

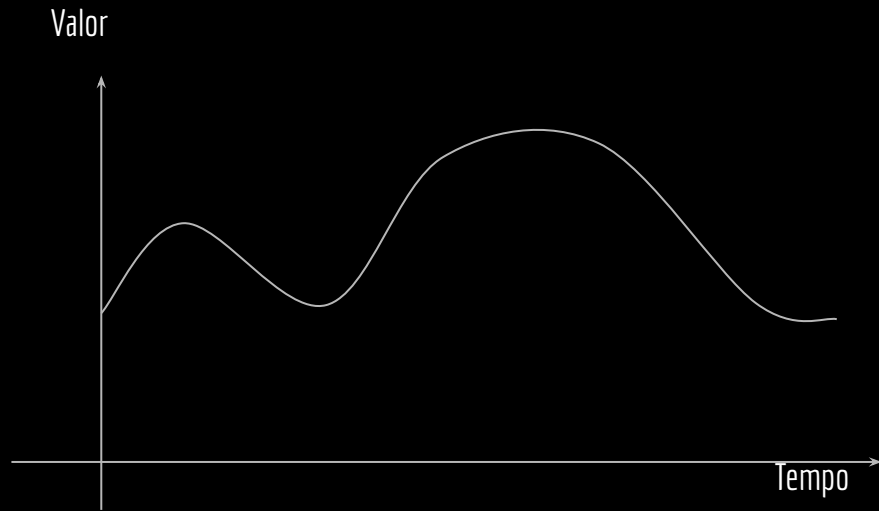
Alexa, faça de conta que você não está escutando.

