

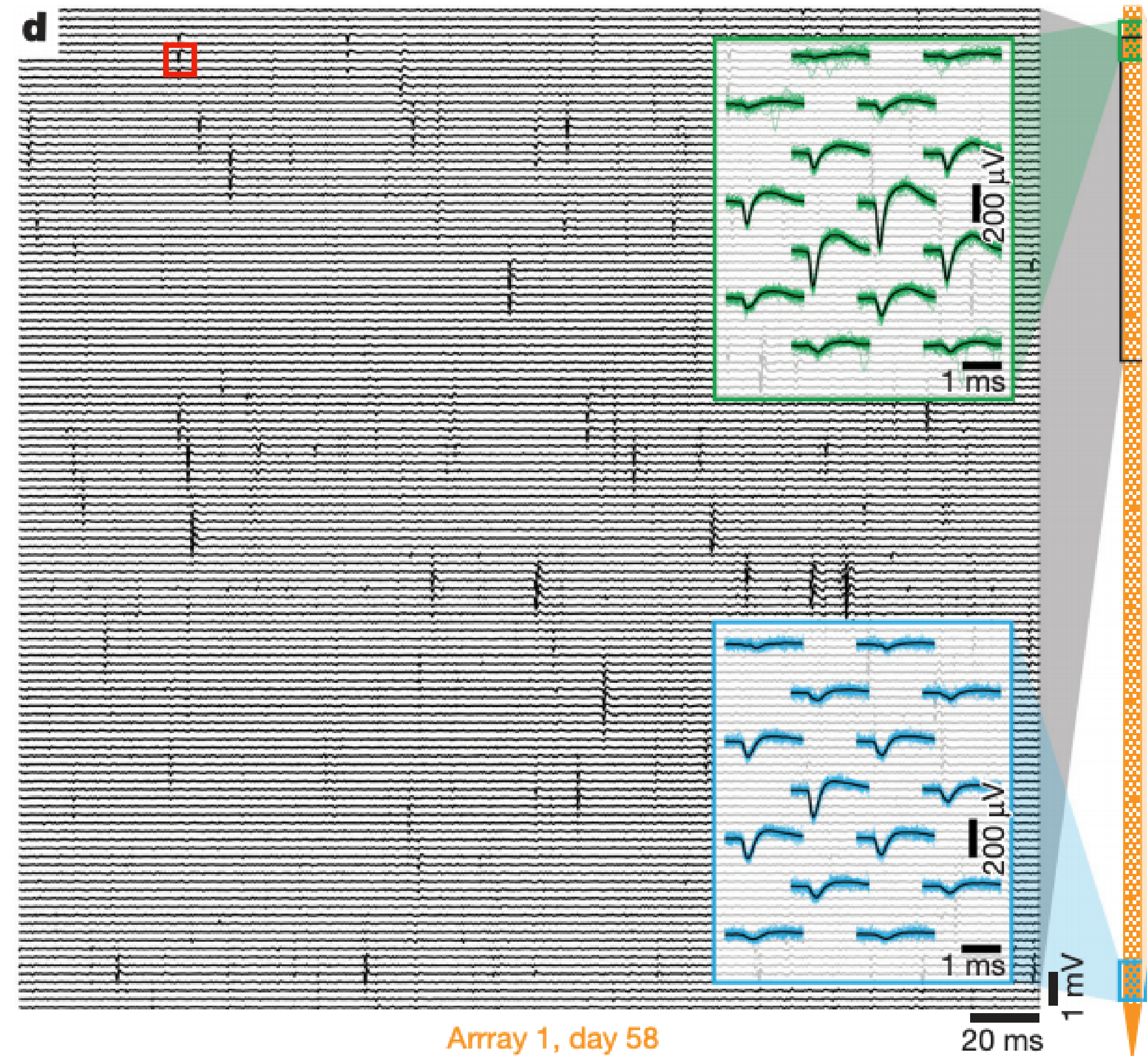
Machine Learning Methods for Neural Data Analysis

Spike Sorting by Deconvolution

Spike Sorting by Deconvolution

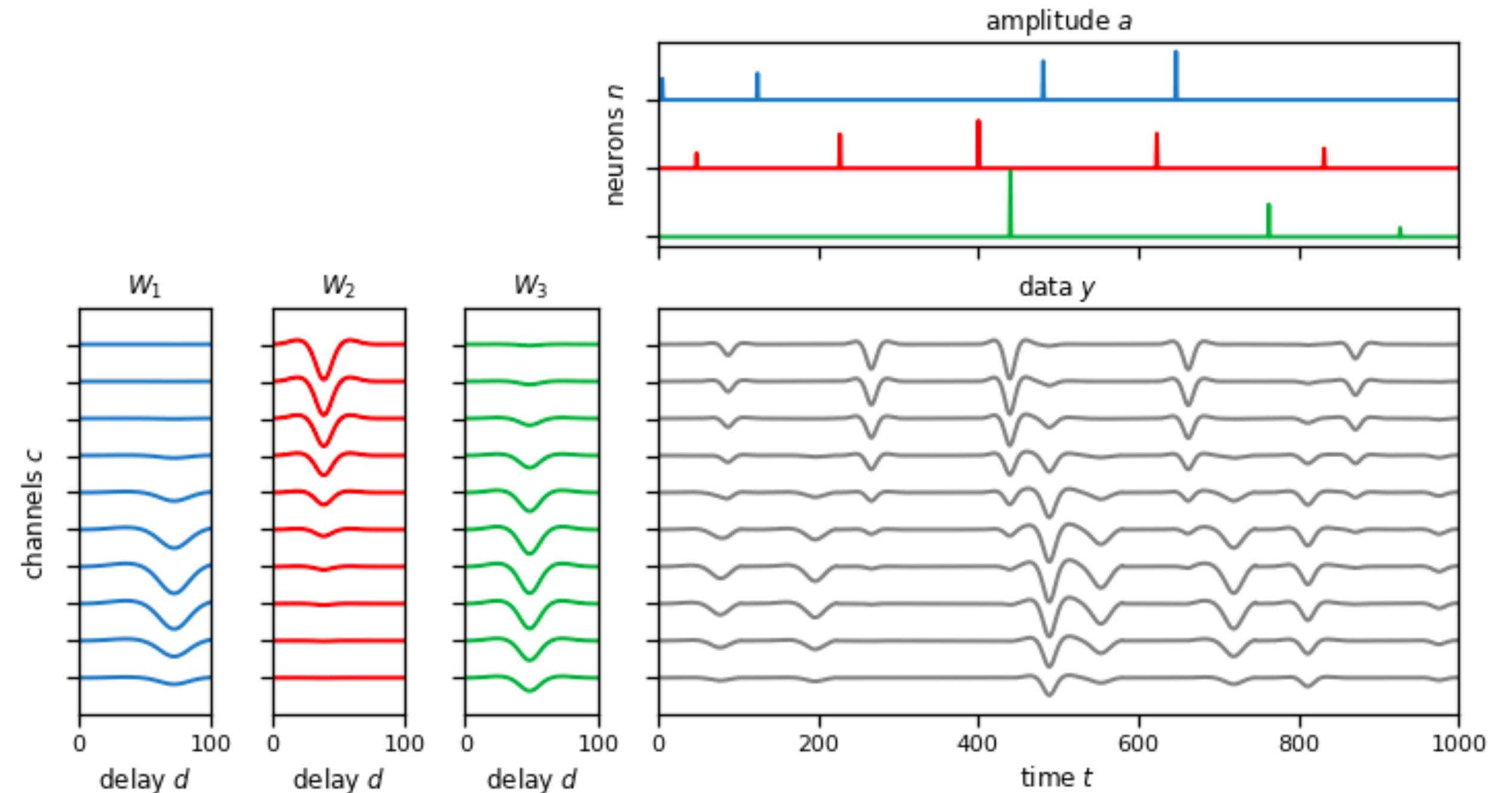
Improving upon the simple model

- Our simple model was a good warm-up, but we made some strong assumptions:
 - Namely, that we could extract windows for each spike and assign them to one neuron.
 - All spikes from the same neuron share the same mean,
- With lots of recording channels, there are likely to be **overlapping spikes**, and they could have **different amplitudes**.
- Next, we'll extend the simple model with a more realistic one using **convolutional matrix factorization**.
- The resulting model will be very similar to **Kilosort** [Pachitariu et al., 2023]



10,000ft view

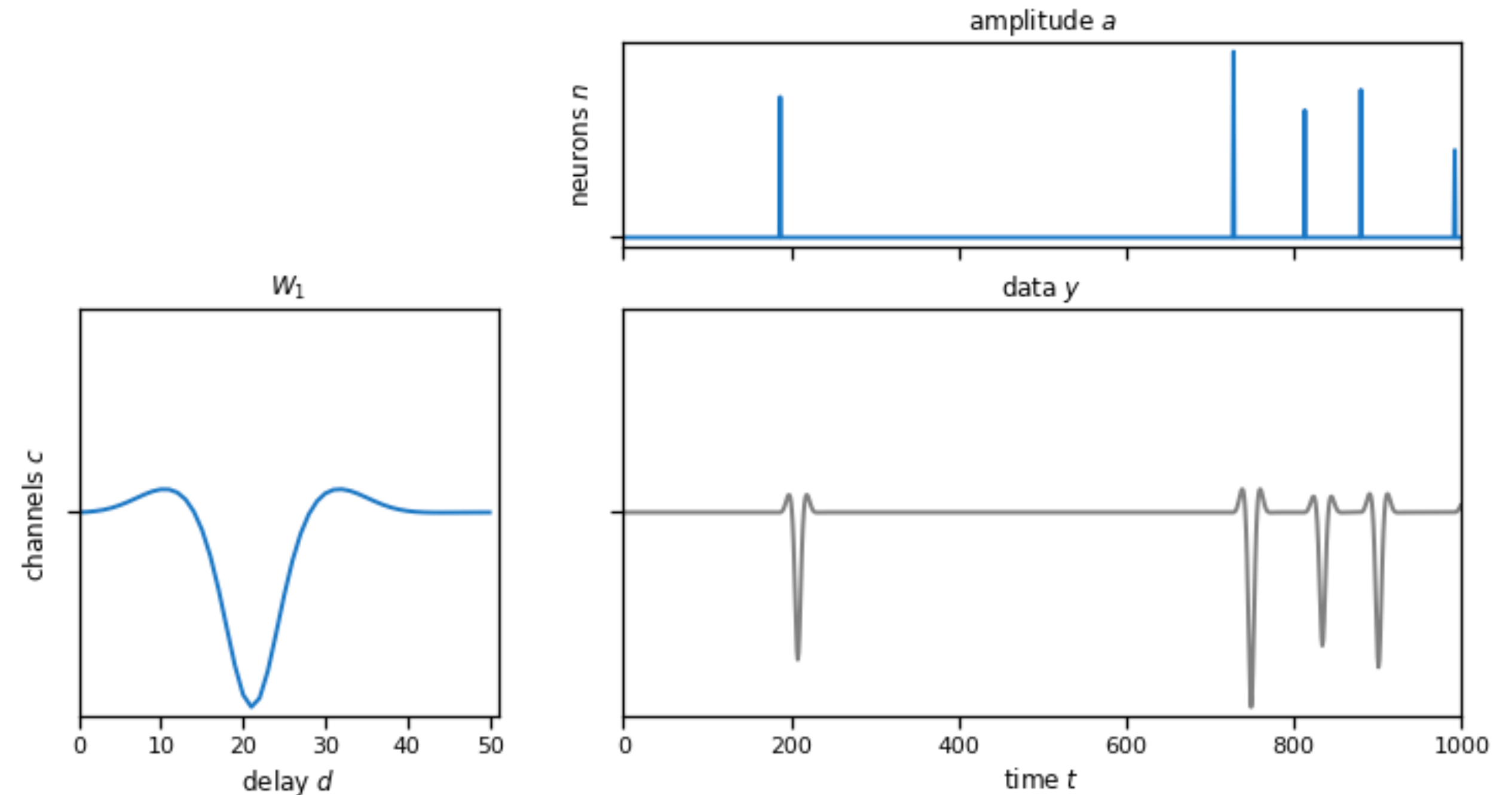
- **Idea:** each time a neuron spikes, it adds a scaled copy of its template to the measured voltage.
- Formally, we model the data as a **sum of convolutions** of templates and amplitudes for each neuron, plus noise.



