



AIDA4Edge



University of Niš
Faculty of Electronic Engineering

Neural Network Quantization

Milan Dinčić

University of Niš, Serbia

2nd Scientific Workshop on Fault-Tolerant Design and Real-Time Reliability Monitoring

July 15, 2026

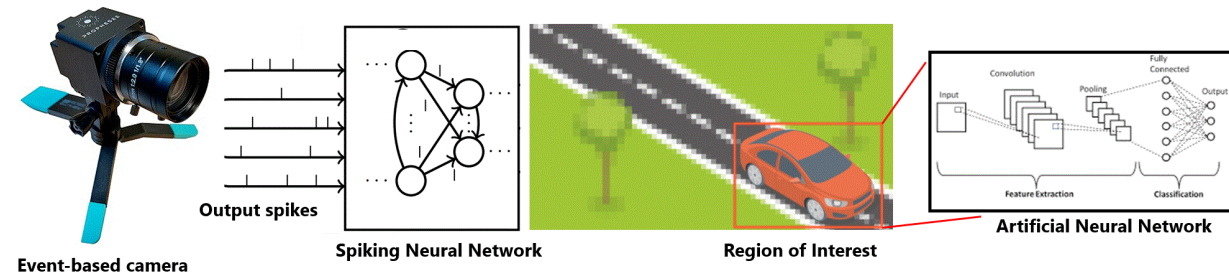
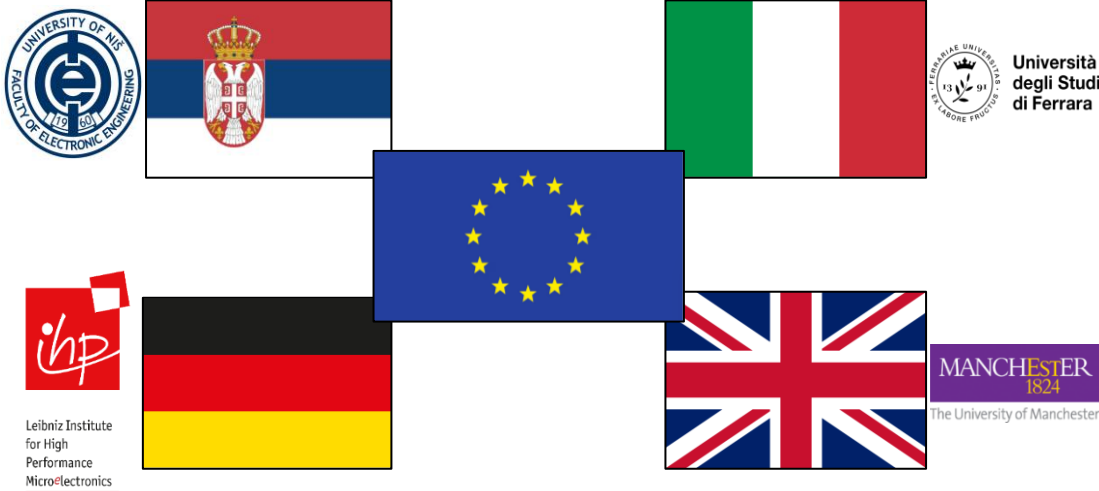
Volos – Greece

AIDA4Edge Horizon Twinning Project



Project title: Twinning for Excellence in Adaptive Edge AI
Time period: 1 October 2024 - 30 September 2027

- To develop low-complexity AI models suitable for deployment on edge devices with limited memory, processing power and energy.
 - Efficient quantization
 - Hardware-aware Neural architecture search
 - Spiking neural networks
 - Hybrid ANN+SNN models
 - Efficient and reliable self-adaptive AI accelerators



Content

- **Quantization**
- **Digital data formats**
 - **Fixed-point formats**
 - **Integer formats**
 - **Floating-point formats**
- **Quantization of Neural networks**
 - **Efficiency**
 - **Reliability**

Content

- **Quantization**
- Digital data formats
 - Fixed-point formats
 - Integer formats
 - Floating-point formats
- Quantization of Neural networks
 - Efficiency
 - Reliability

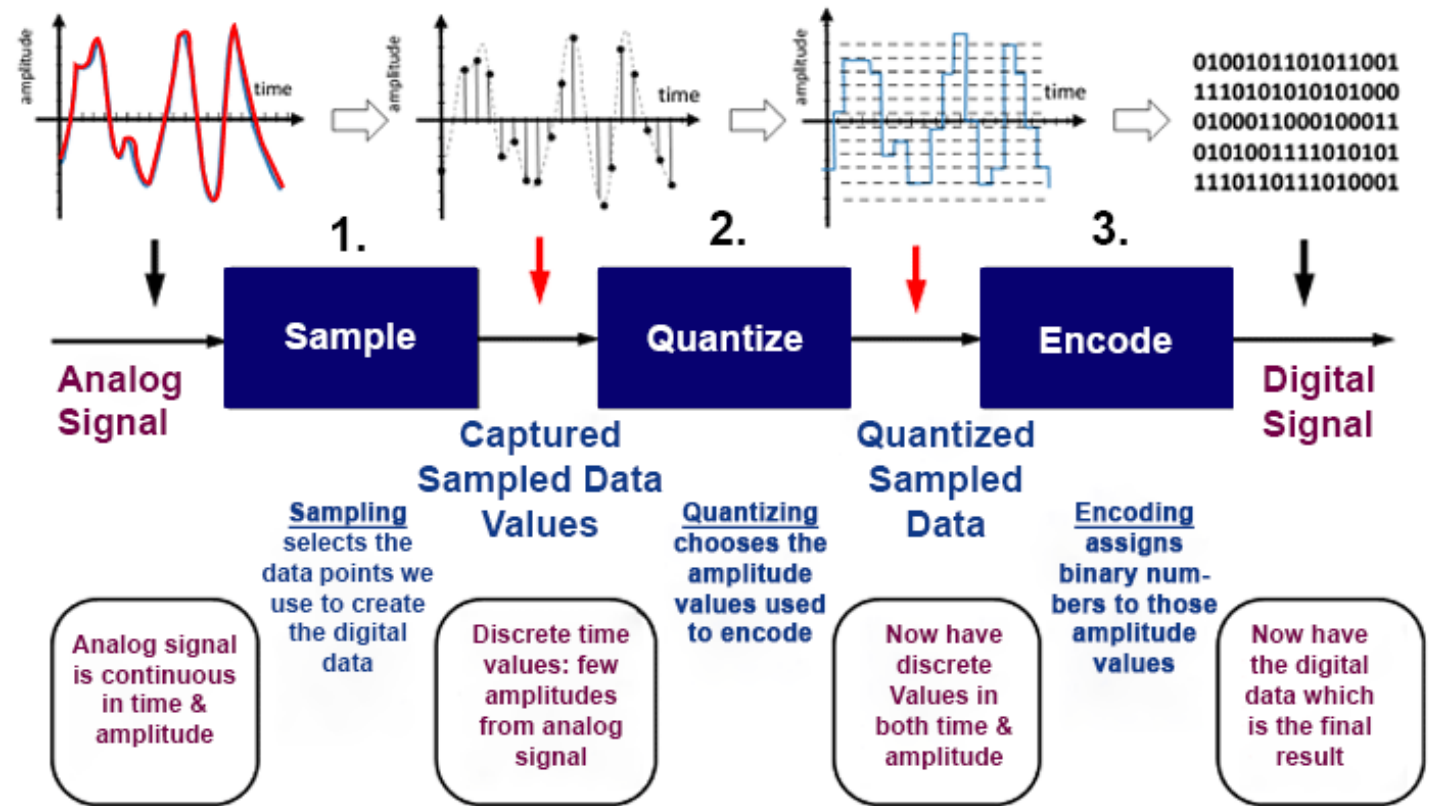
Quantization

- Mapping:

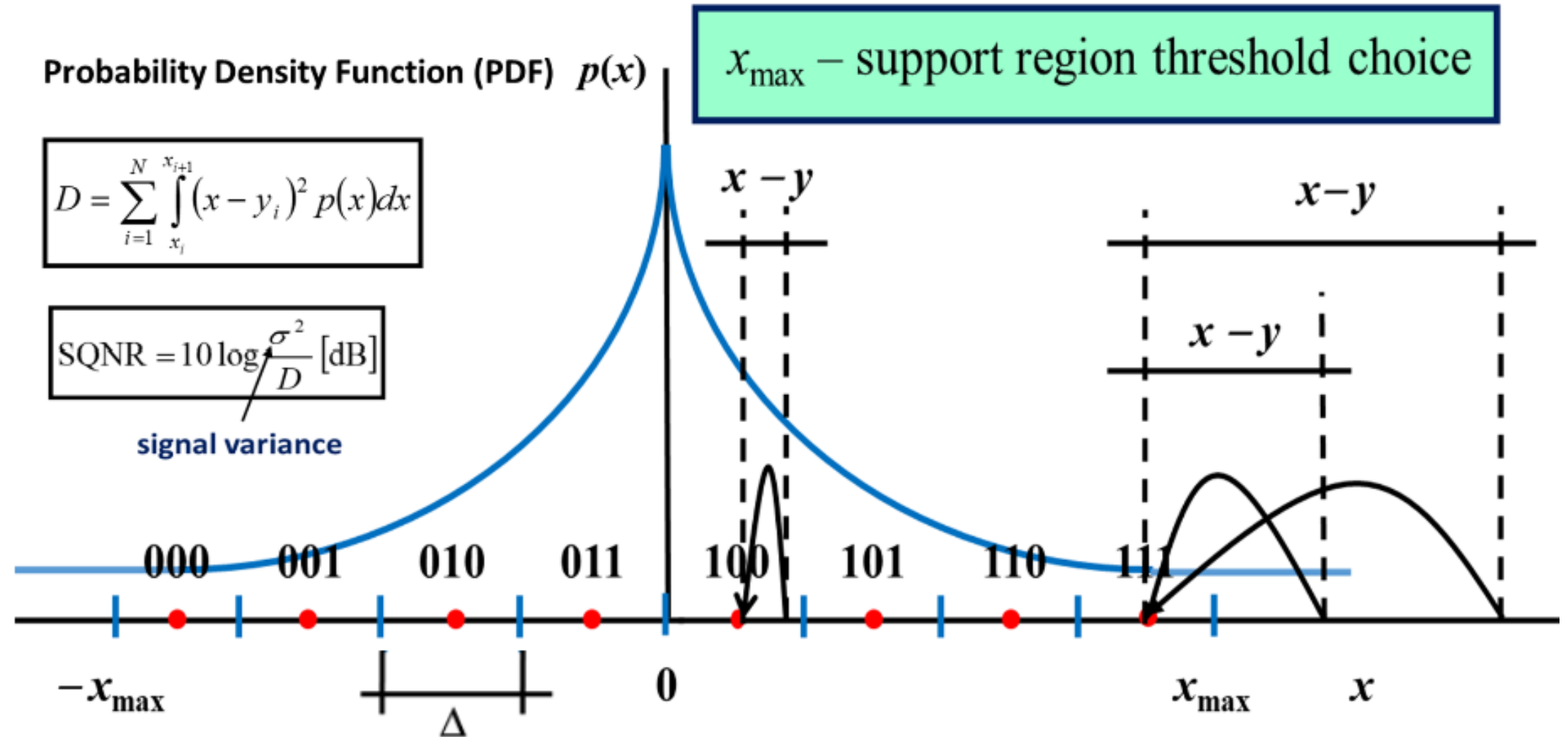
- $\mathbb{R} \rightarrow \{y_1, y_2, \dots, y_N\}$
- $\{y_1, y_2, \dots, y_M\} \rightarrow \{y_1, y_2, \dots, y_N\}, M \gg N$

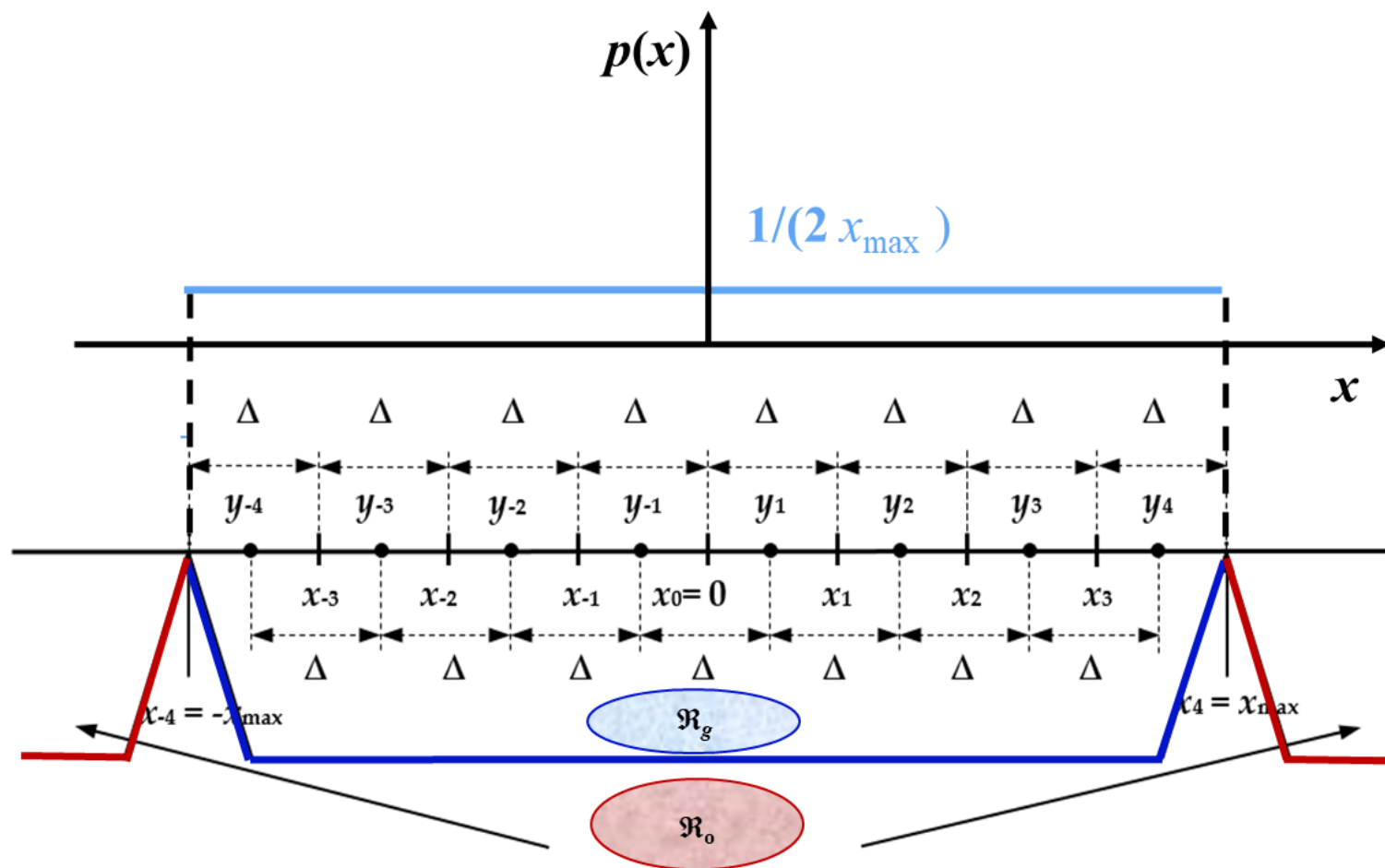
- Types of quantizers:

- Uniform
- Nonuniform
- Piecewise-uniform

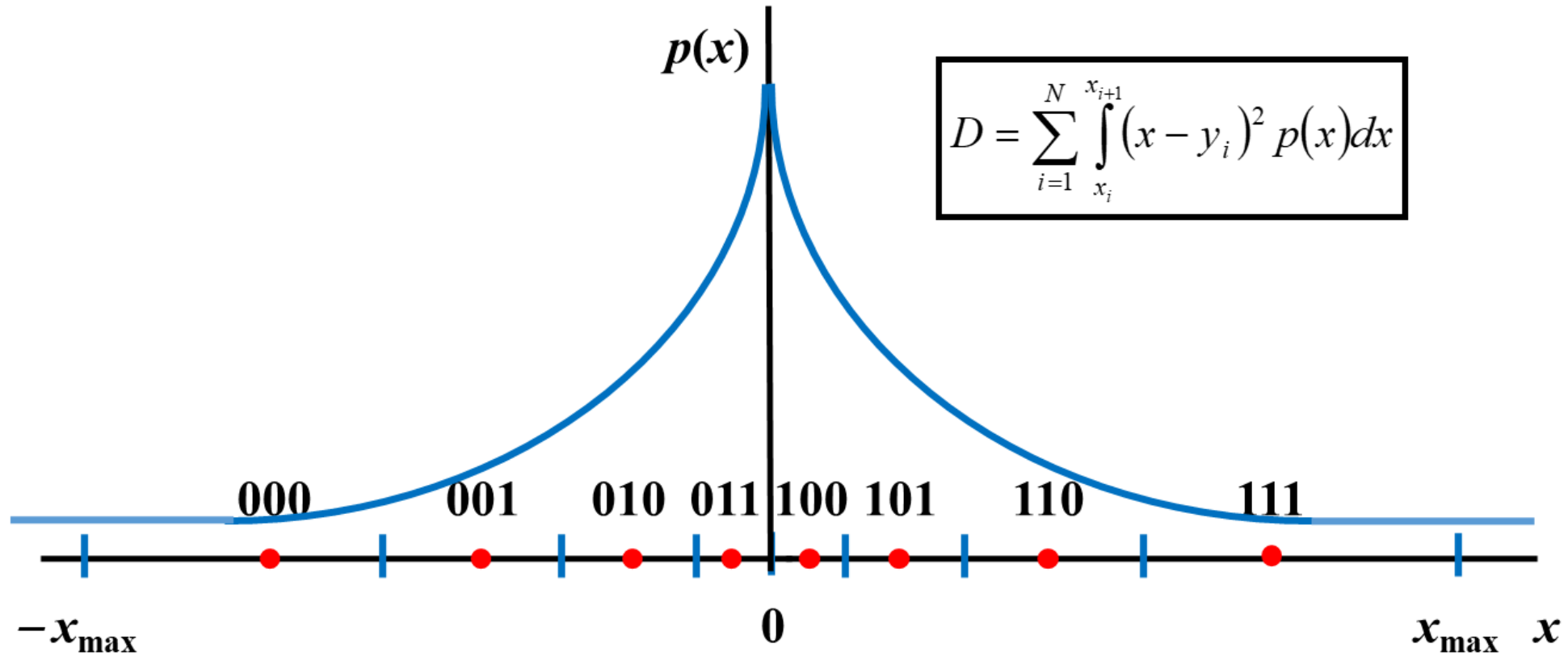


Uniform quantization

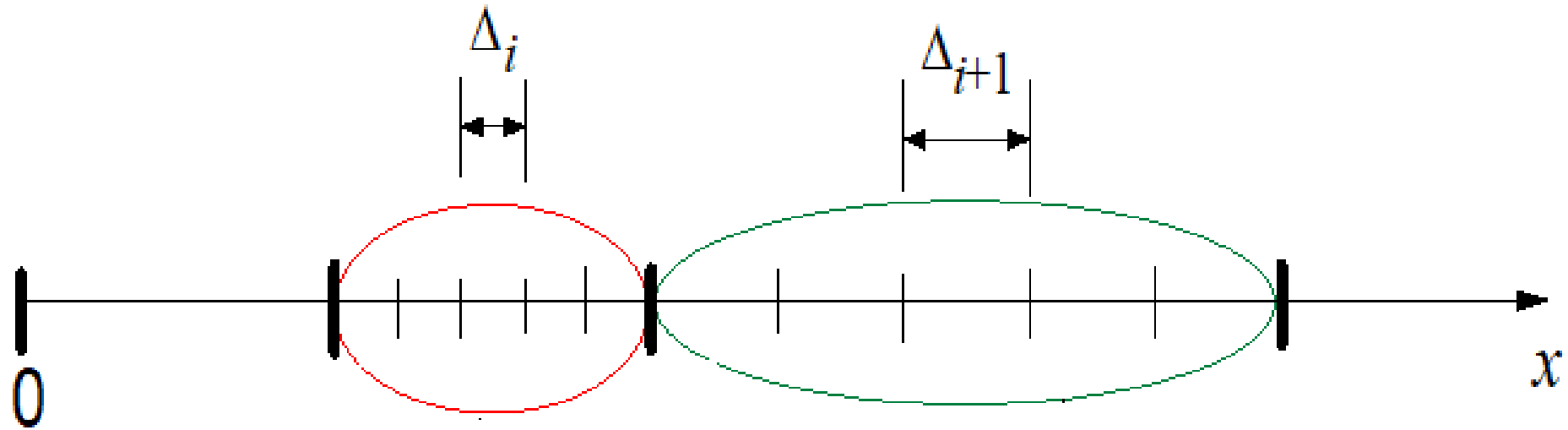




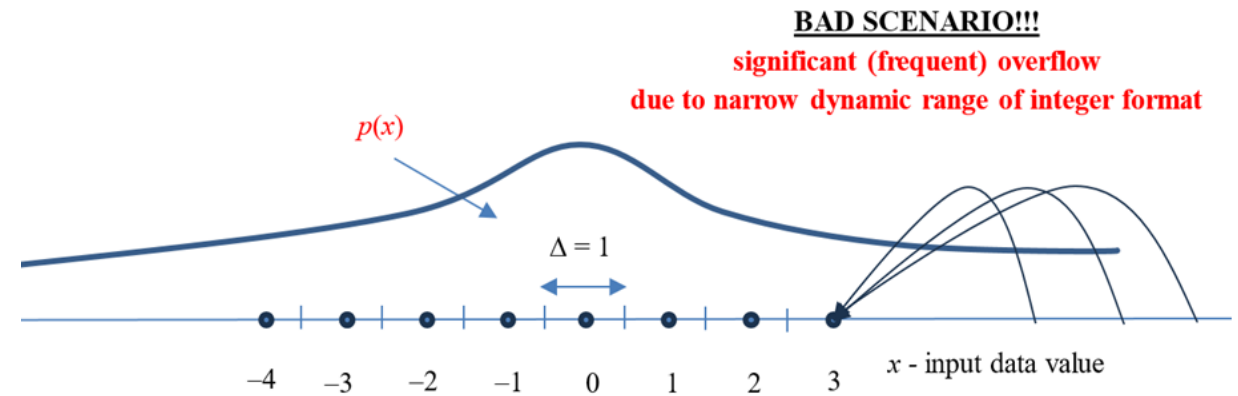
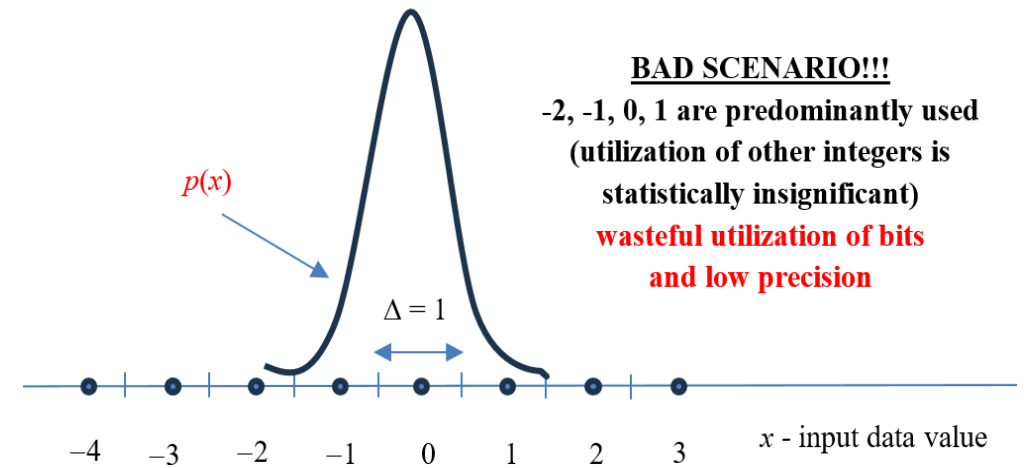
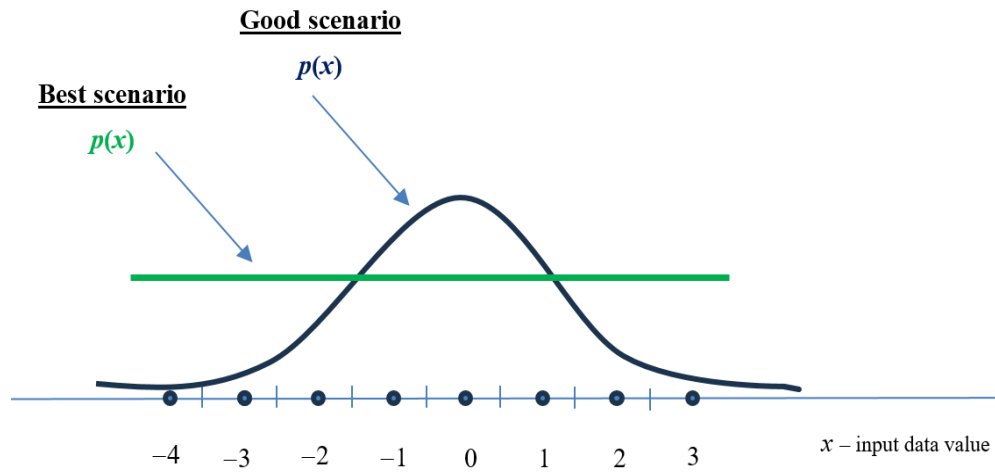
Non-uniform quantization



Piecewise uniform quantization



Quantizer Alignment with the Data PDF



Content

- Quantization
- Digital data formats
 - Fixed-point formats
 - Integer formats
 - Floating-point formats
- Quantization of Neural networks
 - Efficiency
 - Reliability

Fixed-point format

