

What are we about to learn?

How can a trained network reliably execute an addition algorithm?

In the neural tangent kernel (NTK) limit:

$$\underbrace{f_{\infty}(x)}_{\text{trained prediction}} = \underbrace{\mu(x)}_{\text{fixed mean}} + \underbrace{g(x)}_{\text{Gaussian remainder}} .$$

Choose instruction data so the mean gives the correct bit decisions.

Average independent networks so fluctuations rarely change those decisions.

Worked example: learn the updates for $10_2 + 11_2 = 101_2$.

Outline: the three parts of the argument

We will connect training, instruction data, and reliable computation.

1. **Derive the trained prediction.**

Gradient flow $\longrightarrow$ fixed NTK $\longrightarrow$ explicit solution.

Explain why random initialization leaves a Gaussian remainder.

2. **Make the mean execute an instruction.**

Encode input conditions in orthogonal directions.

Use the full output labels to recover the two-bit addition update.

3. **Bound the probability of an incorrect execution.**

Reduce variance by averaging independent trained networks.

Choose the ensemble size to protect every predicted bit.

The endpoint is an error guarantee for fixed inputs and a fixed execution horizon, under the stated conditions.

Gradient flow is continuous gradient descent

For one output: $f_{\theta}(x) = m^{-1/2} W_2 \text{ReLU}(W_1 x)$, with trainable $\theta = (W_1, W_2)$.

$x \in \mathbb{R}^d$, $W_1 \in \mathbb{R}^{m \times d}$, $W_2 \in \mathbb{R}^{1 \times m}$. m is the number of hidden neurons.

On fixed training data (X, y) , follow the squared-loss gradient:

$$L(\theta) = \frac{1}{2} \|f_{\theta}(X) - y\|^2, \quad \frac{d\theta_t}{dt} = -\nabla_{\theta} L(\theta_t).$$

Gradient flow is continuous gradient descent

For one output: $f_{\theta}(x) = m^{-1/2} W_2 \text{ReLU}(W_1 x)$, with trainable $\theta = (W_1, W_2)$.
 $x \in \mathbb{R}^d$, $W_1 \in \mathbb{R}^{m \times d}$, $W_2 \in \mathbb{R}^{1 \times m}$. m is the number of hidden neurons.

On fixed training data (X, y) , follow the squared-loss gradient:

$$L(\theta) = \frac{1}{2} \|f_{\theta}(X) - y\|^2, \quad \frac{d\theta_t}{dt} = -\nabla_{\theta} L(\theta_t).$$

Gradient descent takes one downhill step at a time:

$$\theta^{k+1} = \theta^k - \eta \nabla_{\theta} L(\theta^k), \quad t_k = k\eta.$$

Gradient flow is continuous gradient descent

For one output: $f_{\theta}(x) = m^{-1/2} W_2 \text{ReLU}(W_1 x)$, with trainable $\theta = (W_1, W_2)$.
 $x \in \mathbb{R}^d$, $W_1 \in \mathbb{R}^{m \times d}$, $W_2 \in \mathbb{R}^{1 \times m}$. m is the number of hidden neurons.

On fixed training data (X, y) , follow the squared-loss gradient:

$$L(\theta) = \frac{1}{2} \|f_{\theta}(X) - y\|^2, \quad \frac{d\theta_t}{dt} = -\nabla_{\theta} L(\theta_t).$$

Gradient descent takes one downhill step at a time:

$$\theta^{k+1} = \theta^k - \eta \nabla_{\theta} L(\theta^k), \quad t_k = k\eta.$$
$$\frac{\theta^{k+1} - \theta^k}{\eta} \xrightarrow[\eta \rightarrow 0]{\text{smaller time steps}} \frac{d\theta_t}{dt}.$$

Smaller steps, taken more often, trace a continuous downhill path.

How do the predictions change?

Write $f_t = f_{\theta_t}$. Apply the chain rule, then substitute gradient flow:

$$\frac{d f_t(x)}{d t} = \nabla_{\theta} f_t(x)^{\top} \frac{d \theta_t}{d t}$$

How do the predictions change?

Write $f_t = f_{\theta_t}$. Apply the chain rule, then substitute gradient flow:

$$\frac{d f_t(x)}{d t} = \nabla_{\theta} f_t(x)^{\top} \frac{d \theta_t}{d t} \stackrel{\text{gradient flow}}{=} -\nabla_{\theta} f_t(x)^{\top} \nabla_{\theta} L(\theta_t).$$

How do the predictions change?

Write $f_t = f_{\theta_t}$. Apply the chain rule, then substitute gradient flow:

$$\frac{d f_t(x)}{d t} = \nabla_{\theta} f_t(x)^{\top} \frac{d \theta_t}{d t} \stackrel{\text{gradient flow}}{=} -\nabla_{\theta} f_t(x)^{\top} \nabla_{\theta} L(\theta_t).$$

For squared loss, the remaining gradient is

$$\nabla_{\theta} L(\theta_t) = \sum_i (f_t(x_i) - y_i) \nabla_{\theta} f_t(x_i).$$

How do the predictions change?

Write $f_t = f_{\theta_t}$. Apply the chain rule, then substitute gradient flow:

$$\frac{d f_t(x)}{d t} = \nabla_{\theta} f_t(x)^{\top} \frac{d \theta_t}{d t} \stackrel{\text{gradient flow}}{=} -\nabla_{\theta} f_t(x)^{\top} \nabla_{\theta} L(\theta_t).$$

For squared loss, the remaining gradient is

$$\nabla_{\theta} L(\theta_t) = \sum_i (f_t(x_i) - y_i) \nabla_{\theta} f_t(x_i).$$

$$\frac{d f_t(x)}{d t} = -\Theta_t(x, X) (f_t(X) - y).$$

$$\Theta_t(u, v) = \langle \nabla_{\theta} f_t(u), \nabla_{\theta} f_t(v) \rangle \quad (\text{neural tangent kernel}).$$

Where infinite width enters the calculation

On the training data, set $r_t = f_t(X) - y$. The equation becomes

$$\frac{dr_t}{dt} = -\Theta_t(X, X) r_t.$$

This coefficient depends on the unknown, changing weights.

