

From Correlation to Grammar (Cagnetta)

How NN acquire hierarchical structure from data!

Q. Why does DL work?

□ Curse of dimensionality (poverty of the stimulus)

↓
difficult to sample enough

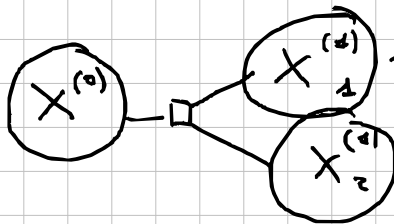
Intelligence becomes experimentally accessible.

□ Introduce Random Hierarchies

Hierarchical Compositionality: complex meaning from simple building blocks

→ Tree-Like Generative Model

Goal: a controlled generator of sentences / sequences



- Draw a root symbol
- Draw a children given the parent

CONTEXT FREE GRAMMARS

$$Q((a,b) | X^{(0)} = r) = Pr(r \rightarrow (a,b))$$

Finite, non recursive : all valid sentences have finite length



LAYER PROBABILISTIC CONTEXT FREE GRAMMAR

RHH :

$$X^{(0)} = a, b, c \dots$$

$$X_1^{(1)} = (a, a)$$

$$X_2^{(1)} = (a, b)$$

$$Q((a,b) | X^{(0)} = c) = ?$$

↓
PRODUCTION RULE

PARENTS

Ⓐ

Ⓑ

Ⓒ

SIBLING PAIRS

(a, a) (a, b) (a, c)

(b, a) (b, b) (b, c)

(c, a) (c, b) (c, c)

RULE BOOK

$a \rightarrow (a, a)$

$a \rightarrow (b, c)$

⋮

Assumption : Tree of depth L , random & independent rules at each layer

Data properties are available analytically and depend on L , vocab size v , m rules

sequence length $d = 2^L$
 # sequences $V^m (d-1)$

Latent variables lower the dimensionality
 $V^d \rightarrow V^{d/2} \rightarrow V^{d/4} \rightarrow \dots \rightarrow V$

Can the learner infer them from
 labelled examples?



$Y=r$

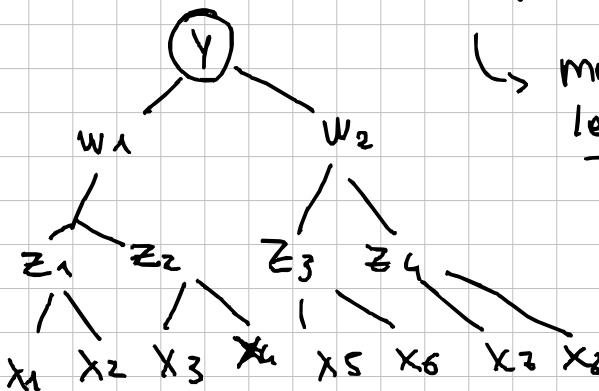
X_1

X_2

$$P(Y=r | X_{1:8}) = 1/V$$

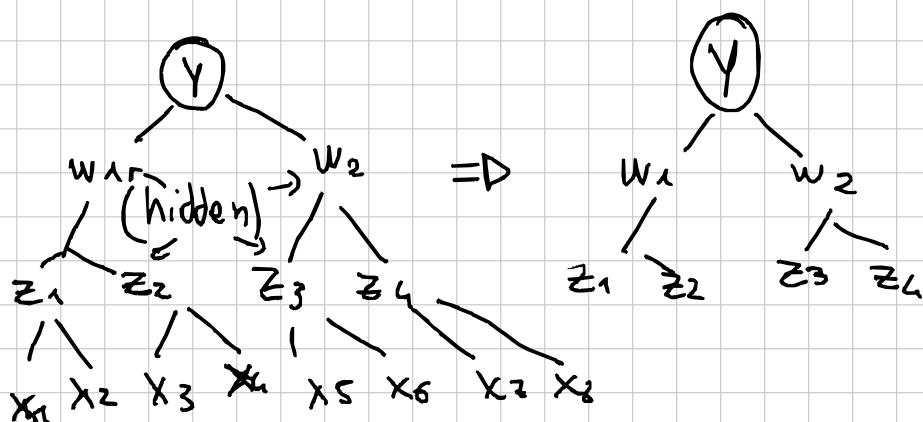
$$P(Y=r | X_{1:2}) = \frac{1}{V} + \epsilon$$

$\hookrightarrow$ mutual info root
leaves



$\hookrightarrow$ i can cluster groups of symbols
 to the latent variables

We can reduce



Training examples $\rightarrow$ low-order models

$\hookrightarrow$ hidden symbol models

sample complexity $\# \{ \text{data} \}$ s.t moments

$\gg$ sampling data

Method of moments

if (x_i, x_{i+1}) are generated by H

$$C_{ab}(y) = \Pr((x_i, x_{i+1}); Y=y) \\ = \sum \underbrace{Q(a, b | H=h)}_{v^2 - by - v} \underbrace{\Pr(H=h; Y=y)}_{v - by - v}$$

