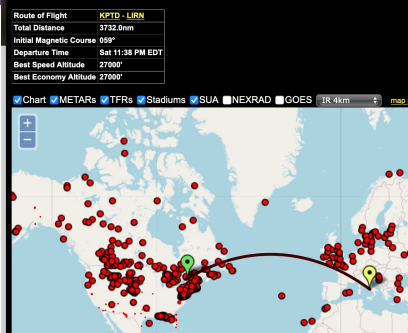




Clarkson
UNIVERSITY



Erik Boltt



C^3S^2

Clarkson Center for Complex Systems Science

-bolitem@clarkson.edu,
-<http://www.clarkson.edu/~bolitem>

How Neural Networks Work: Unraveling the Mystery of Random Projection Networks For Functions and Chaotic Dynamical Systems

Applied and
Computational
Measurable
Dynamics



Erik M. Bollt
Naratip Santitissadeekorn
siam
Mathematical Modeling and Computation

-bolltem@clarkson.edu,
-<http://www.clarkson.edu/~bolltem>

 **mathematics**



Article

Randomized Projection Learning Method for Dynamic Mode Decomposition

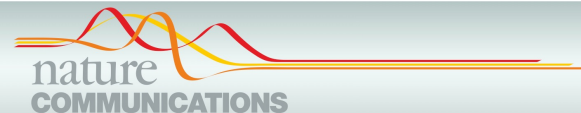
Sudam Surasinghe ^{1,*}, [†], [‡], [§] and Erik M. Bollt ^{2,‡}

Chaos

On explaining the surprising success of
reservoir computing forecaster of chaos?
The universal machine learning dynamical
system with contrast to VAR and DMD 

Cite as: Chaos 31, 013108 (2021); <https://doi.org/10.1063/5.0024890>

Submitted: 16 August 2020, Accepted: 30 November 2020, Published Online: 04 January 2021



ARTICLE

<https://doi.org/10.1038/s41467-021-25801-2>

OPEN

Next generation reservoir computing

Daniel J. Gauthier ^{1,2}, Erik Bollt^{3,4}, Aaron Griffith ¹ & Wendson A. S. Barbosa ¹

So much great work about how well ANN works - WOW they sure do “work” we can all agree.

When I first saw “**multi-layer perceptrons**” over 25 years ago - as a student –
I thought - wow that’s dumb. So much over fitting. So much computation – surely there is a better way.

Somehow along the way it emerged as a leader

Great Salesmanship

– machine learning, neural nets, deep learning, AI, so many variants...

OK OK – let me in – **I want to understand why it works**

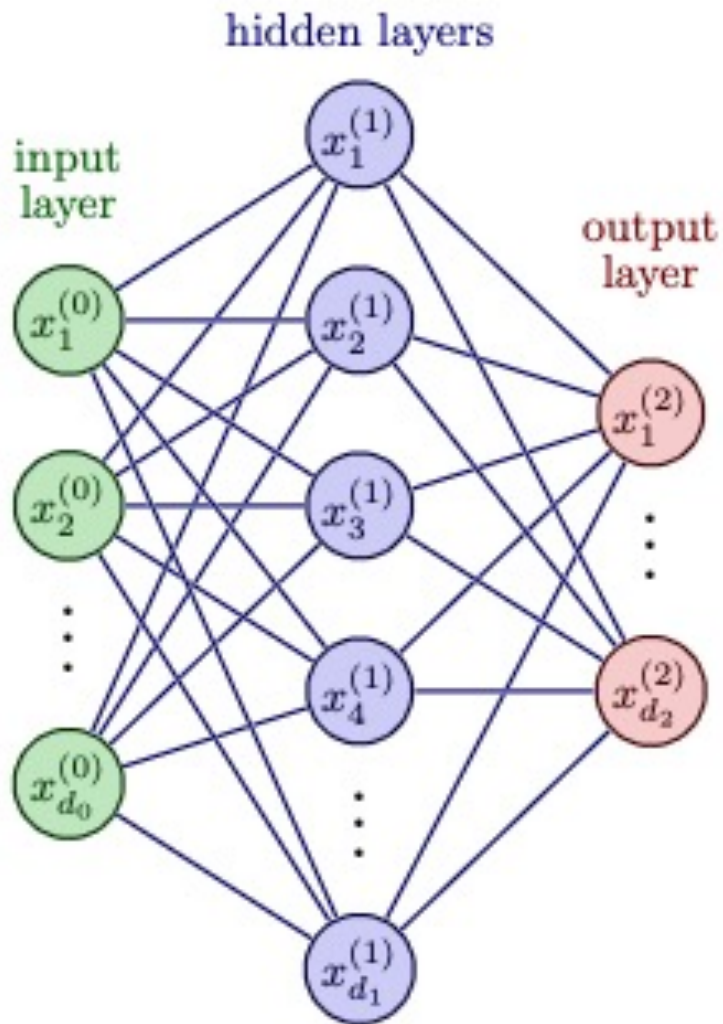
Random ANN variants, RC, and **here ELM**.

All the great things you can do with RC, ELM, Feedforward DNN, etc.

Here focused on how/why does it work – with random parameters - it still works??!

Extreme Learning Machines – ELM

ELM-1



For supervised learning tasks like –

- classification
- regression
- incl forecasting

Here's the crazy part – all but output layer has random parameters

$$\mathcal{D} = \{(X_i, Y_i)\}_{i=1}^N \subset \mathbb{R}^{d_0} \times \mathbb{R}^{d_2} = D.$$

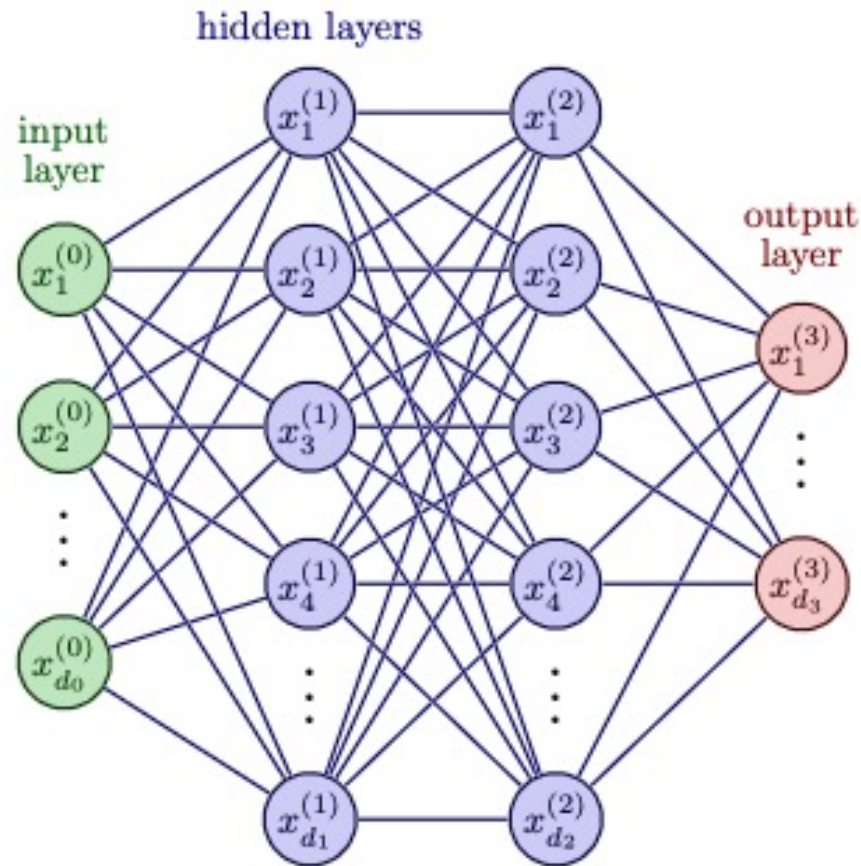
$$f : \mathbb{R}^{d_0} \rightarrow \mathbb{R}^{d_2}$$

$$X^{(1)} = [x_1^{(1)}; x_2^{(1)}; \dots; x_{d_1}^{(1)}] \in \mathbb{R}^{d_1}.$$

$$X^{(2)} = [x_1^{(2)}; x_2^{(2)}; \dots; x_{d_2}^{(2)}] \in \mathbb{R}^{d_2}.$$

I want to explain why it works anyway.

Deep Learning



How to train? Lots of parameters, massive nonlinear optimization problem? Massively parallel computers. Stochastic gradient descent, etc – *not what we are studying here*.

What's old is new again

How Neural Nets Work
 Alan Lapedes
 Robert Farber
 Theoretical Division
 Los Alamos National Laboratory
 Los Alamos, NM 87545

© American Institute of Physics 1988

$$\dot{x} = \frac{ax(t-\tau)}{1+x^{10}(t-\tau)} - bx(t) \quad \text{Glass-Mackey}$$

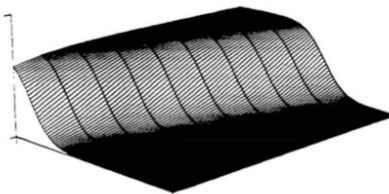
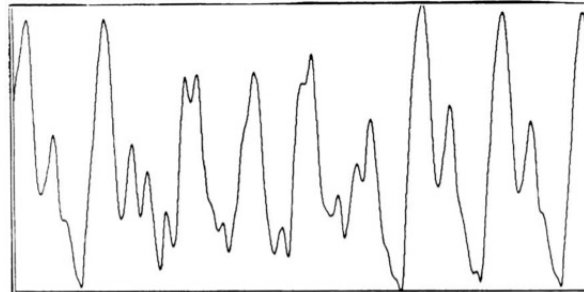


Figure 4. A sigmoidal surface.

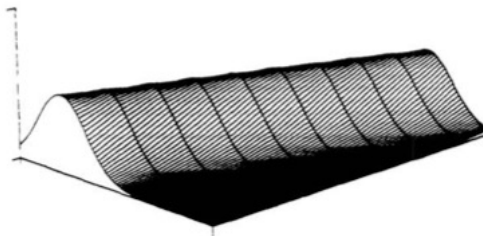


Figure 5. A ridge.

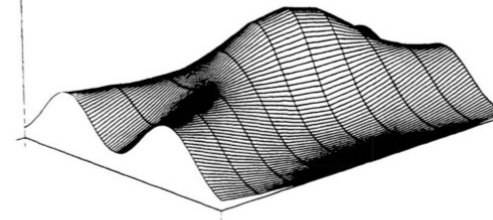


Figure 6. A pseudo-bump.

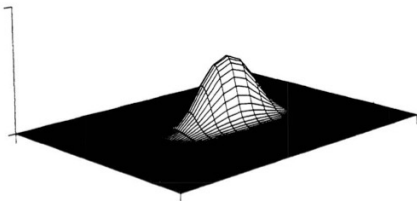


Figure 7. A bump.

