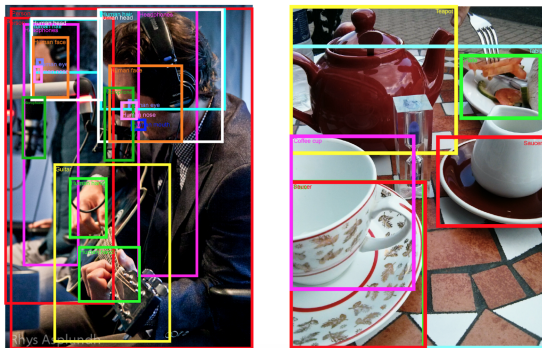


11-695: AI Engineering

Unsupervision

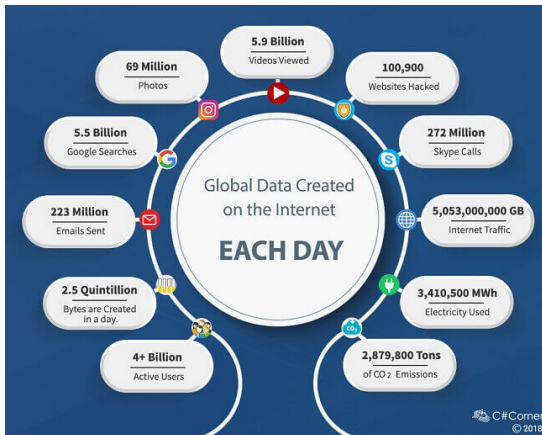
Spring 2019

- ① Data Issue
- ② Methodology
- ③ Auto Encoder
- ④ Data Encoding to Embedded Representation
- ⑤ Decoding from Embedded Space



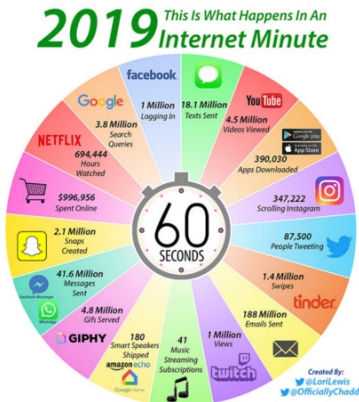
- Data is often carefully chosen,
- carefully cleaned and labelled,
- and usually *balanced*

But, The World's Data is Enormous



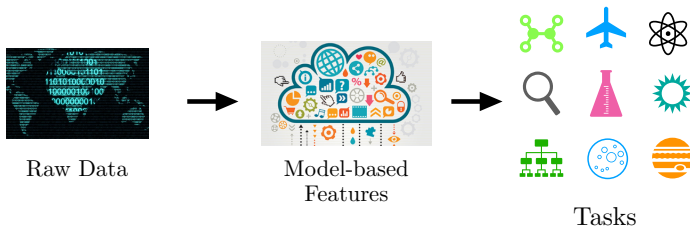
- Data is fast increased every second

Most Data is Unlabelled



- Also unstructured and unclean
- The cost to label them is prohibitively expensive.

- ① Data Issue
- ② Methodology
- ③ Auto Encoder
- ④ Data Encoding to Embedded Representation
- ⑤ Decoding from Embedded Space



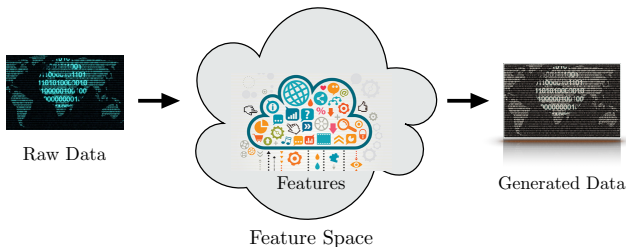
- How about we not only focus on the task, but also put a proper focus on Data itself
- Difference from:
 - Fully labelled data: Supervised
 - Partially labelled data: Semi-supervised
- We have no labels thus no supervision → Unsupervised.