SAMPLE COMPLEXITY REDUCTION

$$|C_{ab}(y)|^2 \gg |C_{ab}(y) - \hat{C}_{ab}(y)|^2 \\ = O(1/\# \{ \text{data} \})$$

Test on NN

Root classification

Given depth to store latent variables

$\downarrow$
 $\exp(d)$

$$\text{Signal } |C_{ab}(z)|^2 \\ = O(1/(vm)^L)$$

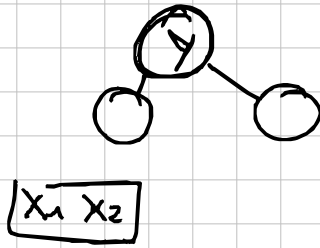
Sample complexity
 $\# \{ \text{data} \} \sim vm^L$

$\text{poly}(d)$
 $(vm)^L$

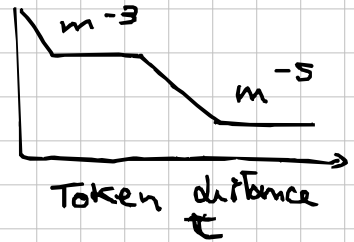
• VARYING TREE TOPOLOGY

|| Correlations reveal latent variables and local topology

Last-token prediction

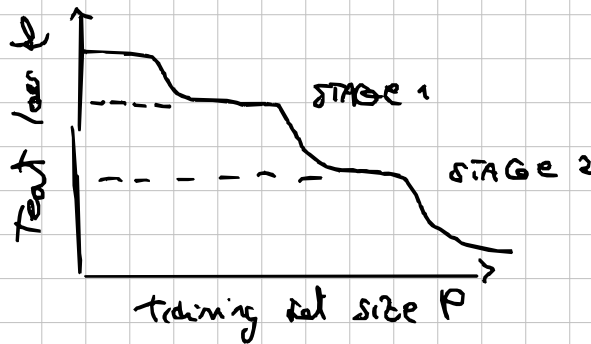


Signal decay with distance



→ Only bottleneck is DATA SET size

Empirical learning curve



[	All sequences	V^d
	Valid sequences	m^d
	sample size	$V m^L \sim d \cdot o(\ln(m))$

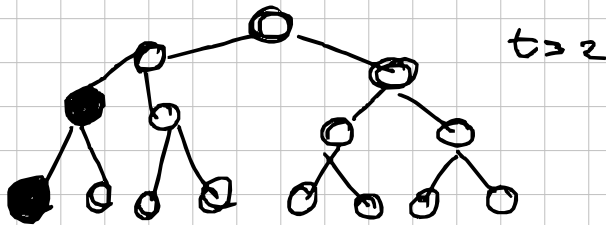
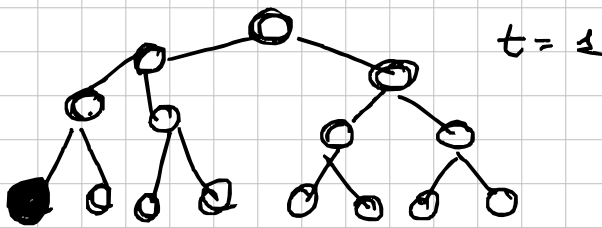
EXPONENTIAL CREATIVITY WITH POLYNOMIAL experience

What has The network learned?

Linear Representation Hypothesis

$$\phi_{k,t}(X) = A \cdot H + \epsilon$$

First-layer activations encode Z_1
 $S_H(\phi) =$ Share of ϕ 's variance due to H 's code



Beyond Synthetic Data

Neural scaling law

Loss

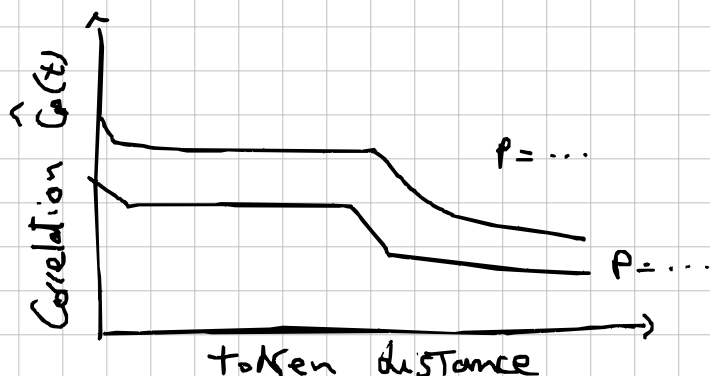
Compute

Dataset size

Parameters

Effective prediction time horizon

From the RHM: $\text{Finite \# of data} = P$



turns out $n^*(P) \sim P^{1/2p}$
 $L_{n^*}(P) \sim P^{-1/2p}$

NO FITTING PARAMS,
PURELY DATA STATISTICS

Persistence of latent variables

$$q_{e,b}(x) = A \cdot H + \epsilon$$

▲ Deeper layers code for deeper latent variables

▲ Deeper latent variables influence longer portions of the context