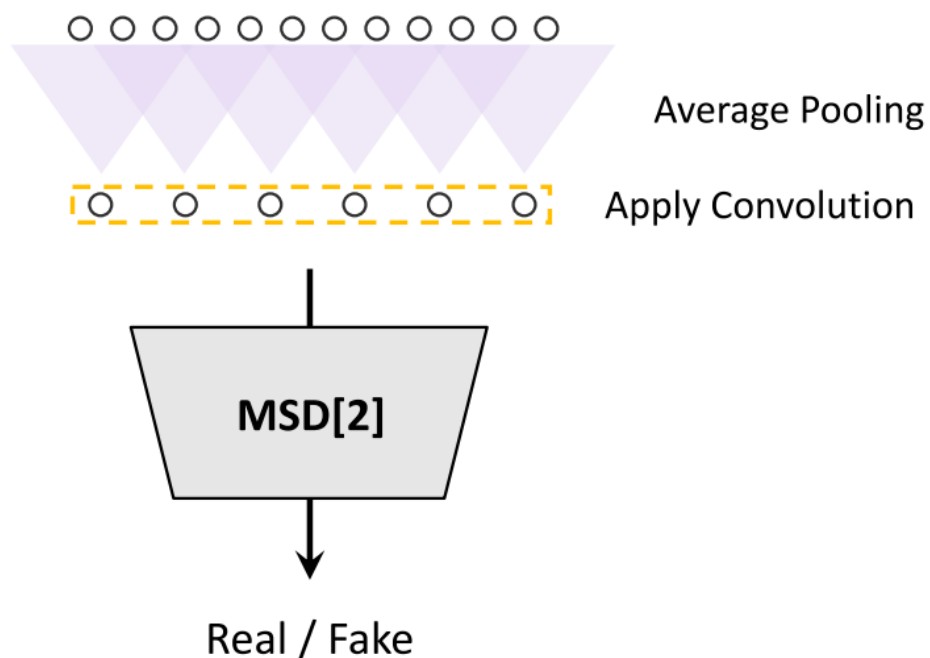
MelGAN (Kumar et al., 2019)



(Source: Kumar et al., 2019)

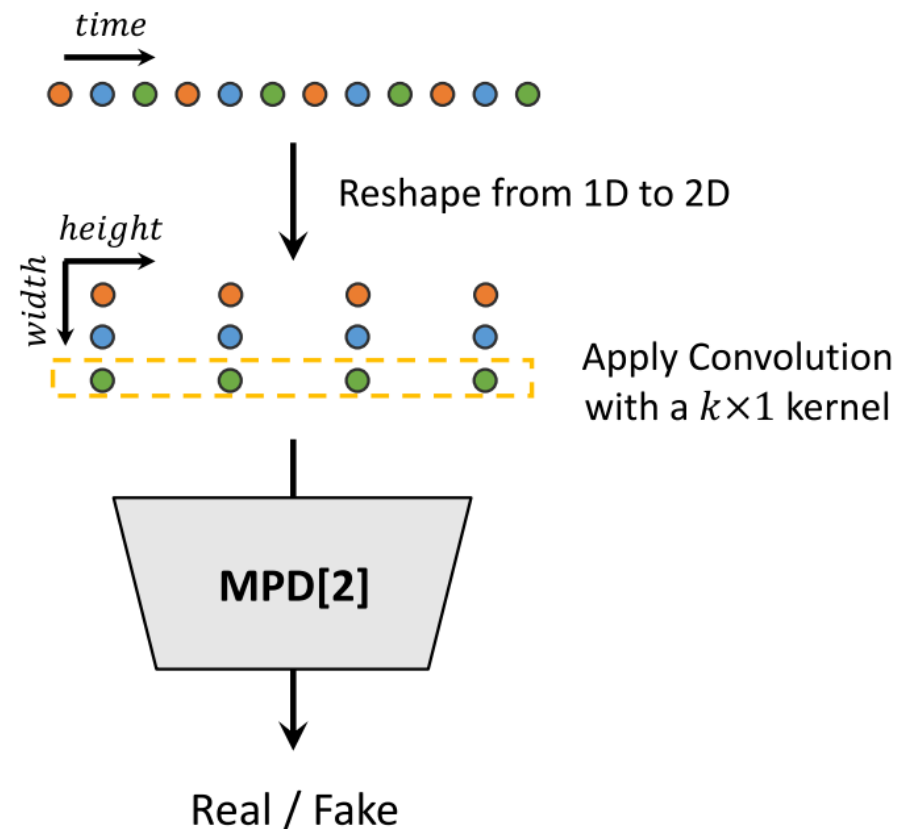
MelGAN (Kumar et al., 2019)

Multi-scale discriminator



(Source: Kong et al., 2019)