Transformers Multimodais

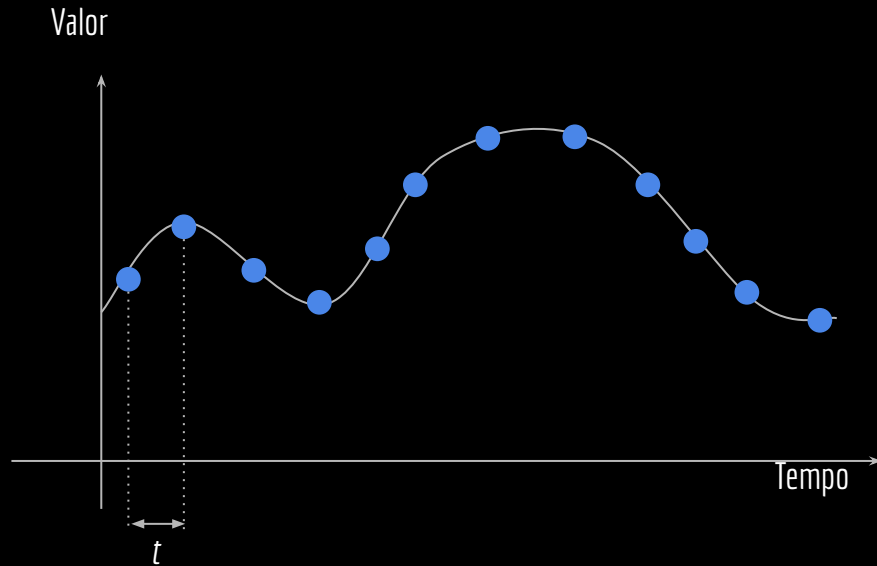
Paulo Ricardo Lisboa de Almeida



Mundo analógico e discreto



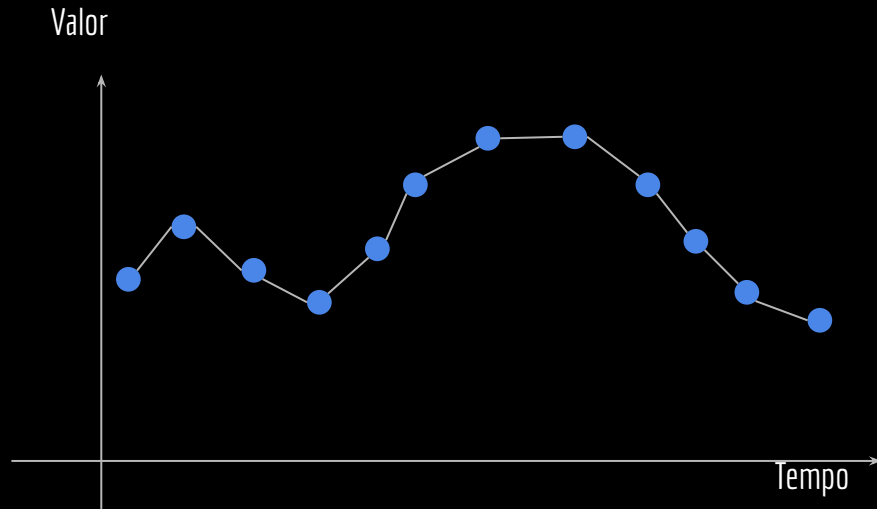
Mundo analógico e discreto



Amostragem.

Podemos discretizar um sinal analógico, verificando o seu valor a cada t instantes de tempo (t é o período).

Mundo analógico e discreto



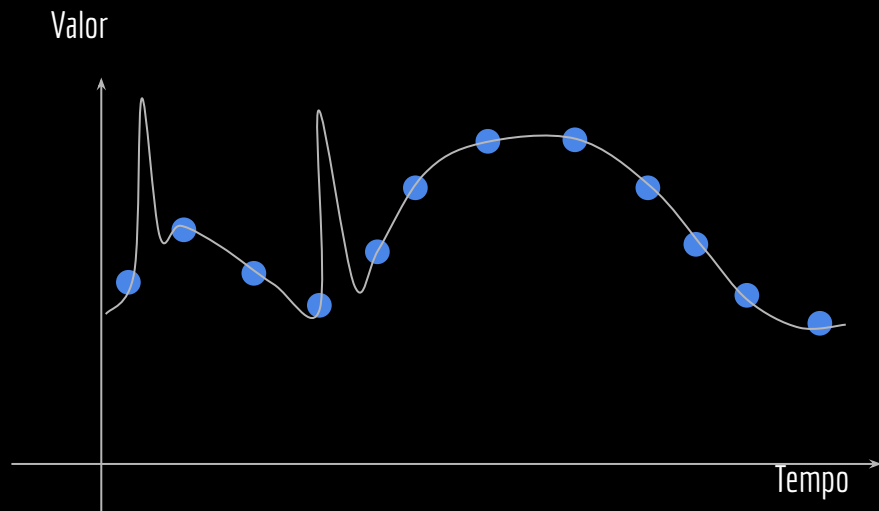
Amostragem.

Podemos discretizar um sinal analógico, verificando o seu valor a cada t instantes de tempo.

Quanto menor for t , mais fiel é o sinal discreto quando comparado ao analógico.

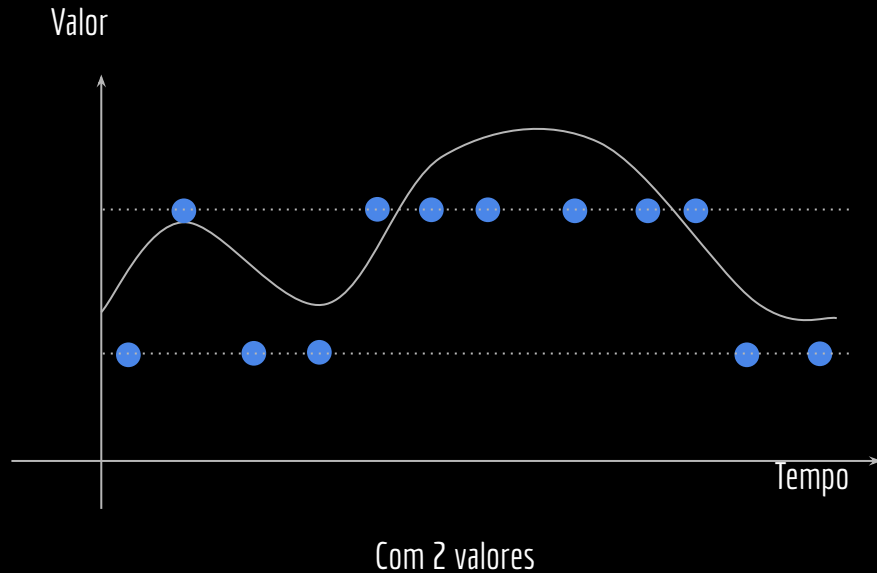
Podemos interpolar o sinal discreto para reconstruir (algo próximo) do sinal original.