$$x = s a_{n-1} \dots a_1 a_0 . a_{-1} a_{-2} \dots a_{-m}$$

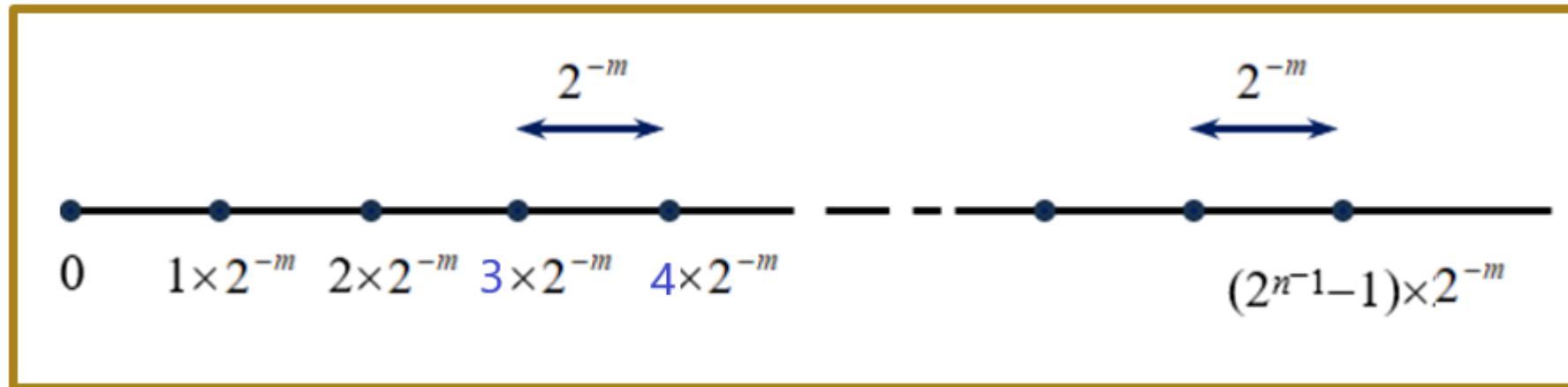
$$x = (-1)^s \sum_{i=-m}^{n-1} a_i 2^i$$

$$(0 \dots 0.0 \dots 01)_2 = 2^{-m}$$

$$(0 \dots 0.0 \dots 010)_2 = 2^{-(m-1)} = 2 \cdot 2^{-m}$$

$$(0 \dots 0.0 \dots 011)_2 = 2^{-(m-1)} + 2^{-m} = 3 \cdot 2^{-m}$$

$$(0 \dots 0.0 \dots 0100)_2 = 2^{-(m-2)} = 4 \cdot 2^{-m}$$



Fixed-point format $\leftrightarrow$ Uniform quantizer

Integer format

fixed-point format: $x = Sa_{n-1} \dots a_1 a_0 . a_{-1} a_{-2} \dots a_{-m}$

for $m = 0$: $x = Sa_{n-1} \dots a_1 a_0$

The integer format is a fixed-point format with $m = 0$.