- ## Feature Space

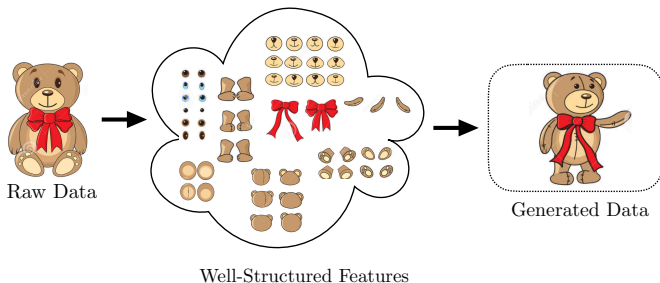


- Spring 2019 8 / 38



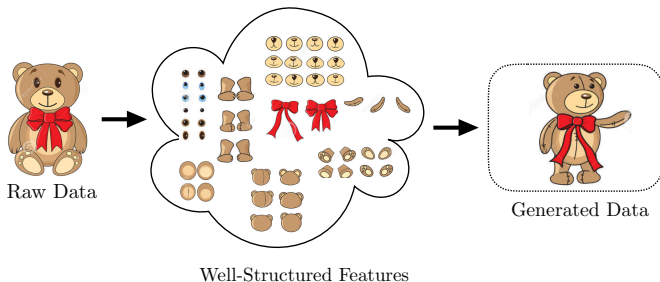
- Use each data sample itself as a label
 - Extract Features from x : $z = f_1(x)$
 - Reconstruct data $\hat{x}$: $\hat{x} = f_2(z)$
 - Minimize the distance $d(x, \hat{x})$, *a.k.a* reconstruction error.
- But, data itself are often highly-dimensional $\rightarrow$ (very) **hard!**
- Inject prior human knowledge/expertise about data, *e.g.*
 - Vision: use CNN as feature extractor
 - Text, sequential data: use RNN instead

Generate a bear to just classify a bear?



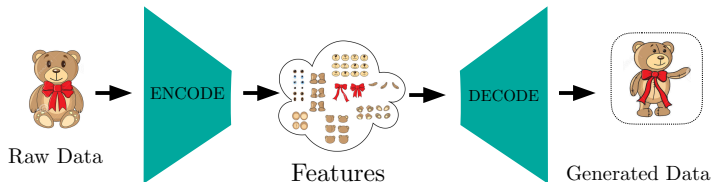
- No
 - But we don't have labels
 - In many cases, we need to actually generate data
 - Or need to have good data features for our tasks
- Remind: for supervised, we don't need to generate data
 - Features needn't be well-structured to predict labels

Generate a bear to just classify a bear?

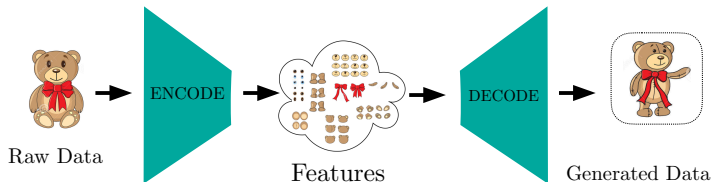


- Remind: for supervised, we don't need to generate data
 - Features needn't be well-structured to predict labels
- Comparison to Unsupervised in the classification/regression:
 - Supervised: we model $p(y|x)$
 - Unsupervised: $p(x, y)$

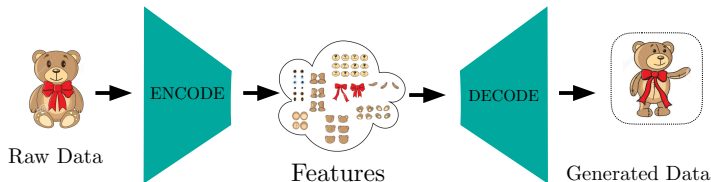
- ① Data Issue
- ② Methodology
- ③ Auto Encoder
- ④ Data Encoding to Embedded Representation
- ⑤ Decoding from Embedded Space



