

**MIT 9.520/6.7910**  
**Statistical Learning Theory and Applications**

**Class 11: Sequential prediction as learning dynamical systems**

Lorenzo Rosasco

## So far

- ▶ Theory and algorithms for learning from i.i.d. data.
- ▶ What about non i.i.d. data?

## Sequential prediction

$$s_0, s_1, \dots, s_T \mapsto s_{T+1} ?$$

# Forecasting Boston temperature

79° 72° 63° 53° 73°  $\mapsto$  ...?

## Forecasting Boston temperature

$$\underbrace{79^\circ}_{s_0} \quad \underbrace{72^\circ}_{s_1} \quad \underbrace{63^\circ}_{s_2} \quad \underbrace{53^\circ}_{s_3} \quad \underbrace{73^\circ}_{s_4} \quad \mapsto \quad \underbrace{\dots}_{s_5}?$$

## Next token prediction

The cat sat on the  $\mapsto$  ...?

## Next token prediction

$\underbrace{\text{The}}_{s_0} \underbrace{\text{cat}}_{s_1} \underbrace{\text{sat}}_{s_2} \underbrace{\text{on}}_{s_3} \underbrace{\text{the}}_{s_4} \mapsto \underbrace{\dots}_{s_5}?$

## Non i.i.d. data

$$s_0, s_1, \dots, s_T \mapsto s_{T+1} ?$$

The i.i.d. assumption breaks down.



# Outline

Data models

Non i.i.d. data

Dynamic learning problems

Autoregressive approaches

Hidden state models

## Discrete time dynamical systems

$$\begin{cases} s_0 \in \mathcal{S}, \\ s_{t+1} = f(s_t). \end{cases}$$

- ▶  $\mathcal{S}$  state space (tokens).
- ▶  $f : \mathcal{S} \rightarrow \mathcal{S}$  evolution map.

## State space

$\mathcal{S}$

- ▶  $\mathcal{S} = \mathbb{R}^d$  continuous state space.
- ▶  $\mathcal{S} = \{w_1, \dots, w_N\}$  finite state space. In this case  $f$  is a "*table*".

## Discrete vs continuous time

$$s_{t+1} = f(s_t)$$

- ▶ Discrete time evolution can arise from discretization of continuous time evolutions

$$\dot{s}(t) = f(s(t)).$$

- ▶ But many problems are inherently discrete in time!

## Non-autonomous systems

$$s_{t+1} = f(s_t)$$

- ▶ The above evolution is called autonomous.
- ▶ For non-autonomous systems, the evolution depends on time

$$s_{t+1} = f(s_t, t).$$

## Higher order systems

$$s_{t+1} = f(s_t)$$

► The above system is first order.

► A system of order  $p$  is

$$s_{t+1} = f(s_t, s_{t-1}, \dots, s_{t-p+1}).$$

In next token prediction, the order  $p$  is called the **context**.

## Stochastic dynamical systems

$$s_{t+1} = f(s_t, \omega_t)$$

- ▶ The variables  $\omega_t$  are i.i.d. random variables, independent of  $s_0$ .
- ▶ The evolution is stochastic.
- ▶ It can be seen as a perturbation of a deterministic evolution, or as an inherently stochastic one.

The above system is equivalent to a Markov process.

## Markov processes

$$\mathbb{P}(s_{t+1} \mid s_t, s_{t-1}, \dots, s_0) = \mathbb{P}(s_{t+1} \mid s_t)$$

- The next state depends only on the current one.



## Stochastic dynamical systems & Markov process

$$s_{t+1} = f(s_t, \omega_t) \iff \mathbb{P}(s_{t+1} \mid s_t, s_{t-1}, \dots, s_0) = \mathbb{P}(s_{t+1} \mid s_t)$$

► Indeed,

$$s_{t+1} = f(s_t, \omega_t) \Rightarrow \mathbb{P}(s_{t+1} \mid s_t) = \mathbb{P}_{\omega_t}(f(s_t, \omega_t))$$

► Conversely,

$$\mathbb{P}(s_{t+1} \mid s_t) = \mathbb{P}(s_{t+1} \mid s_t, \dots, s_0) \Rightarrow s_{t+1} = f(s_t, \omega_t)$$

## Transition kernels

Let  $\mathcal{S} = \mathbb{R}^d$ .

Transition kernel

$$P : \mathcal{S} \times \mathcal{B}(\mathcal{S}) \rightarrow [0, 1].$$

For each  $s \in \mathcal{S}$ ,  $P(s, \cdot)$  is a probability distribution on  $\mathcal{S}$ .