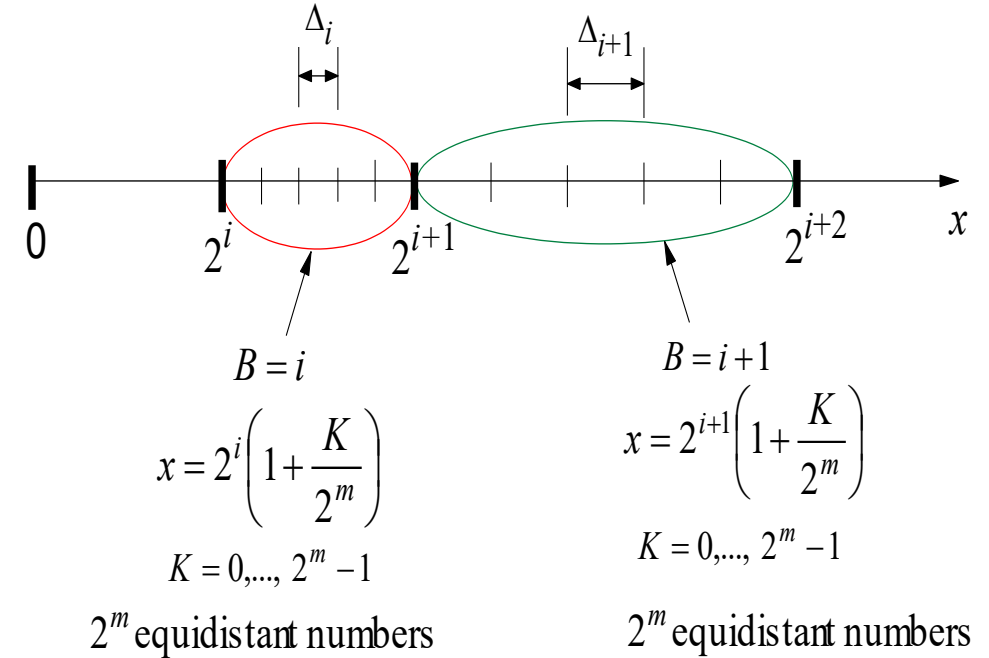
Integer format $\leftrightarrow$ Uniform quantizer

Floating-point format

$$x = (sa_{e-1} \dots a_1 a_0 b_{m-1} \dots b_1 b_0)_2$$

- one bit s to encode the sign of x ,
- e bits to encode the exponent $E = \sum_{i=0}^{e-1} a_i 2^i$
- m bits to represent the significand $K = \sum_{i=1}^m b_{m-i} 2^{m-i}$
- biased exponent: $B = E - bias$

$$x = 2^B \left(1 + \frac{K}{2^m} \right)$$



Floating-point format $\leftrightarrow$ Piecewise uniform quantizer

Conclusion: Different digital data formats represent different types of quantizers.

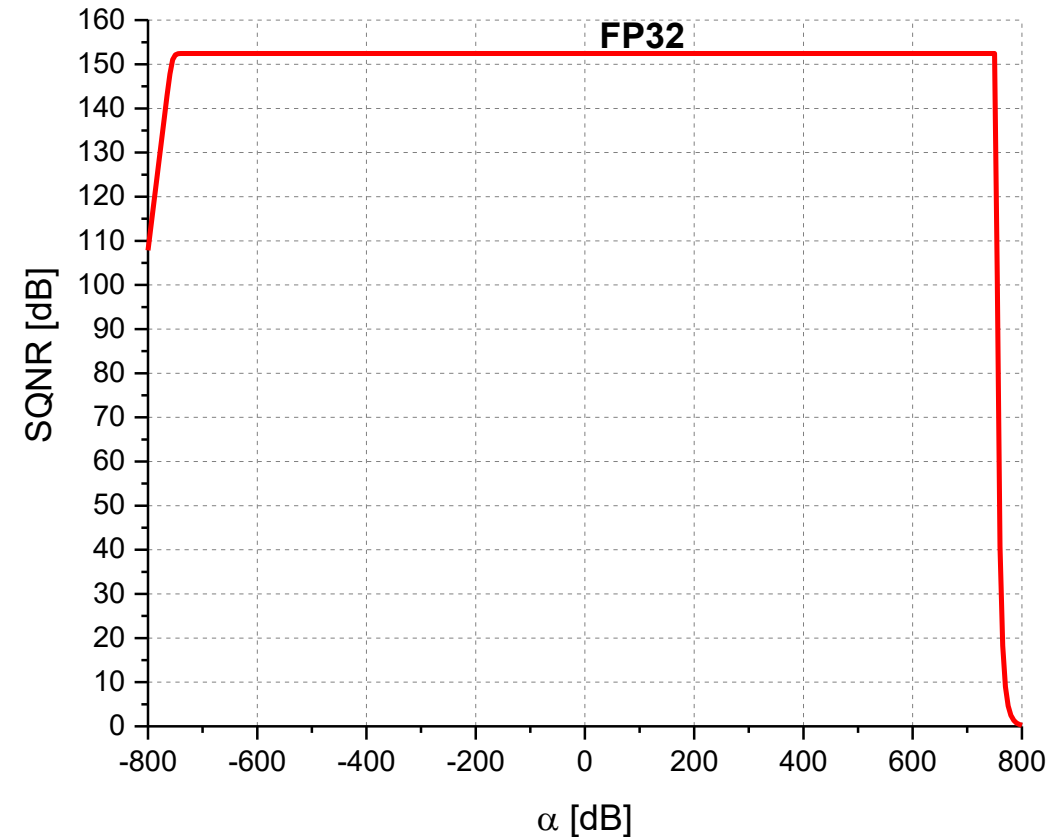
$$\text{SQNR}(\alpha) = -10 \log_{10} \left[\frac{2^{2(B_{\min}-m)}}{12 \cdot 10^{\alpha/10}} \text{erf} \left(\frac{2^{B_{\min}-1/2}}{10^{\alpha/20}} \right) + \sum_{B=B_{\min}}^{B_{\max}} \frac{2^{2(B-m)}}{12 \cdot 10^{\alpha/10}} \left(\text{erf} \left(\frac{2^{B+1/2}}{10^{\alpha/20}} \right) - \text{erf} \left(\frac{2^{B-1/2}}{10^{\alpha/20}} \right) \right) \right. \\ \left. - \frac{x_{\max}}{10^{\alpha/20}} \sqrt{\frac{2}{\pi}} \exp \left(-\frac{x_{\max}^2}{2 \cdot 10^{\alpha/10}} \right) + \left(\frac{x_{\max}^2}{10^{\alpha/10}} + 1 \right) \text{erfc} \left(\frac{x_{\max}}{10^{\alpha/20} \cdot \sqrt{2}} \right) \right]$$

Content

- Quantization
- Digital data formats
 - Fixed-point formats
 - Integer formats
 - Floating-point formats
- Quantization of Neural networks
 - Efficiency
 - Reliability