Where infinite width enters the calculation

On the training data, set $r_t = f_t(X) - y$. The equation becomes

$$\frac{dr_t}{dt} = -\Theta_t(X, X) r_t.$$

This coefficient depends on the unknown, changing weights.

By the law of large numbers and NTK stability:

$$\Theta_t(u, v) \xrightarrow{m \rightarrow \infty} \Theta(u, v).$$

Θ is non-random and constant in training time t .

Where infinite width enters the calculation

On the training data, set $r_t = f_t(X) - y$. The equation becomes

$$\frac{dr_t}{dt} = -\Theta_t(X, X) r_t.$$

This coefficient depends on the unknown, changing weights.

By the law of large numbers and NTK stability:

$$\Theta_t(u, v) \xrightarrow{m \rightarrow \infty} \Theta(u, v).$$

Θ is non-random and constant in training time t .

$$\frac{dr_t}{dt} = -K r_t$$

$$K = \Theta(X, X).$$

Where infinite width enters the calculation

On the training data, set $r_t = f_t(X) - y$. The equation becomes

$$\frac{dr_t}{dt} = -\Theta_t(X, X) r_t.$$

This coefficient depends on the unknown, changing weights.

By the law of large numbers and NTK stability:

$$\Theta_t(u, v) \xrightarrow{m \rightarrow \infty} \Theta(u, v).$$

Θ is non-random and constant in training time t .

$$\frac{dr_t}{dt} = -K r_t \quad \implies \quad r_t = e^{-Kt} r_0, \quad K = \Theta(X, X).$$

Integrate to obtain the prediction

Insert $r_t = e^{-Kt}r_0$ into the prediction equation:

$$\frac{d f_t(x)}{d t} = -\Theta(x, X)e^{-Kt}r_0.$$

Integrate to obtain the prediction

Insert $r_t = e^{-Kt}r_0$ into the prediction equation:

$$\frac{d f_t(x)}{d t} = -\Theta(x, X)e^{-Kt}r_0.$$

Integrate both sides from time 0 to time t :

$$f_t(x) - f_0(x) = - \underbrace{\Theta(x, X) \int_0^t e^{-Ks} ds}_{h_t(x)} r_0.$$

Integrate to obtain the prediction

Insert $r_t = e^{-Kt}r_0$ into the prediction equation:

$$\frac{d f_t(x)}{d t} = -\Theta(x, X) e^{-Kt} r_0.$$

Integrate both sides from time 0 to time t :

$$f_t(x) - f_0(x) = - \underbrace{\Theta(x, X) \int_0^t e^{-Ks} d s}_{h_t(x)} r_0.$$

$$f_t(x) = f_0(x) + h_t(x) (y - f_0(X)).$$

$h_t(x)$ depends on the inputs and time, not on the random initialization.

Use $r_0 = f_0(X) - y$. Lee et al., NeurIPS 2019, Eqs. (9) to (11).

Evaluate the integral: an explicit prediction

The matrix $K = \Theta(X, X)$ is fixed and positive definite, so K^{-1} exists.

$$\frac{d}{ds} \left(-K^{-1} e^{-Ks} \right) = e^{-Ks}.$$

Evaluate the integral: an explicit prediction

The matrix $K = \Theta(X, X)$ is fixed and positive definite, so K^{-1} exists.

$$\frac{d}{ds} \left(-K^{-1} e^{-Ks} \right) = e^{-Ks}.$$
$$\int_0^t e^{-Ks} ds = \left[-K^{-1} e^{-Ks} \right]_{s=0}^{s=t} = K^{-1} (I - e^{-Kt}).$$

Evaluate the integral: an explicit prediction

The matrix $K = \Theta(X, X)$ is fixed and positive definite, so K^{-1} exists.

$$\begin{aligned}\frac{d}{ds} \left(-K^{-1} e^{-Ks} \right) &= e^{-Ks}. \\ \int_0^t e^{-Ks} ds &= \left[-K^{-1} e^{-Ks} \right]_{s=0}^{s=t} = K^{-1} (I - e^{-Kt}).\end{aligned}$$

Substitute into the prediction from the previous slide:

$$f_t(x) = f_0(x) + \underbrace{\Theta(x, X) K^{-1} (I - e^{-Kt})}_{h_t(x)} (y - f_0(X)).$$

Evaluate the integral: an explicit prediction

The matrix $K = \Theta(X, X)$ is fixed and positive definite, so K^{-1} exists.

$$\begin{aligned}\frac{d}{ds} \left(-K^{-1} e^{-Ks} \right) &= e^{-Ks}. \\ \int_0^t e^{-Ks} ds &= \left[-K^{-1} e^{-Ks} \right]_{s=0}^{s=t} = K^{-1} (I - e^{-Kt}).\end{aligned}$$

Substitute into the prediction from the previous slide:

$$f_t(x) = f_0(x) + \underbrace{\Theta(x, X) K^{-1} (I - e^{-Kt})}_{h_t(x)} (y - f_0(X)).$$

$$t \rightarrow \infty : \quad e^{-Kt} \rightarrow 0, \quad f_\infty(x) = f_0(x) + \Theta(x, X) K^{-1} (y - f_0(X)).$$

I is the identity matrix. Infinite-width NTK, squared loss. Evaluation of the preceding integral. Lee et al., 2019, Eqs. (9) to (11).