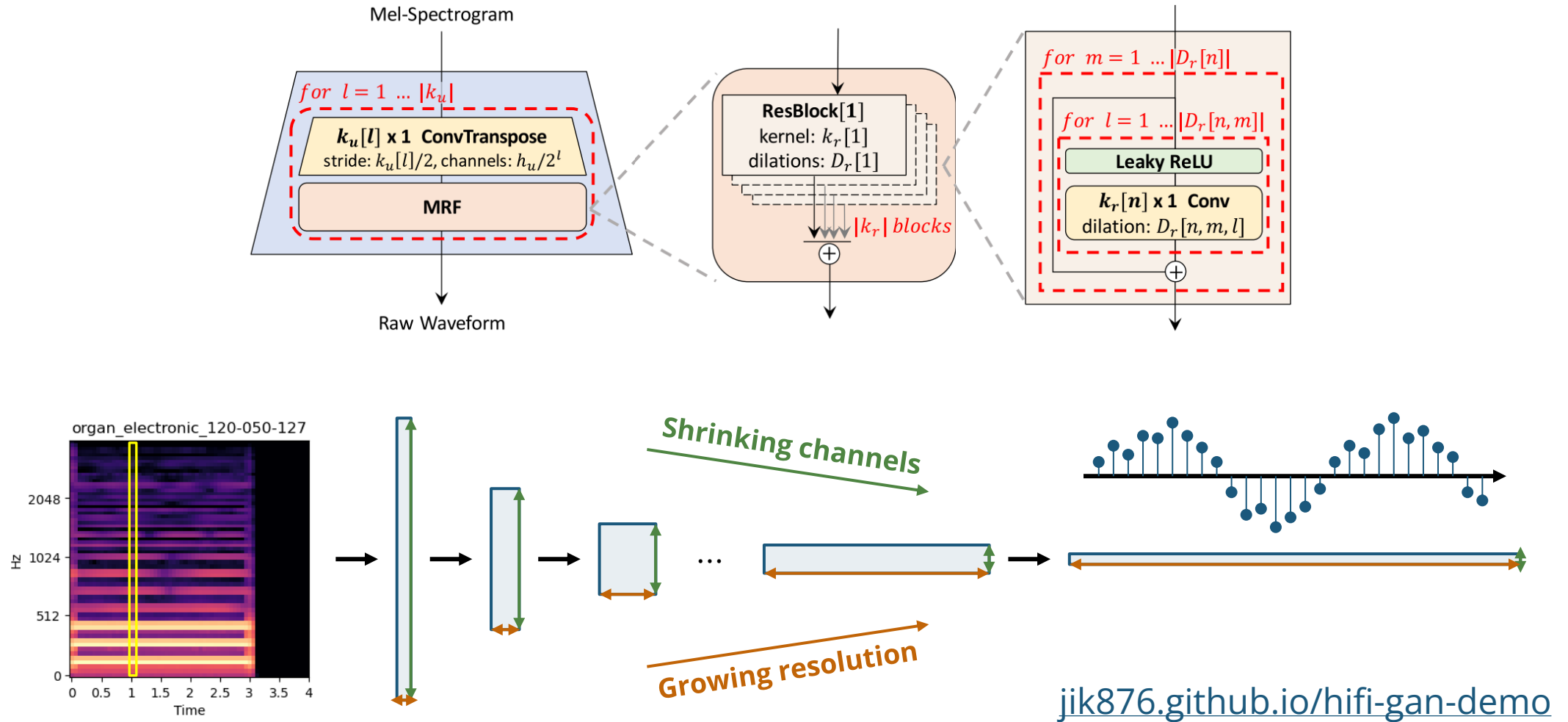
Multi-period discriminator



Kundan Kumar, Rithesh Kumar, Thibault de Boissiere, Lucas Gestin, Wei Zhen Teoh, Jose Sotelo, Alexandre de Brebisson, Yoshua Bengio, and Aaron Courville, "MelGAN: Generative Adversarial Networks for Conditional Waveform Synthesis," *NeurIPS*, 2019.

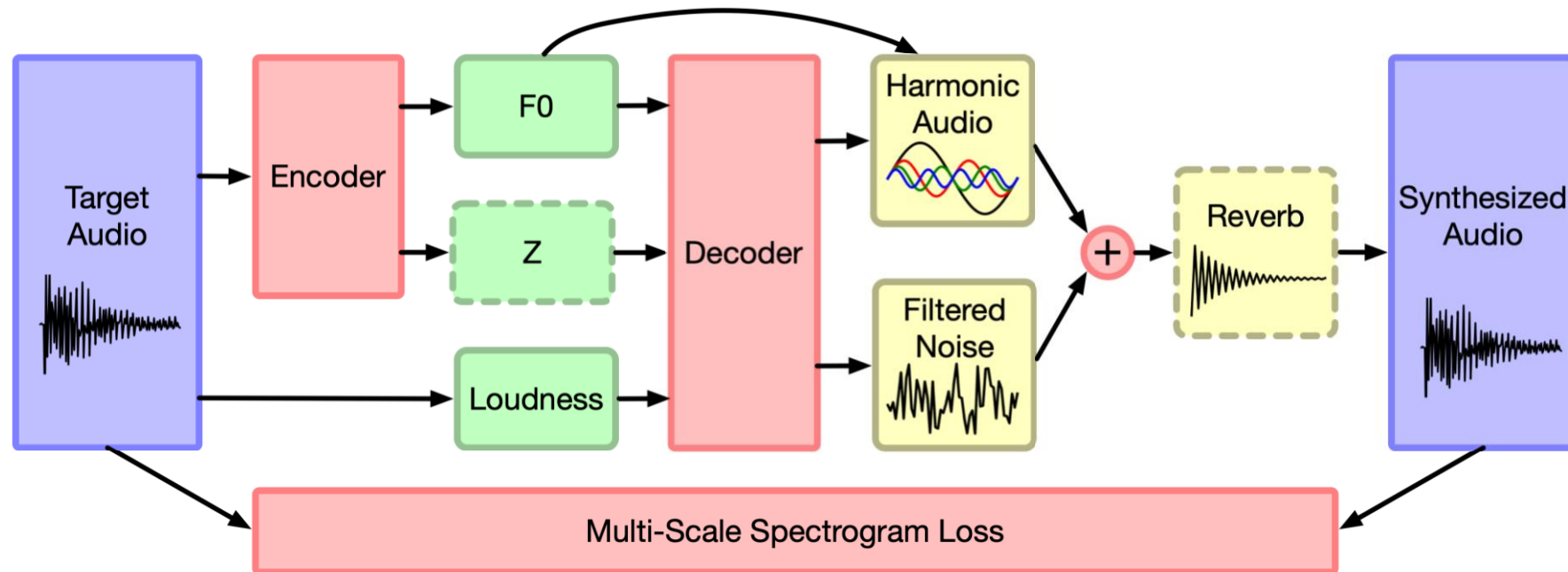
Jungil Kong, Jaehyeon Kim, and Jaekyoung Bae, "HiFi-GAN: Generative Adversarial Networks for Efficient and High Fidelity Speech Synthesis," *NeurIPS*, 2020.

Hifi-GAN (Kong et al., 2020)