Mundo analógico e discreto



Durante a discretização, algumas informações podem ser perdidas.

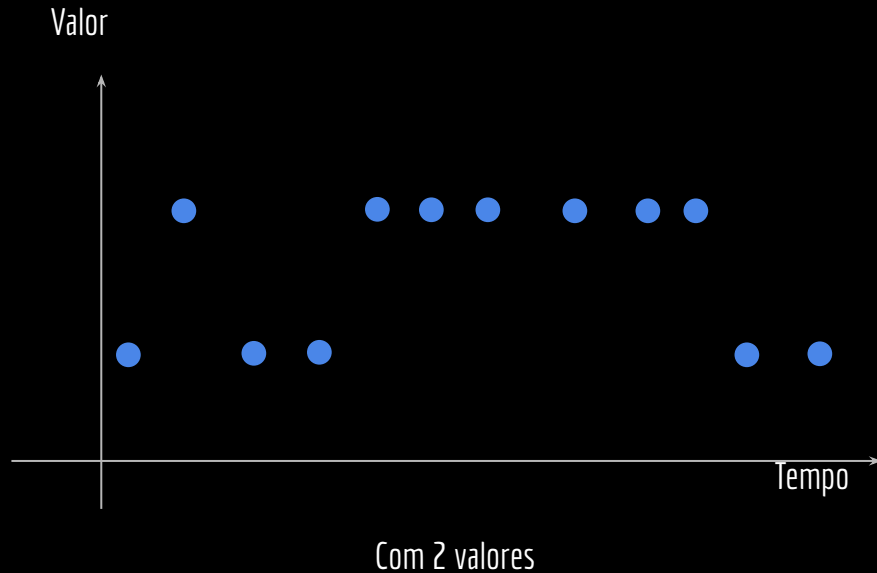
Mundo analógico e discreto



Quantização.

No eixo y, quantos valores podemos representar? Quanto mais valores valores, maior a fidelidade no eixo y com o sinal original, mas o hardware se torna mais complicado.

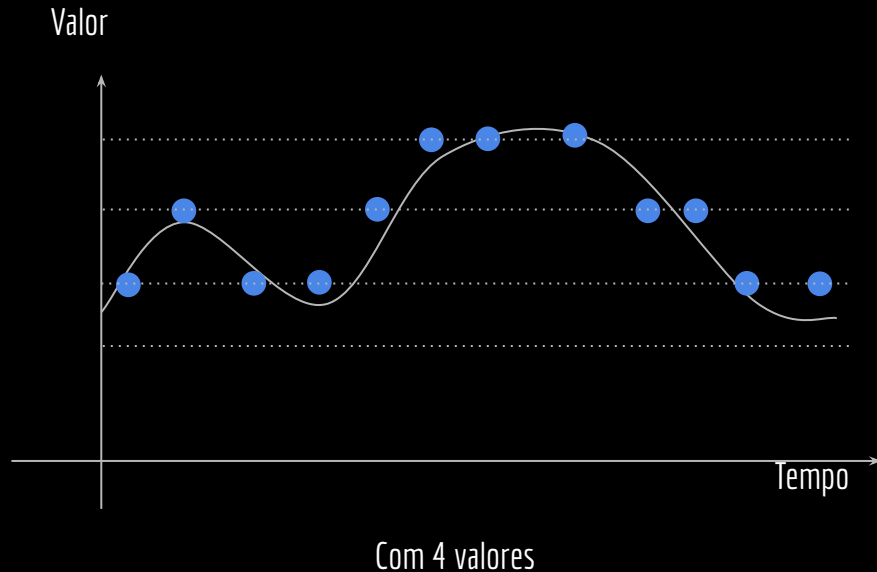
Mundo analógico e discreto



Quantização.

No eixo y, quantos valores podemos representar? Quanto mais valores valores, maior a fidelidade no eixo y com o sinal original, mas o hardware se torna mais complicado.