Main idea:

You can make bump functions and with bump functions
 you can make general functions including a flow.

Thanks Constantinos

Random Projection Neural Networks

- Schmidt et al. (1992): fixing the weights between the input and the hidden layer at random values, and by solving a linear problem for the output weights the approximation accuracy is equivalent to that obtained with back-propagation.
- Random Vector Functional-Link Networks (RVFLNs) were addressed in Pao et al. (1992) in which the input layer is directly connected also to the output layer, the internal weights are chosen randomly in $[-1, 1]$: the output weights are estimated in one step by solving a system of linear equations.
- Igelnik et al. (1995) proved that RVFLNs are universal approximators for continuous functions on bounded finite-dimensional sets.

A Reservoir Computer, RC

Is a kind of neural network - for forecasting dynamical systems
but most of the (millions of) parameters are chosen randomly.

Clearly its cheap

Surprisingly - it actually works!

And surprisingly, it works really well.

(This talk is not about neat things you can do with RC)

(This talk is about how/why/bridge-equivalent to something else)

Last year's talk...

ESN-Jaeger 2001, Jaeger-Haas, 2004.

Reservoir computing – a special case of RNN
spec case ANN, Jaeger-Hass 2004, ESN-Jaeger 2001.

$$\{\mathbf{x}_i\}_{i=1}^N \subset \mathbb{R}^{d_x}$$

$$\mathbf{u}_i = \mathbf{W}^{in} \mathbf{x}_i,$$

$$\mathbf{r}_{i+1} = (1 - \alpha) \mathbf{r}_i + \alpha q(\mathbf{A} \mathbf{r}_i + \mathbf{u}_i + \mathbf{b}),$$

$$\mathbf{y}_{i+1} = \mathbf{W}^{out} \mathbf{r}_{i+1}.$$

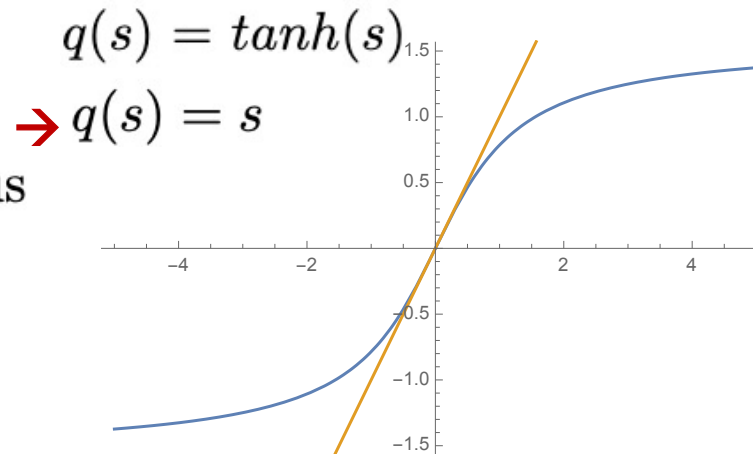
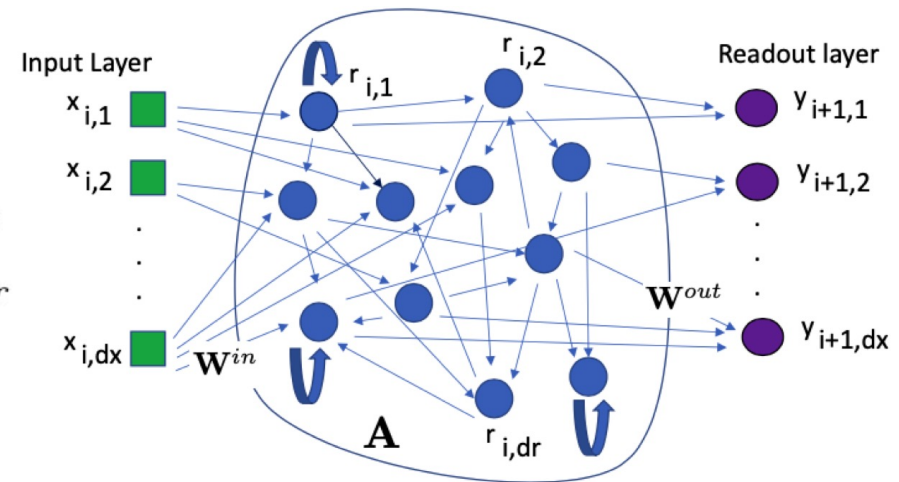
$$\mathbf{W}_{i,j}^{in} \sim U(0, \gamma) \quad d_r \times d_x \text{ read in matrix}$$

$$\mathbf{A}_{i,j} \sim U(-\beta, \beta), \text{ with } \beta \text{ to scale the spectral radius}$$

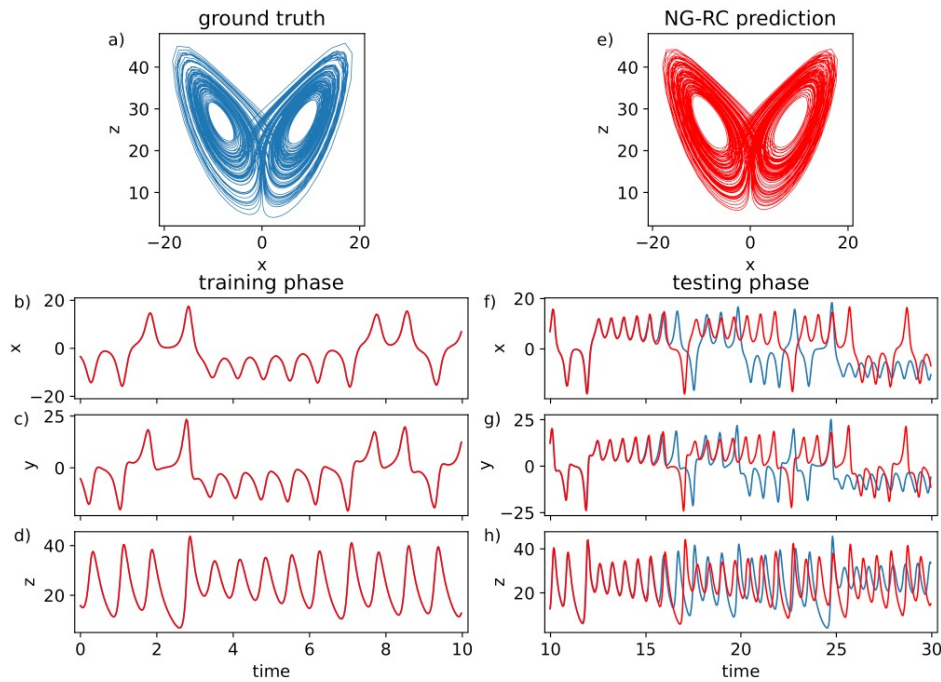
$$d_x \times d_r \text{ trained read-out matrix matrix } \mathbf{W}^{out}$$

Surprise – $\mathbf{A}$ and $\mathbf{W}^{in}$ are random but it still works!

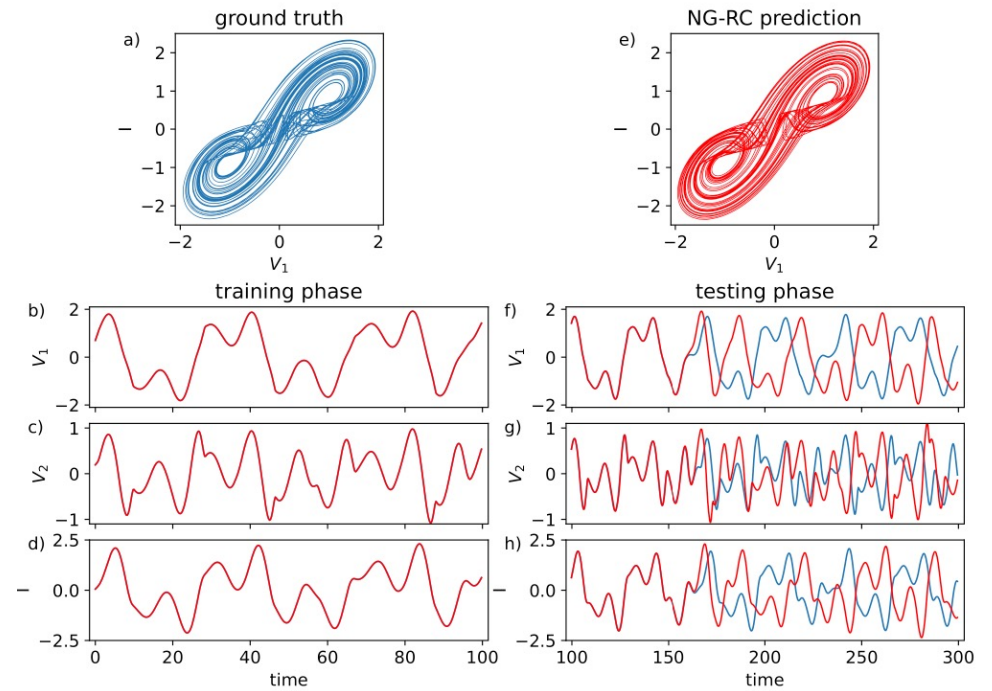
Notable Litt: **Gonon - Ortega 19', 20'** – RC enjoys a universal approximation theorem.
Even if linear with nonlinear readout.