Quantization of neural networks

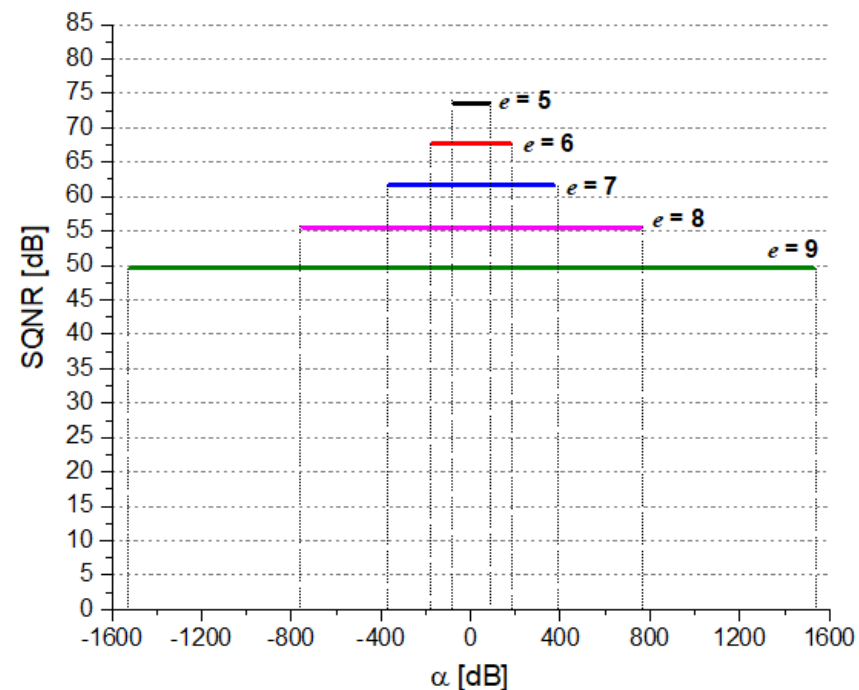
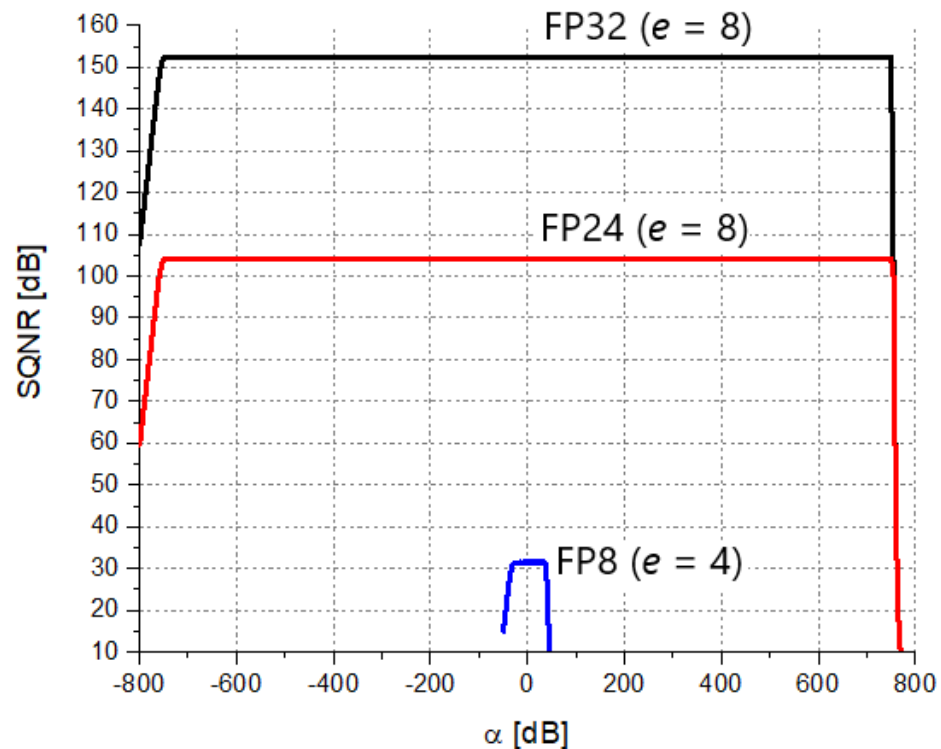
- The 32-bit floating-point format (FP32) is commonly used to represent DNN weights and input data.
- FP32:
 - Advantages: very high precision, very wide dynamic range
 - Disadvantages: very high complexity
- Approaches to reduce complexity:
 - Low-resolution floating-point (FP) formats (e.g., 16-bit, 8-bit, etc.)
 - Fixed-point (FxP) formats
 - Integer formats
 - Power-of-two quantizer



Floating-point formats with lower complexity than FP32

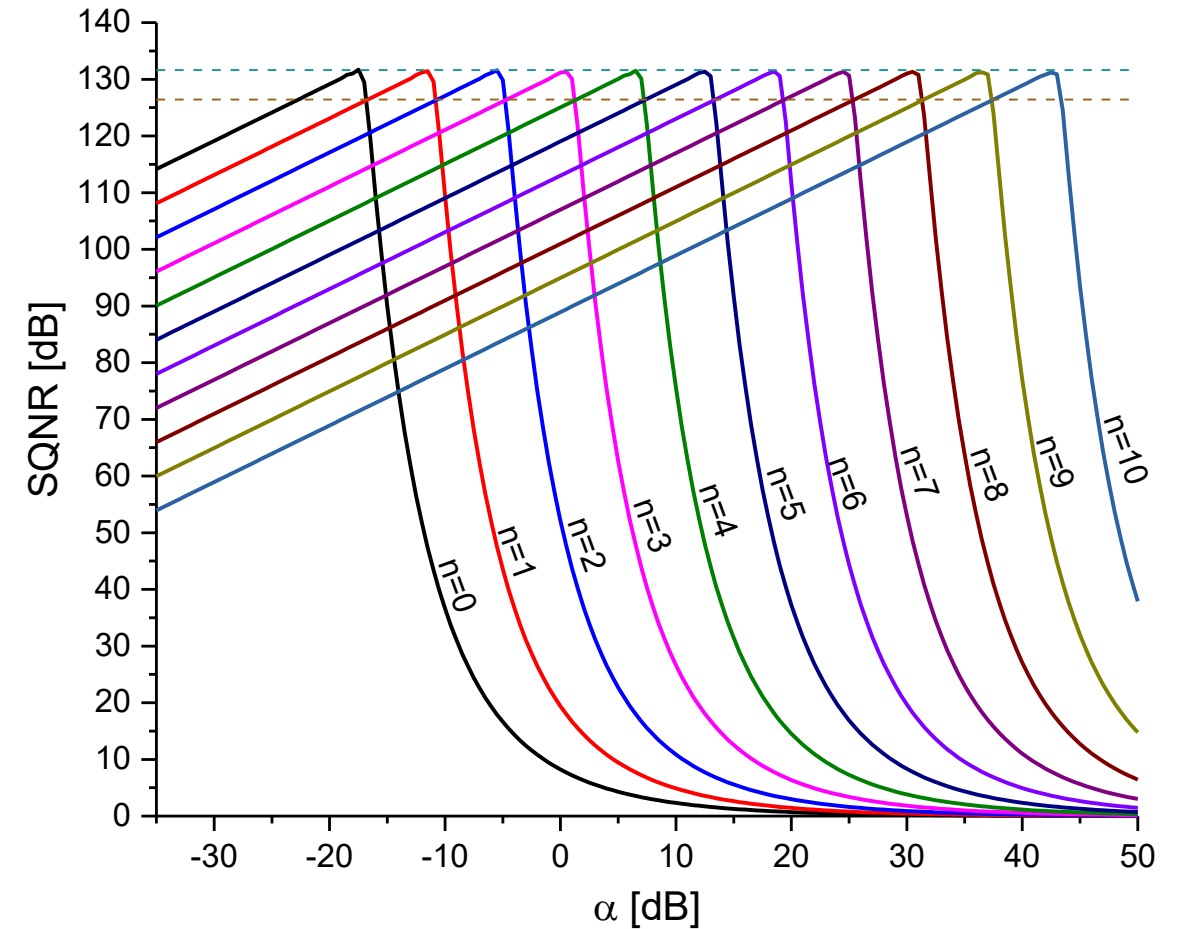
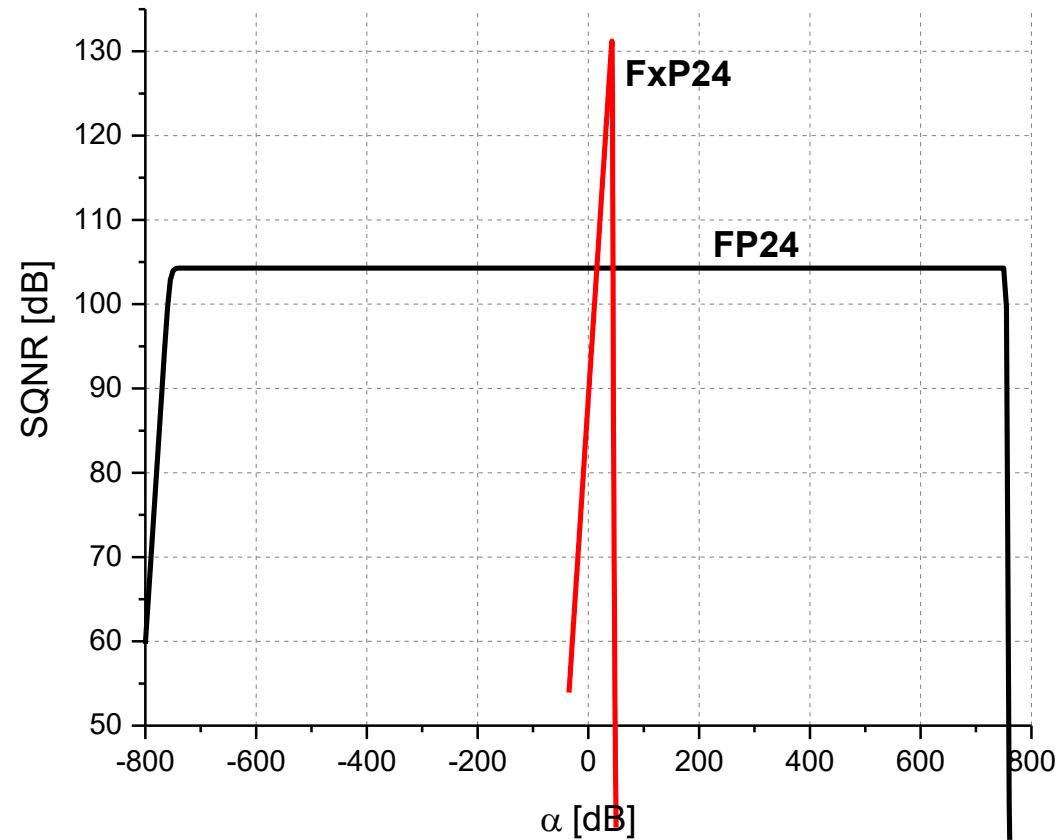
Format	Sign	Exponent	Fraction
FP32	1	8	23
TF32	1	8	10
FP16	1	5	10
bfloat16	1	8	7