Regroup the trained prediction

Write $h(x) = \lim_{t \rightarrow \infty} h_t(x) = \Theta(x, X)K^{-1}$. The previous result is

$$f_{\infty}(x) = f_0(x) + h(x)(y - f_0(X)).$$

Regroup the trained prediction

Write $h(x) = \lim_{t \rightarrow \infty} h_t(x) = \Theta(x, X)K^{-1}$. The previous result is

$$f_{\infty}(x) = f_0(x) + h(x)(y - f_0(X)).$$

$$f_{\infty}(x) = \underbrace{h(x)y}_{\text{fixed part}} + \underbrace{[f_0(x) - h(x)f_0(X)]}_{g(x): \text{ depends on initialization}}$$

Regroup the trained prediction

Write $h(x) = \lim_{t \rightarrow \infty} h_t(x) = \Theta(x, X)K^{-1}$. The previous result is

$$f_{\infty}(x) = f_0(x) + h(x)(y - f_0(X)).$$

$$f_{\infty}(x) = \underbrace{h(x)y}_{\text{fixed part}} + \underbrace{[f_0(x) - h(x)f_0(X)]}_{g(x): \text{ depends on initialization}}$$

Evaluate at all training inputs X :

$$h(X) = \Theta(X, X)K^{-1} = KK^{-1} = I.$$

Regroup the trained prediction

Write $h(x) = \lim_{t \rightarrow \infty} h_t(x) = \Theta(x, X)K^{-1}$. The previous result is

$$f_{\infty}(x) = f_0(x) + h(x)(y - f_0(X)).$$

$$f_{\infty}(x) = \underbrace{h(x)y}_{\text{fixed part}} + \underbrace{[f_0(x) - h(x)f_0(X)]}_{g(x): \text{ depends on initialization}}$$

Evaluate at all training inputs X :

$$h(X) = \Theta(X, X)K^{-1} = KK^{-1} = I.$$

$$f_{\infty}(X) = y + [f_0(X) - f_0(X)] = y.$$

The initial predictions cancel, so $g(X) = 0$.

Why is the remaining term Gaussian?

The coefficients $1, -h_1(x), \dots, -h_n(x)$ are non-random.

$$g(x) = f_0(x) - h_1(x)f_0(x_1) - \dots - h_n(x)f_0(x_n).$$

Why is the remaining term Gaussian?

The coefficients $1, -h_1(x), \dots, -h_n(x)$ are non-random.

$$g(x) = f_0(x) - h_1(x)f_0(x_1) - \dots - h_n(x)f_0(x_n).$$

Initialize $\theta_0 \sim \mathcal{N}(0, I)$. Each f_0 sums independent neuron contributions.

Why is the remaining term Gaussian?

The coefficients $1, -h_1(x), \dots, -h_n(x)$ are non-random.

$$g(x) = f_0(x) - h_1(x)f_0(x_1) - \dots - h_n(x)f_0(x_n).$$

Initialize $\theta_0 \sim \mathcal{N}(0, I)$. Each f_0 sums independent neuron contributions.

Central limit theorem, as $m \rightarrow \infty$: the initial outputs form a Gaussian vector.

$$\left(f_0(x), f_0(x_1), \dots, f_0(x_n)\right)^\top \sim \mathcal{N}(0, C).$$

Why is the remaining term Gaussian?

The coefficients $1, -h_1(x), \dots, -h_n(x)$ are non-random.

$$g(x) = f_0(x) - h_1(x)f_0(x_1) - \dots - h_n(x)f_0(x_n).$$

Initialize $\theta_0 \sim \mathcal{N}(0, I)$. Each f_0 sums independent neuron contributions.

Central limit theorem, as $m \rightarrow \infty$: the initial outputs form a Gaussian vector.

$$\left(f_0(x), f_0(x_1), \dots, f_0(x_n)\right)^\top \sim \mathcal{N}(0, C).$$

The fixed weighted sum in the blue box is therefore Gaussian:

$$g(x) \sim \mathcal{N}\left(0, \Sigma(x)\right), \quad \Sigma(x) = \text{Var}(g(x)).$$

Corollary: the trained prediction is Gaussian

We have proved $g(x) \sim \mathcal{N}(0, \Sigma(x))$. Now add the fixed part:

$$f_{\infty}(x) = \underbrace{h(x)y}_{\text{fixed}} + \underbrace{g(x)}_{\text{Gaussian, mean zero}}$$

Corollary: the trained prediction is Gaussian

We have proved $g(x) \sim \mathcal{N}(0, \Sigma(x))$. Now add the fixed part:

$$f_{\infty}(x) = \underbrace{h(x)y}_{\text{fixed}} + \underbrace{g(x)}_{\text{Gaussian, mean zero}}$$

Adding $h(x)y$ shifts the mean and leaves the variance unchanged.

$$f_{\infty}(x) \sim \mathcal{N}(\mu(x), \Sigma(x)).$$
$$\mu(x) = h(x)y, \quad \Sigma(x) = \text{Var}(g(x)).$$

Corollary: the trained prediction is Gaussian

We have proved $g(x) \sim \mathcal{N}(0, \Sigma(x))$. Now add the fixed part:

$$f_{\infty}(x) = \underbrace{h(x)y}_{\text{fixed}} + \underbrace{g(x)}_{\text{Gaussian, mean zero}}$$