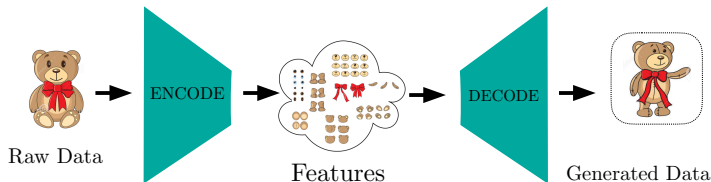
- Extract and Reconstruct *itself* $\rightarrow$ Auto-Encoder + Decoder
- What should the dimensionality of Features be?



- Extract and Reconstruct *itself* $\rightarrow$ Auto-Encoder + Decoder
- What should the dimensionality of Features be?
 - Smaller than Input! Why?
 - So encoding (feature extraction) process is a compression

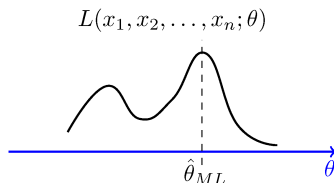
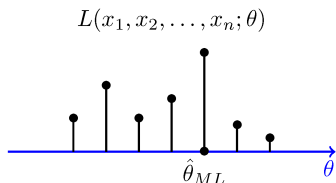


- Extract and Reconstruct *itself* $\rightarrow$ Auto-Encoder + Decoder
- What should the dimensionality of Features be?
 - Smaller than Input! Why?
 - So encoding (feature extraction) process is a compression
- Compression: we need to learn prominent features only
 - Yet should be enough to reconstruct data

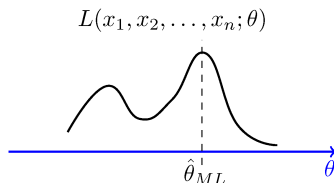
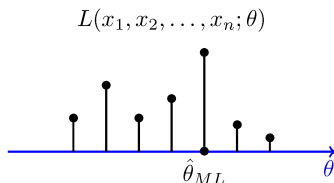


- Extract and Reconstruct *itself* $\rightarrow$ Auto-Encoder + Decoder
- What should the dimensionality of Features be?
 - Smaller than Input! Why?
 - So encoding (feature extraction) process is a compression
- Compression: we need to learn prominent features only
 - Yet should be enough to reconstruct data
 - How? Unless for trivial cases, we need params θ

How to *best* extract features from data?

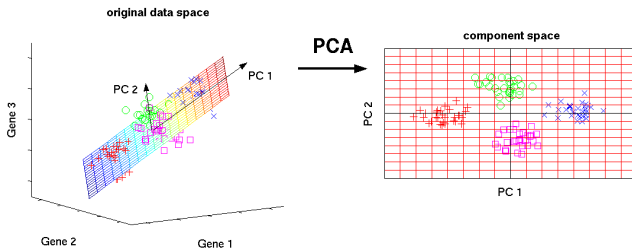


- input x , we need to find an estimator $\hat{\theta}$ best describing data
- Encode/Compress/Extract: $z = f_{\theta}(x)$
- Relation: Maximum Likelihood Estimation (MLE)
 - Likelihood: $\mathcal{L}(\theta; x) = p(x|\theta)$
 - Our job is to find: $\hat{\theta} = \operatorname{argmax}_{\theta} \mathcal{L}(\theta; x)$



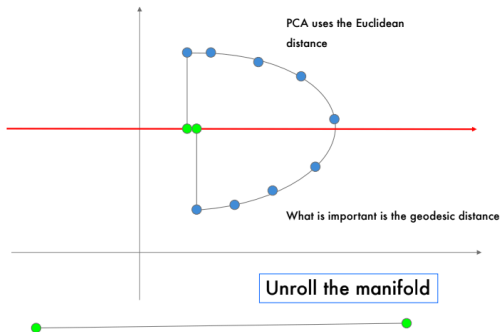
- Trivial case, *e.g.* estimate a quadratic function (convex)
 - Take the first derivative
 - Set to zero and find the estimator from the solution
- Most cases: we need to approximate. How?