Custom FP Formats:



Fixed-point formats

$$x = sa_{n-1} \dots a_1 a_0 . a_{-1} a_{-2} \dots a_{-m}$$



Integer formats

- Direct use ($m = 0$): Usually poor results
- Indirect use: A fixed-point number can be interpreted as an integer multiplied by a fixed scaling factor.

2^{n-f-1}	...	2^1	2^0	2^{-1}	2^{-2}	...	2^{-m}
b_{n-1}	...	b_{f+1}	b_f	b_{f-1}	b_{f-2}	...	b_0

=

2^{n-1}	...	2^{f+1}	2^f	2^{f-1}	2^{f-2}	...	2^0
b_{n-1}	...	b_{f+1}	b_f	b_{f-1}	b_{f-2}	...	b_0

×

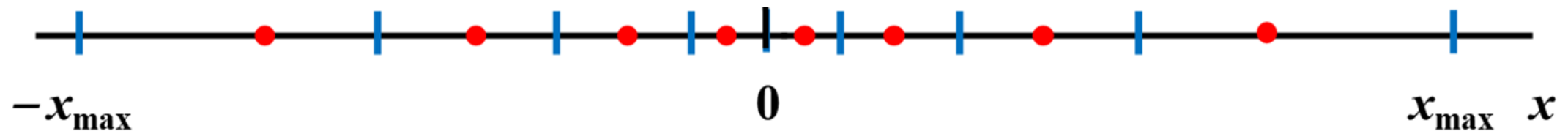
2^{-m}
 scaling factor

integer part of length $n-f$ decimal point fractional part of length m

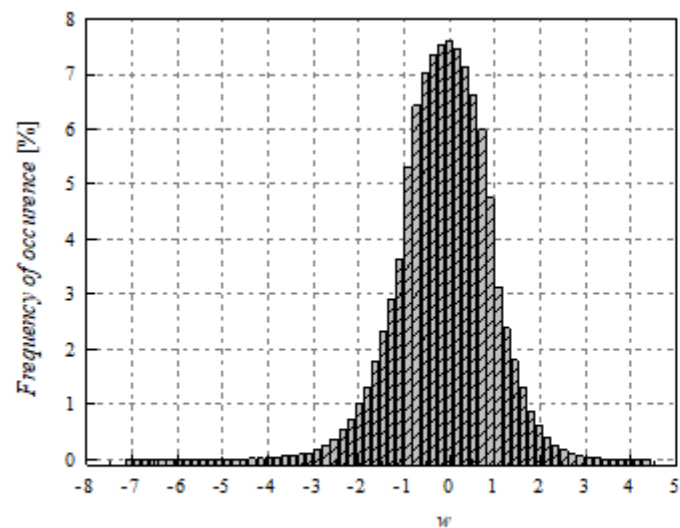
Fixed - point number = Integer × 2^{-m}
scaling factor

Power-of-Two (PoT) quantization

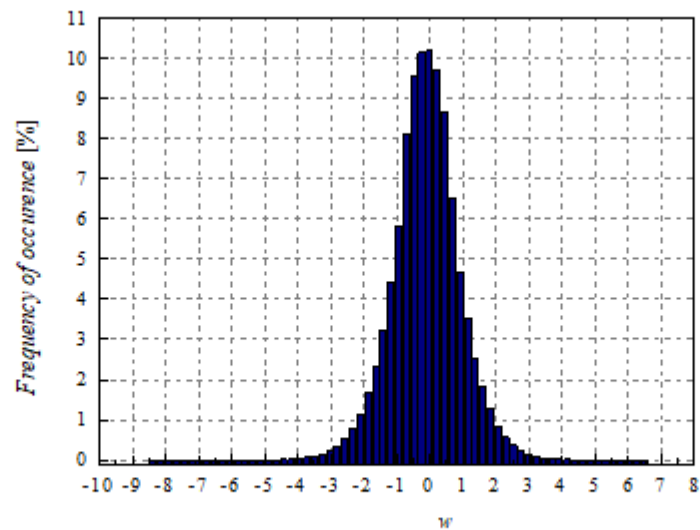
- Non-uniform quantizer
- Representation values are powers of two ($q \pm 2k$)
- Multiplication is replaced by bit-shift operations



