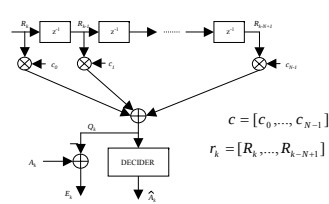


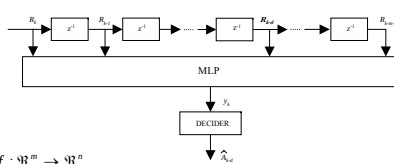
Equalization systems

LTE



$$\text{LMS: } Q_k = cr_k^* \quad E_k = A_k - Q_k \quad c_{k+1} = c_k + \beta E_k r_k$$

Forward MLP

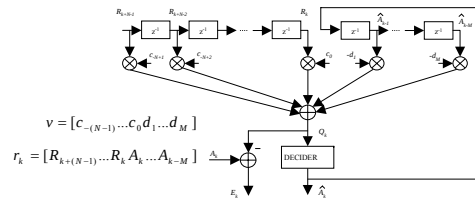


$$f: \mathbb{R}^m \rightarrow \mathbb{R}^n$$

$$y_l = [f(r_k)]_l = \left(\sum_{j=1}^h w_{lj}^2 \left(\sum_{i=1}^m w_{ij}^2 r_i + w_j^1 \right) + w_l^2 \right) \quad l=1, \dots, n$$

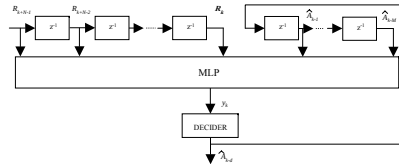
$r_k = [r_{k,1}, \dots, r_{k,m}]$ input vector at period (time) k ,
 $y_k = [y_{k,1}, \dots, y_{k,n}]$ output vector at period (time) k ,
 w_{ij}^1, w_{ij}^2 first and second layer weights

DFE

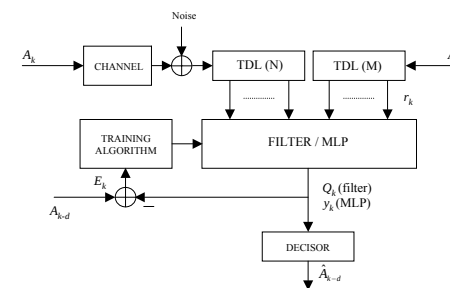


$$\text{LMS: } E_k = A_k - Q_k = A_k - v_k r_k^* \quad v_{k+1} = v_k + \beta E_k r_k$$

Feedback MLP



Learning phase



training set $\rightarrow$ Training sequence
 $tr = [A_1, A_2, \dots, A_k, \dots, A_{L_T}]$

At each learning iteration k :

- output channel sampling r_k
- estimation of the symbol Q_k (filter) or y_k (MLP)
- comparison with the target value A_{k-d}
- output error computation E_k
- weight update

DFE and feedback MLP: feedback outputs

$\rightarrow$ target A_{k-d}

During training: at each learning iteration k , computation of the MSE (E_k)

$$\text{FILTER} \rightarrow \text{MSE}_k = (Q_k - A_k)^2$$

$$\text{MLP} \rightarrow \text{MSE}_k = \frac{(y_k - t_k)(y_k - t_k)^T}{n} \quad (n \text{ uscite})$$

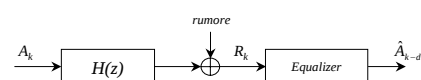
Performance evaluation: BER

$$\text{BER} = \frac{\sum_{i=1}^r d_i / nc}{r}$$

d_i : # of errors (not correctly estimated symbols) at the receiver
 nc : transmitted symbols
 r : # of transmission experiments

Base band 2 levels PAM channel

Discrete time model



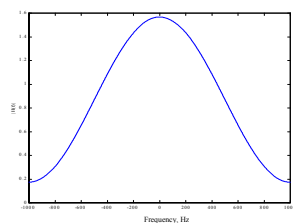
Input stream: random stream of equal-probability and statistically independent symbols $A_k \in \{-1, 1\}$