NG-RC works very well, with very few points, almost no tunable parameters

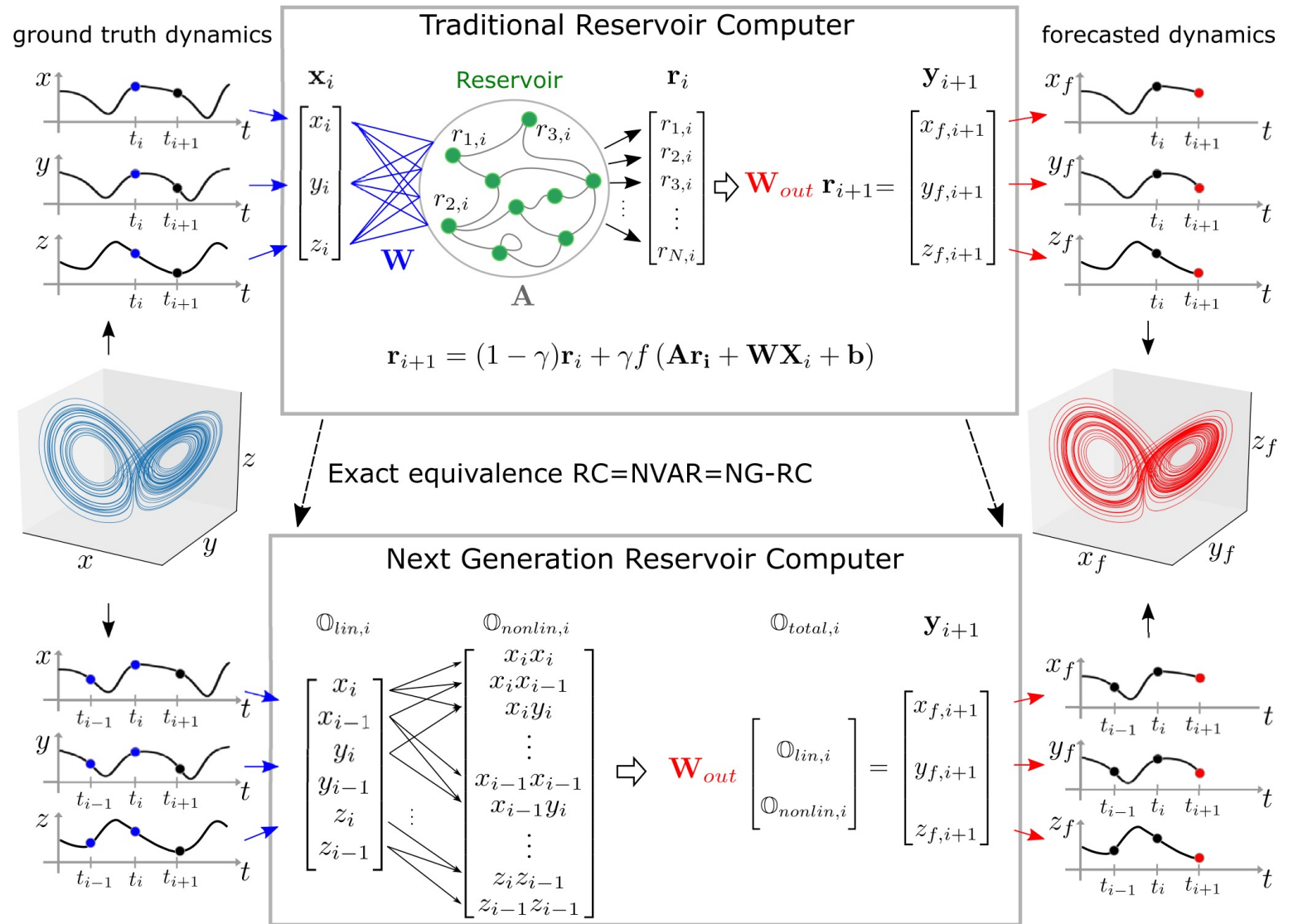


Forecasting a dynamical system using the NG-RC.
Lorenz63 strange attractors.



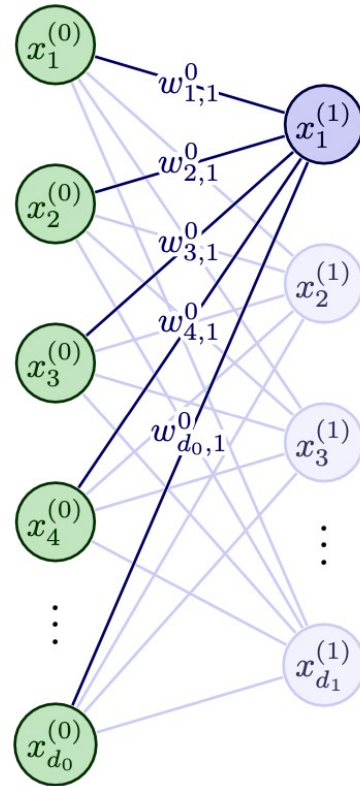
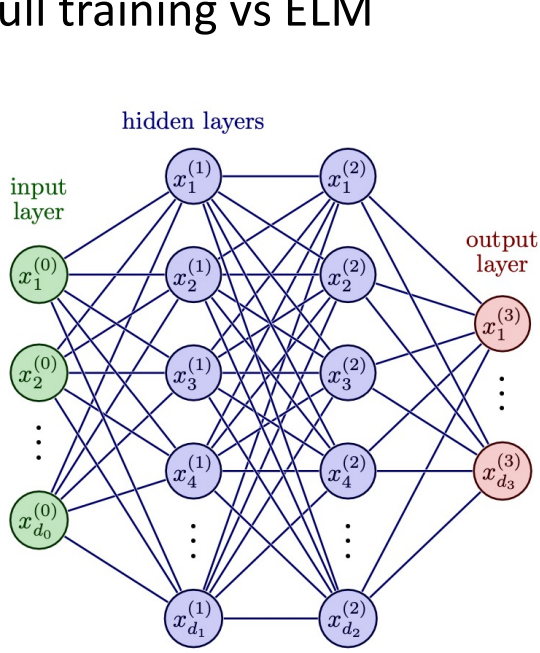
Forecasting the double-scroll system using the NG-RC

A traditional RC
is implicit
in an NG-RC



The “usual” FFNN (Feedforward Neural Network)– DNN (Deep Neural Network)

Full training vs ELM



$$x_1^{(1)} = \sigma \left(w_{1,1}^0 x_1^{(0)} + w_{2,1}^0 x_2^{(0)} + \dots + w_{d_0,1}^0 x_{d_0}^{(0)} + b_1^1 \right)$$

$$= \sigma \left(\sum_{i=1}^{d_0} w_{1,i}^0 x_i^{(0)} + b_1^1 \right)$$

$$\begin{pmatrix} x_1^{(1)} \\ x_2^{(1)} \\ \vdots \\ x_{d_1}^{(1)} \end{pmatrix} = \sigma \left[\begin{pmatrix} w_{1,1}^0 & w_{2,1}^0 & \dots & w_{d_0,1}^0 \\ w_{1,2}^0 & w_{2,2}^0 & \dots & w_{d_0,2}^0 \\ \vdots & \vdots & \ddots & \vdots \\ w_{1,d_1}^0 & w_{2,d_1}^0 & \dots & w_{d_0,d_1}^0 \end{pmatrix} \begin{pmatrix} x_1^{(0)} \\ x_2^{(0)} \\ \vdots \\ x_{d_0}^{(0)} \end{pmatrix} + \begin{pmatrix} b_1^1 \\ b_2^1 \\ \vdots \\ b_{d_1}^1 \end{pmatrix} \right]$$

$$x^{(1)} = \sigma \left(\mathbf{W}^0 x^{(0)} + \mathbf{b}^1 \right)$$

$$\mathcal{L}(\mathcal{D}; \Theta) = \sum_{i=1}^N \|F(X_i, \Theta) - Y_i\|_2 + \lambda \mathcal{R}(\Theta). \quad F(X, \Theta) = F_{q, \Theta_q} \circ F_{q-1, \Theta_{q-1}} \circ \dots \circ F_{0, \Theta_0}(X),$$

$$\Theta_{\text{ANN-}q}^* = \underset{\mathcal{D}; \Theta}{\operatorname{argmin}} \mathcal{L}(\mathcal{D}; \Theta), \quad \text{vs} \quad \Theta_{\text{ELM-}q}^* = \underset{\mathcal{D}; \Theta_q}{\operatorname{argmin}} \mathcal{L}(\mathcal{D}; \Theta_0 \cup \Theta_1).$$

ELM-1 it's a random SLFN

- random parameters except output layer.
- Nonlinear threshold function to hidden layers
- Identify function threshold function output

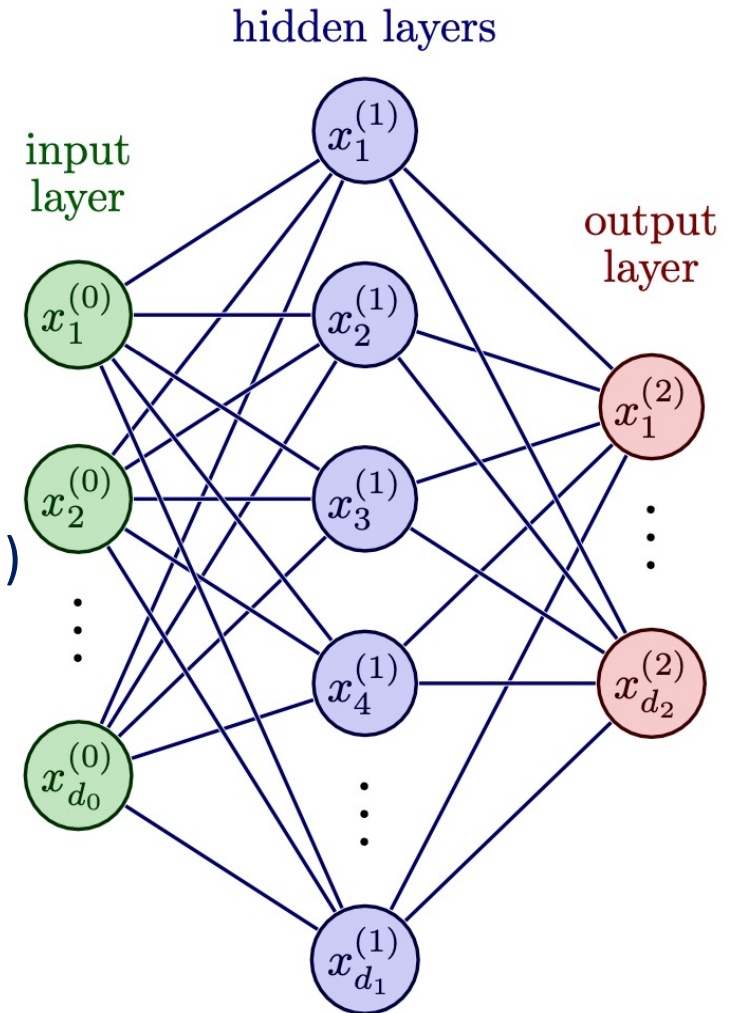
TRAIN only output layer,
Done by simple (linear – OLS regression/regularized)

$$f : \mathbb{R}^{d_0} \rightarrow \mathbb{R}^{d_2}$$

regression of a

$$\mathcal{D} = \{(X_i, Y_i)\}_{i=1}^N \subset \mathbb{R}^{d_0} \times \mathbb{R}^{d_2} = D$$

Based on observed data



Counting Parameters- Full Training.

q hidden layers,

$$\Theta = \cup_{r=0}^q \Theta_r.$$

$$\Theta_r = \{ \{w_{j,k}^r\}_{j=1, k=1}^{d_r, d_{r+1}}, \{b_k^{r+1}\}_{k=1}^{d_{r+1}} \}.$$