Convolution

In one dimension

- **Convolution** is an operation that takes in a signal $a(t)$ and a filter $w(t)$ and outputs
$$y(t) = [a \circledast w](t) = \int a(t - \tau)w(\tau) d\tau.$$
- In **discrete time** this becomes,
$$y_t = [a \circledast w]_t = \sum_{d=-\infty}^{\infty} a_{t-d}w_d$$
- **Causal** filters are constrained so that $w_d = 0$ for $d < 0$. Then y_t is only influenced by $a_{1:t}$.
- Our filters will also have **bounded support** so that $w_d = 0$ for $d \geq D$. Then y_t is only influenced by $a_{t-D+1:t}$.
- In our case, the signal is the time series of spike amplitudes, and the filter is the waveform template. Every time there's a spike, we plop down a scaled template.



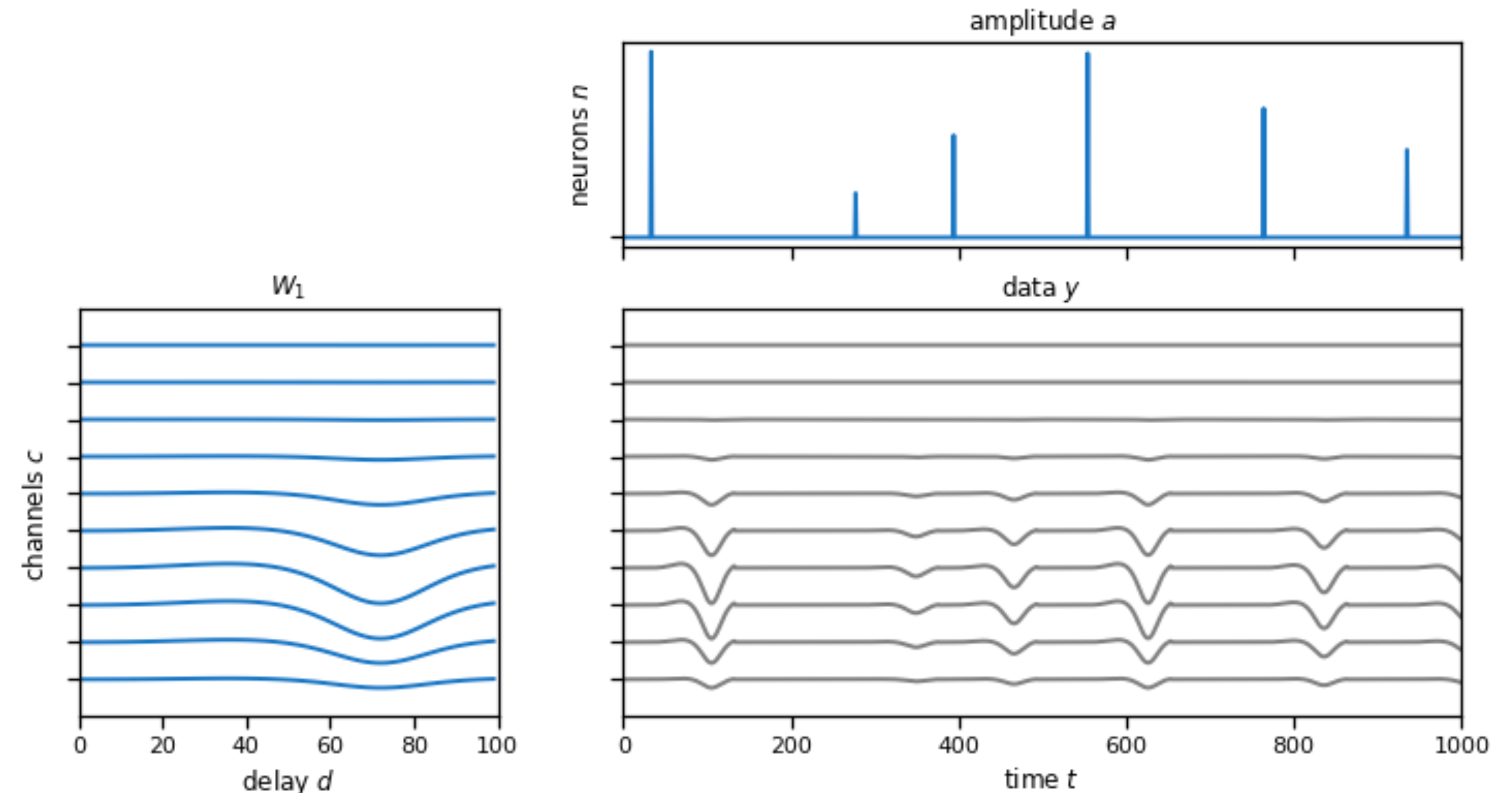
Convolution

With multiple output channels

- We need to convolve the amplitude signal with **multiple filters** in parallel, **one for each channel** of the voltage recording.

$$\mathbf{y}_t = \begin{pmatrix} [a \circledast w_1]_t \\ \vdots \\ [a \circledast w_N]_t \end{pmatrix} = \begin{pmatrix} \sum_{d=1}^D a_{t-d} w_{1,d} \\ \vdots \\ \sum_{d=1}^D a_{t-d} w_{N,d} \end{pmatrix}$$

- (I'm going to index d from $1, \dots, D$ because the notation is a bit simpler.)



Convolution

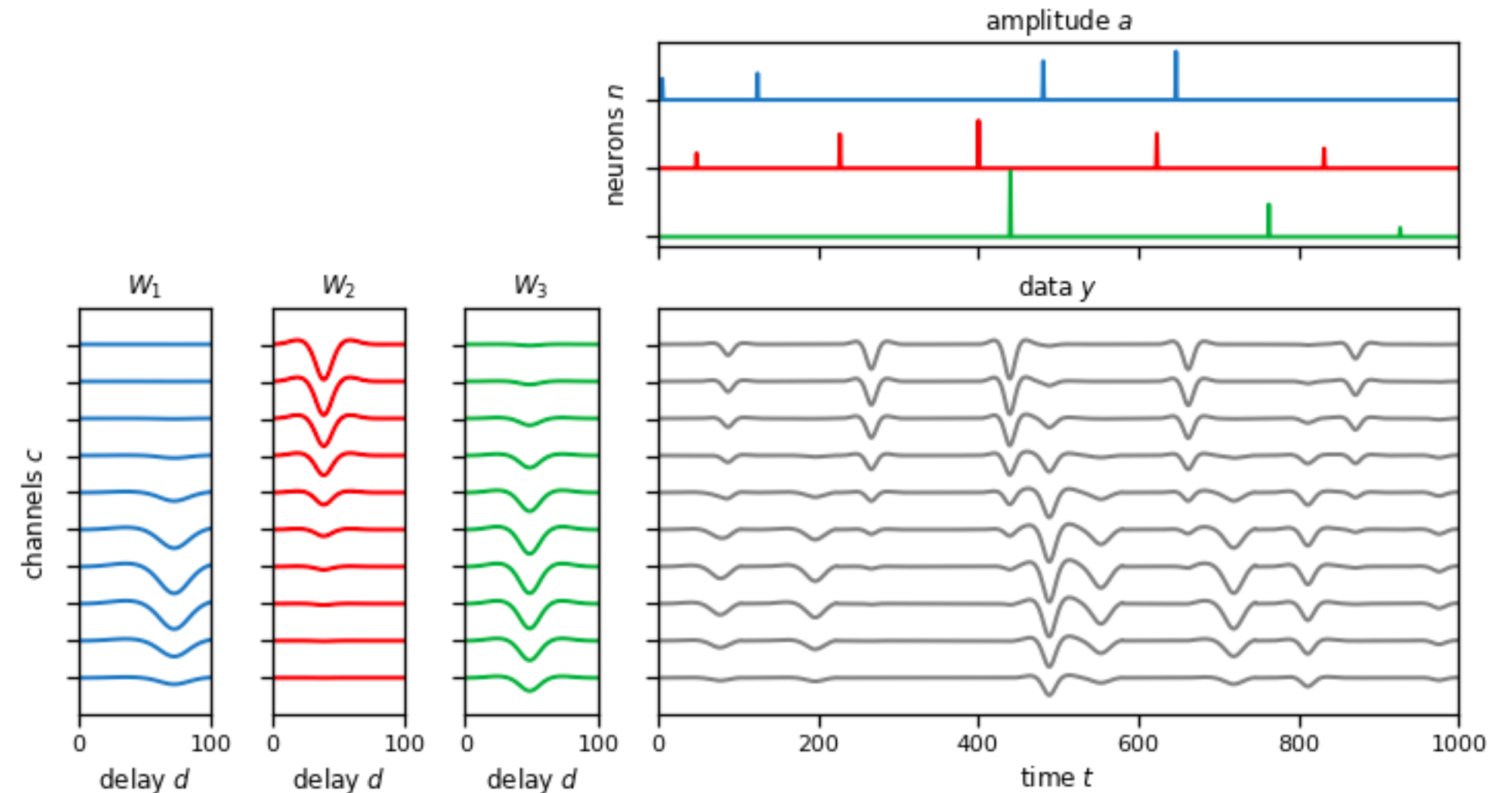
With multiple input & output channels

- Finally, we need to sum convolutions of multiple input signals, **one for each neuron** in the model.

$$\mathbf{X} = \mathbf{A} \circledast \mathbf{W}$$

by which we mean

$$x_{n,t} = \sum_{k=1}^K \sum_{d=1}^D a_{k,t-d} w_{k,n,d}$$



Cross-Correlation

In one dimension

- Cross-correlation essentially goes in reverse.
- In signal processing, the cross-correlation is a sliding dot product of data $y(t)$ and template $w(t)$, which produces a new function $[y \star w](t)$.