## Markov processes and transition kernels

$$P(s, A)$$

- ▶ Each Markov process has an associated kernel

$$P(s, A) = \mathbb{P}(s_{t+1} \in A \mid s_t = s).$$

- ▶ Conversely, given  $P$  and an initial distribution  $\mu_0$ , we can define a Markov process by

$$\mathbb{P}(s_0 \in A) = \mu_0(A), \quad \mathbb{P}(s_{t+1} \in A \mid s_t = s) = P(s, A).$$

## Transition matrices

Let  $\mathcal{S} = \{w_1, \dots, w_N\}$ .

Transition matrix  $P \in \mathbb{R}^{N \times N}$  such that  $P_{ij} \in [0, 1]$ .

For each  $i \in \{1, \dots, N\}$ ,  $P(i, \cdot)$  is a probability distribution on  $\mathcal{S}$ , that is

$$\sum_{j=1}^N P_{ij} = 1.$$

## Markov processes and transition matrices

$$P_{ij}$$

- ▶ Each Markov process with a finite state space has an associated transition matrix

$$P_{ij} = \mathbb{P}(s_{t+1} = w_j \mid s_t = w_i).$$

- ▶ Conversely, given  $P$  and an initial distribution  $\mu_0$ , we can define a Markov process by

$$\mathbb{P}(s_0 = w_i) = \mu_0(w_i), \quad \mathbb{P}(s_{t+1} = w_j \mid s_t = w_i) = P_{ij}.$$

## Invariant distribution

Given a Markov process,  $\pi$  is an *invariant distribution* if <sup>1</sup>

$$\pi P = \pi.$$

There might be no, one, or multiple invariant distributions.

---

<sup>1</sup>For  $\mathcal{S} = \mathbb{R}^d$ , this means  $\int P(s, A) d\pi(s) = \pi(A)$  for all measurable  $A$ .

## Ergodic Markov processes

Let  $\mu_t$  denote the distribution of  $s_t$ . The process is *ergodic* if for any initial distribution  $\mu_0$ ,

$$\mu_t \rightarrow \pi \quad \text{as } t \rightarrow \infty.$$

Then  $\pi$  is unique.

## Mixing and burn-in time

The *mixing time* quantifies how fast the distribution  $\mu_t$  converges to  $\pi$ .

The *burn-in time* refers to the initial transient phase before the process is close to  $\pi$ .



## Dynamical systems & Markov processes

$$s_{t+1} = f(s_t, \omega_t) \iff \mathbb{P}(s_{t+1} \mid s_t) = \mathbb{P}(s_{t+1} \mid s_t, s_{t-1}, \dots, s_0)$$

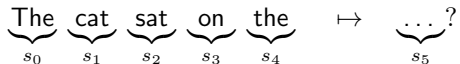
- ▶ Autonomous systems: *time-homogeneous* Markov processes.

- ▶ Higher-order systems: *p-th order* Markov processes

$$\mathbb{P}(s_{t+1} \mid s_t, s_{t-1}, \dots, s_0) = \mathbb{P}(s_{t+1} \mid s_t, \dots, s_{t-p+1}).$$

- ▶ Continuous-time stochastic systems: *stochastic differential equations*, or equivalently, *continuous-time Markov processes*.

## Markov processes and language models



This Markov process perspective on next-token prediction dates back to Shannon<sup>2</sup>.

In language models, it relates to  $p$ -grams ( $n$ -grams).

---

<sup>2</sup>Claude Shannon's 1948 paper "A Mathematical Theory of Communication".

# Outline

Data models

Non i.i.d. data

Dynamic learning problems

Autoregressive approaches

Hidden state models

## Data collection

Data can be seen as sample trajectories.

Single trajectory:

$$(s_0, s_1, \dots, s_T).$$

Multiple trajectories:

$$(s_0^{(1)}, s_1^{(1)}, \dots, s_T^{(1)})$$

$$(s_0^{(2)}, s_1^{(2)}, \dots, s_T^{(2)})$$

$$\vdots$$

$$(s_0^{(N)}, s_1^{(N)}, \dots, s_T^{(N)}).$$

## Dependent data

$$s_0, s_1 \dots, s_T$$

- ▶ We cannot expect independent data.
- ▶ It can be shown that the *mixing time* quantifies dependence.
- ▶ Independence can be reinstated by collecting multiple trajectories, e.g. for  $p = 1$ :

$$(s_0^{(1)}, s_1^{(1)}), (s_0^{(2)}, s_1^{(2)}), \dots, (s_0^{(N)}, s_1^{(N)}).$$

(Not very practical.)

## Non identically distributed data

$$s_t \sim \mu_t$$

- Different distributions at different times.