AWGN noise with variance σ^2

$$\text{SNR} = \frac{\tilde{f} \cdot \tilde{f}^T}{\sigma^2} \quad \text{at the equalizer input}$$

$\tilde{f}$ $H(z)$ coefficients vector

$$H(z) = 0.3482 + 0.8704 z^{-1} + 0.3482 z^{-2}$$



Decoder: signum function

MLP: 10 hidden neurons, 1 output neuron

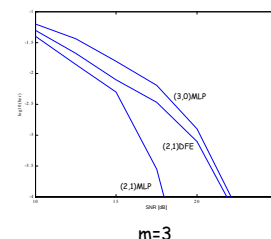
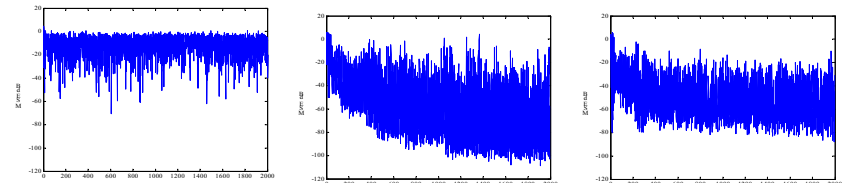
MLP $\rightarrow$ BP: $\eta=0.07, \eta_b=0.05, \alpha=0.3, a_b=0$

DFE $\rightarrow$ LMS: $\beta=0.035$

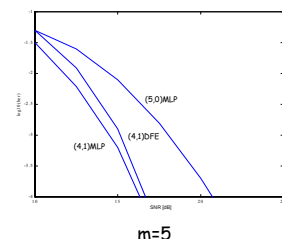
✓ DFE: reference structure for the comparison with the NN

DFE - MLP performance comparison

Training sequence length: 2000 samples



Higher order equalizer

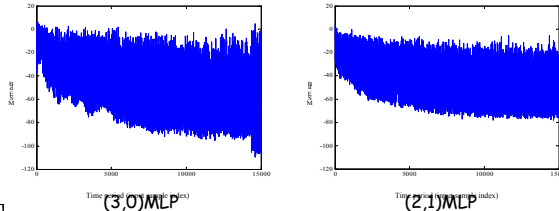
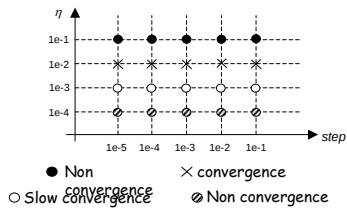


NNs give better performance with respect to DFE filters when the equalizer order is close to the channel filter order

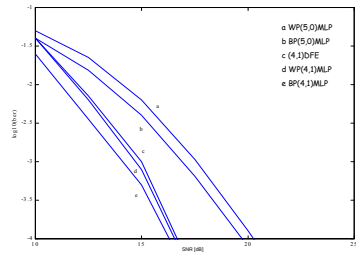
BP - WP comparison

WP: training convergence analysis

step=10⁻³
η=10⁻²
Lt=15000



The performance are comparable unless the training sequence of WP length is 5-fold the one of BP.



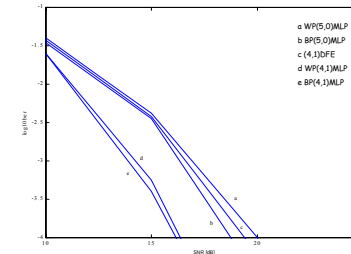
M. Valle

Non-linear channel

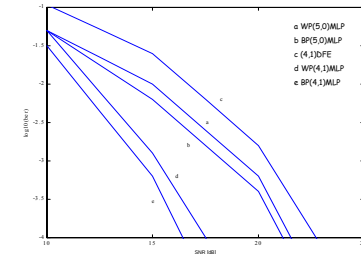


$$H(z) = 0.3482 + 0.8704z^{-1} + 0.3482z^{-2}$$

NL = hyperbolic tangent: $g(x) = \tanh(x)$



NL = exponential: $g(x) = \exp(x) - 1$



In case of non-linear systems/channels, NNs achieve better performance in particular for low SNR values