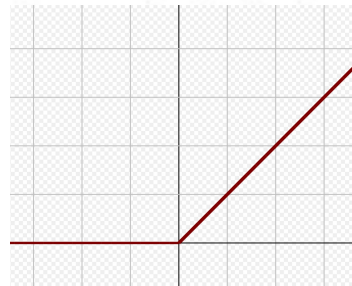
$$|\Theta| = d_0 d_1 + d_1 d_2 + d_2 d_3 + \dots + d_q d_{q+1} + d_q = d_q + \prod_{i=0}^q d_i d_i.$$

$$[W^r]_{j,k} = w_{j,k}^r, \text{ and, } [B^{r+1}]_k = b_k^{r+1}, j = 1, \dots, d_r, k = 1, \dots, d_{r+1},$$

$$F_{r, \Theta_r}(X^r) = \sigma_r(W^r X^r + B^{r+1})$$

Composition

$$F(X, \Theta) = F_{q, \Theta_q} \circ F_{q-1, \Theta_{q-1}} \circ \dots \circ F_{0, \Theta_0}(X)$$



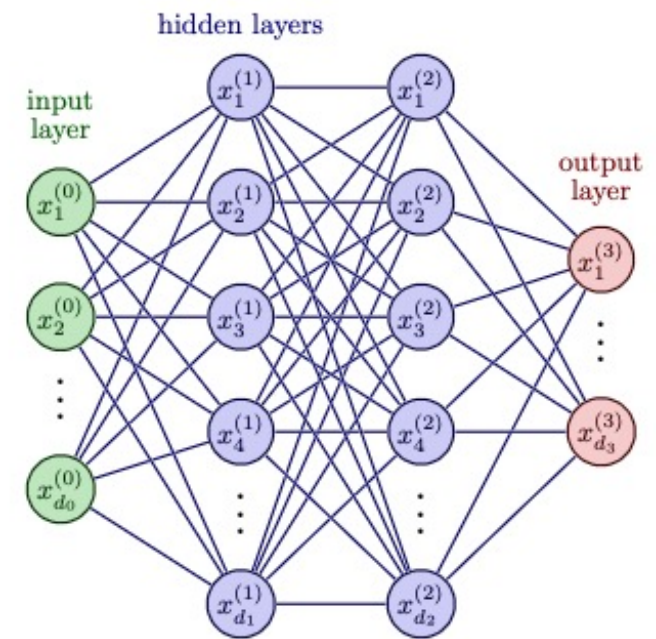
With ReLu Activation Something Simple Happens —configuration of fallen lines

$$\sigma_r(s) = ReLu(s) = \max(s, 0), r < q, \sigma_q(s) = s$$

Vs SLFN

$$|\Theta| = d_0 d_1 + d_1 d_2 + d_1 + d_2,$$

$$|\Theta| = d_0 d_1 + d_1 d_2 + d_1$$



The Joys of not training

SLFN vs ELM-1

$$|\Theta_{out}| = |\Theta_1| = d_1 d_2,$$

$$|\Theta_{out}| = d_1 d_2 \ll |\Theta| = d_0 d_1 + d_1 d_2 + d_1 + d_2,$$

And for DFN vs ELM-q, $q > 1$

$$|\Theta_{out}| = d_q d_{q+1} \ll |\Theta| = d_{q+1} + \prod_{i=0}^q d_i d_{i+1}.$$

Less parameters to train - training easy - OLS - or - Ridge-T

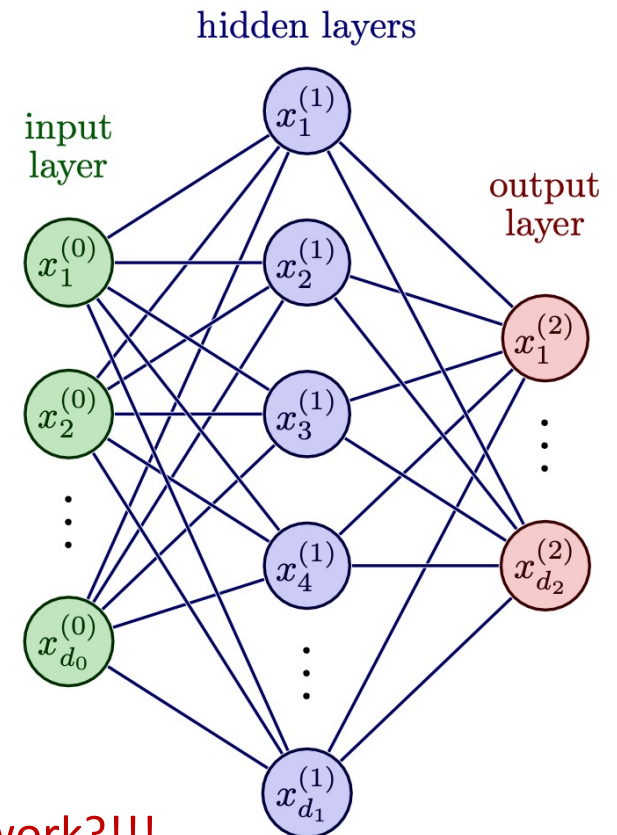
$$\mathbf{X} = [X_1^q | X_2^q | \dots | X_N^q], \mathbf{Y} = [Y_1 | Y_2 | \dots] \quad \mathbf{Y} = \beta \mathbf{X}$$

$$W^q := \underset{\beta}{\operatorname{argmin}} \|\mathbf{Y} - \beta \mathbf{X}\|_F + \lambda \|\beta\|_F,$$

$$\beta^* = \mathbf{Y} \mathbf{X}^T (\mathbf{X} \mathbf{X}^T + \lambda I)^{-1} := \mathbf{Y} \mathbf{X}_\lambda^\dagger.$$

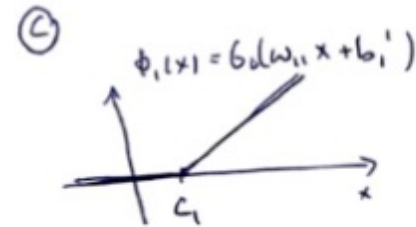
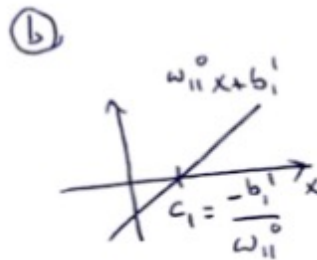
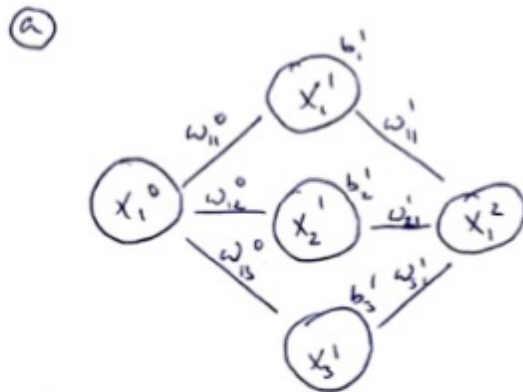
It's Cheap! - But does it work?!!!

Story of configurations of randomly fallen lines in plane
And randomly fallen hyperplanes in space.



ELM-1 - single layer FNN – SLFN

- here simplified to a **really small**, one input dimension, one output dimension

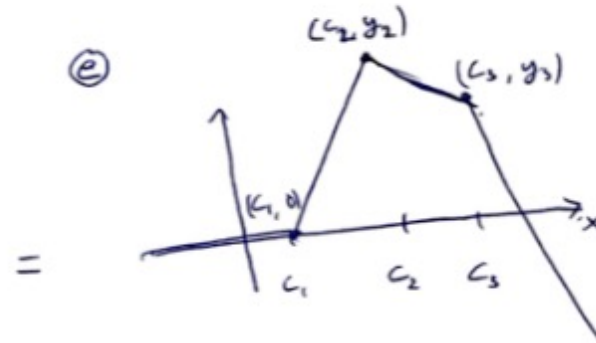
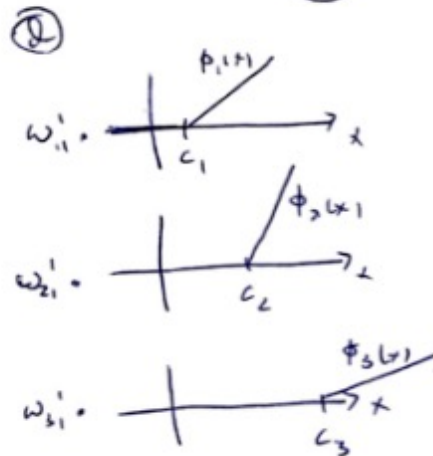


Untrained - Random Parameters.

$$w_{1,i} \sim U([0, 1]), b_i^{(1)} \sim U([-1, 0]),$$

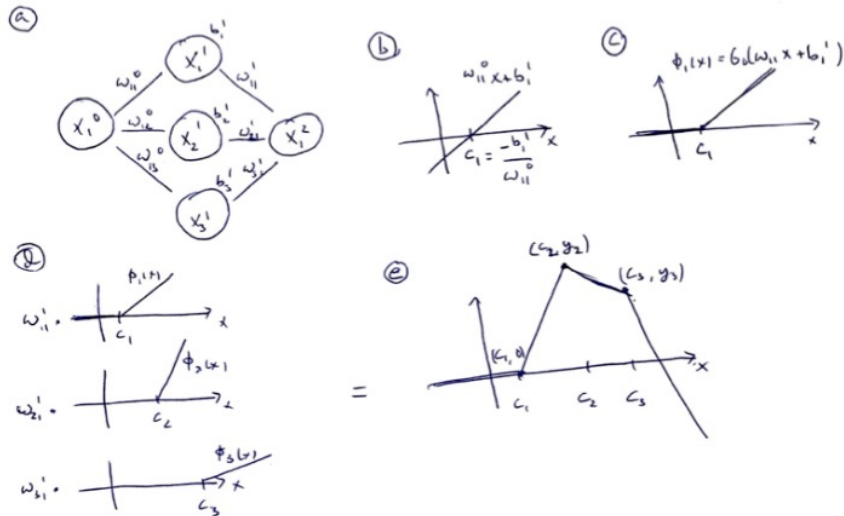
x-intercept

$$c_i^{(0)} = \frac{-b_i^{(1)}}{w_{1,i}^{(0)}},$$



ELM-1 - single layer FNN – SLFN

- here simplified to a **really small**, one input dimension, one output dimension



Untrained - Random Parameters.