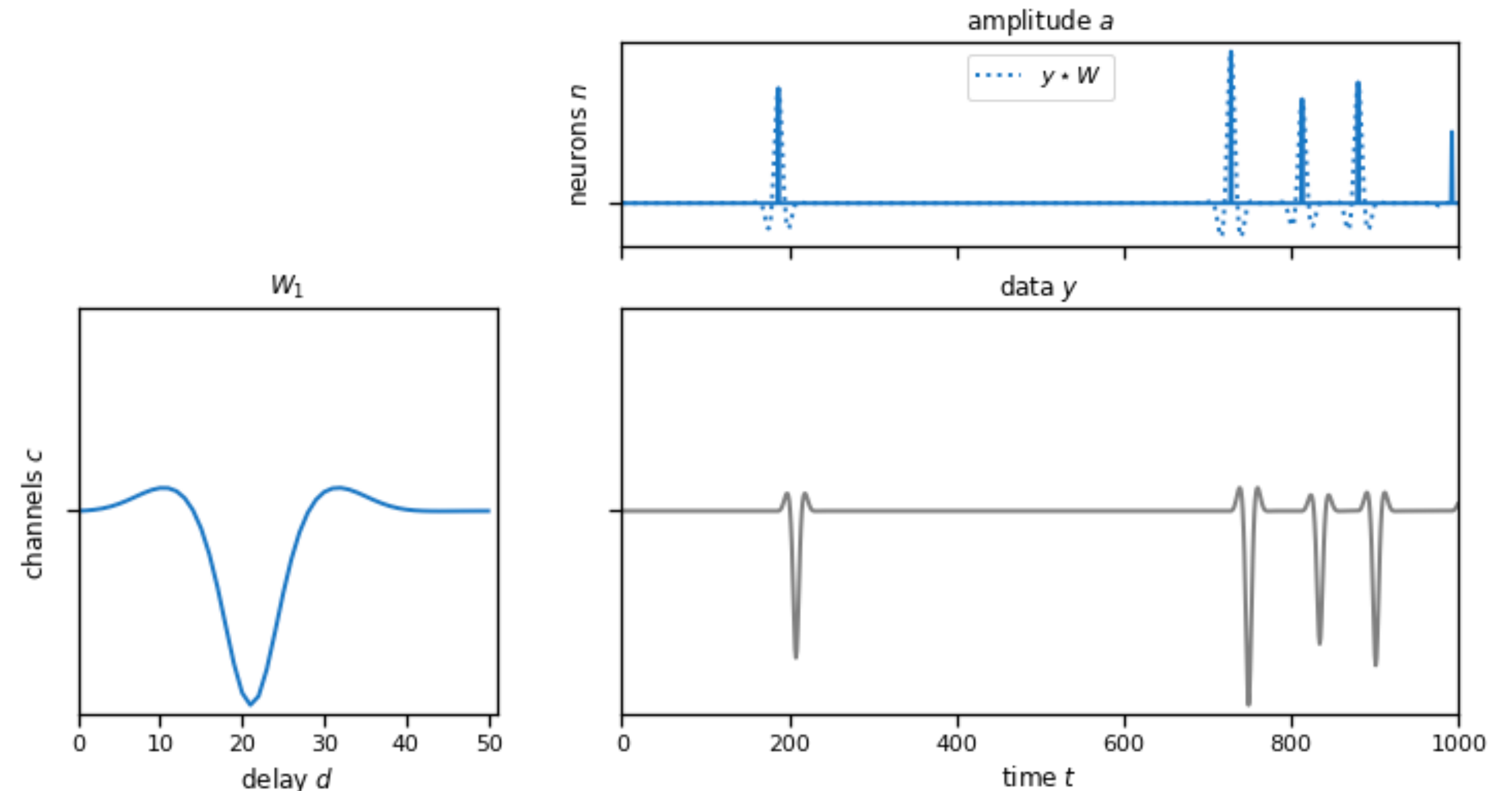
- For discrete time, real-valued inputs,

$$[y \star w]_t = \sum_{d=-\infty}^{\infty} y_{t+d} w_d.$$

- With a change of variables, we see that cross-correlation is equivalent to convolution with a time-reversed filter $\overleftarrow{w}$:

$$[y \star w]_t = \sum_{d=-\infty}^{\infty} y_{t-d} w_{-d} = [y \circledast \overleftarrow{w}]_t.$$

- (Note: the definition of cross-correlation is not unique. This definition is consistent with `np.correlate` but opposite of Wikipedia.)



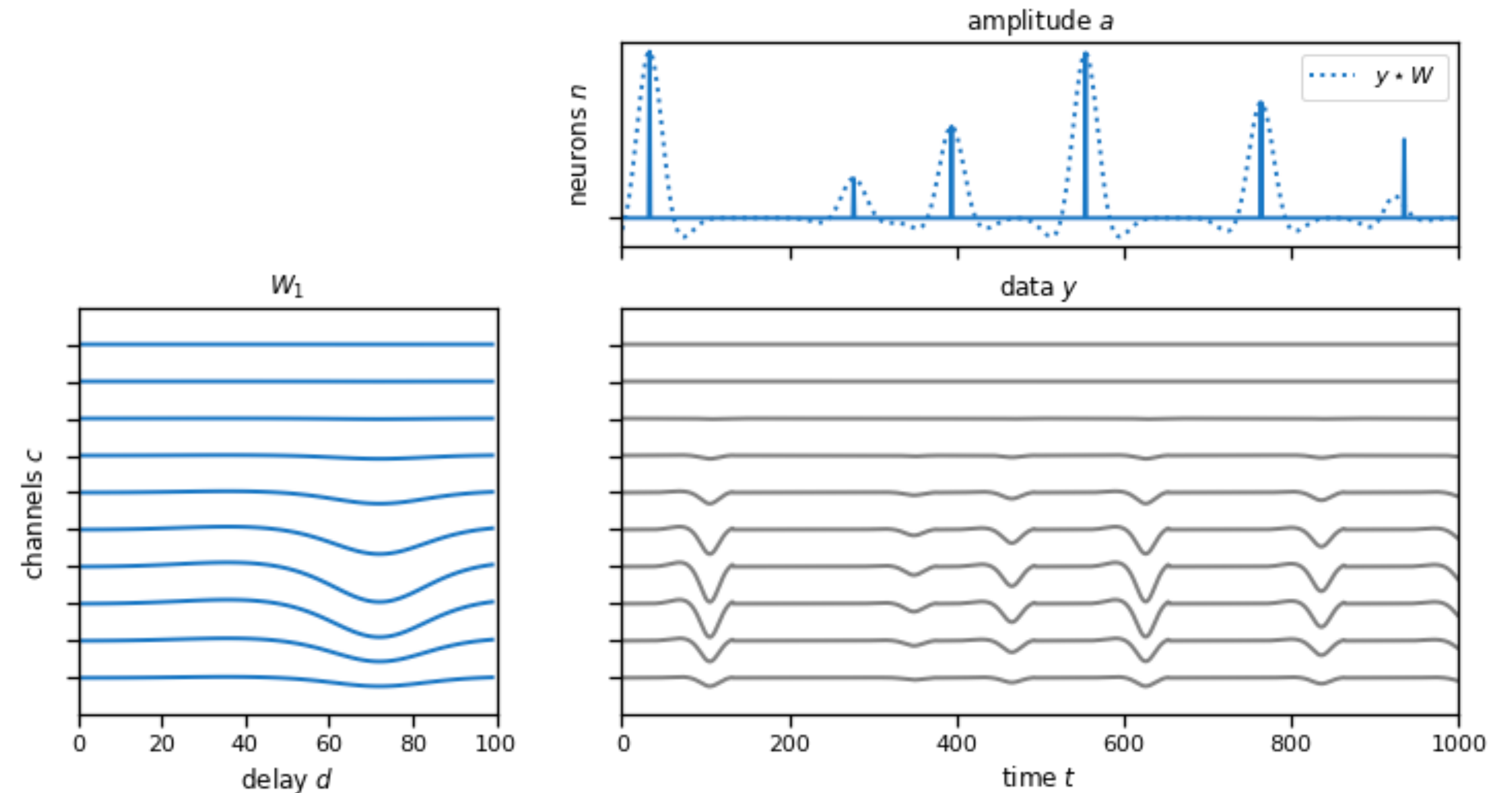
Cross-Correlation

With multiple channels

- As before, we can extend this definition to handle multiple channels

$$[\mathbf{Y} \star \mathbf{W}]_t = \sum_{n=1}^N \sum_{d=-\infty}^{\infty} y_{n,t+d} w_{n,d}$$

- The cross-correlation measures the similarity of the data and the template at each point in time.
- The **auto-correlation** is the cross-correlation of a signal with itself.



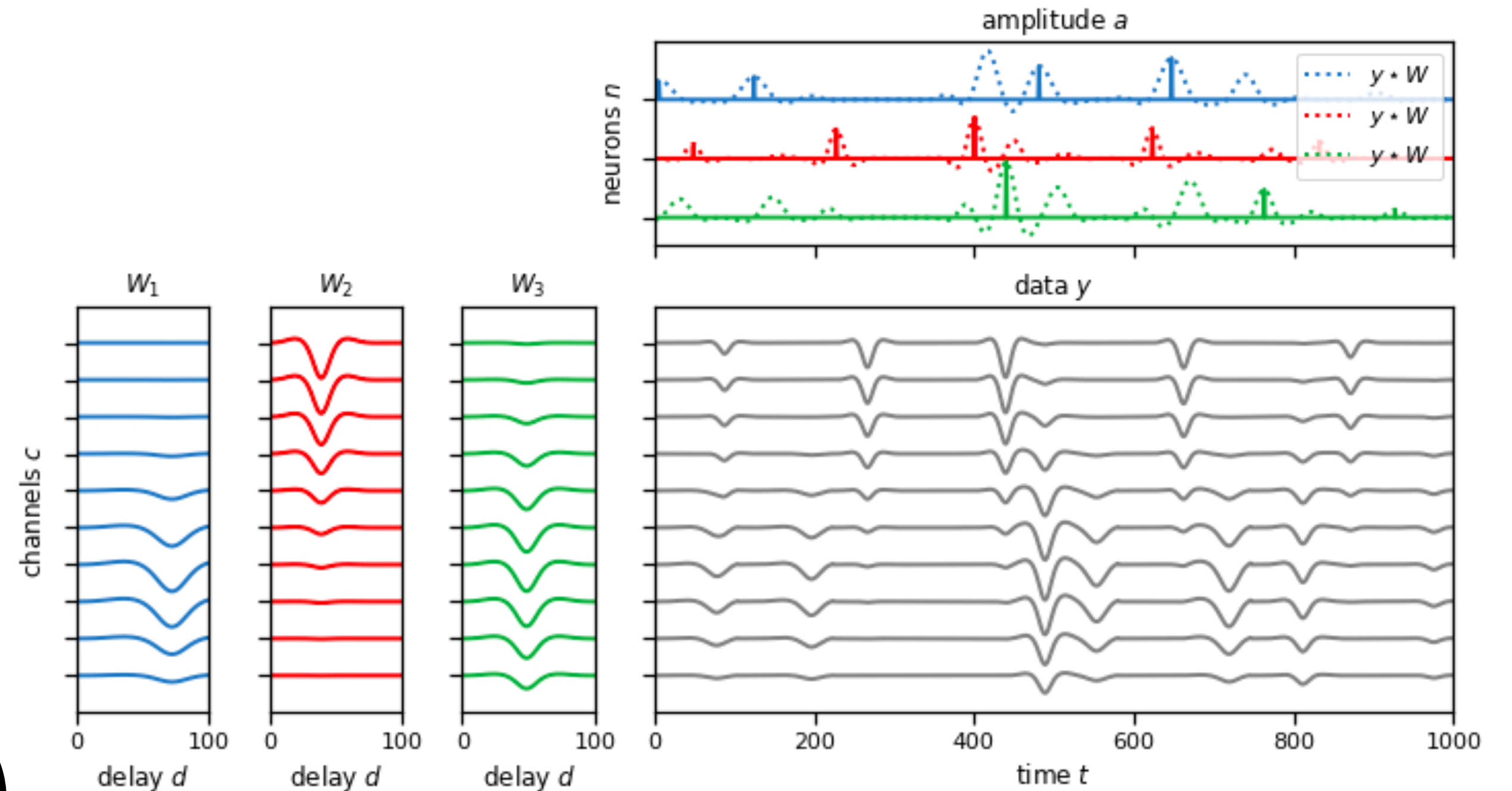
Cross-Correlation

With multiple input & output channels

As before, we can extend this definition to handle multiple channels

$$[Y \star W]_t = \begin{pmatrix} [Y \star W_1]_t \\ \vdots \\ [Y \star W_K]_t \end{pmatrix}$$

$$= \begin{pmatrix} \sum_{n=1}^N \sum_{d=-\infty}^{\infty} y_{n,t+d} w_{1,n,d} \\ \vdots \\ \sum_{n=1}^N \sum_{d=-\infty}^{\infty} y_{n,t+d} w_{K,n,d} \end{pmatrix}$$



Convolution and Cross-Correlation in Pytorch

- PyTorch (and other deep learning libraries) have fast, **GPU-backed implementations** of convolutions.
- **What they call convolution is actually cross-correlation!**
- But remember, we can always get convolution by cross-correlating with the flipped filter.
- For discrete time signals, you have to play with **padding** to handle **edge effects**.
- By default, these functions operate on **mini-batches** of inputs, so you need to add an extra leading dimension to your signal.
- There are **lots of other options** to read about (strides, dilations, groups), but we won't use them this week.

```
torch.nn.functional.conv1d(input, weight, bias=None, stride=1, padding=0, dilation=1, groups=1) → Tensor
```

Applies a 1D convolution over an input signal composed of several input planes.