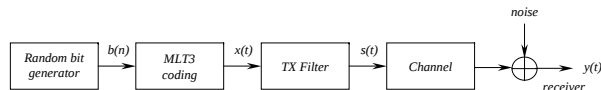
M. Valle

Fast Ethernet: 100BASE-TX

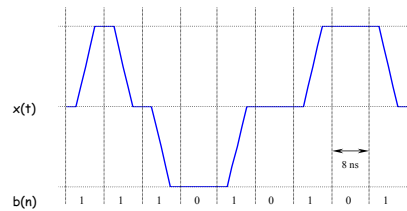
100BASE-TX (IEEE 802.3):

- ✓ 2 twisted pairs CAT5 (1 per TX e 1 per RX)
- ✓ input bit stream @ 100 Mbps, 4b/5b coding
- ✓ input data stream b(n) @ 125 Mbps, MLT3 coding

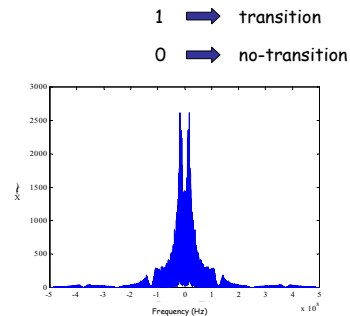
MATLAB model:



- ✓ b(n) random equiprobable bit stream @ 125Mbps
- ✓ MLT3 coding: 3-level (1, 0, -1) baseband signal



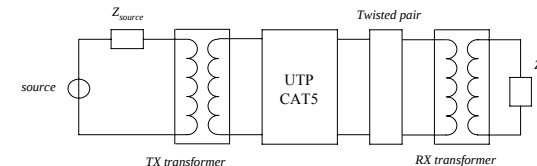
- ✓ TX filter: 4th order BESSEL



M. Valle

Channel model

- ✓ Channel (ANSI/TIA/EIA-568-A):



UTP CAT5

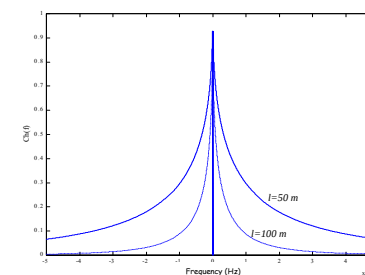
$$H_{dB}(f, l, C) = A_0 \cdot \left(\frac{2j \cdot f}{F_0} \right)^{0.5} \cdot \frac{l}{X_0} \cdot (1 + 0.003 \cdot (C - 20))$$

$$F_0 = 31.25 \cdot 10^6 \text{ Hz} \quad l: \text{twisted pair length}$$

$$A_0 = -36 \text{ dB} \quad C: \text{temperature}$$

$$X_0 = 305 \text{ m}$$

Channel frequency response magnitude



- ✓ @ 100 MHz → 0.3 l = 50 m
0.1 l = 100 m

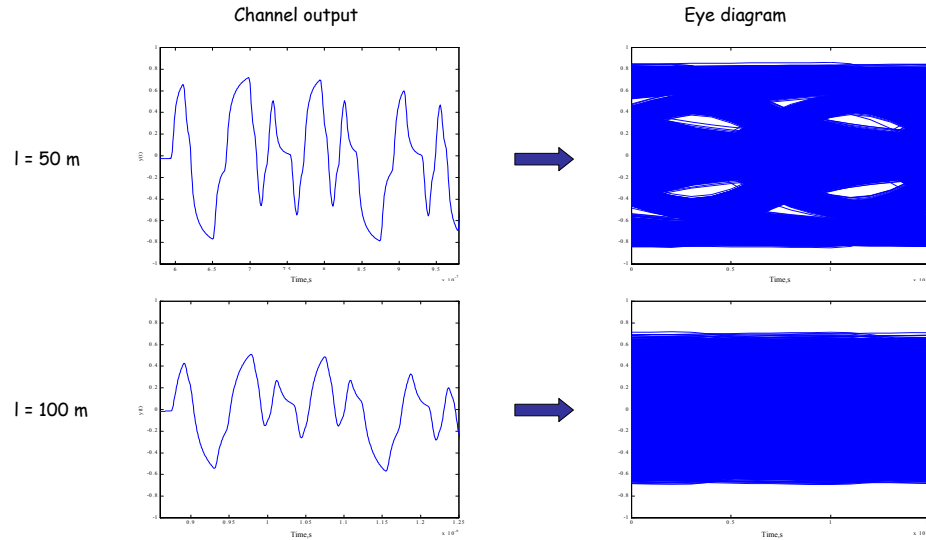
- ✓ transformer → high pass (low cut-off freq = 50 kHz)