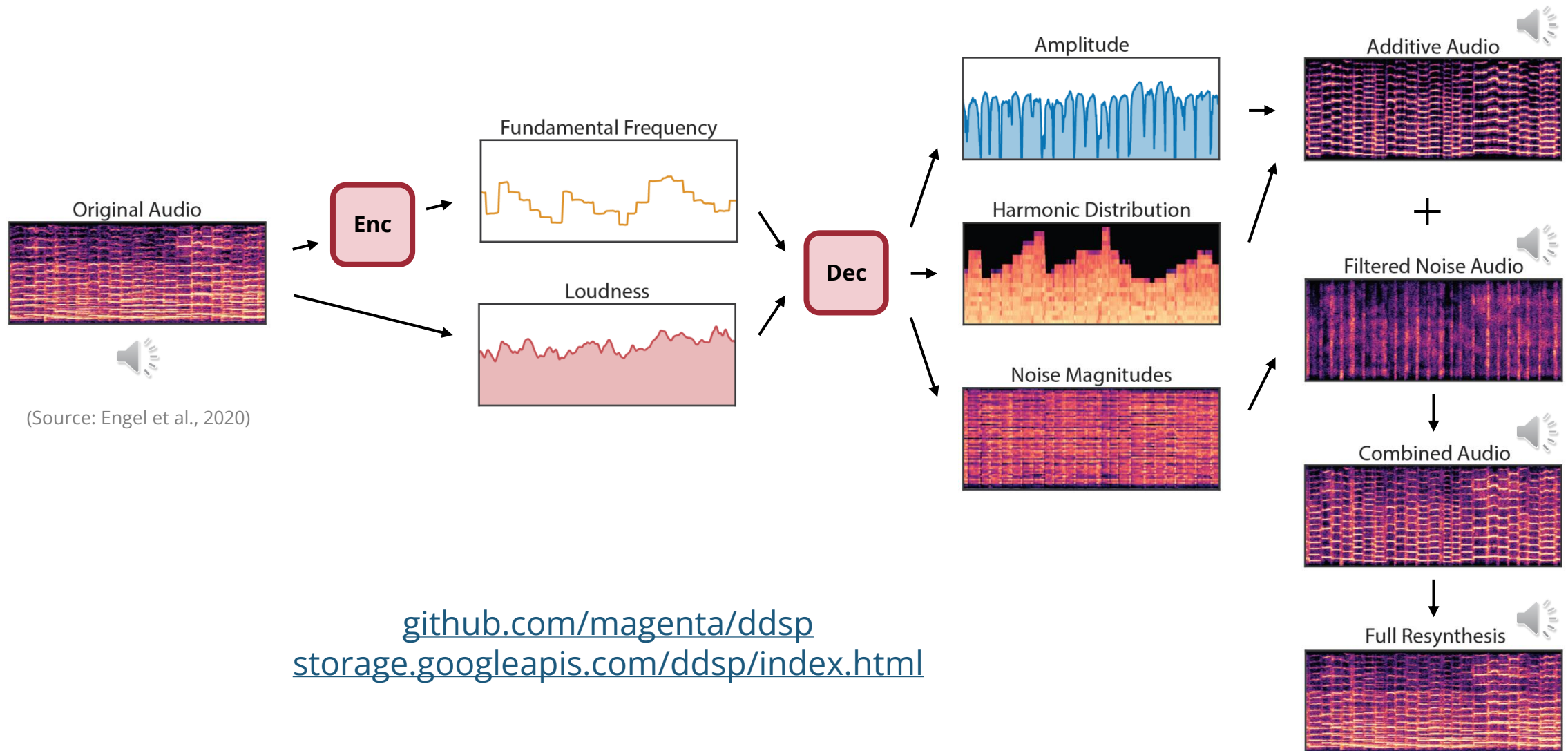
Differentiable Digital Signal Processing (DSP)

Differentiable DSP (DDSP) (Engel et al., 2020)



(Source: Engel et al., 2020)

Differentiable DSP (DDSP) (Engel et al., 2020)



Entering Demons & Gods by Yaboi Hanoi (2022)



youtu.be/PbrRoR3nEVw

soundcloud.com/yaboihanoi/enter-demons-and-gods

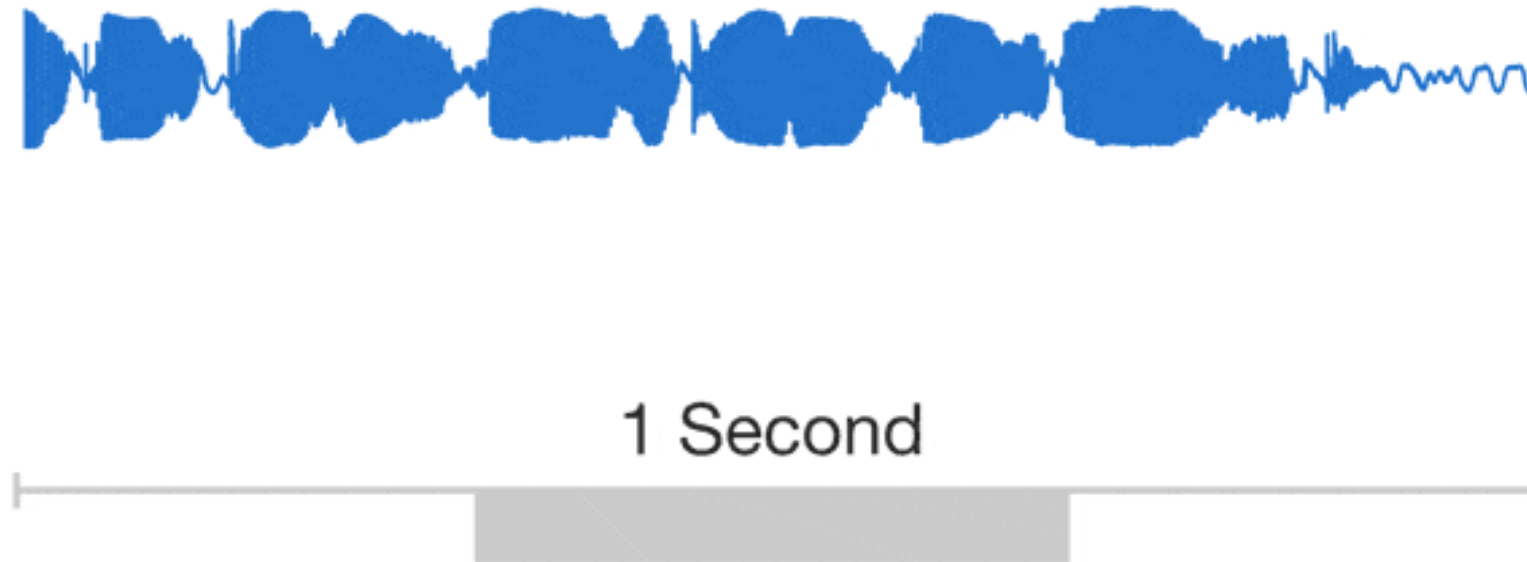


| Optional Reading

- A very nice blog on “**Generating music in the waveform domain**” by Sander Dieleman: sander.ai/2020/03/24/audio-generation

Recap

Generating Waveforms using a Neural Network



(Source: van den Oord et al., 2016)

Autoregressive Models (Mathematically)