This operator supports **TensorFloat32**.

See [Conv1d](#) for details and output shape.

• NOTE

In some circumstances when using the CUDA backend with CuDNN, this operator may select a nondeterministic algorithm to increase performance. If this is undesirable, you can try to make the operation deterministic (potentially at a performance cost) by setting `torch.backends.cudnn.deterministic = True`. Please see the notes on **Reproducibility** for background.

Parameters

- **input** – input tensor of shape (minibatch, in_channels, iW)
- **weight** – filters of shape (out_channels, $\frac{\text{in_channels}}{\text{groups}}$, kW)
- **bias** – optional bias of shape (out_channels). Default: `None`
- **stride** – the stride of the convolving kernel. Can be a single number or a one-element tuple (sW). Default: 1
- **padding** – implicit paddings on both sides of the input. Can be a single number or a one-element tuple ($padW$). Default: 0
- **dilation** – the spacing between kernel elements. Can be a single number or a one-element tuple (dW). Default: 1
- **groups** – split input into groups, in_channels should be divisible by the number of groups. Default: 1

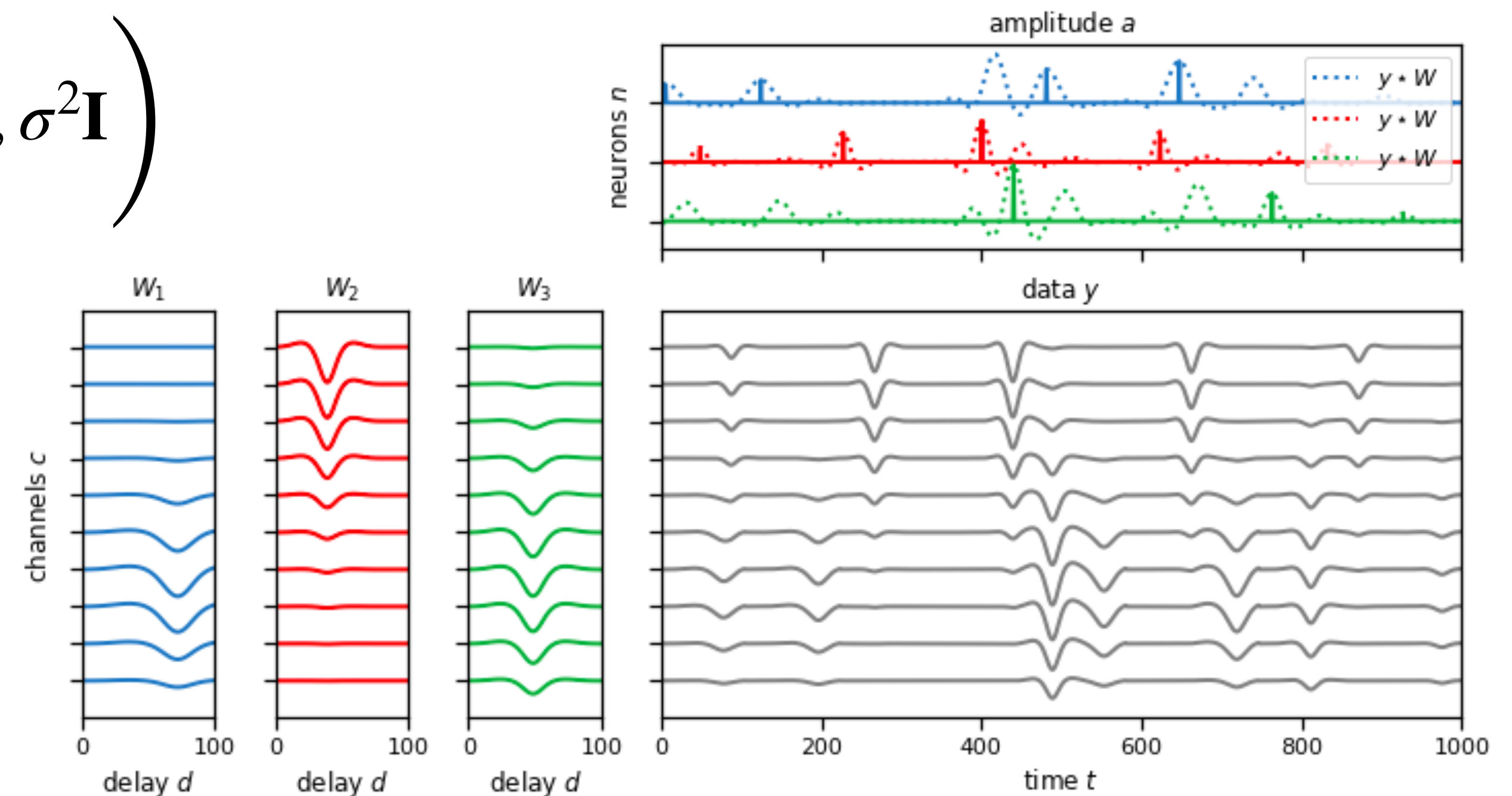
Spike sorting by deconvolution

Probabilistic Model

Likelihood

- Assume each spike is a noisy, scaled version of the template of the neuron that generated it.

$$p(\mathbf{X} \mid \mathbf{A}, \mathbf{W}) = \prod_{t=1}^T \mathcal{N} \left(\mathbf{x}_t \mid \sum_{k=1}^K [\mathbf{a}_k \circledast \mathbf{W}_k]_t, \sigma^2 \mathbf{I} \right)$$



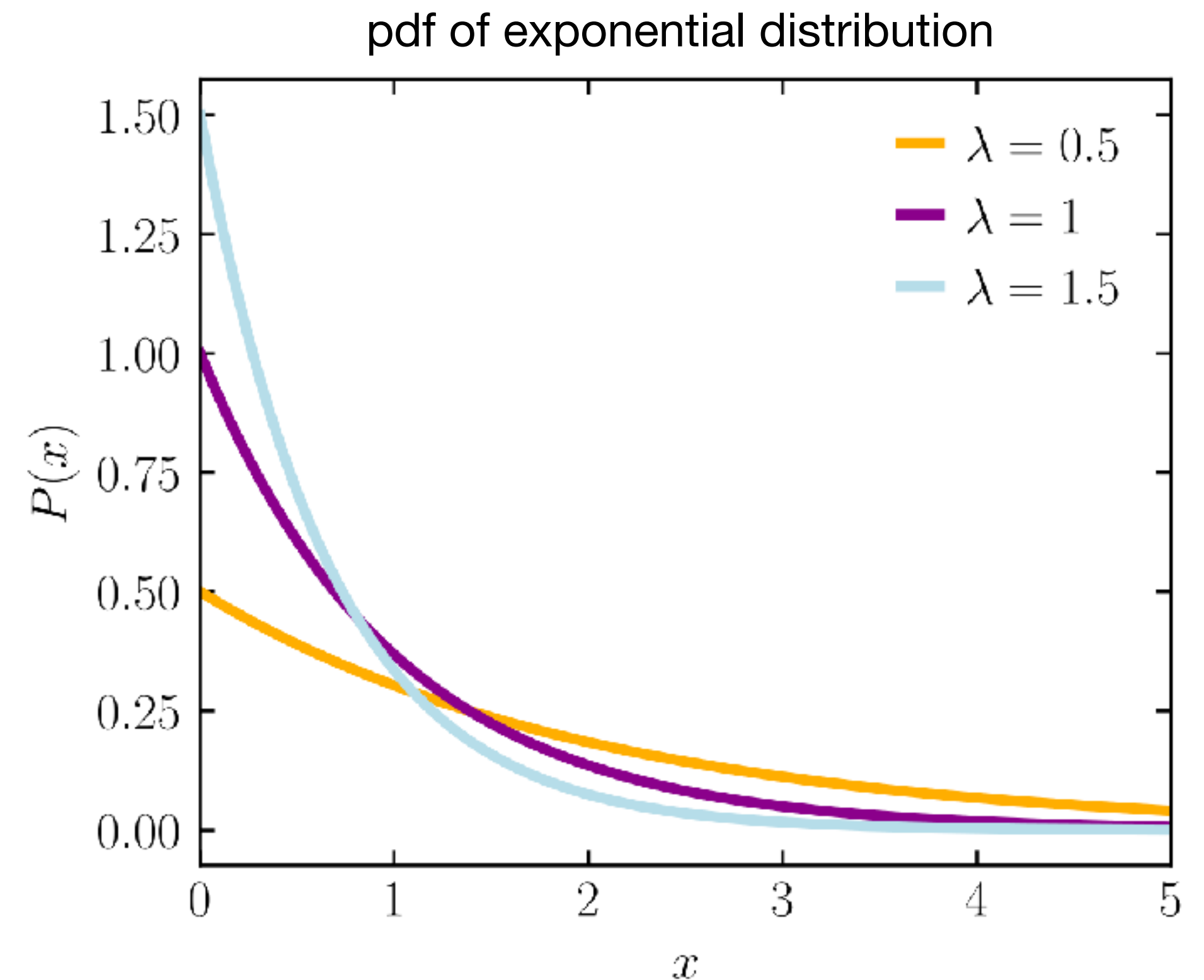
Probabilistic Model

Prior on spike amplitudes

- Assume the spike amplitudes are drawn from an exponential distribution.

$$a_{k,t} \sim \text{Exp}(\lambda)$$

- This simple prior will lead to sparse amplitudes, but it does not encode any dependencies between time steps.
- Ideally, we would also like to prohibit a neuron from producing two spikes within D samples of each other.



Probabilistic Model

Scale invariance via Frobenius norm constraint

- For the amplitudes to be meaningful, we need to **constrain the scale of the templates**.
- Assume the matrix $\mathbf{W}_k \in \mathbb{R}^{N \times D}$ has unit **Frobenius norm** $\|\mathbf{W}_k\|_F = 1$.