- After the burn-in time  $t_b$ ,

$$\mu_t \approx \pi \quad \text{for } t \geq t_b.$$

- Stationary processes often assumed in time series analysis, i.e.

$$(s_t, s_{t+1}, \dots) \stackrel{d}{=} (s_{t+\tau}, s_{t+\tau+1}, \dots), \quad \forall \tau \geq 1.$$

# Outline

Data models

Non i.i.d. data

Dynamic learning problems

Autoregressive approaches

Hidden state models

## Different problems

Dynamical systems are studied by different communities with different focus.

- **Prediction / forecasting:** find  $\hat{f}$  s.t.

$$\hat{f}(s_t) \approx s_{t+1}.$$

- **Estimation:** find  $\hat{p}$  s.t.

$$\hat{p}(s_{t+1} \mid s_t) \approx p(s_{t+1} \mid s_t).$$

- **Long-time behaviour:** stability, chaos, attractors.



## Loss functions

$$\hat{f}(s_t) \approx s_{t+1}.$$

How shall we measure accuracy of  $\hat{f} : \mathcal{S} \rightarrow \mathcal{S}$ ?

Pick a loss:

► If  $\mathcal{S} = \mathbb{R}^d$ :

$$\ell(s, \hat{f}(s)) = \|s - \hat{f}(s)\|^2.$$

► If  $\mathcal{S} = \{w_1, \dots, w_N\}$ :

$$\ell(s', \hat{f}(s)) = \mathbf{1}_{\{s' \neq \hat{f}(s)\}}.$$

## Empirical risk

$$\hat{L}(f) = \frac{1}{T} \sum_{t=0}^{T-1} \ell(s_{t+1}, f(s_t))$$

## How good is a solution?

$$\hat{f}$$

- **Practice: test error**<sup>3</sup>

$$\hat{L}_{\text{Test}}(\hat{f}) = \frac{2}{T} \sum_{t=T/2}^{T-1} \ell(s_{t+1}, \hat{f}(s_t))$$

- **Excess risk** (with ergodicity)

$$L(f) = \mathbb{E}_{s_t \sim \pi}[\ell(s_{t+1}, f(s_t))] = \int \ell(s', f(s)) P(s, ds') d\pi(s).$$

This might not be the natural definition; other error metrics can be considered, e.g. the regret.

---

<sup>3</sup>Use  $T/2$  points for  $\hat{f}$ .

## Target functions

$$f_*(s) = \arg \min_{\bar{s} \in \mathcal{S}} \int \ell(s', \bar{s}) P(s, ds')$$

Analog of the inner risk.

- For  $\mathcal{S} = \mathbb{R}^d$  and  $\ell(s, \hat{f}(s)) = \|s - \hat{f}(s)\|^2$ :

$$f_*(s) = \mathbb{E}[s_{t+1} \mid s_t = s].$$

- For  $\mathcal{S} = \{1, \dots, N\}$  and  $\ell(s, \hat{f}(s)) = \mathbf{1}_{\{s \neq \hat{f}(s)\}}$ :

$$f_*(i) = \arg \max_{j=1, \dots, N} P_{ij}.$$

## Remarks

- ▶ Next, we focus on  $\mathcal{S} = \mathbb{R}^d$ .
- ▶ The ideas can be extended to  $\mathcal{S} = \{w_1, \dots, w_N\}$  using ideas from multiclass learning.
- ▶ Also we could consider predicting  $s_{t+\tau}$  rather than just  $s_{t+1}$ , e.g. by

$$\hat{f}^\tau(s_t) \approx s_{t+\tau}.$$

# Outline

Data models

Non i.i.d. data

Dynamic learning problems

**Autoregressive approaches**

Hidden state models

## Basic idea

$$x_t = s_t \quad y_t = s_{t+1}$$

This is a vector valued regression problem  $y \in \mathbb{R}^d$ .

Train your favorite model  $f_\theta$  on

$$(x_0, y_0), \dots, (x_{T-1}, y_{T-1}).$$

## Adding context

$$x_t = (s_t, \dots, s_{t-p+1}) \quad y_t = s_{t+1}.$$

Still a regression problem, but the input is  $x \in \mathcal{S}^p$ , that is  $x \in \mathbb{R}^{d \times p}$ .

Each  $x$  is a collection of states (tokens).