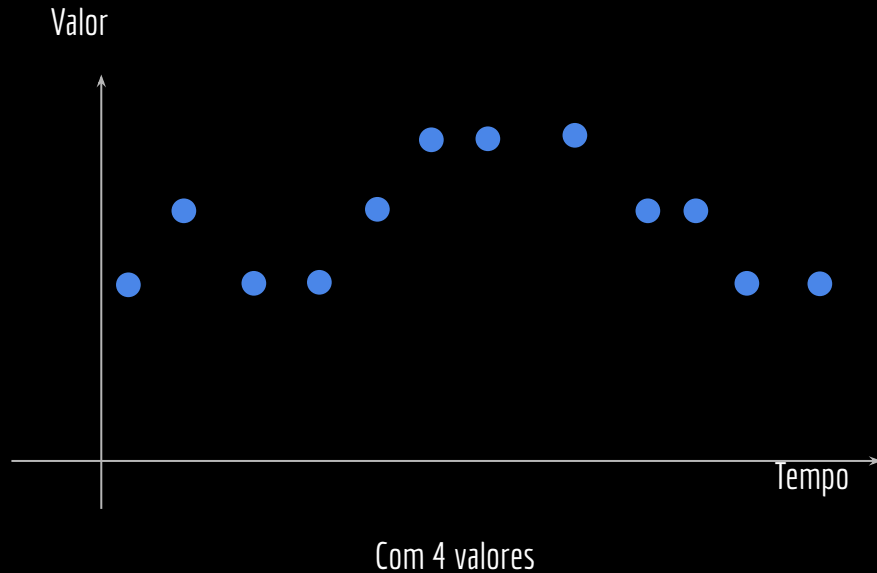
Mundo analógico e discreto



Quantização.

No eixo y, quantos valores podemos representar? Quanto mais valores valores, maior a fidelidade no eixo y com o sinal original, mas o hardware se torna mais complicado.

Mundo analógico e discreto



Quantização.

No eixo y, quantos valores podemos representar? Quanto mais valores valores, maior a fidelidade no eixo y com o sinal original, mas o hardware se torna mais complicado.

Exemplo

Geralmente as músicas que ouvimos possuem uma amostragem de 44.100Hz.

Isso significa que a cada $t = 1/44100 \approx 0,000023$ segundos um sinal é amostrado.

Exemplo

Geralmente as músicas que ouvimos possuem uma amostragem de 44.100Hz.

Isso significa que a cada $t = 1/441000 \approx 0,000023$ segundos um sinal é amostrado.

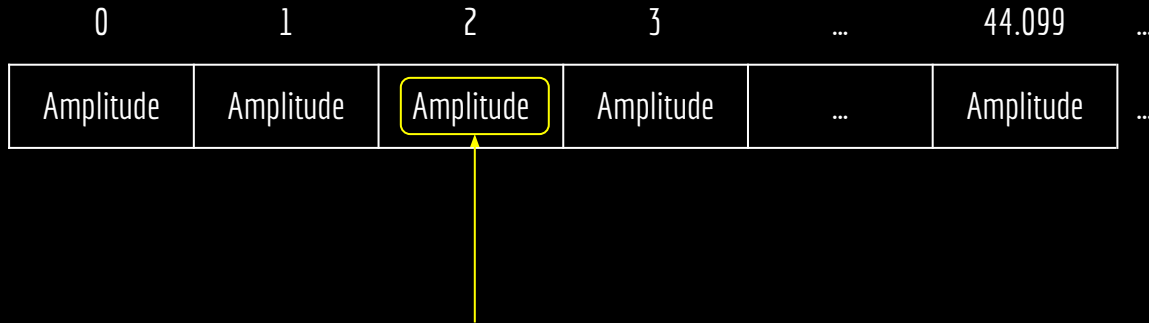
$f = 1/t$, onde f é a frequência, e t o período.

Exemplo - CD



Técnica chamada de Pulse-Code Modulation.

Exemplo - CD



Cada amostra de amplitude é armazenada como um valor de 16 bits.

Técnica chamada de Pulse-Code Modulation.

Exemplo - CD

Podemos adicionar mais fluxos para ter múltiplos canais de áudio. Por exemplo, em um CD com dois canais (stereo).

0	1	2	3	...	44.099	...
Amplitude	Amplitude	Amplitude	Amplitude	...	Amplitude	...