Probabilistic Model

Aside: The Frobenius norm and the SVD

- The Frobenius norm of a matrix is induced by the Frobenius inner product $\langle \mathbf{A}, \mathbf{B} \rangle_F = \text{Tr}(\mathbf{A}^\top \mathbf{B})$. With this definition,

$$\|\mathbf{W}\|_F^2 = \langle \mathbf{W}, \mathbf{W} \rangle_F = \text{Tr}(\mathbf{W}^\top \mathbf{W}).$$

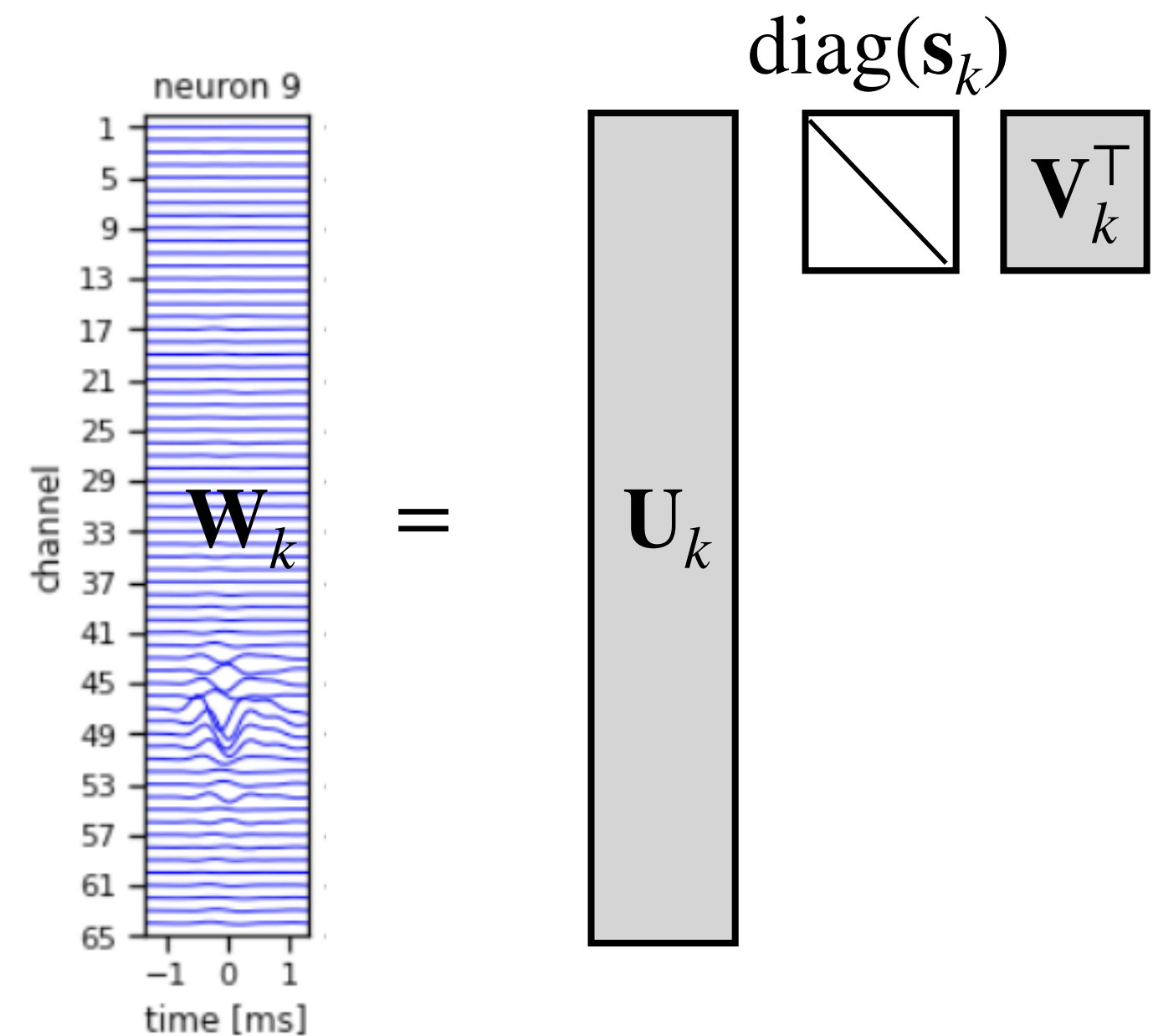
- It is equal to the ℓ_2 norm of the flattened matrix,

$$\|\mathbf{W}\|_F^2 = \sum_{n=1}^N \sum_{d=1}^D w_{n,d}^2 = \text{vec}(\mathbf{W})^\top \text{vec}(\mathbf{W}) = \|\text{vec}(\mathbf{W})\|_2^2.$$

- We can also write it in terms of the singular values of $\mathbf{W}$,

$$\|\mathbf{W}\|_F = \text{Tr}(\mathbf{V} \mathbf{S} \mathbf{U}^\top \mathbf{U} \mathbf{S} \mathbf{V}^\top) = \text{Tr}(\mathbf{S}^2) = \|\mathbf{s}\|_2,$$

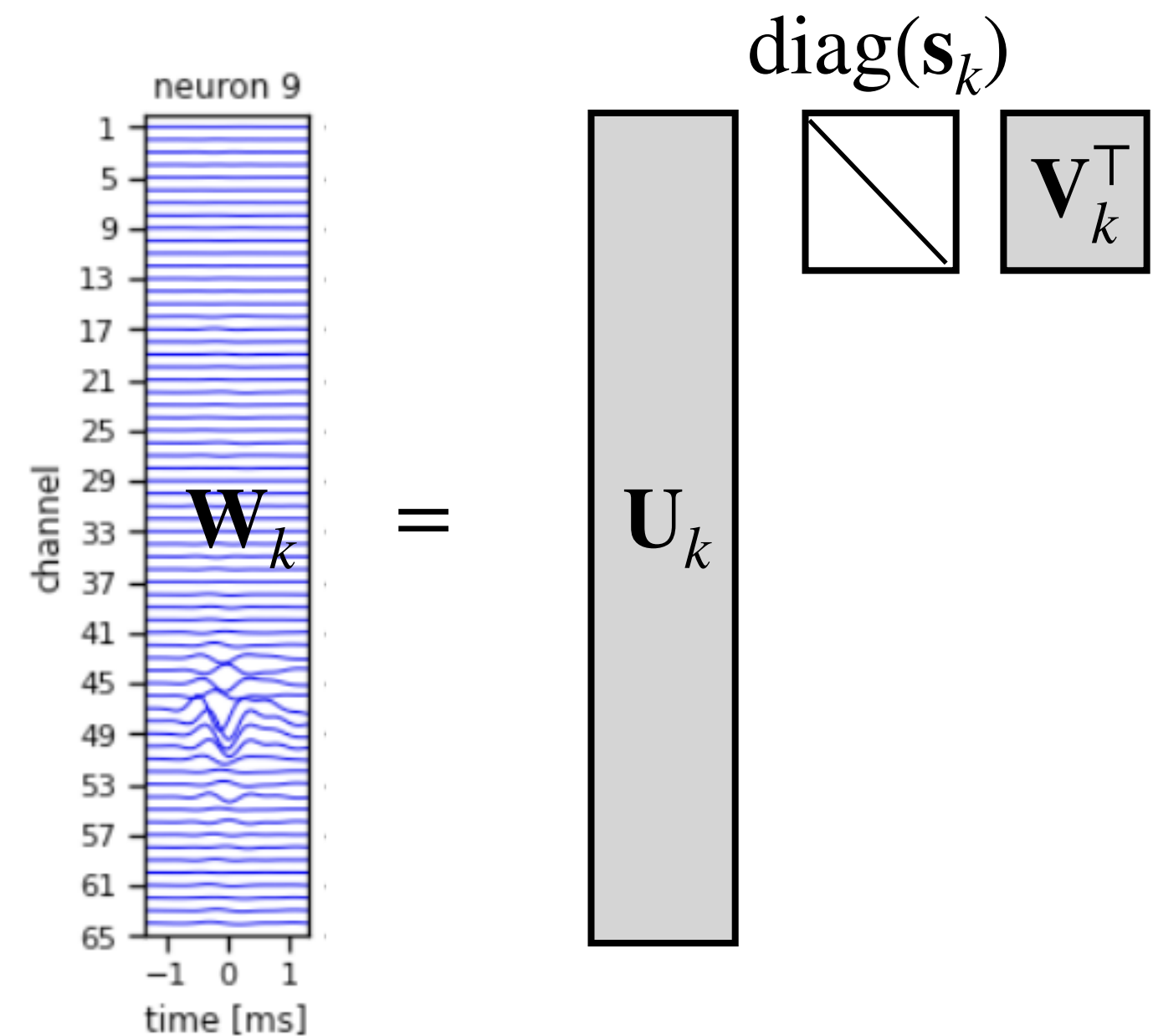
- where $\mathbf{W} = \mathbf{U} \mathbf{S} \mathbf{V}^\top$ with $\mathbf{S} = \text{diag}(\mathbf{s})$ is the **singular value decomposition (SVD)**.



Probabilistic Model

Scale invariance via Frobenius norm constraint

- For the amplitudes to be meaningful, we need to **constrain the scale of the templates**.
- Assume the matrix $\mathbf{W}_k \in \mathbb{R}^{N \times D}$ has unit **Frobenius norm** $\|\mathbf{W}_k\|_F = 1$.
- This is equivalent to constraining the **singular values** to be normalized $\|\mathbf{s}_k\|_2 = 1$.



Probabilistic Model

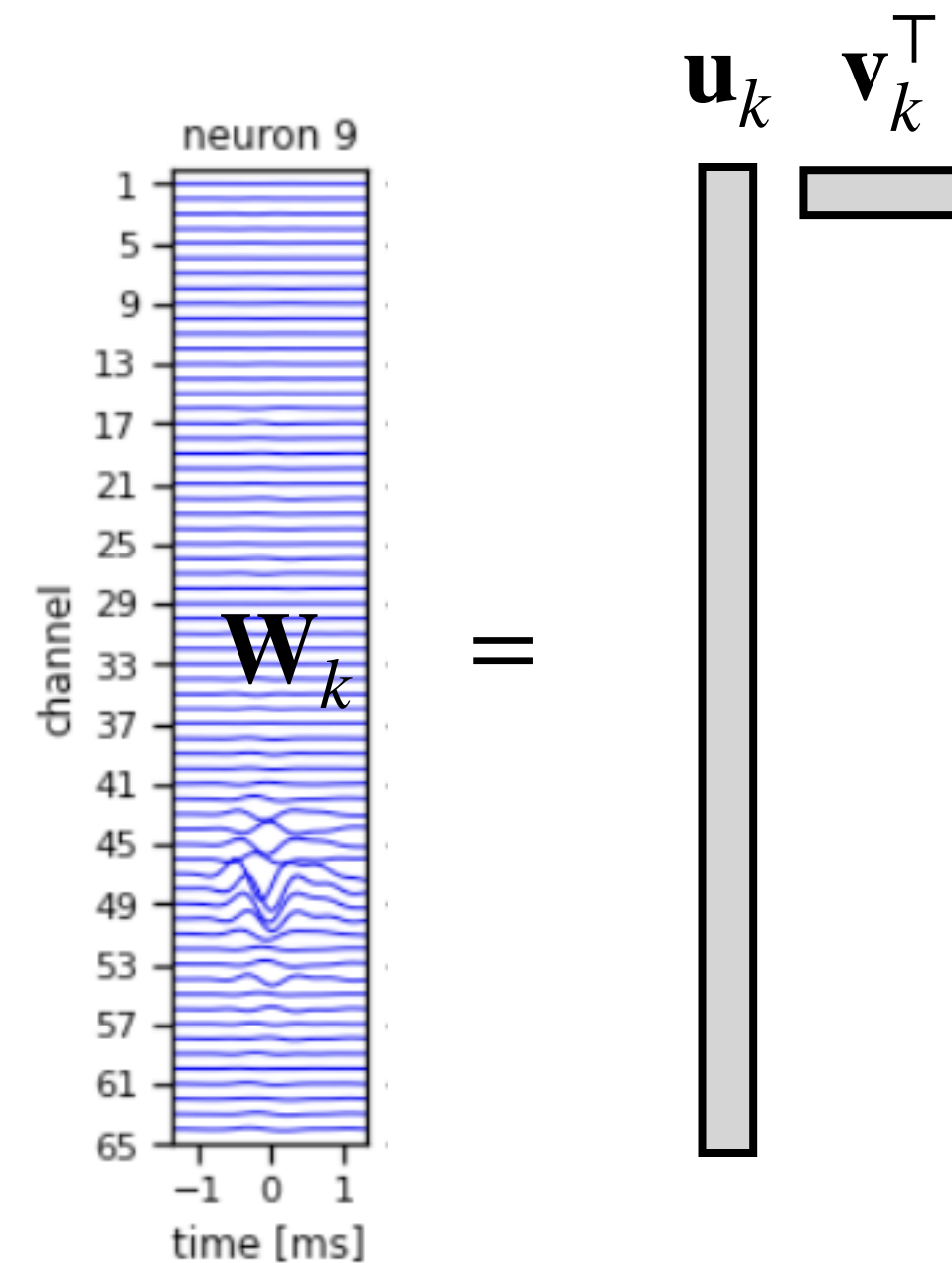
Low-rank constraint

- This view suggests a further assumption: constraint the **rank** of the templates as well.
- If we constrain it to be rank 1 (i.e., only one nonzero singular value), then