- ✓ AWGN noise

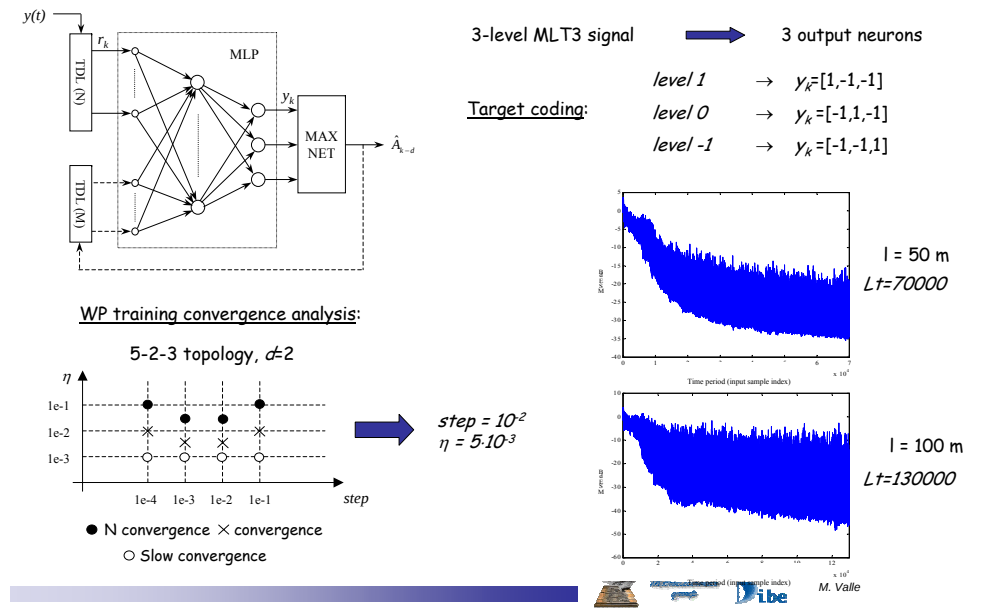


M. Valle

Waveforms



Equalizer and training phase



Performance

BER value averaged on 50 different 10000-symbol sequences

Analysis on the equalizer structure

- ✓ input neurons (i.e. equalizer order)(m) $\rightarrow$ Non convergence if $m > 10$
- ✓ hidden neurons (h) $\rightarrow$ No performance enhancement if $h > 2$
- ✓ estimation delay (d) $\rightarrow$ Performance enhancement when d approaches 0

Forward MLP equalizer

- ✓ TDL 5 taps with 8 ns delay
- ✓ MLP topology $5 \times 2 \times 3$
- ✓ Estimation delay $d=0$

Training

- ✓ Weight Perturbation $\eta = 5 \cdot 10^{-3}$ step $= 10^{-2}$
- ✓ Training sequence ≥ 70000 $l = 50 \text{ m}$ and training sequence ≥ 130000 $l = 100 \text{ m}$

Performance

- ✓ BER $< 10^{-4}$ @ SNR $\geq 30 \text{ dB}$, $l \in [0, 100] \text{ m}$
- ✓ BER $= 10^{-3}$ @ SNR $= 20 \text{ dB}$, $l = 50 \text{ m}$
- ✓ BER $= 10^{-3}$ @ SNR $= 25 \text{ dB}$, $l = 100 \text{ m}$

Non-ideal effects: baseline wander and jitter