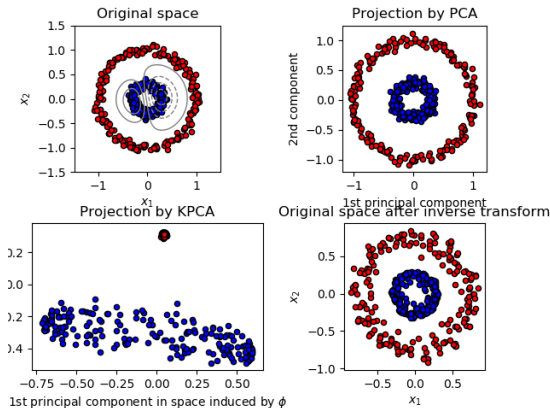
- ① Data Issue
- ② Methodology
- ③ Auto Encoder
- ④ Data Encoding to Embedded Representation
- ⑤ Decoding from Embedded Space



- Each Principal Component is a feature, and orthogonal to the rest
- Find a s.t. $x*a = \lambda a$ where λ is a scalar
 - One solution: perform SVD for x
 - Take only first few eigenvectors (hence approximation)

Principal Components Analysis (PCA)





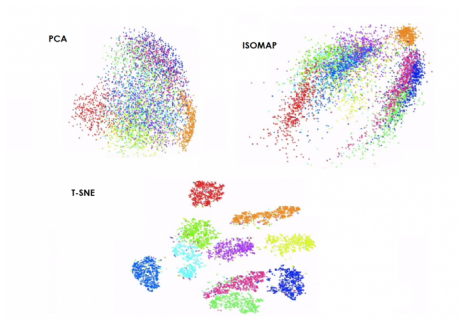
- Recall: why do we use a kernel?
- Still a PCA, and have to carefully choose the kernel

Decoded PCA (Generated Data)

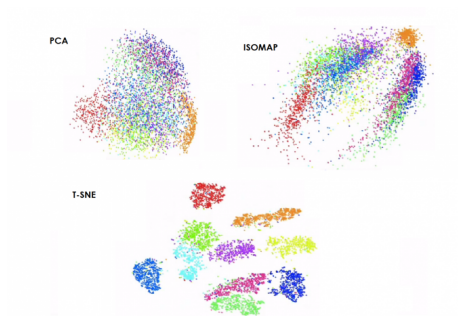


- Number of eigenvectors taken: 2000, 1000, 500, 100, 50, 10, 2, 1

Image credit: Gene Kogan



- Take into account relative distances between data points
- Is a non-linear method
- Both PCA and t-SNE suffer in high dimensional space.



- Take into account relative distances between data points
- Is a non-linear method
- Both PCA and t-SNE suffer in high dimensional space.
- *It's time for Deep NN to be used!*

- Original space:

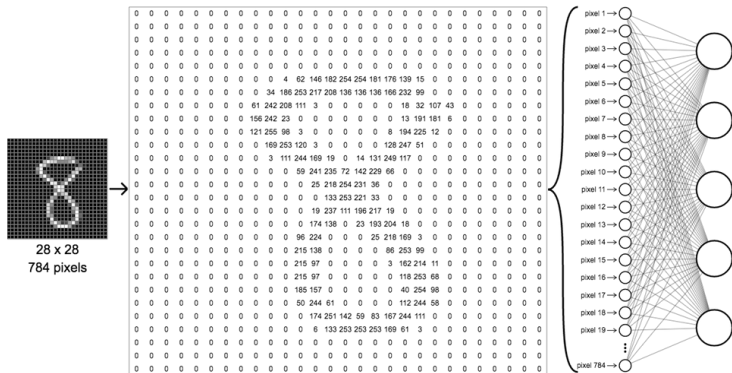
- one-sided distance: $p_{i|j} = \frac{\exp(-\|x_i - x_j\|^2 / 2\sigma_i^2)}{\sum_{k \neq i} \exp(-\|x_i - x_k\|^2 / 2\sigma_i^2)}$