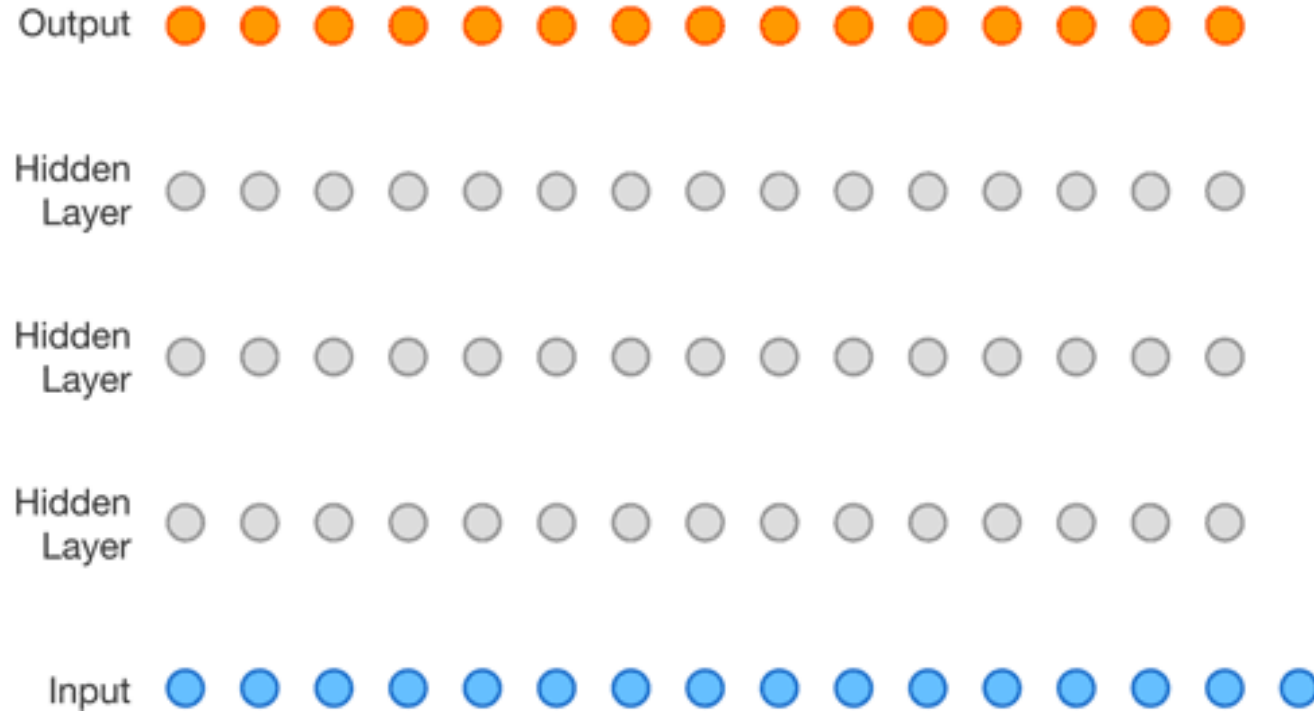
- A class of machine learning models that **learn** the probability of the next value given previous values

$$P(\underbrace{x_i}_{\text{Next number}} \mid \underbrace{x_1, x_2, \dots, x_{i-1}}_{\text{Previous numbers}})$$

$$\begin{array}{l} P(0.1 \mid 0.5, 0.4, 0.3, 0.2) \uparrow \\ P(0.09 \mid 0.5, 0.4, 0.3, 0.2) \uparrow \\ P(0.11 \mid 0.5, 0.4, 0.3, 0.2) \uparrow \\ P(99 \mid 0.5, 0.4, 0.3, 0.2) \downarrow \\ P(-1 \mid 0.5, 0.4, 0.3, 0.2) \downarrow \end{array}$$

The term “autoregressive” has different definitions in machine learning and signal processing.
In signal processing, an autoregressive model needs to be a linear model.

WaveNet (van den Oord et al., 2016)

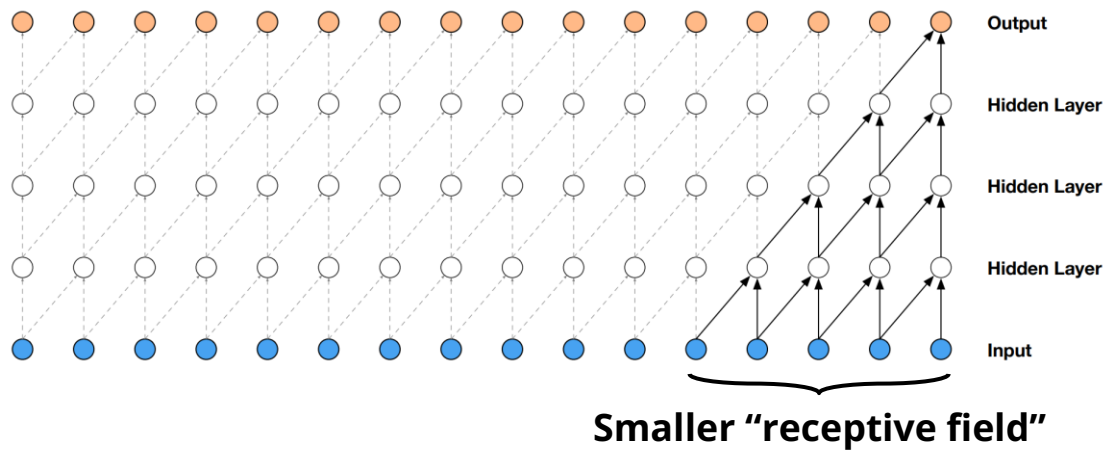


(Source: van den Oord et al., 2016)

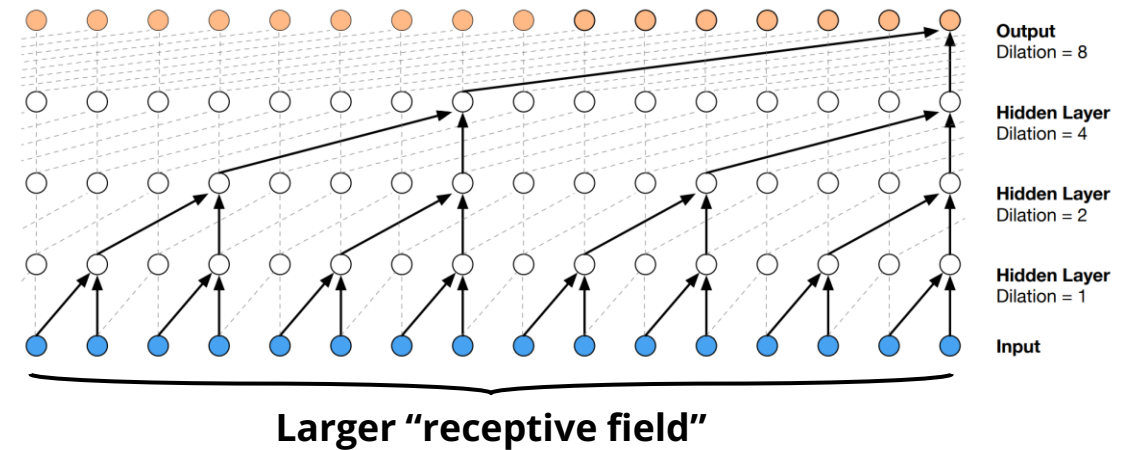
A convolutional neural network for raw waveform generation