$$\mathbf{W}_k = \mathbf{u}_k \mathbf{v}_k^\top$$

where $\mathbf{u}_k \in \mathbb{S}_{N-1}$ is the **spatial footprint** and $\mathbf{v}_k \in \mathbb{S}_{D-1}$ is the **temporal profile**,

and $\mathbb{S}_{N-1} = \{ \mathbf{u} : \mathbf{u} \in \mathbb{R}^N \text{ and } \|\mathbf{u}\|_2 = 1 \}$ is the set of unit vectors in $\mathbb{R}^N$.



MAP estimation

Maximum a posteriori estimation

Coordinate ascent

- Initialize templates $\mathbf{W}$ and set $\mathbf{A} = 0$.
- Iterate until convergence:
 - For neuron $k = 1, \dots, K$:
 - a. Optimize **amplitudes** a_k for neuron k .
 - b. Optimize **templates** $\mathbf{W}_k$ for neuron k .

[In each case, maximize log joint probability wrt one variable, holding others fixed.]

Maximum a posteriori estimation

Deriving the log likelihood

The likelihood is,

$$\begin{aligned} p(\mathbf{X} \mid \mathbf{W}, \mathbf{A}) &= \prod_{n=1}^N \prod_{t=1}^T \mathcal{N} \left(x_{n,t} \mid \sum_{k=1}^K \sum_{d=1}^D a_{k,t-d} w_{k,n,d}, \sigma^2 \right) \\ &= \prod_{n=1}^N \prod_{t=1}^T \mathcal{N} \left(x_{n,t} \mid \sum_{k=1}^K [\mathbf{a}_k \circledast \mathbf{w}_{k,n}]_t, \sigma^2 \right) \\ &= \prod_{t=1}^T \mathcal{N} \left(\mathbf{x}_t \mid \sum_{k=1}^K [\mathbf{a}_k \circledast \mathbf{W}_k]_{:,t}, \sigma^2 \mathbf{I} \right) \end{aligned}$$

Maximum a posteriori estimation

Deriving the log likelihood

Taking the log of both sides and expanding the Gaussian density yields,

$$\begin{aligned}\log p(\mathbf{X} \mid \mathbf{W}, \mathbf{A}) &= -\frac{1}{2\sigma^2} \sum_{t=1}^T \left\| \mathbf{x}_t - \sum_{k=1}^K [\mathbf{a}_k \circledast \mathbf{W}_k]_{:,t} \right\|_2^2 + c \\ &= -\frac{1}{2\sigma^2} \left\| \mathbf{X} - \sum_{k=1}^K [\mathbf{a}_k \circledast \mathbf{W}_k] \right\|_F^2 + c\end{aligned}$$

Maximum a posteriori estimation

Deriving the log likelihood

As a function of $\mathbf{a}_k$ and $\mathbf{W}_k$, fixing the other neurons' amplitudes and templates,

$$\log p(\mathbf{X} \mid \mathbf{A}, \mathbf{W}) = -\frac{1}{2\sigma^2} \|\mathbf{R}_k - \mathbf{a}_k \circledast \mathbf{W}_k\|_F^2$$

where $\mathbf{R}_k \in \mathbb{R}^{N \times T}$ is the residual for neuron k , defined as

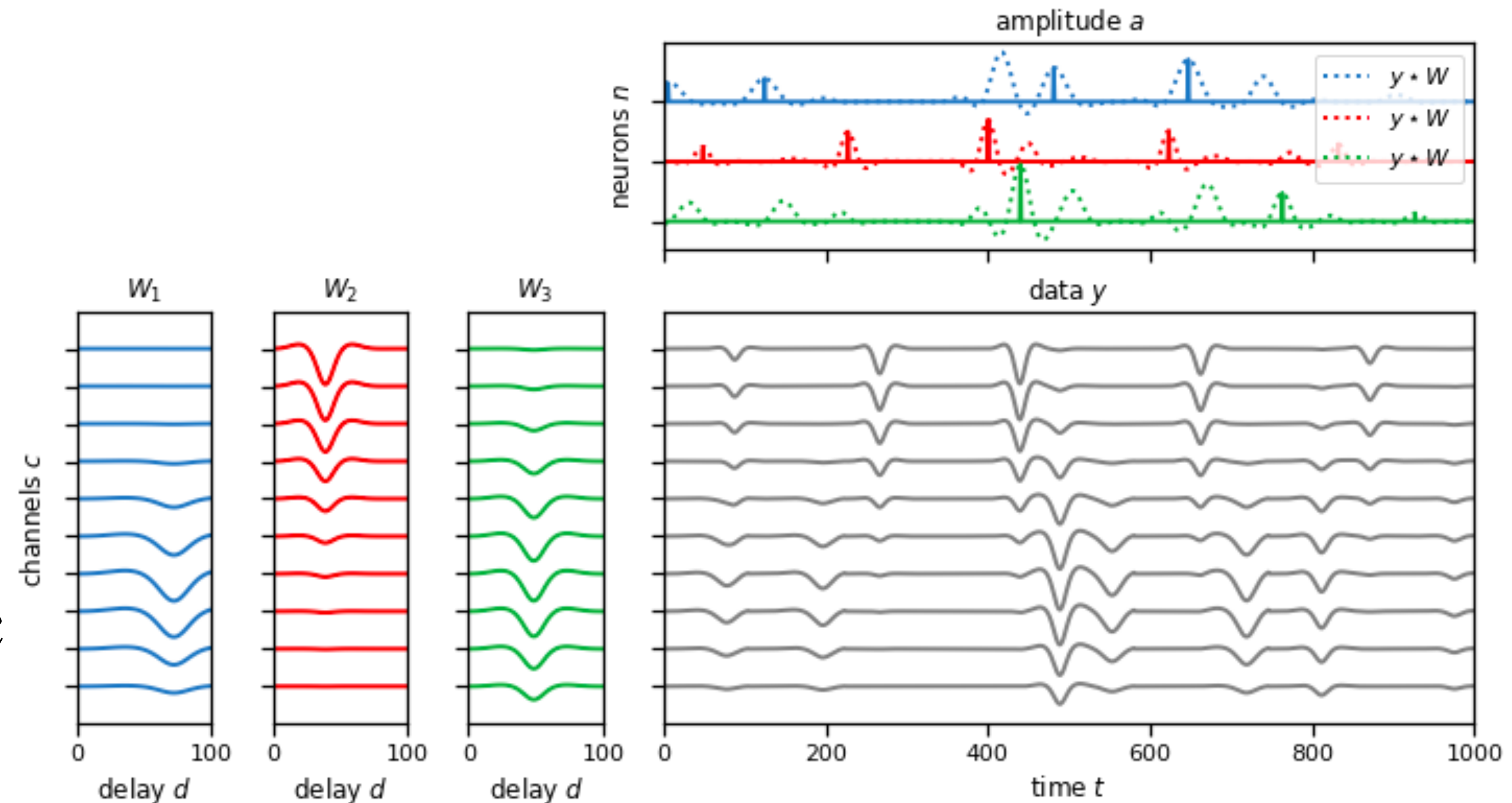
$$\mathbf{R}_k = \mathbf{X} - \sum_{j \neq k} [\mathbf{a}_j \circledast \mathbf{W}_j].$$

Maximum a posteriori estimation

Optimizing the amplitudes

As a function of $\mathbf{a}_k$, the log *joint* probability is the log *likelihood* plus the log *prior* probability,

$$\begin{aligned}\log p(\mathbf{X}, \mathbf{W}, \mathbf{A}) &= \\ &= \log p(\mathbf{X} \mid \mathbf{A}, \mathbf{W}) + \log p(\mathbf{a}_k; \lambda) \\ &= -\frac{1}{2\sigma^2} \|\mathbf{R}_k - \mathbf{a}_k \circledast \mathbf{W}_k\|_F^2 + \log p(\mathbf{a}_k) + c\end{aligned}$$



Maximum a posteriori estimation

Optimizing the amplitudes

Now, expand the square in the Frobenius norm,

$$\begin{aligned}\log p(\mathbf{X}, \mathbf{W}, \mathbf{A}) &= -\frac{1}{2\sigma^2} \|\mathbf{R}_k - \mathbf{a}_k \circledast \mathbf{W}_k\|_F^2 + \log p(\mathbf{a}_k) + c \\ &= \underbrace{-\frac{1}{2\sigma^2} \|\mathbf{a}_k \circledast \mathbf{W}_k\|_F^2}_{\mathcal{L}_2(\mathbf{a}_k)} + \underbrace{\frac{1}{\sigma^2} \langle \mathbf{R}_k, \mathbf{a}_k \circledast \mathbf{W}_k \rangle_F}_{\mathcal{L}_1(\mathbf{a}_k)} + \log p(\mathbf{a}_k) + c.\end{aligned}$$

Maximum a posteriori estimation

Optimizing the amplitudes

Further expanding the quadratic term,

$$\begin{aligned}\mathcal{L}_2(\mathbf{a}_k) &= -\frac{1}{2\sigma^2} \|\mathbf{a}_k \circledast \mathbf{W}_k\|_F^2 \\ &= -\frac{1}{2\sigma^2} \sum_{t=1}^T \sum_{n=1}^N \left(\sum_{d=1}^D a_{k,t-d}^2 w_{k,n,d}^2 + 2 \sum_{d=1}^D \sum_{d'=1}^{d-1} a_{k,t-d} a_{k,t-d'} w_{k,n,d} w_{k,n,d'} \right) \\ &\approx -\frac{1}{2\sigma^2} \sum_{t=1}^T a_{k,t}^2 \|\mathbf{W}_k\|_F^2 \\ &= -\frac{1}{2\sigma^2} \sum_{t=1}^T a_{k,t}^2\end{aligned}$$

with equality when nonzero entries (i.e. “spikes”) of $\mathbf{a}_k$ are separated by at least D samples.