- two-sided (gaussian centered density): $p_{ij} = \frac{p_{i|j} + p_{j|i}}{2}$

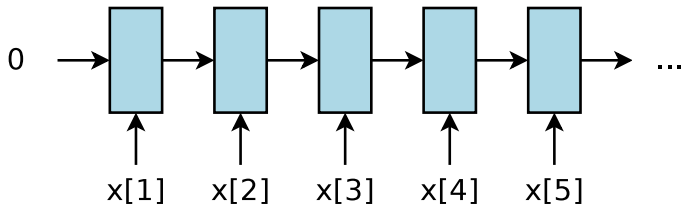
- Embedded space: similar distance: $q_{ij} = \frac{(1 + \|y_i - y_j\|^2)^{-1}}{\sum_{k \neq i} (1 + \|y_i - y_k\|^2)^{-1}}$

- Target: by minimizing the KL distance:

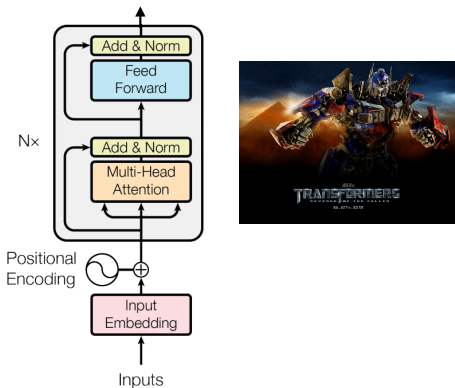
$$y_1, \dots, y_D = \operatorname{argmin} KL(P \| Q) = \operatorname{argmin} \sum_{i \neq j} p_{ij} \log \frac{p_{ij}}{q_{ij}}$$



- Caveat: need to flatten the input
- With data knowledge, we can do a better job



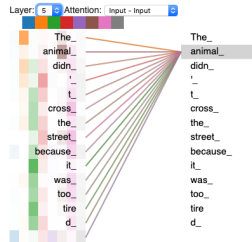
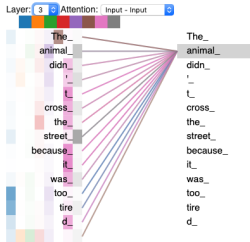
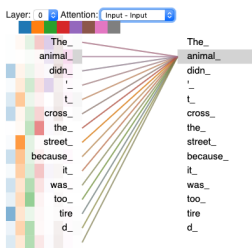
- Representation has sequential correlations



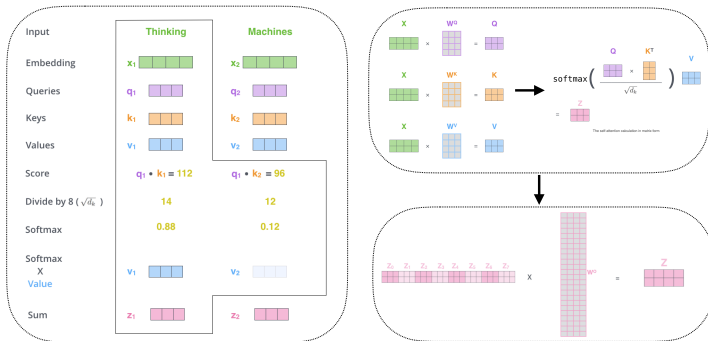
- We don't rely on RNN function anymore
- Instead, rely on N^* (multi-head attentions + MLP) blocks
- And more ...