WaveNet (van den Oord et al., 2016)

Standard convolution



Dilated convolution



[deepmind.google/discover/
blog/wavenet-a-generative-
model-for-raw-audio](https://deepmind.google/discover/blog/wavenet-a-generative-model-for-raw-audio)

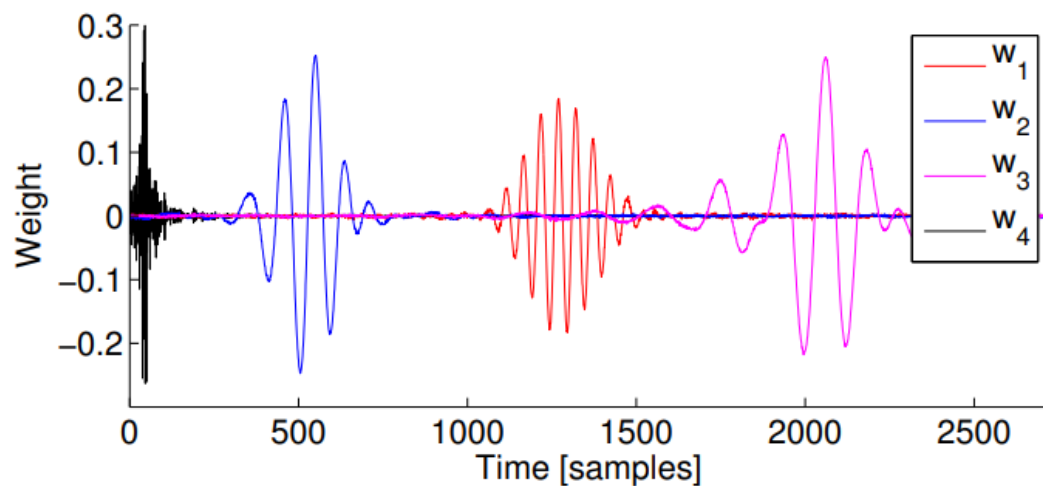
Example of generated music



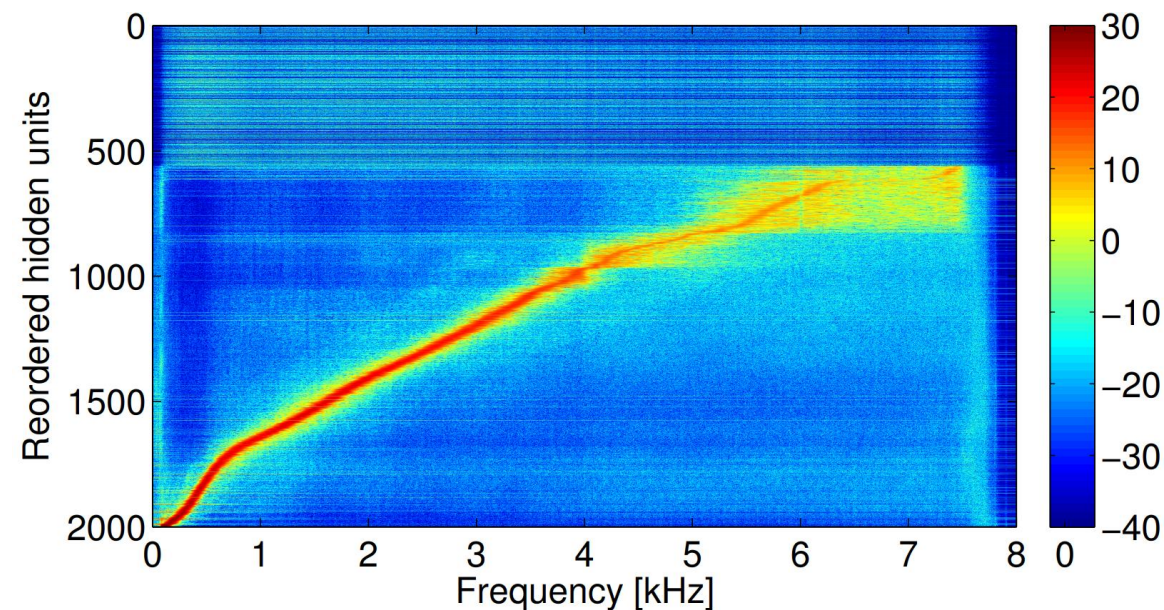
(Source: van den Oord et al., 2016)