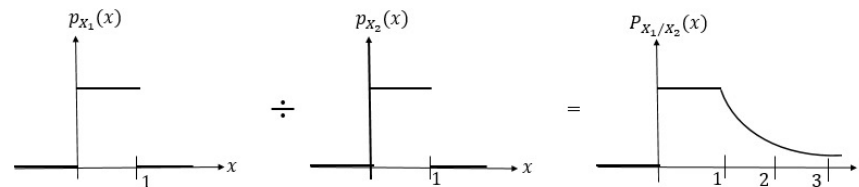
$$w_{1,i} \sim U([0, 1]), b_i^{(1)} \sim U([-1, 0]),$$

x-intercept

$$c_i^{(0)} = \frac{-b_i^{(1)}}{w_{1,i}^{(0)}},$$

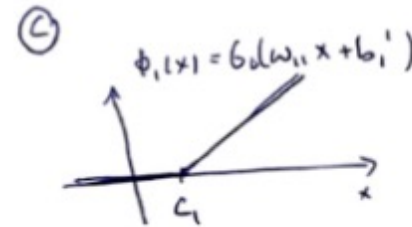
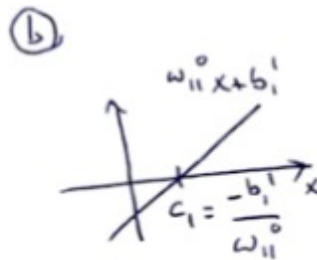
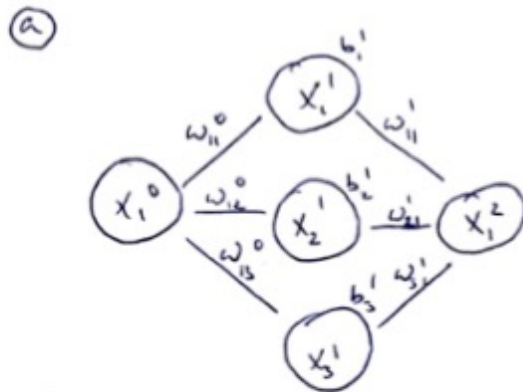
Quotient distribution

$$c_i \sim \frac{(H(1-x)) - H(x - \frac{1}{2})}{2}, \text{ where } H(x) = \begin{cases} 0, & x < 0, \\ \frac{1}{2}, & x = 0, \\ 1, & x > 0, \end{cases}$$



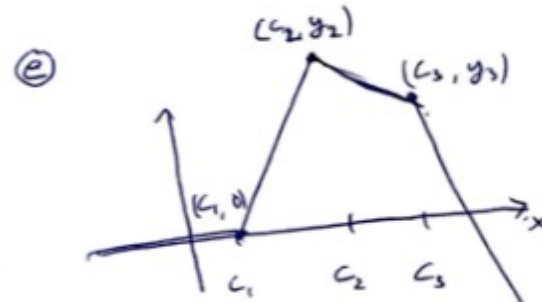
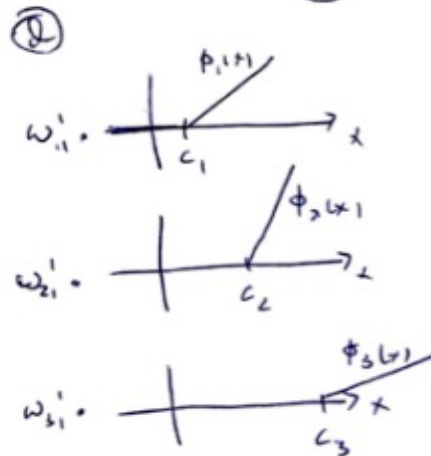
ELM-1 - single layer FNN – SLFN

- here simplified to a **really small, one input dimension, one output dimension**



x-intercept

$$c_i = \frac{-b_i^{(1)}}{w_{1,i}^{(0)}}.$$



Produce piece-wise linear spline in 3 knots – more with larger d_1

$$\hat{f}(x) = w_{1,1}^{(1)}\phi_1(x) + w_{2,1}^{(1)}\phi_2(x) + w_{3,1}^{(1)}\phi_3(x).$$

$$\hat{f} \in \Delta_{d_1} = \text{span}(\{\phi_1(x), \phi_2(x), \phi_3(x)\})$$

Universal Approximation Theorem ELM as a picture

Given lots of data,

$$\mathcal{D} = \{(X_i, Y_i)\}_{i=1}^N \subset \mathbb{R}^{d_0} \times \mathbb{R}^{d_2} = D$$

With a large hidden layer of the SLFN

My take:

- random hidden layer

- (the c_i are placed randomly)

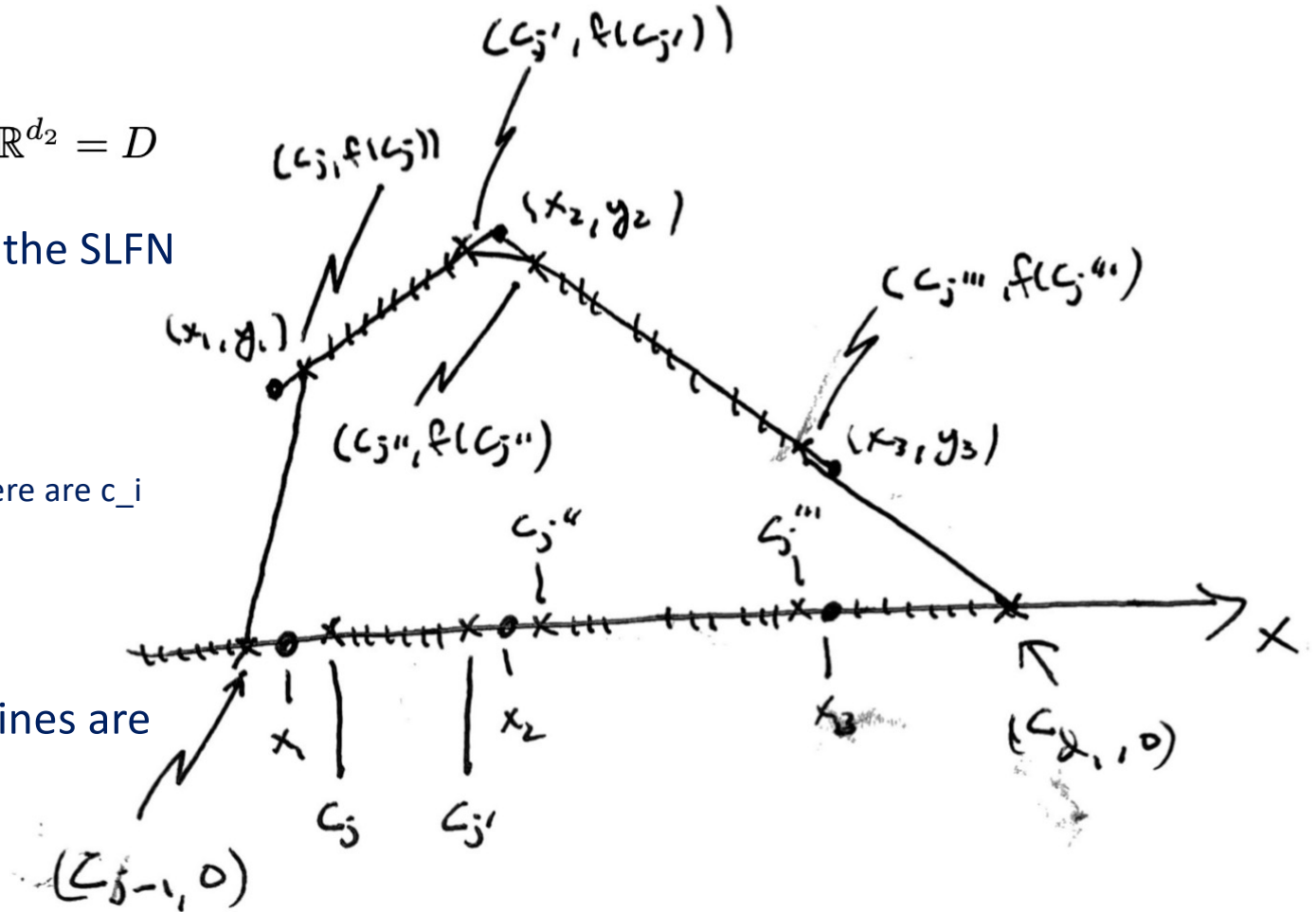
- between data points x_i there are c_i

- a linear spline can run through data.

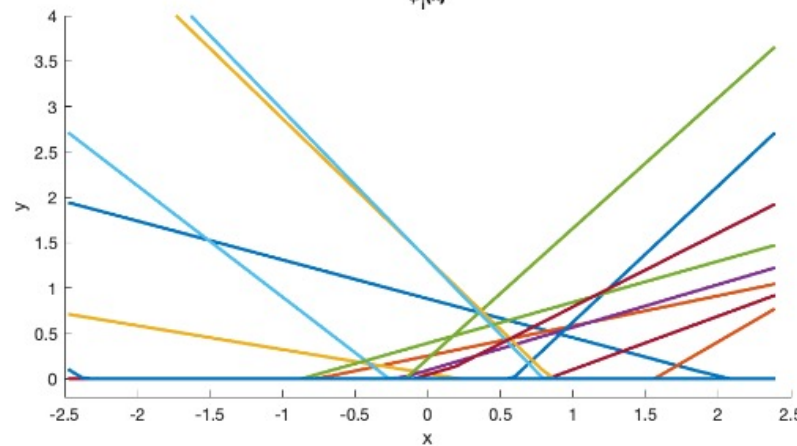
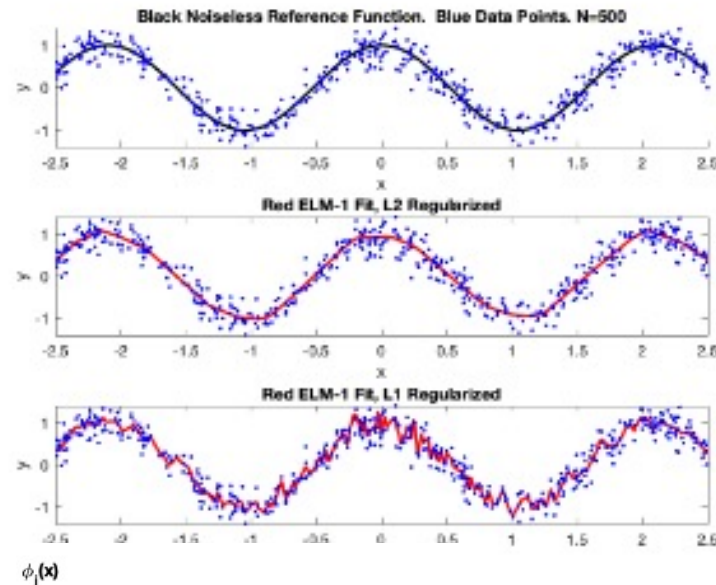
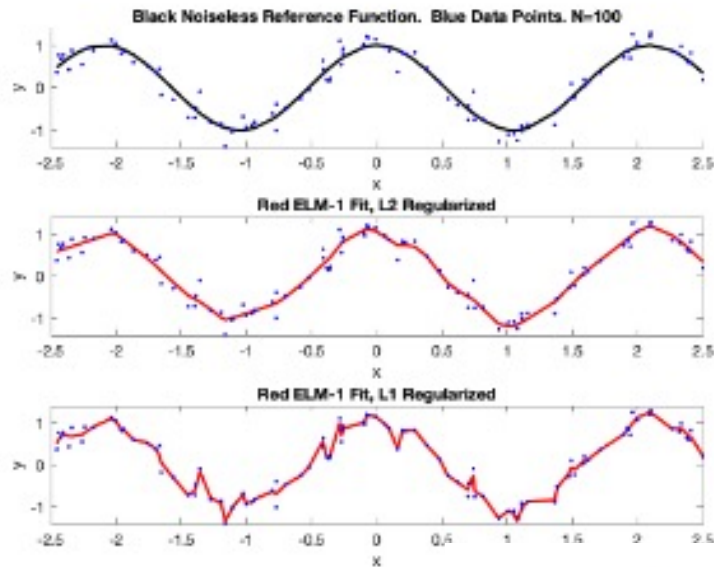
General piecewise linear splines are
Dense in C^0 .