Image credit: Ashish Vaswani *et al.*

Self-Attention (of *it* to *itself*)

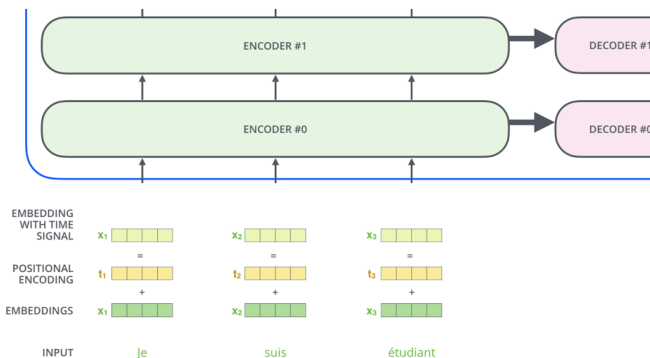


- Instead of encoding the sequence into a single hidden vector
- Now we calculate the relationship of each word vs. *every* another word
- Maybe a better idea for languages in which the order is very flexible



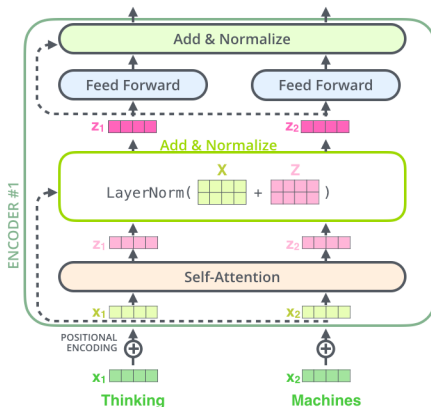
- Instead of using a single FC, use a triple (q, k, v) (query, key, value)
- Normalized with input length and with softmax
- Merge all head using concatenation and a FC with weight W^O

Multi-head Attn Able to Replace RNN?



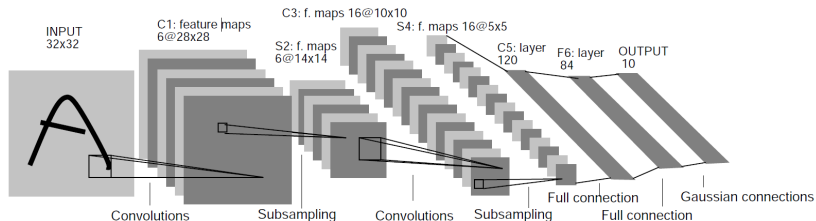
- Not yet, needs an alternative way to model sequence
- *Positional Encoding* is added to word embedding

Putting together: Transformer Encoder



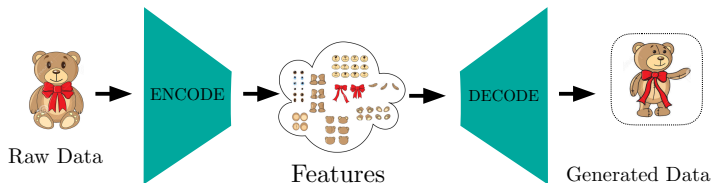
- Residual connection is used at self-attention and MLP
- Batch Norm is the last layer.
- Note: it also stacks multiple encoders (and decoders)