Canal 1.

0	1	2	3	...	44.099	...
Amplitude	Amplitude	Amplitude	Amplitude	...	Amplitude	...

Canal 2.

Técnica chamada de Pulse-Code Modulation.

Convoluções

0	1	2	3	...	44.099	...
Amplitude	Amplitude	Amplitude	Amplitude	...	Amplitude	...

Podemos aplicar uma CNN.

Convoluções de $I \times N \times C$, onde C é o número de canais, e N é o tamanho do Kernel.

Transformer

0	1	2	3	...	44.099	...
Amplitude	Amplitude	Amplitude	Amplitude	...	Amplitude	...

Ou podemos aplicar um transformer, que vai ser capaz de capturar o contexto global.
Mas possui um custo quadrático de acordo com o tamanho da entrada.

Ideia

A mesma ideia aplicada para modelos de linguagem pode ser aplicada.

Dado um conjunto de tokens (de áudio) prever, por exemplo, qual o token (de texto) pertence a esse trecho.

Nesse exemplo, o modelo se torna um *speech to text*.

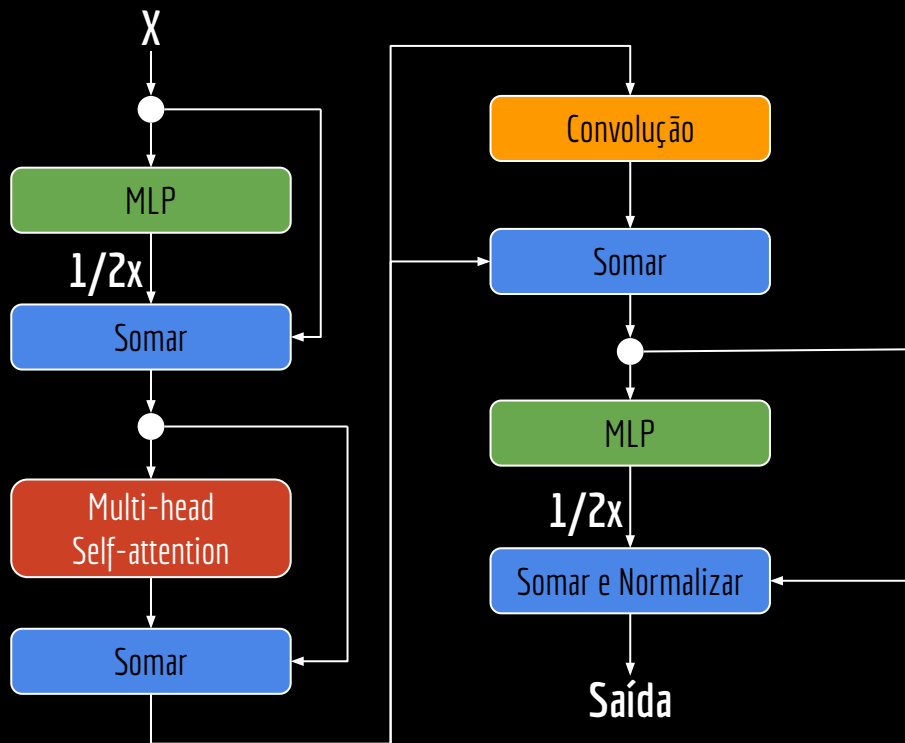
Estado da arte

Técnicas do estado da arte combinam convoluções com transformers.

Convoluções para capturar as correlações locais e reduzir a dimensionalidade dos dados.

Transformers para capturar as correlações globais.