Vs Stone-Weierstrass

Cite Huang, ELM as universal approximation theorem.

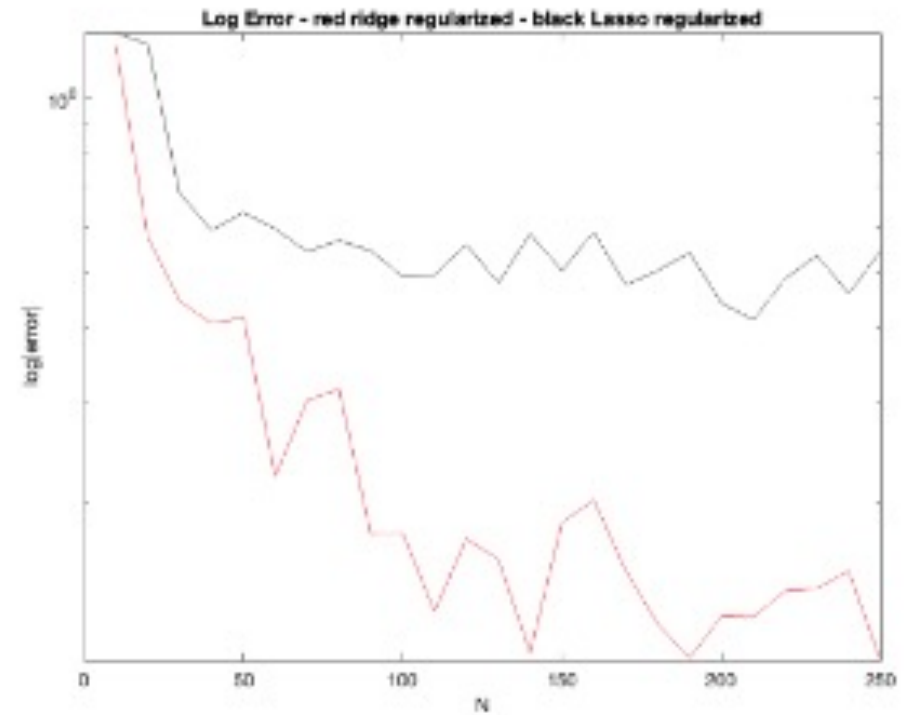
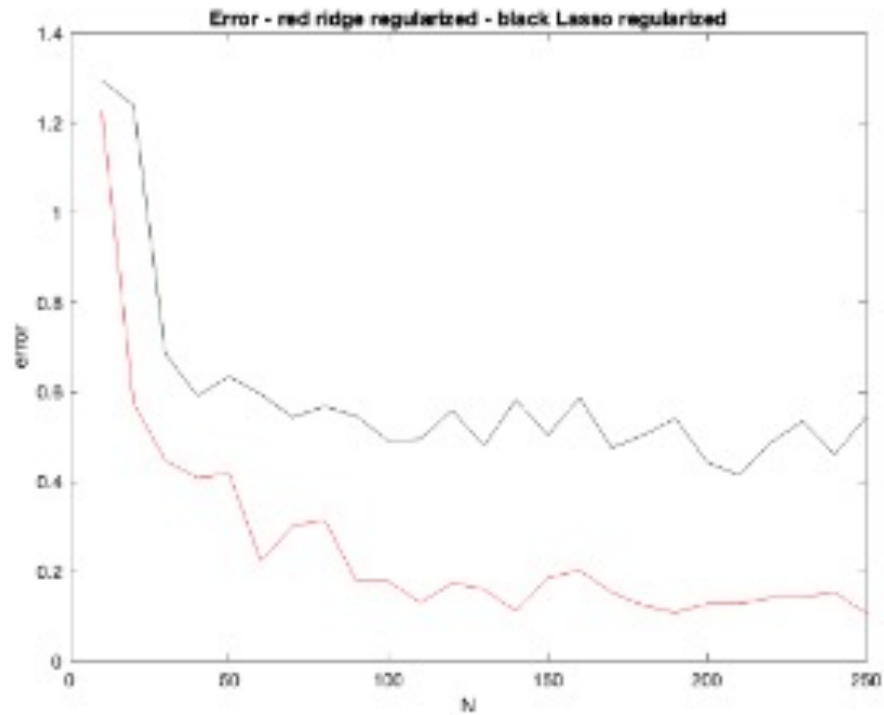


A One-D ELM-1 experiment – with noise – and increasing data set size



Some of the $\phi_i(x)$

A One-D ELM-1 experiment – with noise – and increasing data set size



ELM-1 - single layer FNN – SLFN

- here simplified to a **really small, TWO input dimension, one output dimension**

Consider the i^{th} neuron in the hidden layer, but before the threshold is applied,

$$z = W_{i,:}^0 \cdot \mathbf{X}^{(0)} + b_i^{(1)}, i = 1, 2, \dots, d_1.$$

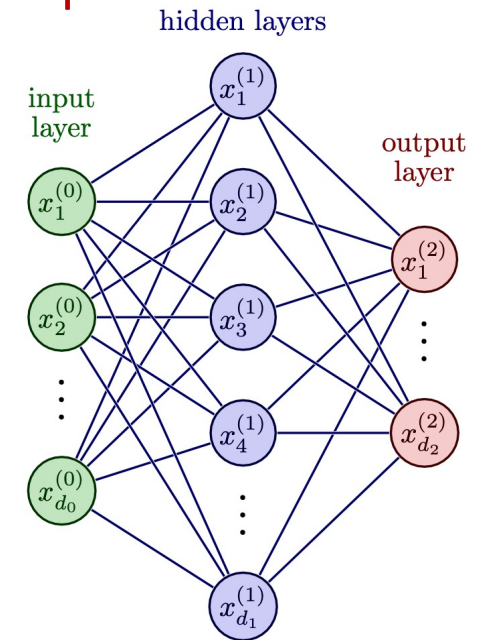
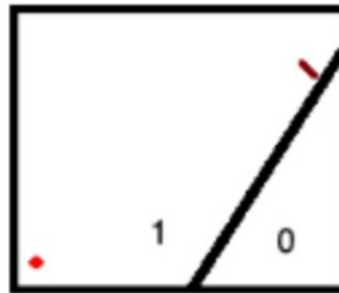
!!equation of a line in 2-d! Or a plane in 3-d, or a hyperplane in n-d.

Still ReLu activation function:

$$\sigma_r(s) = \text{ReLU}(s) = \max(s, 0), r < q, \sigma_q(s) = s$$

half plane ReLu does nothing, and for other half plane, you get zero.

$$z \geq W_{i,:}^0 \cdot \mathbf{X}^{(0)} + b_i^{(1)}.$$



ELM-1 - single layer FNN – SLFN

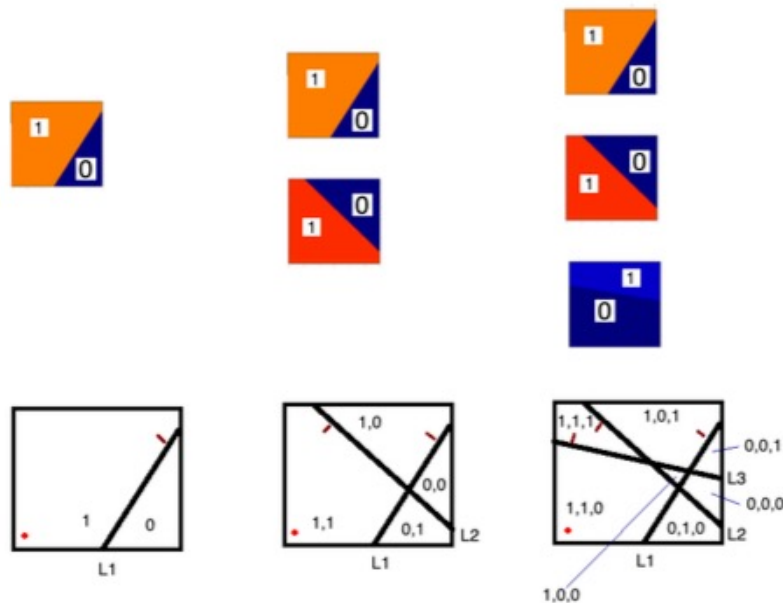
- here simplified to a **really small, one input dimension, one output dimension**

Consider the i^{th} neuron in the hidden layer, but before the threshold is applied,

$$z = W_{i,:}^0 \cdot \mathbf{X}^{(0)} + b_i^{(1)}, i = 1, 2, \dots, d_1.$$

For half plane ReLu, you get zero.