Image credit: Jay Alammar

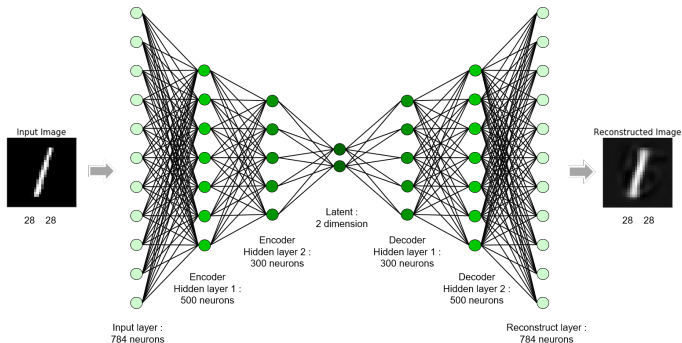


- Encoded representation has spatial correlations

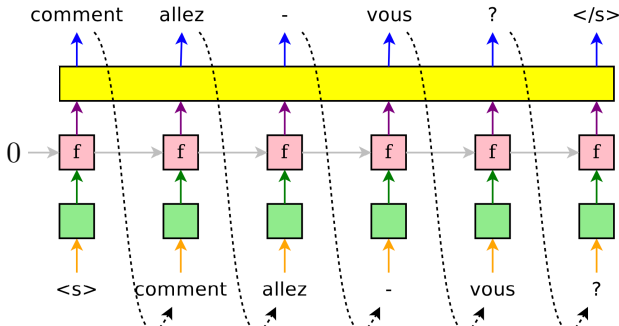
- ① Data Issue
- ② Methodology
- ③ Auto Encoder
- ④ Data Encoding to Embedded Representation
- ⑤ Decoding from Embedded Space



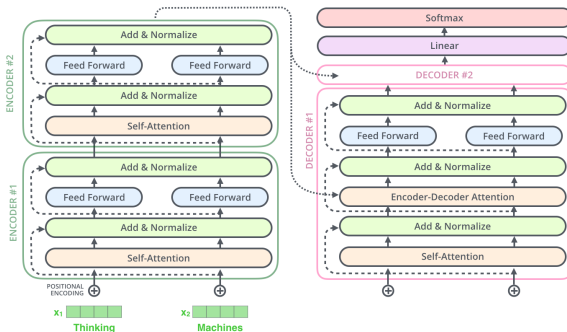
- Perform the inverse process of decoder
 - Encoder: compress from dimension N to D for $N > D$
 - Decoder: uncompress from D back to N



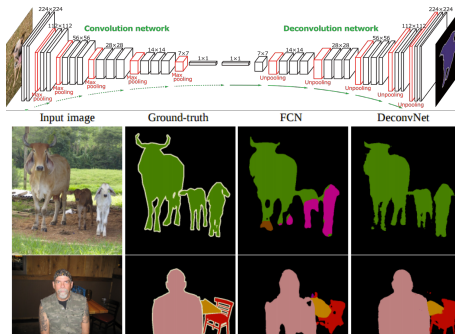
- Perform the inverse process of decoder
 - Encoder: MLP with decreasing dimensions
 - Decoder: MLP with increasing dimensions to the original
- Tutorial: [▶ MNIST AutoEncoder Tensorflow](#)



- Perform the inverse process of Encoder
 - Using previous output and previous hidden
 - It models $p(x_t | x_{<t})$ so sometimes it's called an *auto-regressive model*.
- Tutorial: [▶ RNN Text Generation with Eager Execution](#)

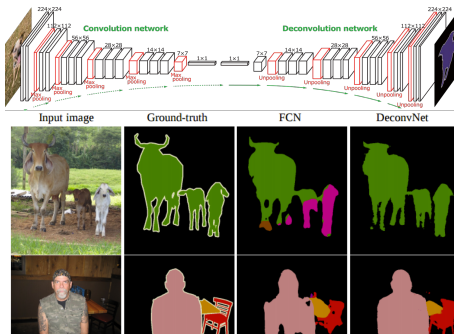


- Self-Attention is used, but only for the seen sequence to that point
- Then normal Attention is used as in Seq2Seq
- Decoding as usual to find a word in vocabulary.
- Tutorial: [▶ NMT-Transformer with pre-trained weights](#)



- Deconvolution is only good enough for such tasks as segmentation where detail is not the first priority
- Maybe we need to add some constraint into embedded space?
- Or maybe we should not use MLE-based method?

Image credit: Hyeonwoo Noh et al.



- Maybe we need to add some constraint into embedded space?
- Or maybe we should not use MLE-based method?
- *Next lectures will help provide some appealing solutions to those questions.*