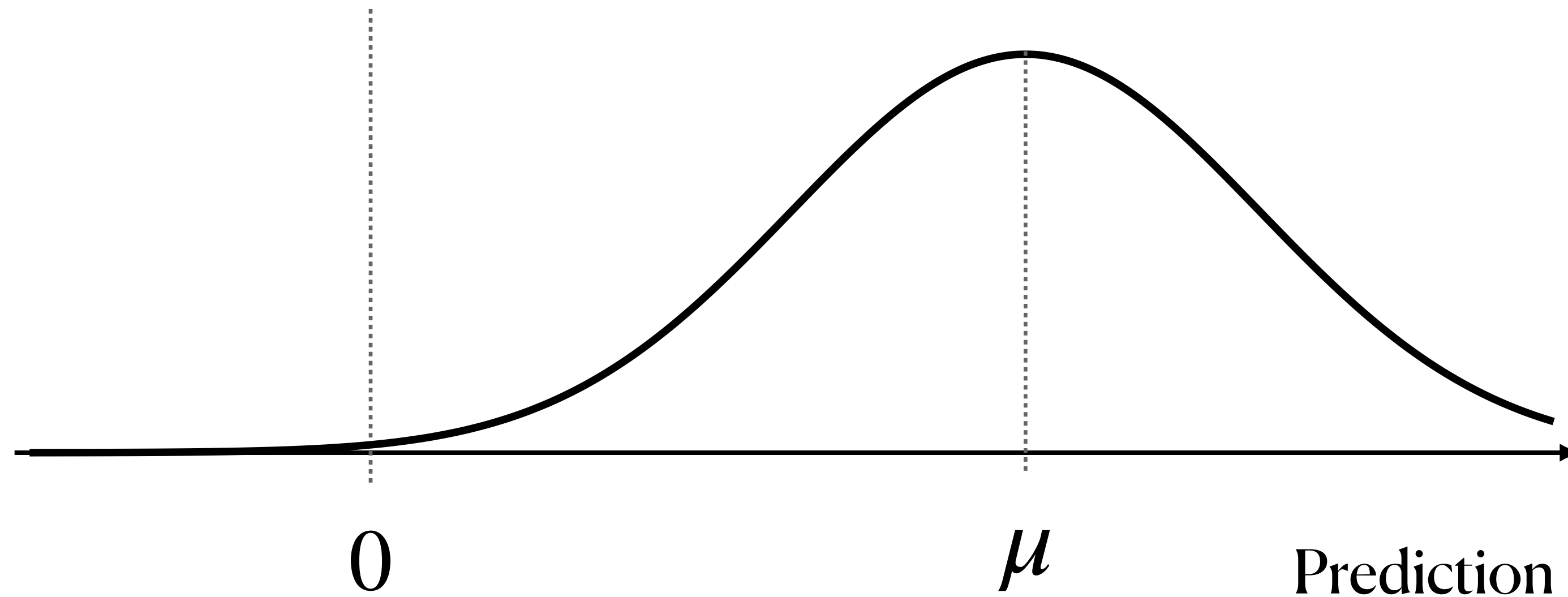
Adding $h(x)y$ shifts the mean and leaves the variance unchanged.

$$f_{\infty}(x) \sim \mathcal{N}(\mu(x), \Sigma(x)).$$
$$\mu(x) = h(x)y, \quad \Sigma(x) = \text{Var}(g(x)).$$

At training inputs, $\Sigma(x_i) = 0$. At new inputs, the variance can remain.

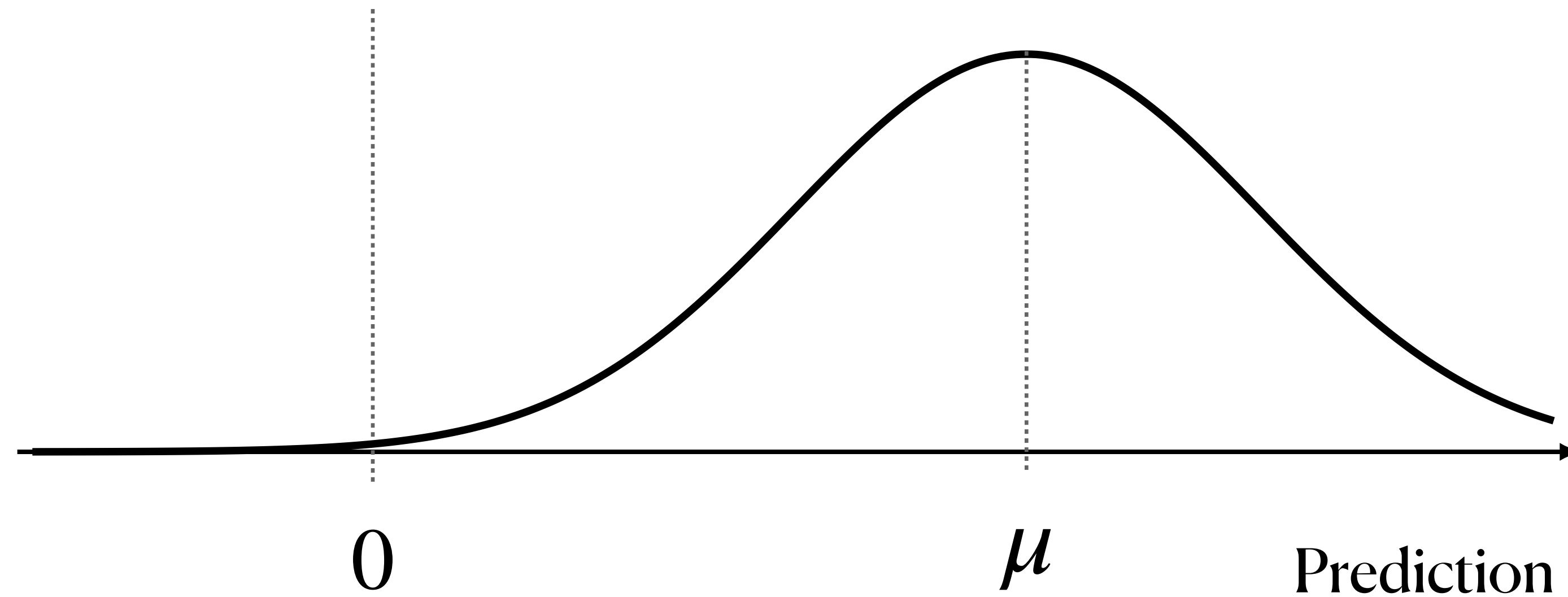
A Gaussian prediction is not enough

For a bit that should be 1, the prediction must stay above zero.



A Gaussian prediction is not enough

For a bit that should be 1, the prediction must stay above zero.