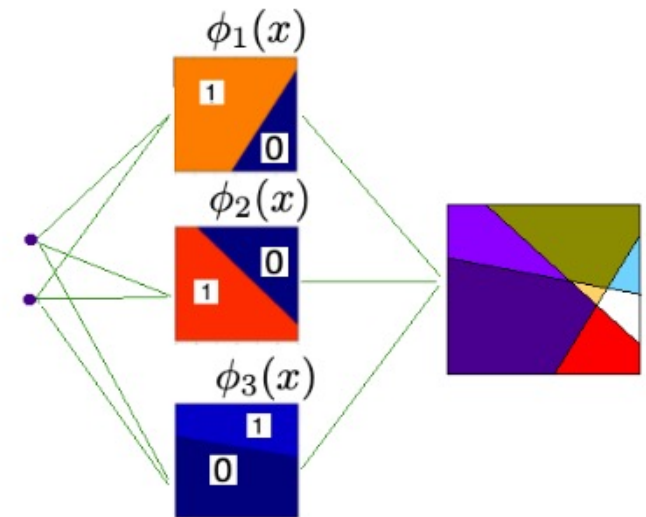
$$z \geq W_{i,:}^0 \cdot \mathbf{X}^{(0)} + b_i^{(1)}.$$



a)

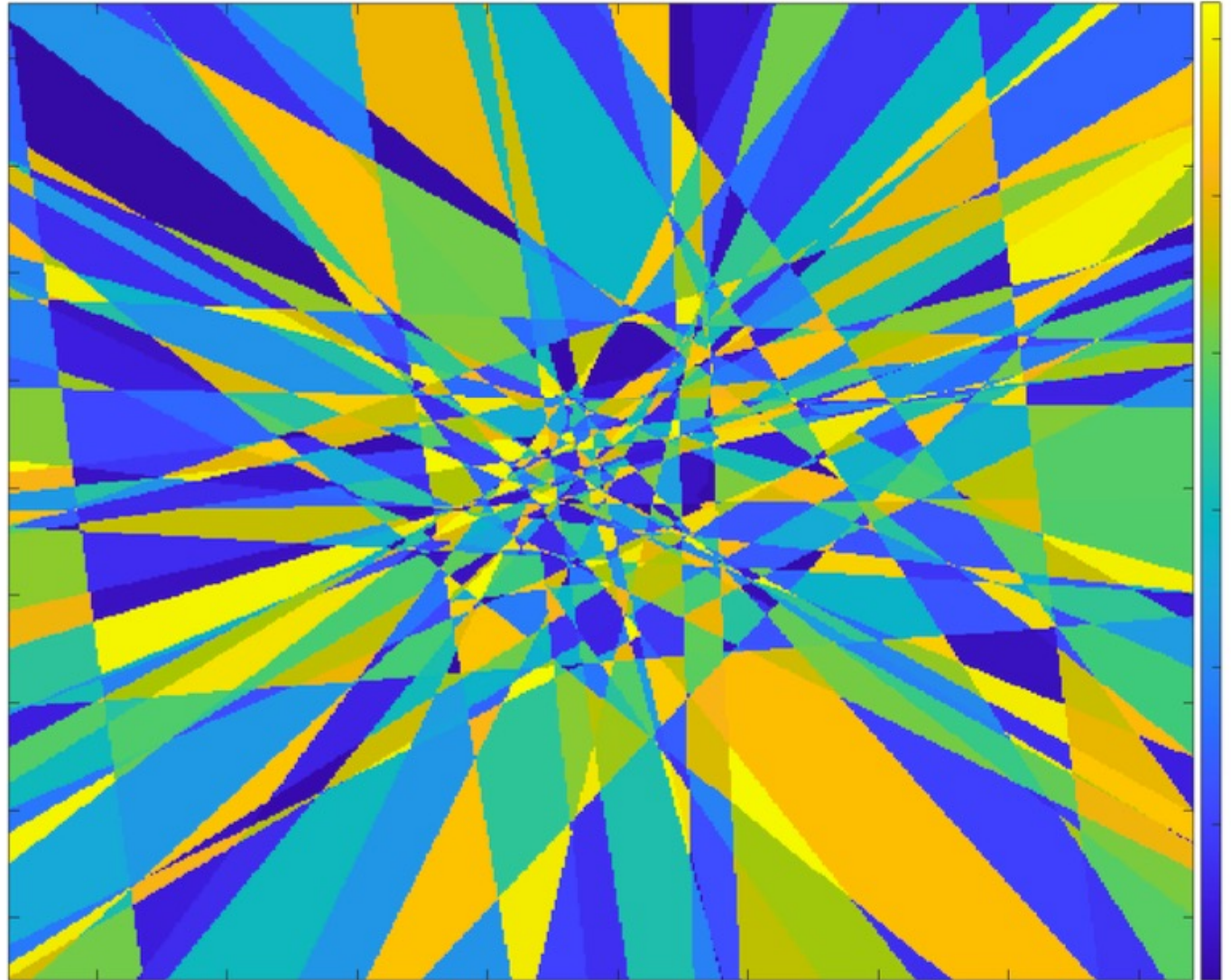
Notice Symbology

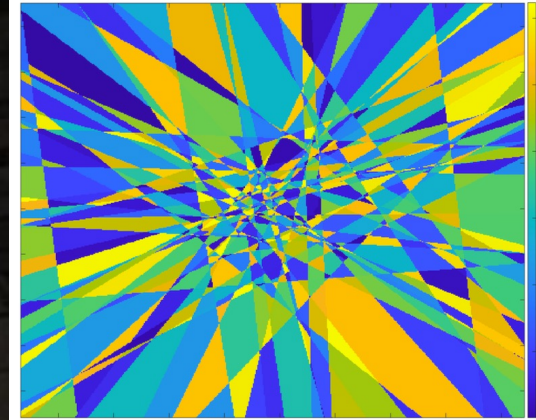
**Still end up with
Linear combinations of $\phi_i(x)$**



b)

ELM-1 in
2-input dimensions
Lots of hidden layer nodes

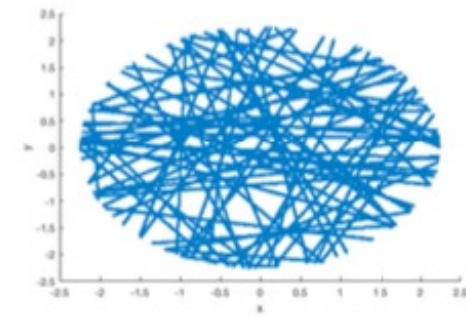
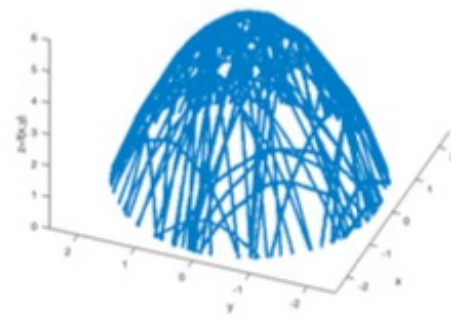
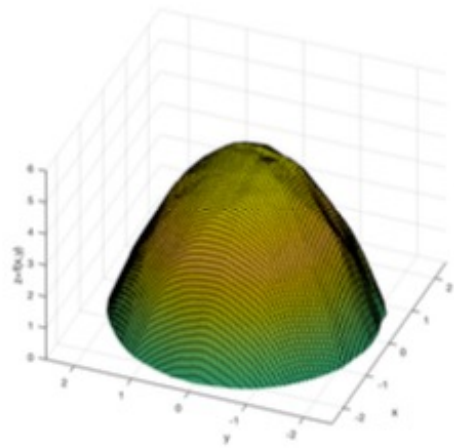
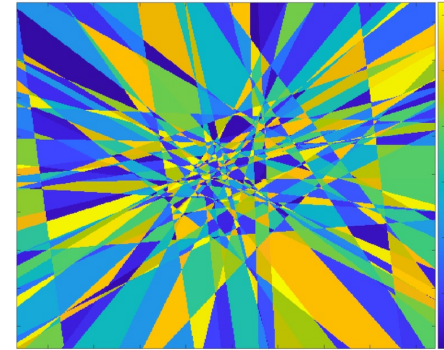
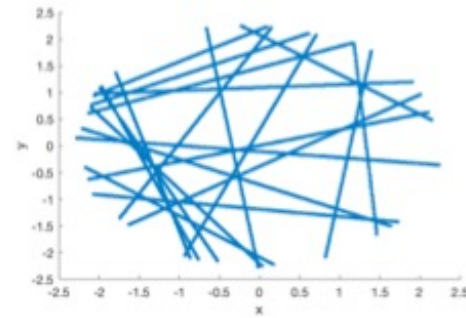
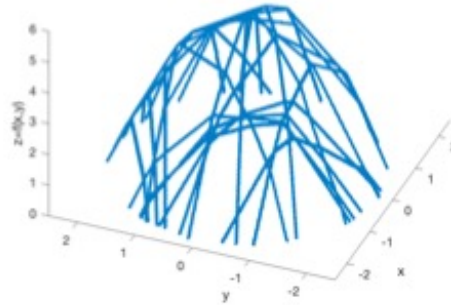
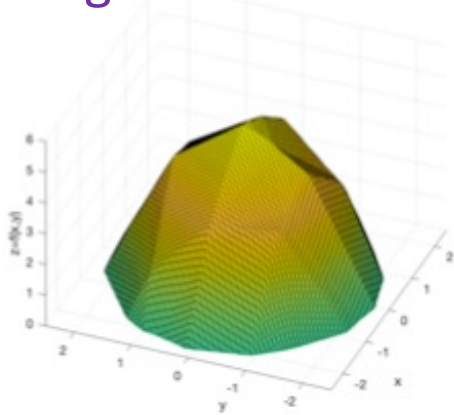