## Linear models aka AR(p)

For  $d = 1$ , that is  $\mathcal{S} = \mathbb{R}$ ,

$$x = (s_1, \dots, s_p) \in \mathbb{R}^p$$

Consider,

$$f_{\theta}(x) = \sum_{j=1}^p a_j s_j, \quad \theta = (a_1, \dots, a_p) \in \mathbb{R}^p.$$

## Linear models aka VAR(p)

For  $d > 1$ , that is  $\mathcal{S} = \mathbb{R}^d$ ,

$$x = (s_1, \dots, s_p), \quad s_j \in \mathbb{R}^d.$$

Consider,

$$f_{\theta}(x) = \sum_{j=1}^p A_j s_j, \quad \theta = (A_1, \dots, A_p), \quad A_j \in \mathbb{R}^{d \times d}.$$

## Beyond linear models aka NARX

$$f_{\theta}(x) = \sum_{j=1}^p A_j s_j, \quad x = (s_1, \dots, s_p), \quad s_j \in \mathbb{R}^d.$$

► Feature maps

$$f_{\theta}(x) = \sum_{j=1}^p A_j \Phi(s_j), \quad A_j \in \mathbb{R}^{d \times m}, \quad \Phi(s_j) \in \mathbb{R}^m, \quad \theta = (A_1, \dots, A_p).$$

► Neural networks, as above, but

$$\Phi(s_j) = \sigma(W_j s_j + b_j), \quad \theta = (A_1, \dots, A_p, W_1, \dots, W_p, b_1, \dots, b_p).$$

# Transformers

$$f_{\theta}(x) = \text{MLP}(A(s_1, \dots, s_p)), \quad x = (s_1, \dots, s_p).$$

- ▶ MLP is a fully connected neural network, as before.
- ▶  $A$  is the attention layer.

## Attention

$$A(x) = (A_1(x), \dots, A_p(x)), \quad x = (s_1, \dots, s_p).$$

Where  $A_j(x)$

$$A_j(x) = \sum_{i=1}^p \alpha_{ji}(x) s_i, \quad \alpha_{ji}(x) = \frac{\exp(s_j^\top s_i)}{\sum_{k=1}^p \exp(s_j^\top s_k)}.$$

Each token is represented as a weighted combination of all tokens.

## Attention and embeddings

The inner product may not capture tokens similarity.

$$A_j(x) = \sum_{i=1}^p \alpha_{ji}(x) V s_i, \quad \alpha_{ji}(x) = \frac{\exp((Q s_j)^\top (K s_i))}{\sum_{k=1}^p \exp((Q s_j)^\top (K s_k))}.$$

- ▶ Tokens are embedded into so-called queries  $Q s_j$ , keys  $K s_i$ , and values  $V s_i$ .
- ▶  $Q$ ,  $K$ , and  $V$  are parameters to be learned.

# Outline

Data models

Non i.i.d. data

Dynamic learning problems

Autoregressive approaches

Hidden state models

## Non-autonomous systems

$$s_{t+1} = f(s_t, t)$$

- ▶ Can be approximated using **high**-order autonomous systems
- ▶ Next, we describe models based on hidden states.



## Hidden state models

$$\begin{cases} h_{t+1} = g(h_t, s_t) \\ s_{t+1} = q(h_{t+1}) \end{cases}$$

- ▶ The evolution of the pair  $(s_t, h_t)$  is autonomous.
- ▶ The evolution of  $s_t$  alone depends on the entire past (non-autonomous).
- ▶ Also called state space models (SSM).

## Linear state space models

$$\begin{cases} h_{t+1} = Ah_t + Bs_t, \\ s_{t+1} = Ch_{t+1}. \end{cases}$$

- ▶ Evolution maps are linear.
- ▶ Related to Kalman filtering, where noise is modeled explicitly.<sup>4</sup>

---

<sup>4</sup>Note that here we are only interested in predicting the observable state.

## Recurrent neural networks (RNN)

$$\begin{cases} h_{t+1} = \sigma(Ah_t + Bs_t) \\ s_{t+1} = Ch_{t+1} \end{cases}$$

- ▶ Evolution maps are nonlinear.
- ▶  $\sigma$  is an activation function.

## Recurrent neural networks (RNN) (cont.)

$$\begin{cases} h_{t+1} = \sigma(Ah_t + Bs_t) \\ s_{t+1} = Ch_{t+1} \end{cases}$$

- ▶ Unrolling: RNNs are DNNs with  $\# \text{ layers} = \# \text{time steps}$ , and same weights in each layer.
- ▶ Optimized through so-called backpropagation through time (BPTT).
- ▶ Vanishing gradients lead to variants, e.g. LSTM and GRU, etc.

## Summary

- ▶ A dynamical system tour of sequential prediction.
- ▶ Focus on a simple next token prediction task.
- ▶ Basic algorithmic principles.

## Forecasting Boston temperature

$$\underbrace{79^\circ}_{s_1} \underbrace{72^\circ}_{s_2} \underbrace{63^\circ}_{s_3} \underbrace{53^\circ}_{s_4} \underbrace{73^\circ}_{s_5} \mapsto \underbrace{55^\circ}_{s_6}$$

Time to go to Italy.