The mean:

$$\mu > 0$$

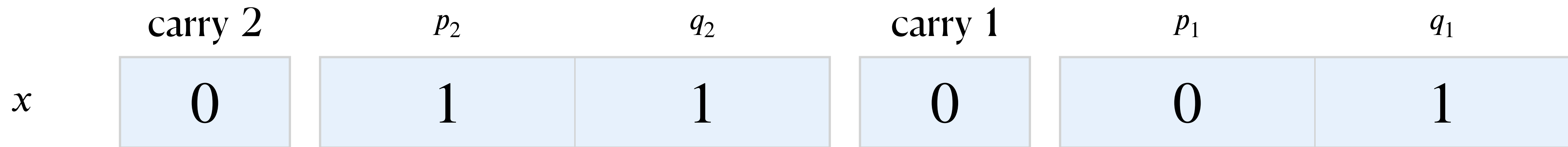
The spread:

$$\sigma \ll |\mu|$$

The data control the mean. Averaging controls the remaining spread.

The two-bit state has four blocks

Follow the paper's example: $p = 10_2$, $q = 11_2$, with both carries initially zero.



$$x = (0 \mid 1 \ 1 \mid 0 \mid 0 \ 1)^\top \in \mathbb{R}^6.$$

Two blocks hold summand bits. Two blocks hold carries.

The eight data points as full input–output vectors

Each training example isolates one instruction in the six-bit state.

Data point	Raw input $u_r^\top$ carry 2 p_2 q_2 carry 1 p_1 q_1	Complete output $y^{(r)\top}$ carry 2' p'_2 q'_2 carry 1' p'_1 q'_1
1	$(1 \mid 0 \ 0 \mid 0 \mid 0 \ 0)$	$(1 \mid 0 \ 0 \mid 0 \mid 0 \ 0)$
2	$(0 \mid 1 \ 1 \mid 0 \mid 0 \ 0)$	$(1 \mid 0 \ 0 \mid 0 \mid 0 \ 0)$
3	$(0 \mid 0 \ 1 \mid 0 \mid 0 \ 0)$	$(0 \mid 1 \ 0 \mid 0 \mid 0 \ 0)$
4	$(0 \mid 1 \ 0 \mid 0 \mid 0 \ 0)$	$(0 \mid 1 \ 0 \mid 0 \mid 0 \ 0)$
5	$(0 \mid 0 \ 0 \mid 1 \mid 0 \ 0)$	$(0 \mid 0 \ 1 \mid 0 \mid 0 \ 0)$
6	$(0 \mid 0 \ 0 \mid 0 \mid 1 \ 1)$	$(0 \mid 0 \ 0 \mid 1 \mid 0 \ 0)$
7	$(0 \mid 0 \ 0 \mid 0 \mid 0 \ 1)$	$(0 \mid 0 \ 0 \mid 0 \mid 1 \ 0)$
8	$(0 \mid 0 \ 0 \mid 0 \mid 1 \ 0)$	$(0 \mid 0 \ 0 \mid 0 \mid 1 \ 0)$

The eight data points as full input–output vectors

Each training example isolates one instruction in the six-bit state.

Data point	Raw input $u_r^\top$ carry 2 p_2 q_2 carry 1 p_1 q_1	Complete output $y^{(r)\top}$ carry 2' p'_2 q'_2 carry 1' p'_1 q'_1
1	$(1 \mid 0 \ 0 \mid 0 \mid 0 \ 0)$	$(1 \mid 0 \ 0 \mid 0 \mid 0 \ 0)$
2	$(0 \mid 1 \ 1 \mid 0 \mid 0 \ 0)$	$(1 \mid 0 \ 0 \mid 0 \mid 0 \ 0)$
3	$(0 \mid 0 \ 1 \mid 0 \mid 0 \ 0)$	$(0 \mid 1 \ 0 \mid 0 \mid 0 \ 0)$
4	$(0 \mid 1 \ 0 \mid 0 \mid 0 \ 0)$	$(0 \mid 1 \ 0 \mid 0 \mid 0 \ 0)$
5	$(0 \mid 0 \ 0 \mid 1 \mid 0 \ 0)$	$(0 \mid 0 \ 1 \mid 0 \mid 0 \ 0)$
6	$(0 \mid 0 \ 0 \mid 0 \mid 1 \ 1)$	$(0 \mid 0 \ 0 \mid 1 \mid 0 \ 0)$
7	$(0 \mid 0 \ 0 \mid 0 \mid 0 \ 1)$	$(0 \mid 0 \ 0 \mid 0 \mid 1 \ 0)$
8	$(0 \mid 0 \ 0 \mid 0 \mid 1 \ 0)$	$(0 \mid 0 \ 0 \mid 0 \mid 1 \ 0)$

State $x^\top$ $(0 \mid 1 \ 1 \mid 0 \mid 0 \ 1)$

Matches data points 2 and 7.
Outside-block zeros are padding.

Recover the sum by matching data rows

$\text{Match}(x)$ returns the rows whose condition holds in that row's block.

$$\text{Match}(x) = \{2, 7\}.$$

$$f_{\text{match}}(x) = \bigvee_{r \in \text{Match}(x)} y^{(r)}.$$

Recover the sum by matching data rows

$\text{Match}(x)$ returns the rows whose condition holds in that row's block.

$$\text{Match}(x) = \{2, 7\}.$$

$$f_{\text{match}}(x) = \bigvee_{r \in \text{Match}(x)} y^{(r)}.$$

$$f_{\text{match}}(x) = y^{(2)} + y^{(7)} = (1 \mid 0 \ 0 \mid 0 \mid 1 \ 0)^{\top}.$$

Here the two writes are disjoint, so bitwise OR equals their sum.

$$\text{Read carry } 2, \text{ then } p_2, p_1: \quad 10_2 + 11_2 = 101_2.$$

How does the network use these data points?

Data points 2 and 7 match our state. Can the mean recover their outputs?

$$\mu_j(x) = \underbrace{\Theta(x, X)}_{\text{similarities}} \underbrace{K^{-1}}_{\text{reweighting}} \underbrace{y_j}_{\text{output labels}} .$$

$K = \Theta(X, X)$ compares the training inputs with each other.

y_j lists one output bit's labels across the eight training data points.