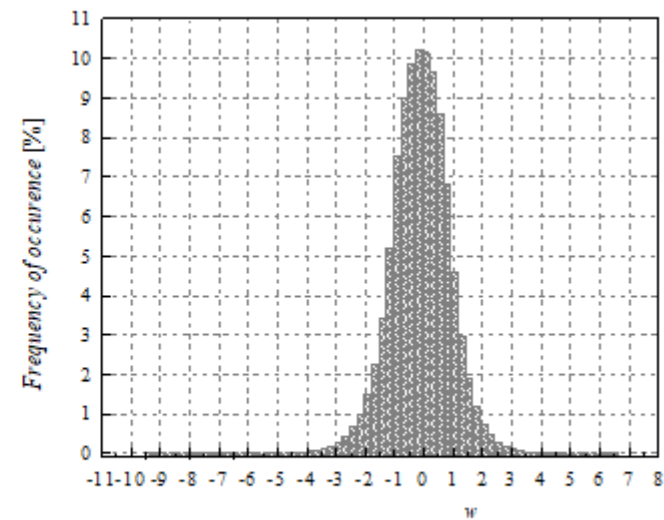
Weights PDF



MLP+MNIST



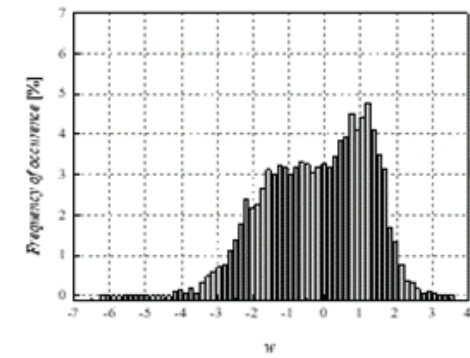
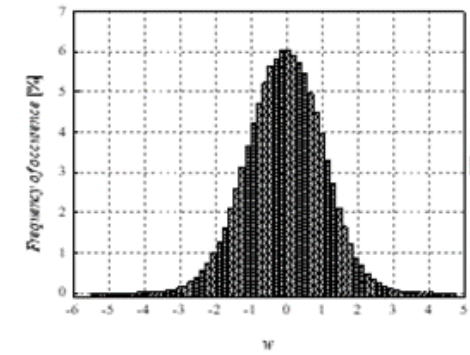
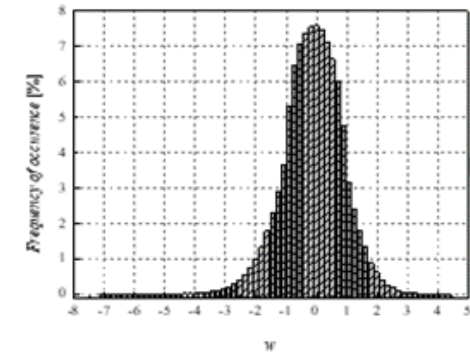
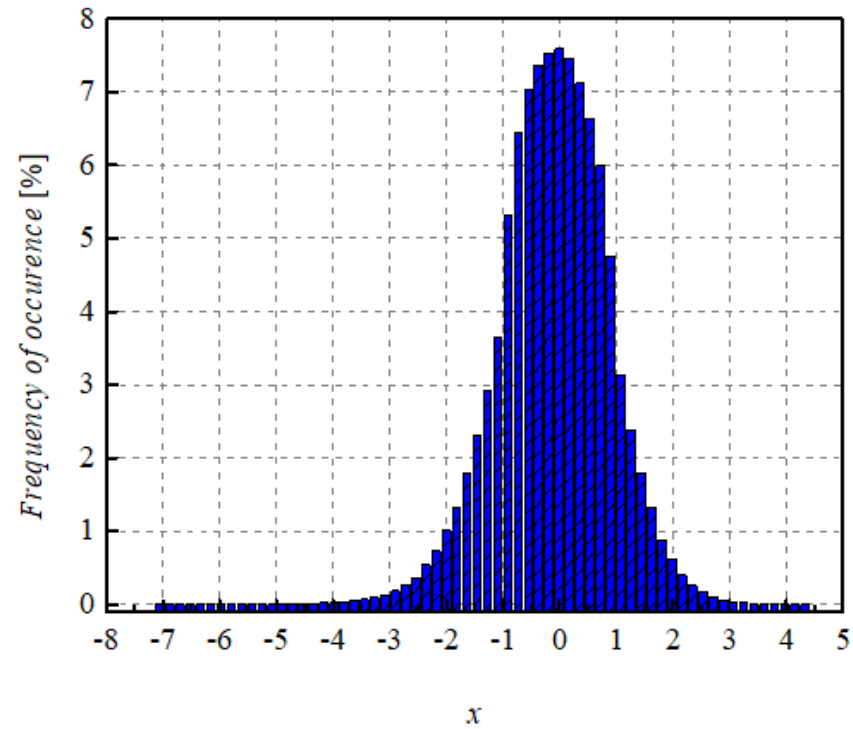
CNN+MNIST



MLP + Fashion-MNIST

Weights across layers

All weights



Quantization and Reliability

- Although the impact of quantization on neural network reliability remains an active area of research and many open questions still exist, quantization has the potential to significantly improve system reliability:
- Reduces the memory area required for storing network parameters.
- Enables memory redundancy (e.g., storing each weight multiple times) without increasing the overall memory footprint.
- Fixed-point representations are inherently more robust to hardware faults than floating-point formats.
- In low-resolution fixed-point formats, the effect of a single-bit error on the represented value is typically smaller, reducing the impact of memory faults on inference accuracy.

Conclusion

- Quantization plays a crucial role in determining both the computational efficiency and the reliability of neural networks.
- The numerical formats used to represent neural network weights, activations, and input data can be interpreted as different types of quantizers.
- Although FP32 remains the standard representation, its high computational and hardware complexity makes it inefficient for Edge AI applications.
- Lower-precision numerical formats offer a more efficient alternative while significantly reducing implementation cost.
- Low-bit fixed-point formats are particularly attractive; however, their parameters must be carefully optimized. Otherwise, variance mismatch can severely degrade network accuracy.
- Quantization theory provides a principled framework for designing application-specific numerical formats, with the statistical characteristics of network weights and activations playing a key role in selecting the optimal representation.

International Symposium on Efficient and Reliable Intelligent Edge Systems

- **Date:** 25–26 May 2027
- **Place:** Niš, Serbia
- **Submission deadline:** 25 January 2027
- **General Co-Chairs:**
 - Milan Dinčić (University of Niš)
 - Davide Bertozzi (University of Manchester)
 - Miloš Krstić (IHP Germany)
- **Program Committee Co-Chairs:**
 - Zoran Perić (University of Niš)
 - Riccardo Zese (University of Ferrara)
 - Oliver Rhodes (University of Manchester)
 - Marko Andjelković (IHP Germany)
- **Submission options:** Authors may submit either a full paper or an abstract only.
- **Journal publications:**
 - up to **5 selected papers** published free of charge in **JLPEA**,
 - up to **10 selected papers** published in **Facta Universitatis**.
- **Registration fee:** €150 for regular participants (both with a full paper or an abstract only) and €100 for PhD students.
- **Topics:**
 - Optimization techniques for Edge AI models
 - Neural architecture search for Edge AI models
 - Quantization and compression of neural networks
 - Approximate computing
 - Dynamic neural network models and implementations
 - Neuromorphic computing for Edge AI
 - Neuromorphic vision
 - Spiking neural networks
 - Hybrid ANN–SNN models
 - In-memory and near-memory computing
 - Reliable electronics for Edge AI
 - Hardware accelerators for Edge AI
 - Hardware/software co-exploration for neural architectures
 - Edge AI applications

THANKS!