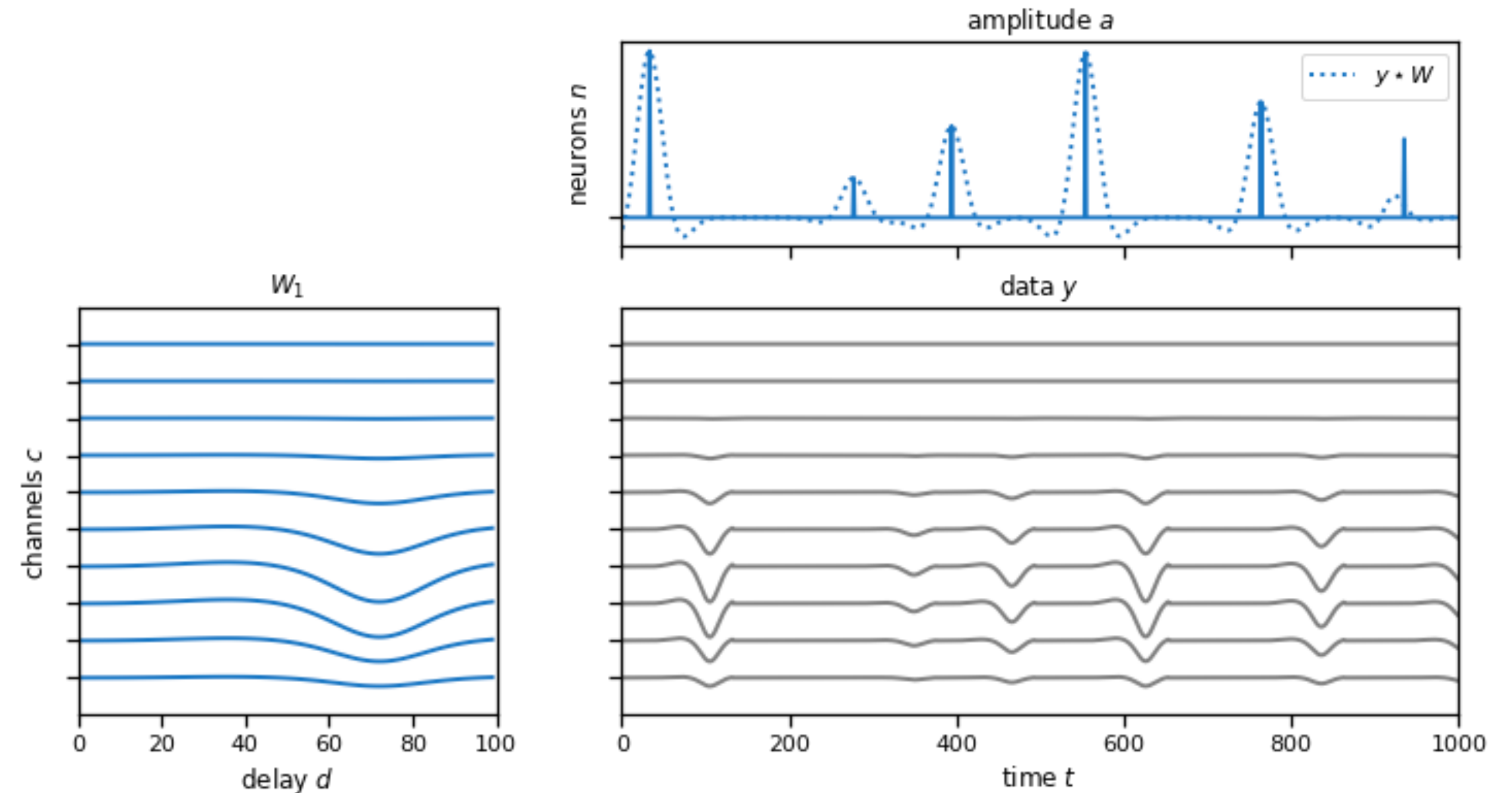
Maximum a posteriori estimation

Optimizing the amplitudes

Now take the linear term...

$$\begin{aligned}
 \mathcal{L}_1(\mathbf{a}_k) &= \frac{1}{\sigma^2} \langle \mathbf{R}_k, \mathbf{a}_k \circledast \mathbf{W}_k \rangle \\
 &= \frac{1}{\sigma^2} \sum_{t=1}^T \sum_{n=1}^N r_{k,n,t} [\mathbf{a}_k \circledast \mathbf{w}_{k,n}]_t \\
 &= \frac{1}{\sigma^2} \sum_{t=1}^T \sum_{n=1}^N \sum_{d=1}^D a_{k,t-d} r_{k,n,t} w_{k,n,d} \\
 &= \frac{1}{\sigma^2} \sum_{t=1}^T a_{k,t} \sum_{n=1}^N \sum_{d=1}^D r_{k,n,t+d} w_{k,n,d} \\
 &= \frac{1}{\sigma^2} \sum_{t=1}^T a_{k,t} [\mathbf{R}_k \star \mathbf{W}_k]_t
 \end{aligned}$$

where $[\mathbf{R} \star \mathbf{W}_k]_t$ is the cross-correlation of the residual and the template for neuron k .



Maximum a posteriori estimation

Optimizing the amplitudes

Putting it all together

$$\mathcal{L}(\mathbf{a}_k) = \sum_{t=1}^T \left[-\frac{1}{2\sigma^2} a_{k,t}^2 + \frac{1}{\sigma^2} \mu_{k,t} a_{k,t} - \lambda a_{k,t} \right] + c,$$

which separates into a sum of quadratic objective functions for each time t .

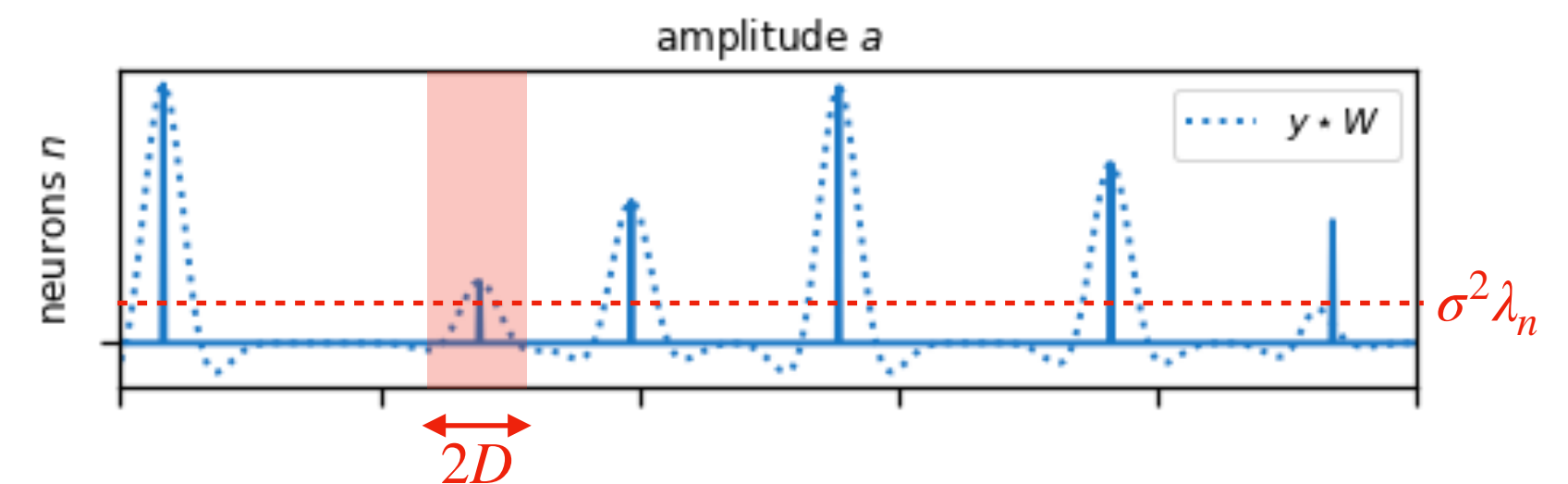
Maximum a posteriori estimation

Completing the square and solving for the optimal amplitudes

- The maximum of this quadratic objective, subject to non-negativity constraints, is obtained at

$$a_{k,t} = \max \{0, \mu_{k,t} - \sigma^2 \lambda\}$$

- However, we also want spikes to be well-separated; i.e. $a_{k,t} > 0 \implies a_{k,t+d} = 0$ for $d = 1, \dots, D$.
- We'll enforce this with a **simple heuristic**: use `find_peaks` to select local maxima of this “score” signal.



Maximum a posteriori estimation

Optimizing the templates

As a function of $\mathbf{W}_k$

$$\begin{aligned}\log p(\mathbf{X}, \mathbf{A}, \mathbf{W}) &= -\frac{1}{2\sigma^2} \sum_{t=1}^T \langle a_{k,t} \mathbf{R}_{k,t}, \mathbf{W}_k \rangle + c' \\ &= -\frac{1}{2\sigma^2} \left\langle \sum_{t=1}^T a_{k,t} \mathbf{R}_{k,t}, \mathbf{W}_k \right\rangle + c'\end{aligned}$$

where

$$\mathbf{R}_{k,t} = \begin{bmatrix} r_{k,1,t} & \cdots & r_{k,1,t+D} \\ \vdots & & \vdots \\ r_{k,N,t} & \cdots & r_{k,N,t+D} \end{bmatrix}$$

is a slice of the residual matrix ($\mathbf{R}[\mathbf{k}, :, \mathbf{t} : \mathbf{t} + \mathbf{D}]$ in code).

Maximum a posteriori estimation

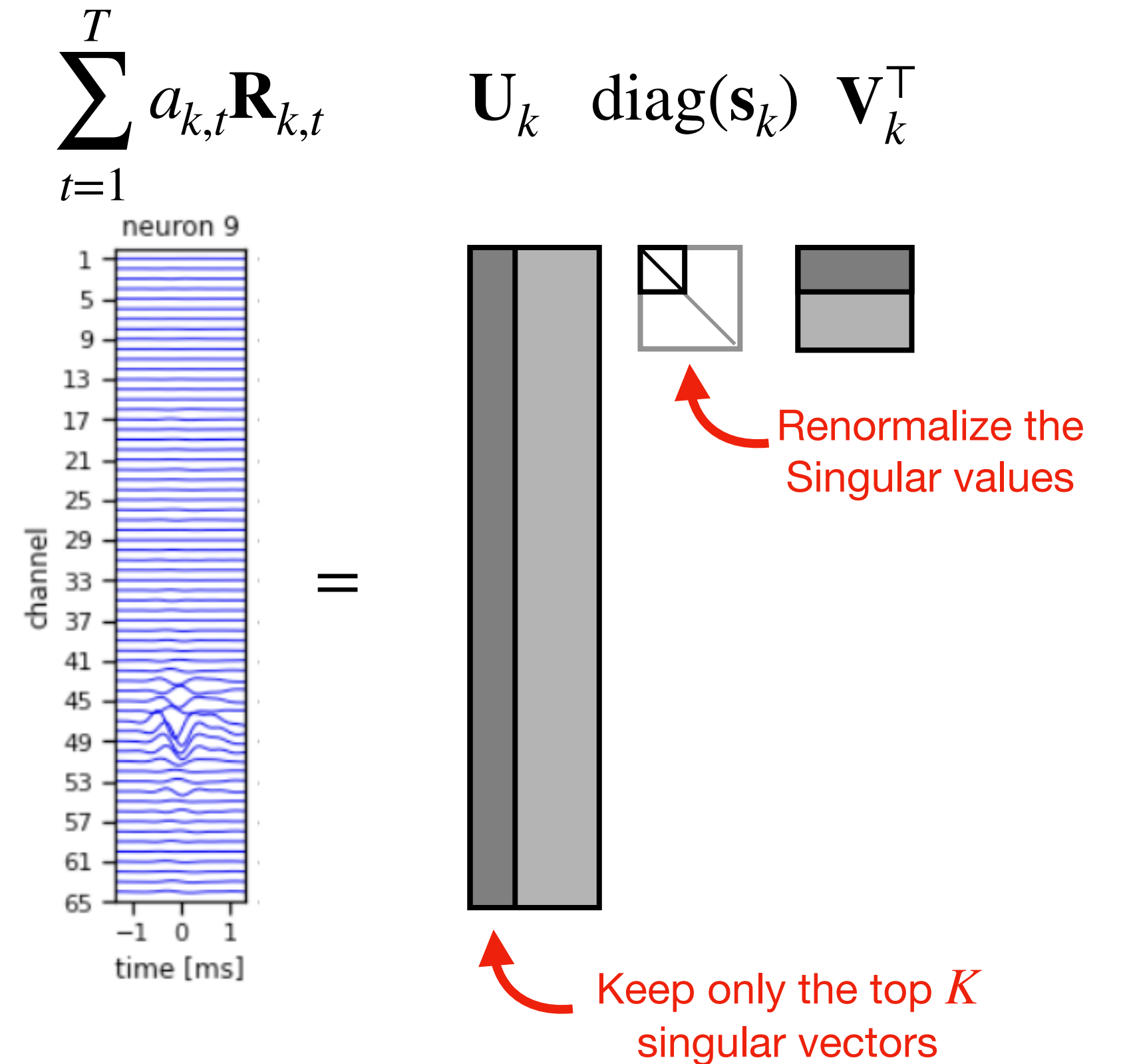
Optimizing the templates

We want to maximize this log joint probability over the space of low-rank, unit-norm matrices,