Encode the whole condition in its own direction

Make similarity recognize the whole condition, not just shared 1s.

$$\begin{array}{ccc} (1, 1) & \mapsto & (1, 0, 0) \\ (0, 1) & \mapsto & (0, 1, 0) \\ (1, 0) & \mapsto & (0, 0, 1) \end{array}$$

Use $(0, 0, 0)$ for $(0, 0)$. Carry blocks keep their one coordinate.

Encode the whole condition in its own direction

Make similarity recognize the whole condition, not just shared 1s.

$$\begin{array}{ccc} (1, 1) & \mapsto & (1, 0, 0) \\ (0, 1) & \mapsto & (0, 1, 0) \\ (1, 0) & \mapsto & (0, 0, 1) \end{array}$$

Use $(0, 0, 0)$ for $(0, 0)$. Carry blocks keep their one coordinate.

Different conditions now have zero overlap, even when their raw bits overlap.

Apply the encoding to the full vectors

Data points 1–4: the high carry and high summand block. Every zero is shown.

Data point	Raw input $u_r^\top$	Encoded input $v_r^\top$
	carry 2 p_2 q_2 carry 1 p_1 q_1	carry 2 (11, 01, 10) carry 1 (11, 01, 10)
1	(1 0 0 0 0 0)	(1 0 0 0 0 0 0 0)
2	(0 1 1 0 0 0)	(0 1 0 0 0 0 0 0)
3	(0 0 1 0 0 0)	(0 0 1 0 0 0 0 0)
4	(0 1 0 0 0 0)	(0 0 0 1 0 0 0 0)

The other four data points

Data points 5–8: the low carry and low summand block. The row order stays fixed.

Data point	Raw input $u_r^\top$	Encoded input $v_r^\top$
	carry 2 p_2 q_2 carry 1 p_1 q_1	carry 2 (11, 01, 10) carry 1 (11, 01, 10)
5	(0 0 0 1 0 0)	(0 0 0 0 1 0 0 0)
6	(0 0 0 0 1 1)	(0 0 0 0 0 1 0 0)
7	(0 0 0 0 0 1)	(0 0 0 0 0 0 1 0)
8	(0 0 0 0 1 0)	(0 0 0 0 0 0 0 1)

Orthogonality simplifies the inverse in the mean

The eight encoded training inputs are orthonormal.

$$\mu_j(\hat{x}) = \Theta(\hat{x}, X) K^{-1} y_j.$$

$$K^{-1} = \alpha I_8 - \beta \mathbf{1}\mathbf{1}^\top, \quad \alpha, \beta > 0.$$

Orthogonality simplifies the inverse in the mean

The eight encoded training inputs are orthonormal.

$$\mu_j(\hat{x}) = \Theta(\hat{x}, X) K^{-1} y_j.$$

$$K^{-1} = \alpha I_8 - \beta \mathbf{1}\mathbf{1}^\top, \quad \alpha, \beta > 0.$$

The constants α and β are fixed by the training kernel.

The identity rescales $\Theta(\hat{x}, X)$.

The all-ones term subtracts the same correction from every entry.

What multiplies each data point's output?

For our two matches, $\hat{x} = (v_2 + v_7)/\sqrt{2}$. Write the mean as $\mu_j = h y_j$.

$$h := \Theta(\hat{x}, X)K^{-1}.$$

What multiplies each data point's output?

For our two matches, $\hat{x} = (v_2 + v_7)/\sqrt{2}$. Write the mean as $\mu_j = h y_j$.

$$h := \Theta(\hat{x}, X)K^{-1}.$$

$$h_r = \underbrace{\alpha \Theta(\hat{x}, v_r)}_{\text{rescale data point } r\text{'s similarity}} - \underbrace{\beta \sum_{s=1}^8 \Theta(\hat{x}, v_s)}_{\text{same subtraction for every data point}}.$$

This is the previous inverse multiplied into the similarity row.

Evaluate that same calculation for our input

The paper's kernel gives the following rounded values.

Matching data point: 2 or 7

Rescaled
similarity

0.76

Shared sub-
traction

— 0.20

Any nonmatching data point

Rescaled
similarity

0.19

Shared sub-
traction

— 0.20

Evaluate that same calculation for our input

The paper's kernel gives the following rounded values.

Matching data point: 2 or 7

Rescaled
similarity 0.76

Shared sub-
traction -0.20

$$h_2 = h_7 \simeq +0.56$$

Any nonmatching data point

Rescaled
similarity 0.19

Shared sub-
traction -0.20

$$h_r \simeq -0.01$$

These are the coefficients in $\mu_j = h y_j$, not hand-chosen data point scores.

Low sum bit: multiply by the full label vector

y_{p_1} contains the eight labels for p'_1 in the dataset on slide 11.

$$\mu_{p_1}(\hat{x}) = h \, y_{p_1}.$$

Data point r	1	2	3	4	5	6	7	8
h (rounded)	−0.01	+0.56	−0.01	−0.01	−0.01	−0.01	+0.56	−0.01
$y_{p_1}^\top$	0	0	0	0	0	0	1	1

Multiply each coefficient by the label directly below it, then add.