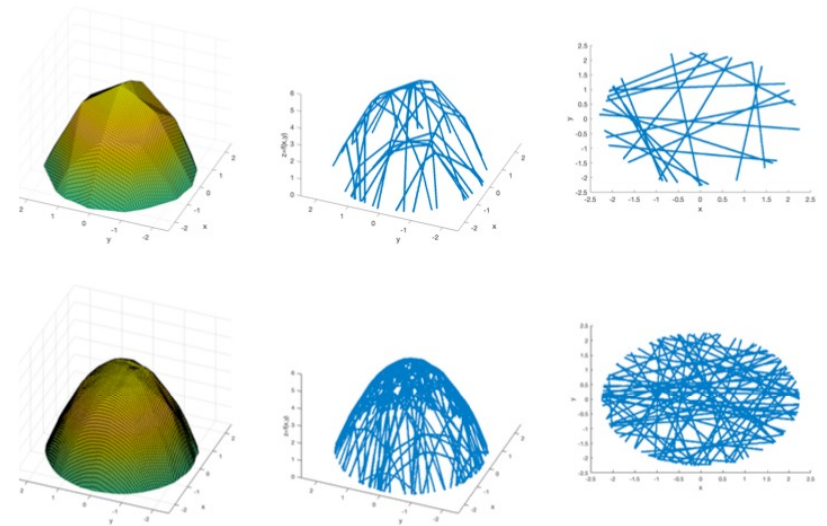
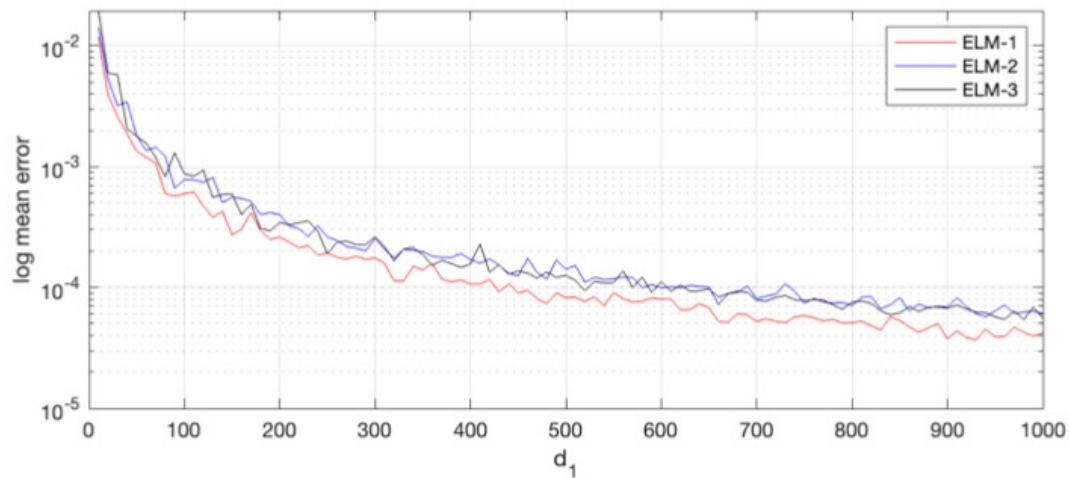
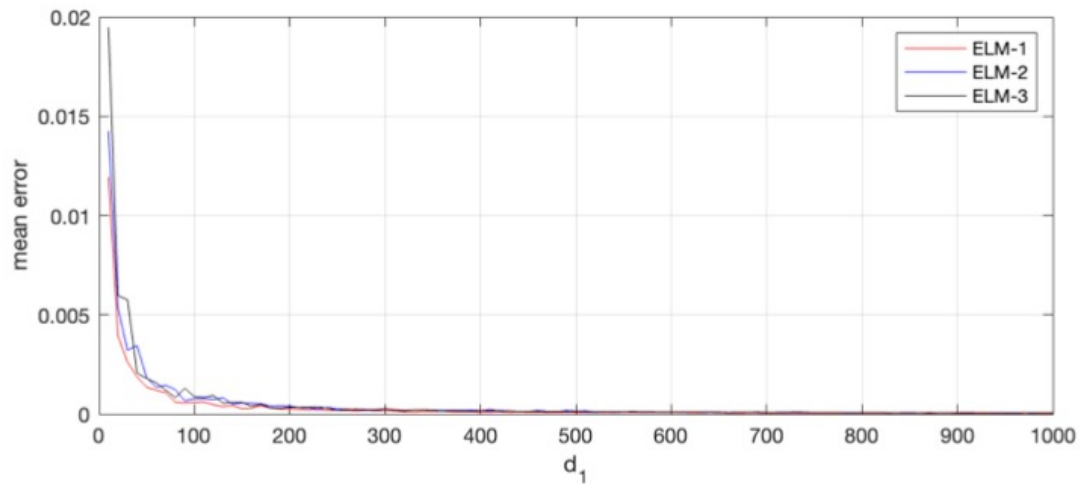
Church Windows!

Neural Nets in
The Middle Ages!

Fitting a function with an ELM-1 and increasing hidden layer size



Larger d_1 means more randomly placed lines, - finer random partition => more precision of piecewise linear cts fn.



Story summary –

ELM-1 – bigger hidden layer,
 -more random lines
 -finer fitting of piecewise linear fn
 -piecewise linear dense in C^0

ELM-q – $q > 1$

-even though for fully trained – deep
 Is good. More expressive

-for ELM not any better

Thank you!



Theory of Random Configurations of Lines in the Plane

Of (hyper)Planes in Space.

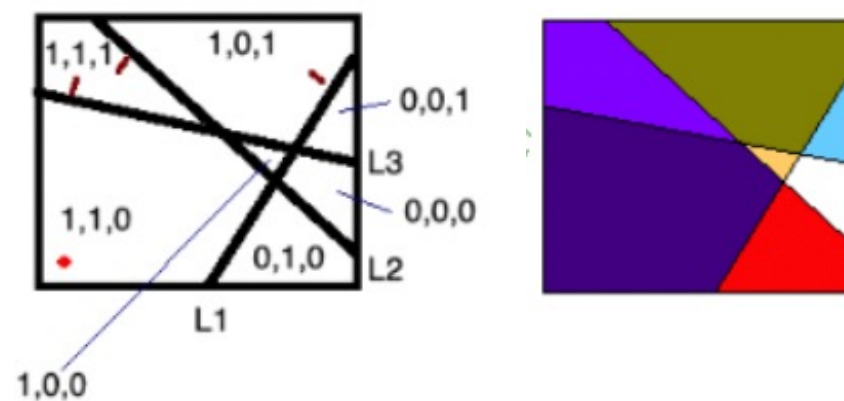
$$M_n \leq \frac{n(n+1)}{2} + 1 = \# \text{ of cells in an arrangement of } n\text{-lines in 2-dimensional space.}$$

Notice:

$$M_1 = 2, M_2 = 4, \text{ but } M_3 = 7.$$

Fewer than you might guess – due to occlusion

$$\frac{n(n+1)}{2} + 1 \ll 2^n.$$



011 is missing.

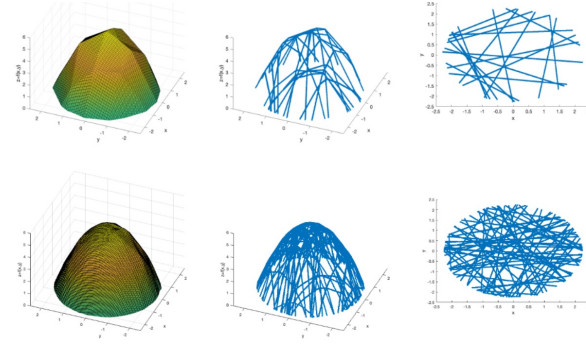
$$M_{m+1} \leq M_m + (m + 1), M_1 = 2.$$

Theory of Random Configurations of Lines in the Plane Of (hyper)Planes in Space.

$$M(d_0, n) \leq \sum_{i=0}^{d_0} \binom{n}{i} = \# \text{ of cells in an arrangement of } n\text{-planes (hyperplanes) in } d_0\text{-dimensional space.}$$

recursion relationship. This time,

$$M(d_0, n) \leq M(d_0, n-1) + M(d_0, n-1).$$

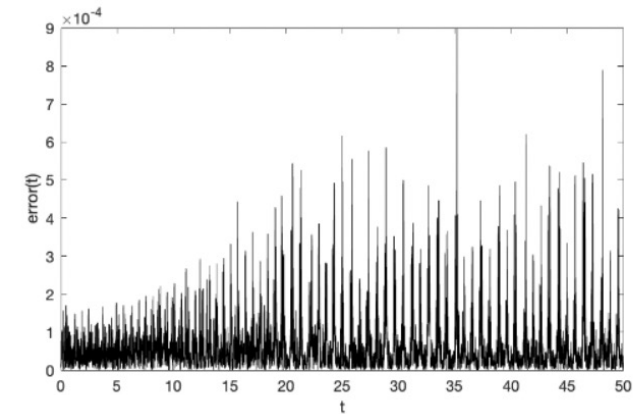
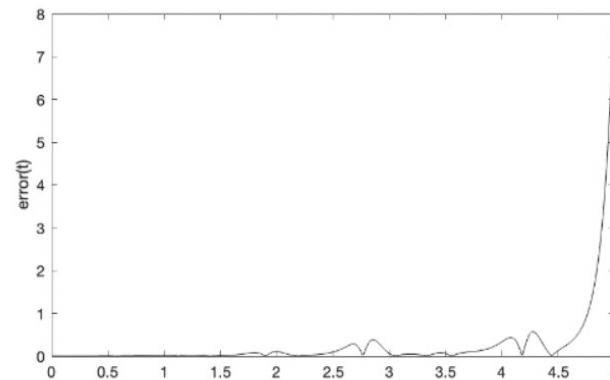
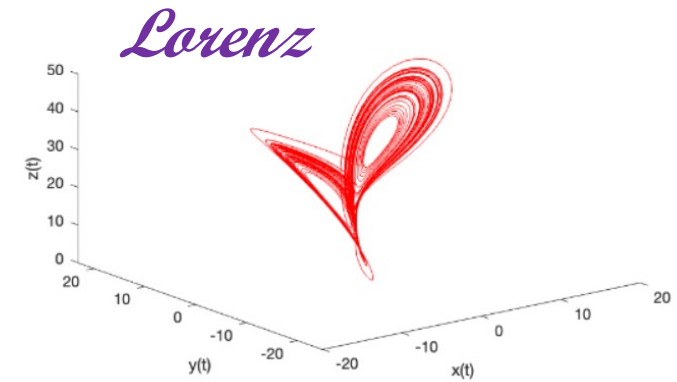
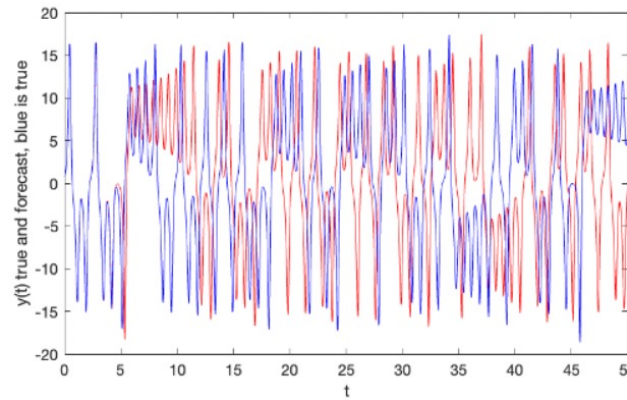


$$B(d_0, n) \leq \binom{n-1}{d_0} = \# \text{ of bounded cells in an arrangement of } n\text{-hyperplanes in } d_0\text{-dimensional space,}$$

$$B(d_0, n) \leq \binom{n-1}{d_0} = \# \text{ of bounded cells in an arrangement of } n\text{-hyperplanes in } d_0\text{-dimensional space,}$$

- As the size of the hidden layer increases, $d_1 \uparrow$, then the number of linear domains increases.
- As $d_1 \uparrow$, the size of these domains decreases. This is a random refinement.
- As $d_1 \uparrow$, the number of basis elements in Δ_{d_1} increases, and the fit accuracy improves.

ELM for Forecasting Chaos



Differential Equation

$\dot{z} = f(z)$, for a given initial condition in its phase space $z(t_0) = z_0 \in \mathcal{M} \subset \mathbb{R}^{d_0}$,

Flow – FUNCTION $\varphi : \mathcal{M} \times \mathbf{R} \rightarrow \mathcal{M}$
 $(t, z_0) \mapsto z(t) = \varphi(t; z_0).$

Orbit DATA
 $\mathcal{D} = \{(z_k, z_{k+1})\}_{k=1}^N.$