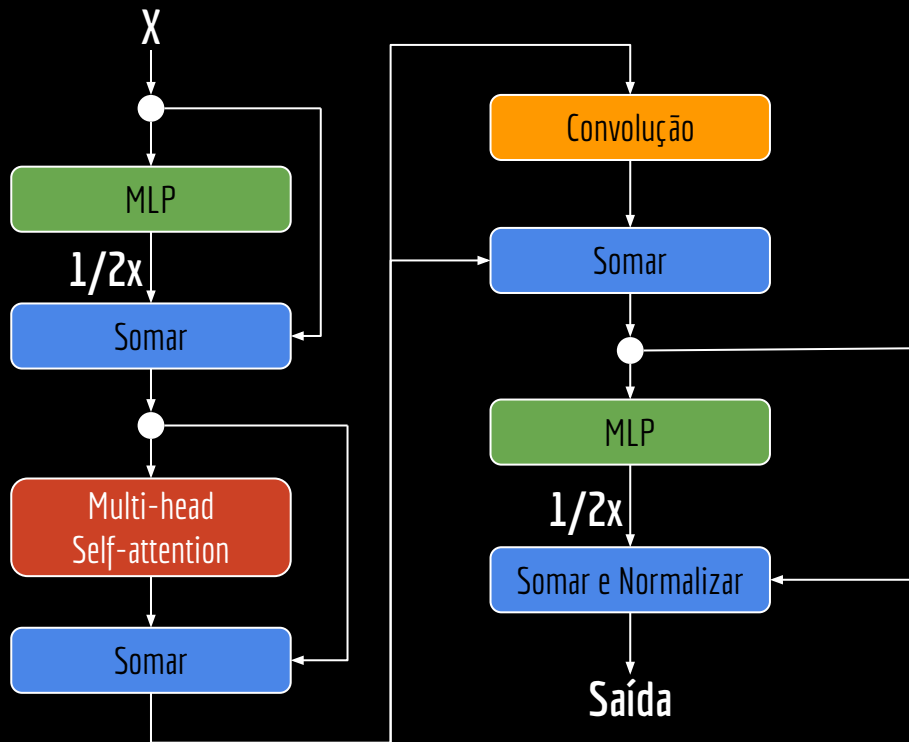
Conformer



GULATI, Anmol et al. Conformer: Convolution-augmented transformer for speech recognition. 2020.



Conformer



GULATI, Anmol et al. Conformer: Convolution-augmented transformer for speech recognition. 2020.

$$\begin{aligned}\tilde{x} &= x + \frac{1}{2}MLP(x) \\ x' &= \tilde{x} + MHSA(\tilde{x}) \\ x'' &= x' + Conv(x') \\ y &= normalizar(x'' + \frac{1}{2}MLP(x''))\end{aligned}$$



Conformer

Cabeças de atenção ficam em um “sanduíche” de MLPs.

O primeiro MLP projeta os dados para uma dimensão 4x menor. O segundo desfaz essa operação.

GULATI, Anmol et al. Conformer: Convolution-augmented transformer for speech recognition. 2020.

$$\tilde{\mathbf{x}} = \mathbf{x} + \frac{1}{2}MLP(\mathbf{x})$$

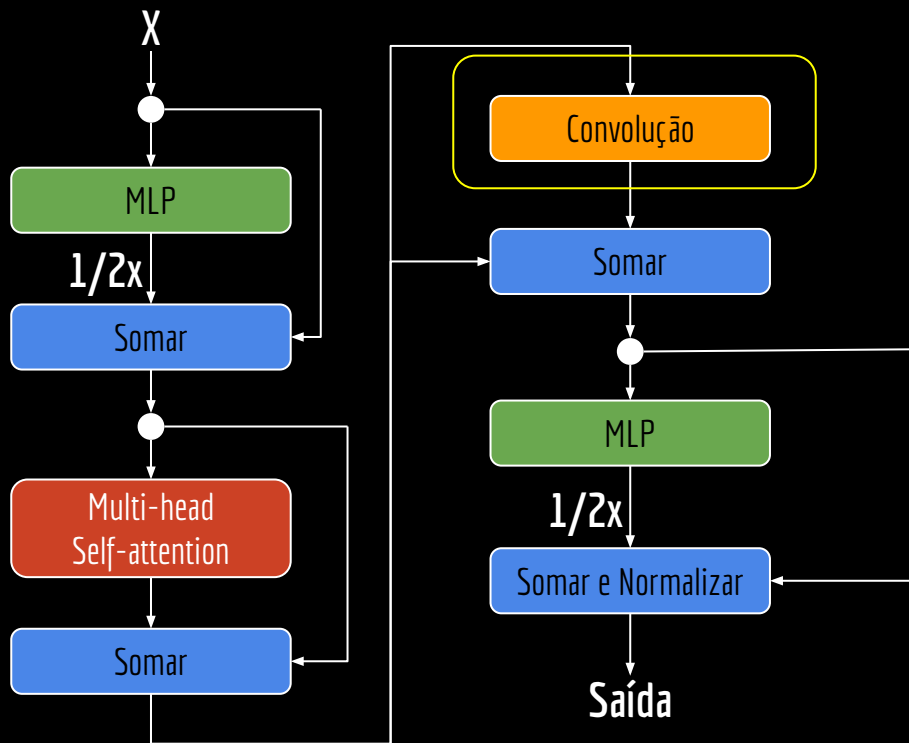
$$\mathbf{x}' = \tilde{\mathbf{x}} + MHSA(\tilde{\mathbf{x}})$$