Low sum bit: multiply by the full label vector

y_{p_1} contains the eight labels for p'_1 in the dataset on slide 11.

$$\mu_{p_1}(\hat{x}) = h y_{p_1}.$$

Data point r	1	2	3	4	5	6	7	8
h (rounded)	−0.01	+0.56	−0.01	−0.01	−0.01	−0.01	+0.56	−0.01
$y_{p_1}^\top$	0	0	0	0	0	0	1	1

Multiply each coefficient by the label directly below it, then add.

$$\mu_{p_1} = h_1 \cdot 0 + h_2 \cdot 0 + h_3 \cdot 0 + h_4 \cdot 0 + h_5 \cdot 0 + h_6 \cdot 0 + h_7 \cdot 1 + h_8 \cdot 1.$$

Low sum bit: multiply by the full label vector

y_{p_1} contains the eight labels for p'_1 in the dataset on slide 11.

$$\mu_{p_1}(\hat{x}) = h y_{p_1}.$$

Data point r	1	2	3	4	5	6	7	8
h (rounded)	−0.01	+0.56	−0.01	−0.01	−0.01	−0.01	+0.56	−0.01
$y_{p_1}^\top$	0	0	0	0	0	0	1	1

Multiply each coefficient by the label directly below it, then add.

$$\mu_{p_1} = h_1 \cdot 0 + h_2 \cdot 0 + h_3 \cdot 0 + h_4 \cdot 0 + h_5 \cdot 0 + h_6 \cdot 0 + h_7 \cdot 1 + h_8 \cdot 1.$$

$$\mu_{p_1} = h_7 + h_8 \simeq 0.55 > 0 \quad \Longrightarrow \quad p'_1 = 1.$$

Even matching data point 2 contributes zero here: its label for this bit is 0.

High sum bit: change only the label vector

Keep the same input and coefficients h . Now use the eight labels for p'_2 .

$$\mu_{p_2}(\hat{x}) = h y_{p_2}.$$

Data point r	1	2	3	4	5	6	7	8
h (rounded)	−0.01	+0.56	−0.01	−0.01	−0.01	−0.01	+0.56	−0.01
$y_{p_2}^\top$	0	0	1	1	0	0	0	0

Multiply each coefficient by the label directly below it, then add.

High sum bit: change only the label vector

Keep the same input and coefficients h . Now use the eight labels for p'_2 .

$$\mu_{p_2}(\hat{x}) = h y_{p_2}.$$

Data point r	1	2	3	4	5	6	7	8
h (rounded)	−0.01	+0.56	−0.01	−0.01	−0.01	−0.01	+0.56	−0.01
$y_{p_2}^\top$	0	0	1	1	0	0	0	0

Multiply each coefficient by the label directly below it, then add.

$$\mu_{p_2} = h_1 \cdot 0 + h_2 \cdot 0 + h_3 \cdot 1 + h_4 \cdot 1 + h_5 \cdot 0 + h_6 \cdot 0 + h_7 \cdot 0 + h_8 \cdot 0.$$

High sum bit: change only the label vector

Keep the same input and coefficients h . Now use the eight labels for p'_2 .

$$\mu_{p_2}(\hat{x}) = h y_{p_2}.$$

Data point r	1	2	3	4	5	6	7	8
h (rounded)	−0.01	+0.56	−0.01	−0.01	−0.01	−0.01	+0.56	−0.01
$y_{p_2}^\top$	0	0	1	1	0	0	0	0

Multiply each coefficient by the label directly below it, then add.

$$\mu_{p_2} = h_1 \cdot 0 + h_2 \cdot 0 + h_3 \cdot 1 + h_4 \cdot 1 + h_5 \cdot 0 + h_6 \cdot 0 + h_7 \cdot 0 + h_8 \cdot 0.$$

$$\mu_{p_2} = h_3 + h_4 \simeq -0.02 < 0 \quad \Longrightarrow \quad p'_2 = 0.$$

Data points 2 and 7 match the input, but both have label 0 for this bit.

The mean gives the desired update

Decode each mean coordinate: positive $\rightarrow$ 1, otherwise $\rightarrow$ 0.

carry 2' | p'_2 q'_2 | carry 1' | p'_1 q'_1

$$x' = (1 \mid 0 \ 0 \mid 0 \mid 1 \ 0)^\top.$$

The mean gives the desired update

Decode each mean coordinate: positive $\rightarrow$ 1, otherwise $\rightarrow$ 0.

carry 2' | p'_2 q'_2 | carry 1' | p'_1 q'_1

$$x' = (1 \mid 0 \ 0 \mid 0 \mid 1 \ 0)^\top.$$

Read carry 2, then p_2, p_1 : 101_2 .

The decoded mean is correct. A single network can still fluctuate.

Now bring back the random initialization

The mean is fixed by the instruction data; individual trained networks can differ.

$$F_{\infty,j}(\hat{x}) = \mu_j(\hat{x}) + g_j(\hat{x}).$$

Now bring back the random initialization

The mean is fixed by the instruction data; individual trained networks can differ.

$$F_{\infty,j}(\hat{x}) = \mu_j(\hat{x}) + g_j(\hat{x}).$$

$$g_j(\hat{x}) \sim \mathcal{N}\left(0, \sigma^2(\hat{x})\right).$$

Correct mean and small random spread are two separate requirements.

What controls the remaining variance?

In the paper's model normalization, Lemma 6.1 gives

$$\sigma^2(\hat{x}) = O(1/D).$$

D is the encoded dimension. In our example, $D = 8$.

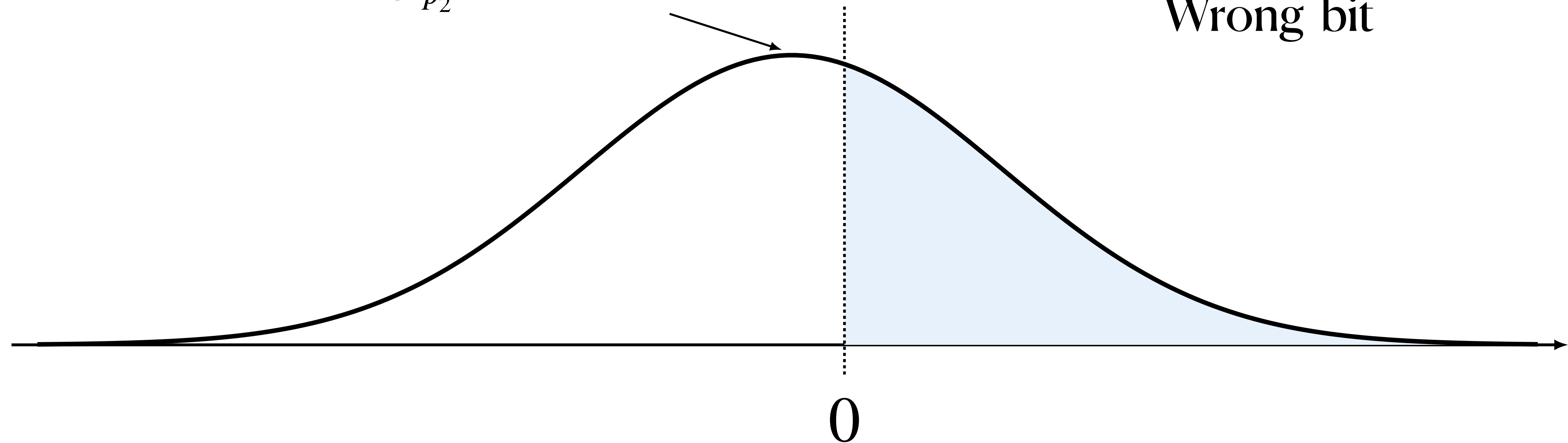
Small variance alone is not enough: compare the spread with the distance to zero.

