

From outcomes to predictive distributions

A model should say which values are possible and how plausible each one is.

A neural network maps an input to the parameters of a distribution for the target.

1 • Start with the random variable

The **support** $\mathcal{S}_Y$ is the set of values the target can take. The support should match the task before you choose a loss.

BINARY	$Y \in \{0, 1\}$
CLASS	$Y \in \{1, \dots, K\}$
COUNT	$Y \in \{0, 1, 2, \dots\}$
REAL	$Y \in \mathbb{R}$

Impossible targets should receive zero probability or density. A support mismatch is a modelling error, not an optimizer problem.

2 • Mass for discrete, density for continuous

PMF	$p_{Y(y)} = P(Y = y) \text{ and } \sum_{y \in \mathcal{S}_Y} p_{Y(y)} = 1.$
PDF	$p_{Y(y)} \geq 0 \text{ and } \int_{\mathcal{S}_Y} p_{Y(y)} dy = 1.$

For continuous Y ,

$$P(a \leq Y \leq b) = \int_a^b p_{Y(y)} dy.$$

A density value can exceed 1; it is the **area** that is a probability.

3 • Four distribution families

Family	Parameters	Support
Bernoulli	$p \in [0, 1]$	$\{0, 1\}$
Categorical	$p_1, \dots, p_K$	$\{1, \dots, K\}$
Poisson	$\lambda > 0$	nonnegative integers
Gaussian	$\mu \in \mathbb{R}, \sigma > 0$	$\mathbb{R}$

4 • Read conditional model notation

	$p_{\theta}(y \mid x)$
x	one observed input
y	one possible or observed target
θ	all learned weights and biases
$p_{\theta(y x)}$	probability or density assigned to y , given x

During supervised training the observed inputs are treated as fixed; the model describes the conditional target distribution.

5 • The network predicts distribution parameters

REG	$x \rightarrow (\mu_{\theta(x)}, \sigma_{\theta(x)})$
BINARY	$x \rightarrow p_{\theta(x)}$
K-CLASS	$x \rightarrow (p_{\theta,1}(x), \dots, p_{\theta,K}(x))$

A point prediction is only a summary: often the conditional mean or mode. Training scores the **full distribution** at the observed target.

6 • Prediction and likelihood ask different questions

Prediction: fix θ and x ; vary the possible outcome y .

Likelihood: fix the observed data; vary candidate parameters θ .

The same expression $p_{\theta(y|x)}$ is read differently because a different quantity is allowed to vary.

Choose the target variable first, then its support, then a distribution family, then the neural outputs that parameterize it.

Build a dataset likelihood

Likelihood ranks parameter settings by how well they explain the outcomes that were actually observed.

1 • Score one fixed observation

For one pair (x^*, y^*) , the likelihood contribution is

$$L_{i(\theta)} = p_{\theta}(y^* | x^*).$$

A larger value means that candidate θ explains this fixed observation better.

For a continuous target this is a density. Comparing candidates is valid even though the density is not itself an event probability.

2 • Independence gives a product

For conditionally independent examples,

$$L(\theta) = p(\mathcal{D} | \theta) = \prod_{i=1}^n p_{\theta}(y_i | x_i).$$

Independent means multiply the per-example factors.

Identically distributed means reuse the same conditional model form.

These are separate assumptions; either can fail without the other.

3 • Tiny Bernoulli comparison

Compare $\theta = 0.5$ and $\theta = 0.8$ where $P(H | \theta) = \theta$.

Data	$L(0.5)$	$L(0.8)$
H, T	$0.5(0.5) = 0.25$	$0.8(0.2) = 0.16$
H, H	$0.5^2 = 0.25$	$0.8^2 = 0.64$

The candidate ranking changes when the fixed observed data change.

4 • MLE names the winner

$$\hat{\theta}_{\text{MLE}} = \arg \max_{\theta} L(\theta).$$

Maximum likelihood estimation chooses the parameter vector that assigns the largest likelihood to the observed dataset.

5 • Logs preserve the ranking

Because log is strictly increasing on positive values,

$$\arg \max_{\theta} L(\theta) = \arg \max_{\theta} \log L(\theta).$$

For independent examples, the product becomes a sum:

$$\log L(\theta) = \sum_i \log p_{\theta}(y_i | x_i).$$

This avoids underflow and makes derivatives easier to combine.

6 • Negative log-likelihood is a loss

Optimizers conventionally minimize, so negate the log-likelihood:

$$\ell_{i(\theta)} = -\log p_{\theta}(y_i | x_i).$$

$$\hat{R}(\theta) = \frac{1}{n} \sum_i \ell_{i(\theta)} = -\frac{1}{n} \log L(\theta).$$

The factor $\frac{1}{n}$ changes gradient scale, not the minimizer.

7 • Support becomes a hard constraint

If the model assigns $p_{\theta}(y_i | x_i) = 0$ to an observed target, then

$$\ell_i = -\log 0 = +\infty.$$

A zero-likelihood observation is an infinite penalty. Check support before blaming numerical instability.

8 • Final audit

- ✓ Target support matches the chosen distribution.
- ✓ Neural outputs obey parameter constraints.
- ✓ Each observation contributes one log-density or log-probability.
- ✓ Independence assumptions are stated, not silently assumed.
- ✓ Likelihood and prediction are not confused.

Probability models possible outcomes. Likelihood scores parameters using the outcome that occurred.