1D CNNs & Fourier Transform

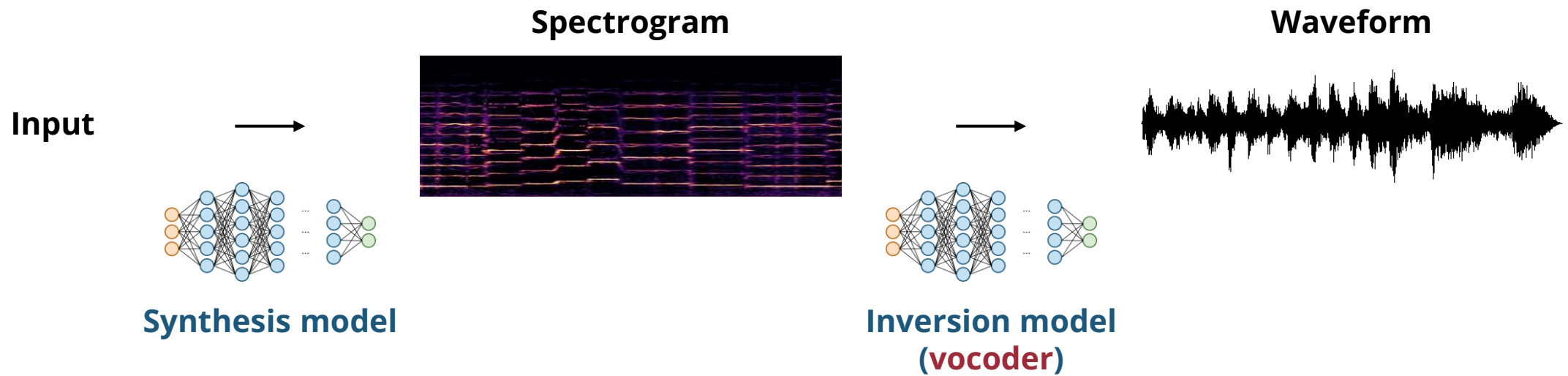
Convolution kernels learned



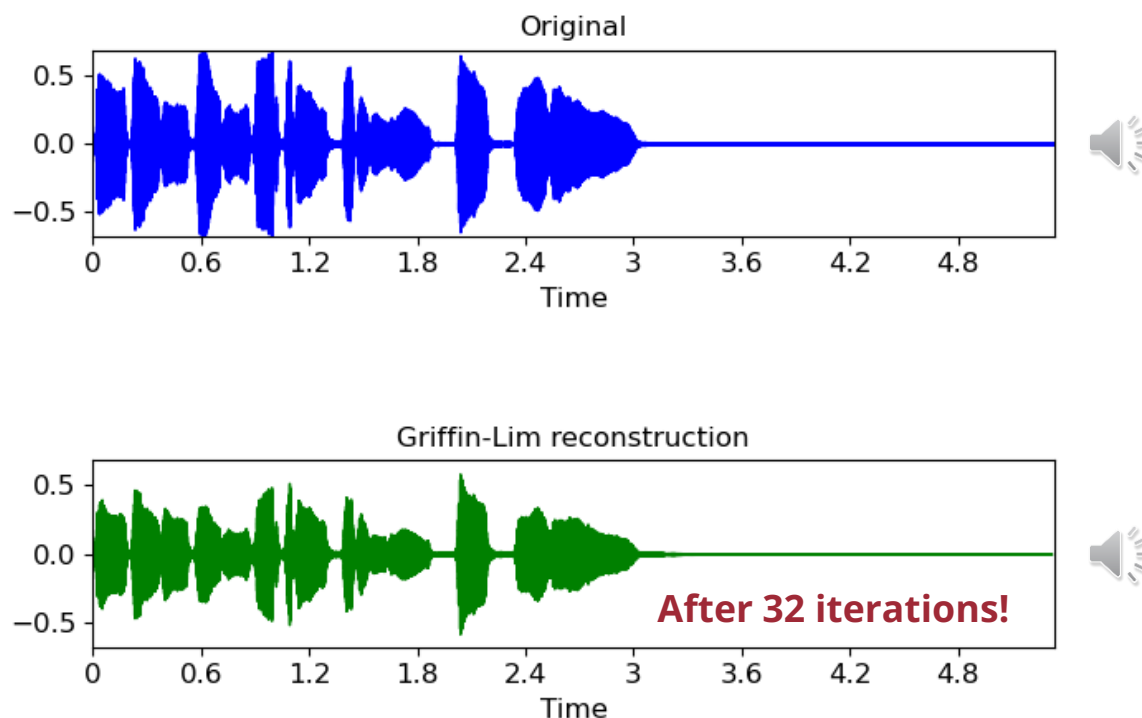
Peak frequency detected by the learned kernels



Frequency-domain Audio Synthesis



Griffin-Lim Algorithm (Griffin & Lim, 1984)



(Source: librosa documentation)

Given a magnitude-only STFT matrix



Randomly initialize the phase



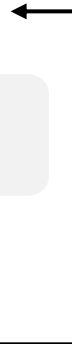
$$y' = \arg \min_y (M - \text{STFT}(y))^2$$