A correct mean can still give a wrong prediction

The high sum bit should be 0. Its mean is negative, but a prediction can cross zero.

Mean $\mu_{p_2} < 0$

Prediction > 0
Wrong bit



The spread must be small compared with the distance to zero.

Average predictions from independent models

Train N networks on the same instructions, with independent initializations.

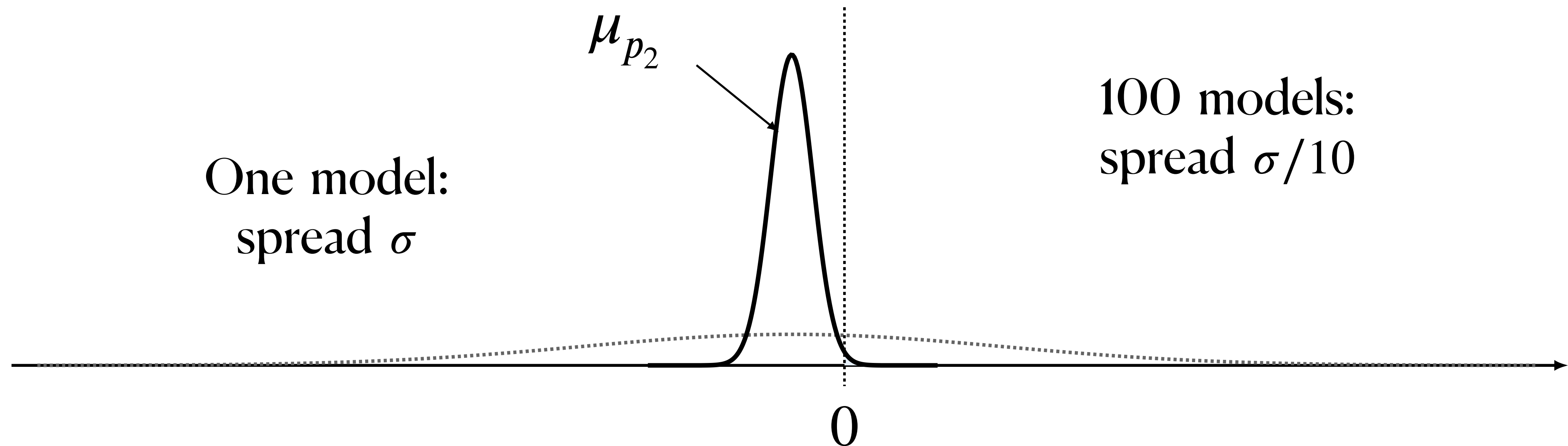
$$\overline{F}_j = \frac{1}{N} \sum_{n=1}^N F_j^{(n)}.$$

$$\mathbb{E}[\overline{F}_j] = \mu_j, \quad \text{Var}(\overline{F}_j) = \frac{\sigma^2}{N}.$$

The mean stays fixed. The standard deviation shrinks by a factor of $\sqrt{N}$.

Same bit, same mean, much less spread

Return to the same high sum bit p'_2 , now averaging $N = 100$ networks.



Averaging narrows the distribution without moving its center.

Substitute the mean and variance bounds

For addition, use the weakest nonzero-mean rate in Lemma 6.1.

Variance upper bound

$$\sigma^2(\hat{x}) = O(D^{-1})$$

Mean magnitude lower bound

$$|\mu_j(\hat{x})| = \Omega(D^{-1})$$

$$\sigma^2(\hat{x}) / \mu_j(\hat{x})^2 \leq CD.$$

$C > 0$ is independent of D and N ; use the same C for every bit being protected.

$$\Pr(\text{wrong bit } j) \leq 2 \exp\left(-\frac{N \mu_j^2}{8 \sigma^2}\right) \leq 2 \exp\left(-\frac{N}{8CD}\right).$$

What does M count?

One decision means rounding one predicted next-state bit to 0 or 1.

Each update of our two-bit state predicts five nonconstant bits:

carry $2'$, p'_2 , q'_2 , carry $1'$, p'_1 .

The known-zero q'_1 is set to 0 directly (the added decoder convention in the notes).

Which computation is protected?

Number of bit predictions M

Our $10_2 + 11_2$ example: one update

$$M = 5$$

One two-bit addition: up to 4 updates

$$M \leq 5 \times 4 = 20$$

All 16 two-bit additions: up to 4 updates each

$$M \leq 16 \times 5 \times 4 = 320$$

M counts output-bit predictions, not training data points or networks.

Choose N to protect every predicted bit

Use the same ensemble at every update. Allow total failure probability at most δ .

$$\Pr(\text{at least one wrong bit}) \leq 2M \exp\left(-\frac{N}{8CD}\right).$$

Add the one-bit error bounds over the M predictions. They need not be independent.

To make that total at most δ , it is sufficient to choose

$$N \geq 8CD \ln(2M/\delta).$$

$$\Pr(\text{every predicted bit is correct}) \geq 1 - \delta.$$

Correct state bits at every update give a correct execution.