$$\mathbf{W}_k^\star = \arg \max_{\mathbf{W}_k \in \mathbb{S}_R^{N,D}} \left\langle \sum_{t=1}^T a_{k,t} \mathbf{R}_{k,t}, \mathbf{W}_k \right\rangle$$

The solution is to set the waveform matrix "proportional to" the weighted sum of residual matrices by taking its SVD and renormalizing the singular values.

$$\mathbf{W}_k^\star = \sum_{r=1}^R \bar{s}_r \mathbf{u}_r \mathbf{v}_r^\top \quad \text{where} \quad \bar{s}_r = \frac{s_r}{\sqrt{\sum_{r'=1}^R s_{r'}^2}}.$$



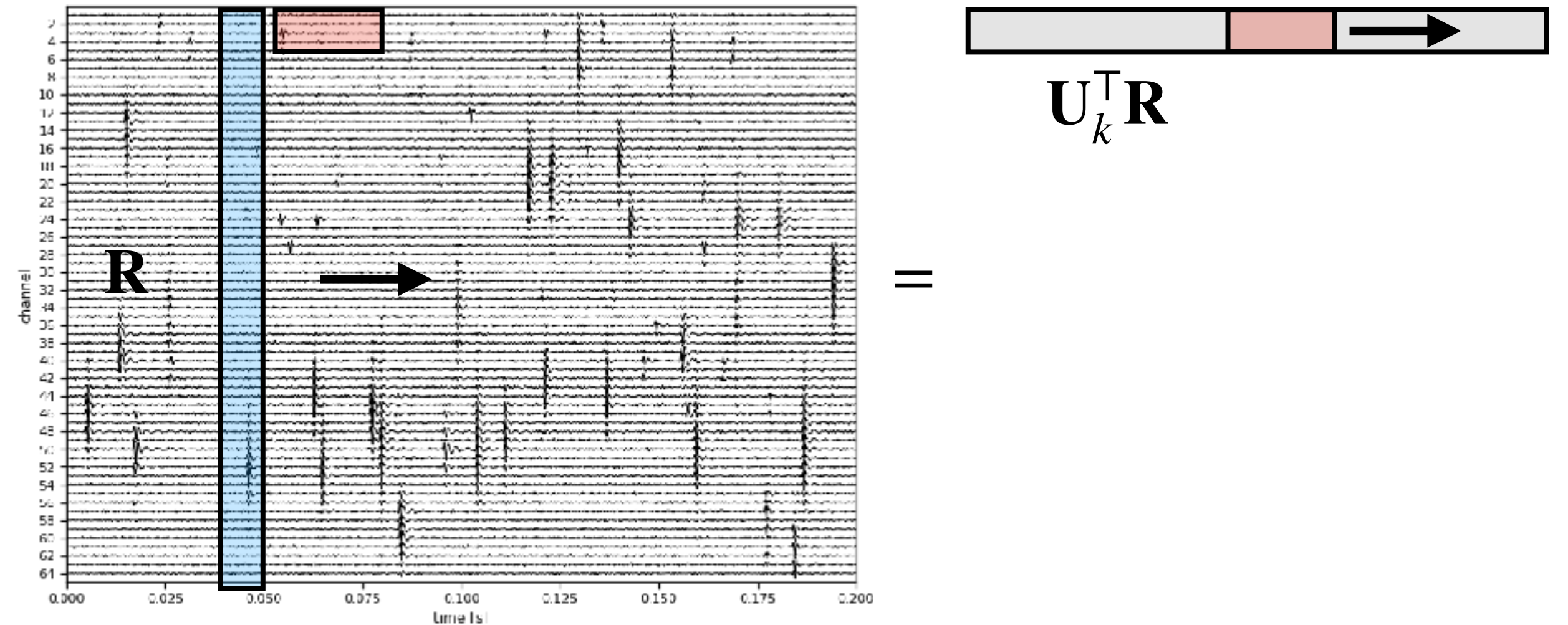
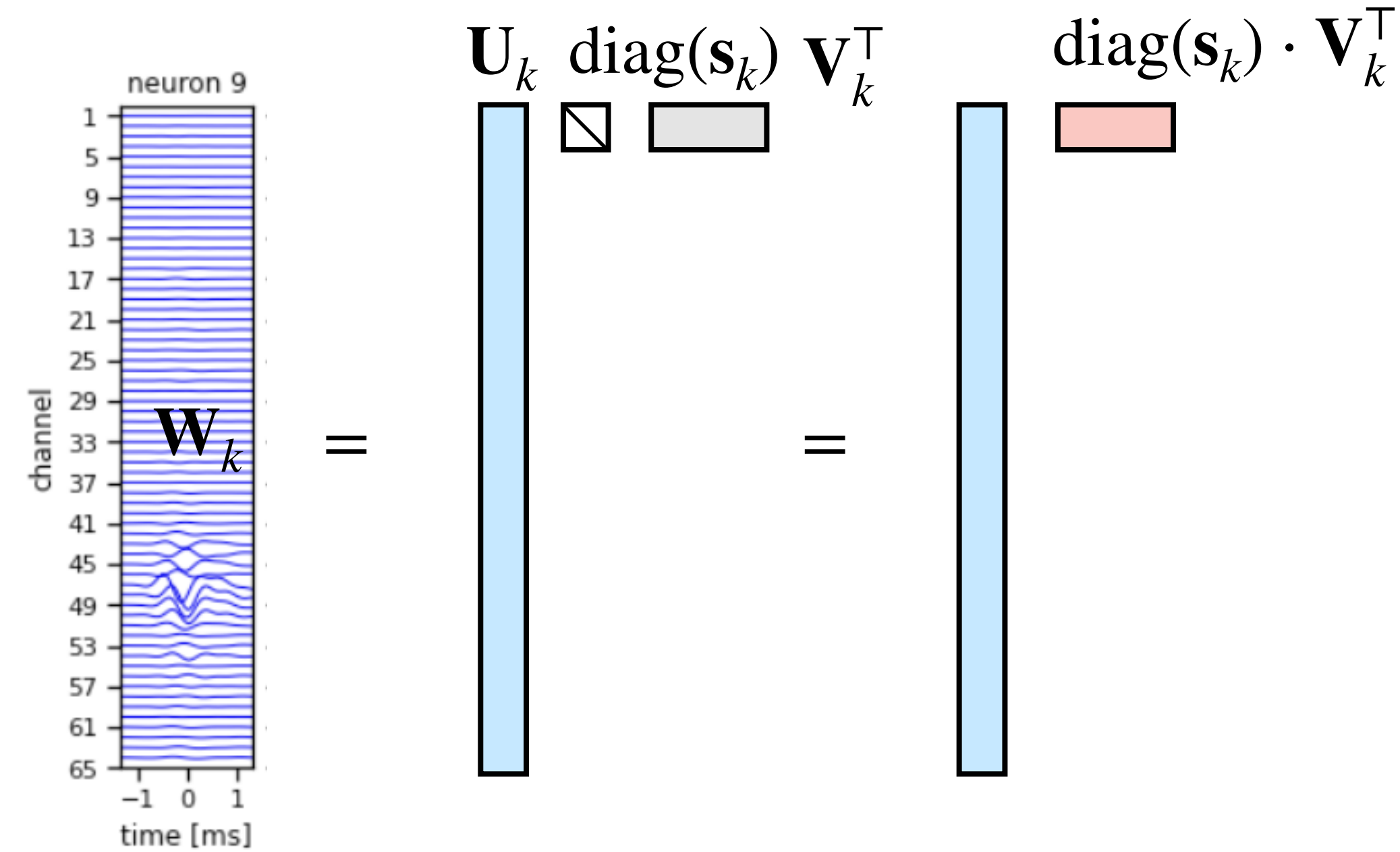
More efficient computation

Leveraging the low-rank templates

We can compute the “scores” for amplitude updates more efficiently by leveraging the low-rank templates,

$$\begin{aligned}
 [\mathbf{R}_k \star \mathbf{W}_k]_t &= \sum_{n=1}^N \sum_{d=1}^D r_{k,n,t+d} w_{k,n,d} \\
 &= \sum_{d=1}^D \mathbf{r}_{k,:,t+d}^\top \mathbf{w}_{k,:,d} \\
 &= \sum_{d=1}^D \mathbf{r}_{k,:,t+d}^\top \mathbf{U}_k \mathbf{S}_k \mathbf{v}_{k,:,d} \\
 &= \sum_{d=1}^D (\mathbf{U}_k^\top \mathbf{r}_{k,:,t+d})^\top [\mathbf{S}_k \mathbf{V}_k^\top]_{:,d} \\
 &= [(\mathbf{U}_k^\top \mathbf{R}_k) \star (\mathbf{S}_k \mathbf{V}_k^\top)]_t
 \end{aligned}$$

In other words, we cross-correlate the projected residual.



Conclusion

- We developed a basic spike sorting model that was good for building intuition, but not very practical.
- We developed a new model for the voltage in terms of a superposition of templates convolved with spike amplitudes for each neuron.
 - Along the way, we learned about convolution and cross-correlation.
- We derived a **coordinate ascent algorithm** for *maximum a posteriori* (MAP) inference.
- **Next time:** you'll implement the algorithm in lab! You'll learn a bit of PyTorch for implementing the convolutions and cross-correlations, then test it out on the GPU.

Further reading

- **Simple Spike Sorting** and **Spike Sorting by Deconvolution** course notes.
- Convolution and cross-correlation:
 - Chapter 9 of *The Deep Learning Book* (deeplearningbook.org/contents/convnets.html)
 - Start reading up on PyTorch convolutions! <https://pytorch.org/docs/stable/generated/torch.nn.functional.conv1d.html>
- Spike sorting:
 - Pachitariu, Marius, Shashwat Sridhar, and Carsen Stringer. "Solving the spike sorting problem with Kilosort." bioRxiv (2023).
 - The model we presented is a slightly modified version of *Kilosort*