Find the signal y that minimize the MSE between the input and $\text{STFT}(y)$

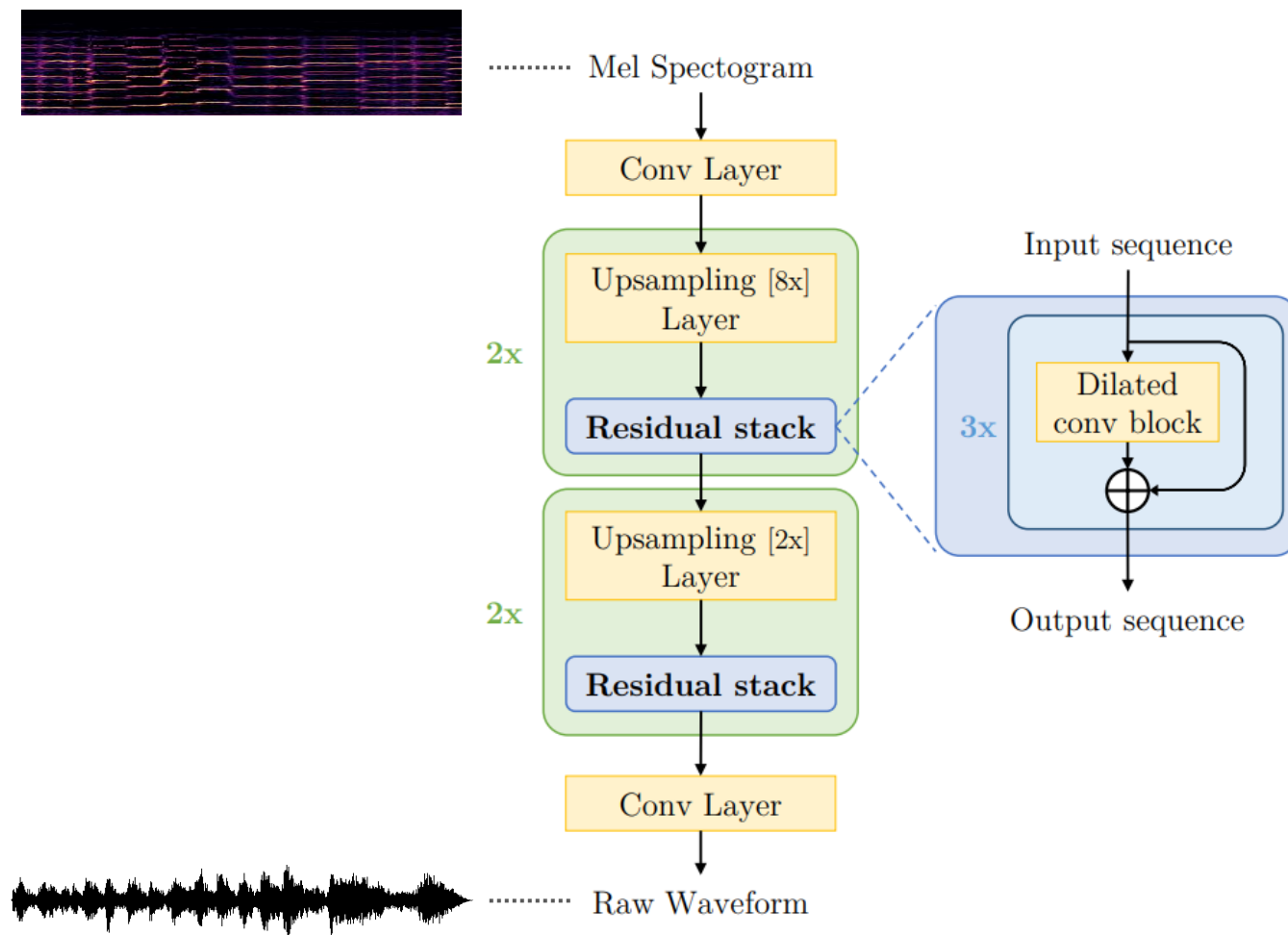


$$M' = \text{STFT}(y')$$

Find the STFT of the signal y

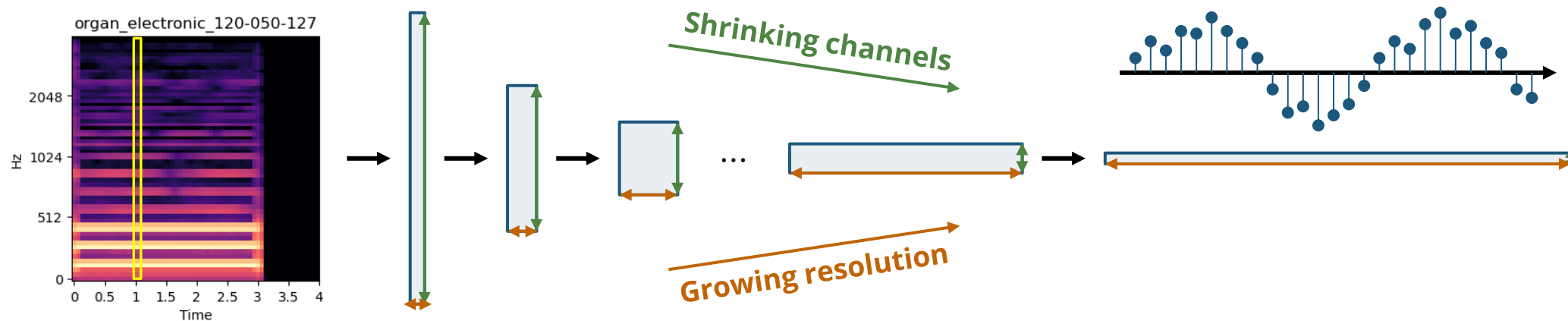


MelGAN (Kumar et al., 2019)

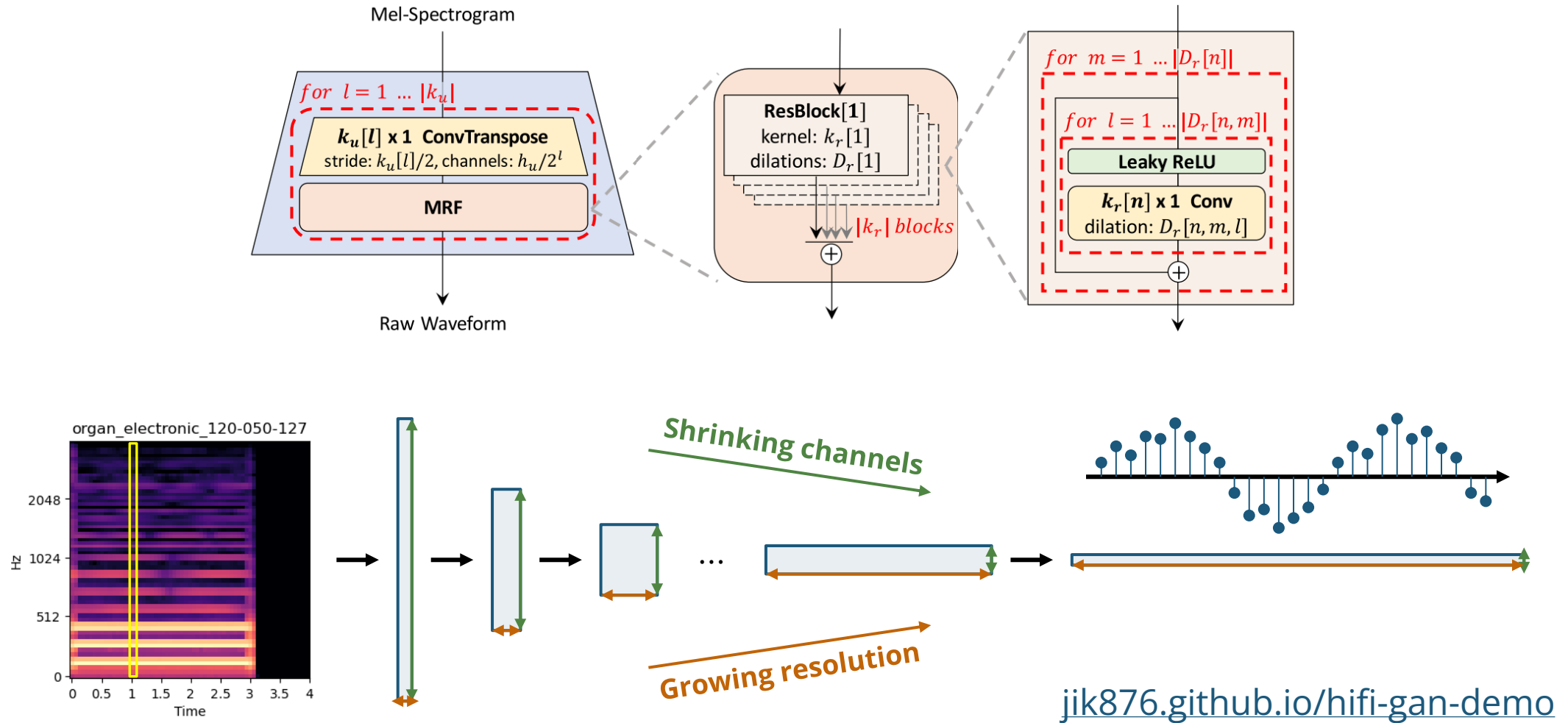


(Source: Kumar et al., 2019)