$$\mathbf{x}'' = \mathbf{x}' + Conv(\mathbf{x}')$$

$$\mathbf{y} = \text{normalizar}(\mathbf{x}'' + \frac{1}{2}MLP(\mathbf{x}''))$$



Conformer



GULATI, Anmol et al. Conformer: Convolution-augmented transformer for speech recognition. 2020.

Camada convolucional aplica convolução de tamanho 1 (pointwise) para combinar os canais, e depois aplica uma convolução de tamanho 31 para capturar as correlações locais.



Fast Conformer

Em 2023, os Fast Conformers foram propostos.

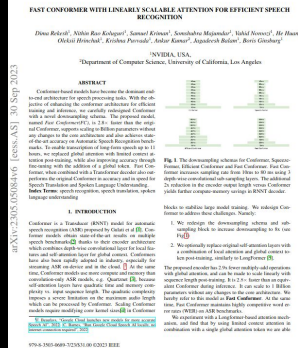
Mesma ideia básica do Conformer.

Troca do tokenizador.

Possibilitou fazer o downsampling de 4x para 8x.

Kernels de tamanho 9.

REKESH, Dima et al. Fast conformer with linearly scalable attention for efficient speech recognition. IEEE Automatic Speech Recognition and Understanding Workshop. 2023.



[illegible]

Outras modalidades

Dados que podemos projetar em um espaço (*embedding*) de tokens, e onde há uma dependência sequencial entre os dados, podem ser processados através de transformers.

O processamento de áudio foi um exemplo.

Outras modalidades

Dados que podemos projetar em um espaço (*embedding*) de tokens, e onde há uma dependência sequencial entre os dados, podem ser processados através de transformers.

O processamento de áudio foi um exemplo.

Outro exemplo: processamento de imagens.

Transformers para visão

Em visão, podemos utilizar cada pixel individual como um *token*.

Vai ser computacionalmente custoso, mas funciona.

Transformers para visão

Em visão, podemos utilizar cada pixel individual como um *token*.

Vai ser computacionalmente custoso, mas funciona.

Para reduzir o custo podemos, por exemplo:

- Utilizar patches de $N \times N$ das imagens, e projetar cada patch como um token, ou
- Utilizar filtros de convolução para reduzir a dimensionalidade das imagens originais.

Ideia do ViT: *Vision Transformer*

Transformers para visão

O embedding final gerado pelos transformers pode ser usado para, por exemplo, passar por um MLP para classificar a imagem.

Ou podemos montar um modelo baseado em auto regressão:

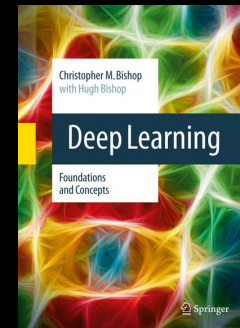
Dado o início de uma imagem, prever os próximos pixels dela (terminar de desenhar).

Multimodalidade

Ao transformar as entradas e saídas em tokens, é relativamente trivial unir múltiplos modais de informação em uma única rede.

Exemplos:

- Uma rede que recebe ao um prompt de texto, e uma imagem de base, e a modifica de acordo com o prompt.
- Uma rede que recebe um prompt, e gera uma imagem de saída (na forma de auto regressão, na próxima iteração, recebe o prompt, e o primeiro pixel já gerado, ...).



Licença

Esta obra está licenciada com uma Licença [Creative Commons Atribuição 4.0 Internacional](https://creativecommons.org/licenses/by/4.0/).