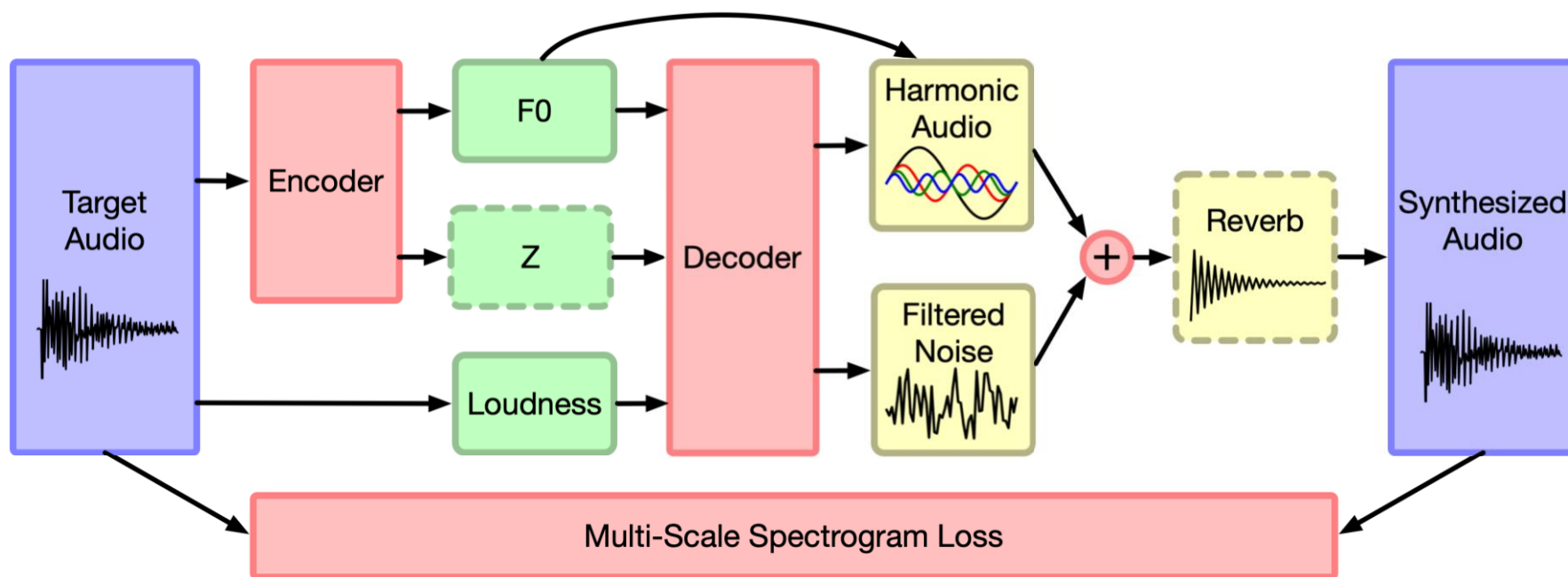
Upsampling for Vocoders



Hifi-GAN (Kong et al., 2020)

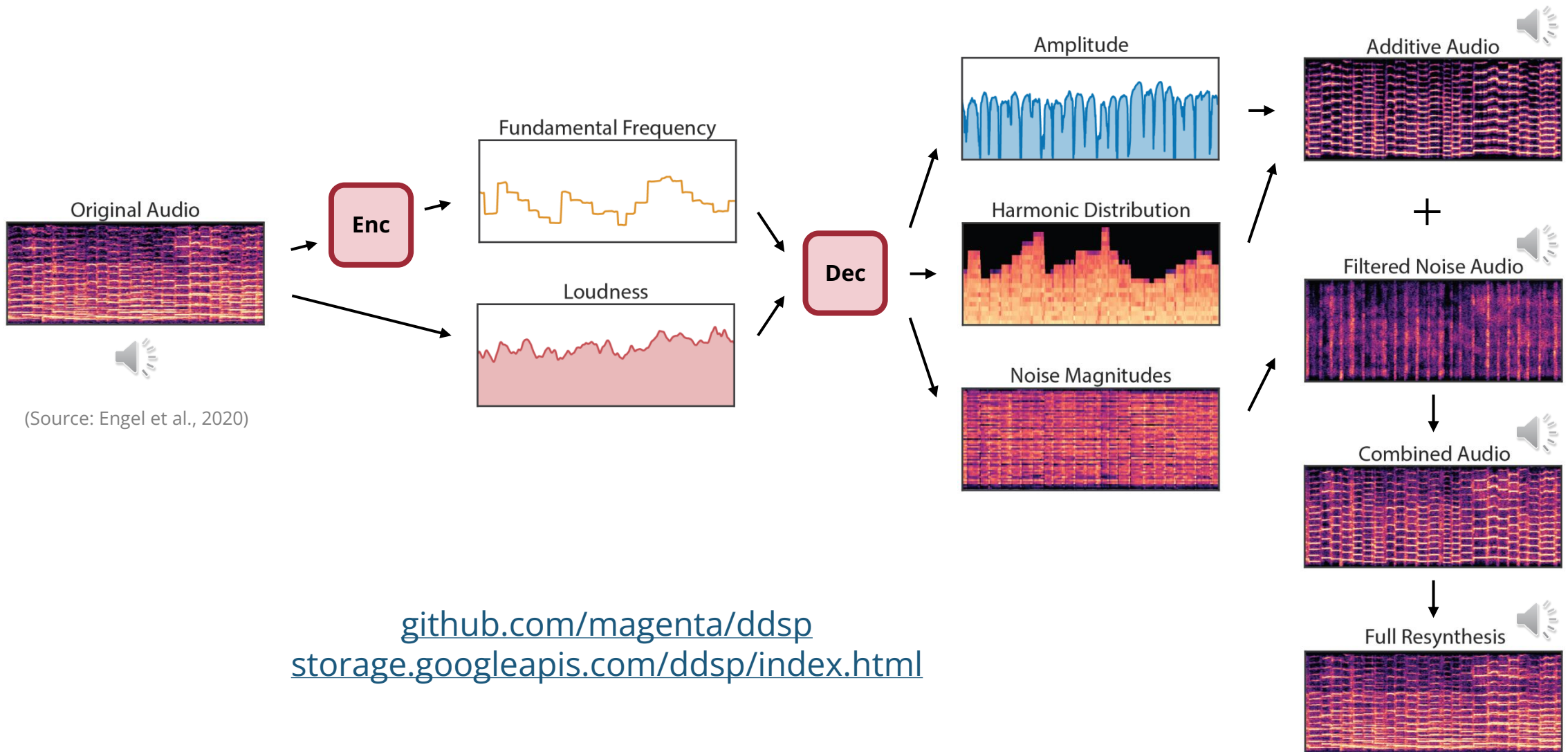


Differentiable DSP (DDSP) (Engel et al., 2020)



(Source: Engel et al., 2020)

Differentiable DSP (DDSP) (Engel et al., 2020)



Entering Demons & Gods by Yaboi Hanoi (2022)



youtu.be/PbrRoR3nEVw

soundcloud.com/yaboihanoi/enter-demons-and-gods



Next Lecture

Latent-based Music Generation

