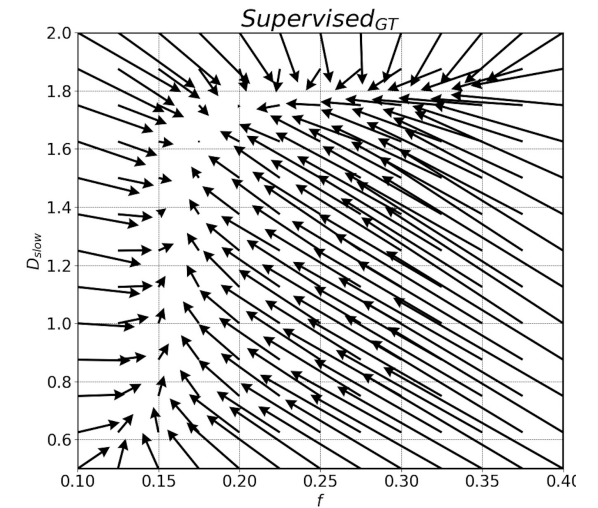
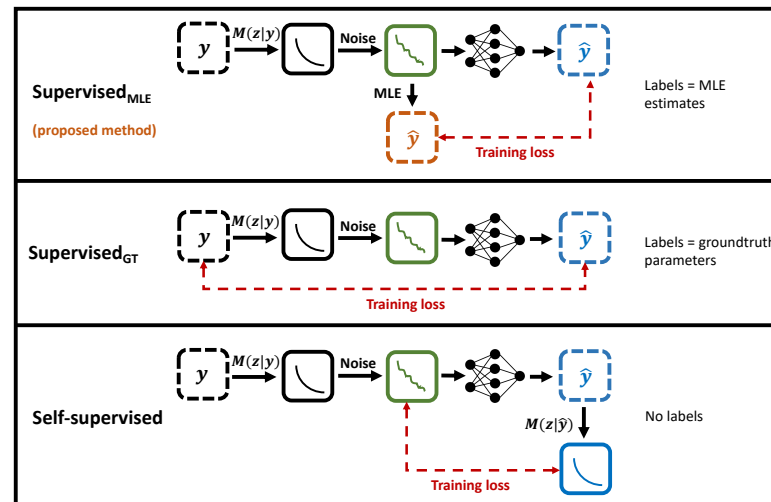
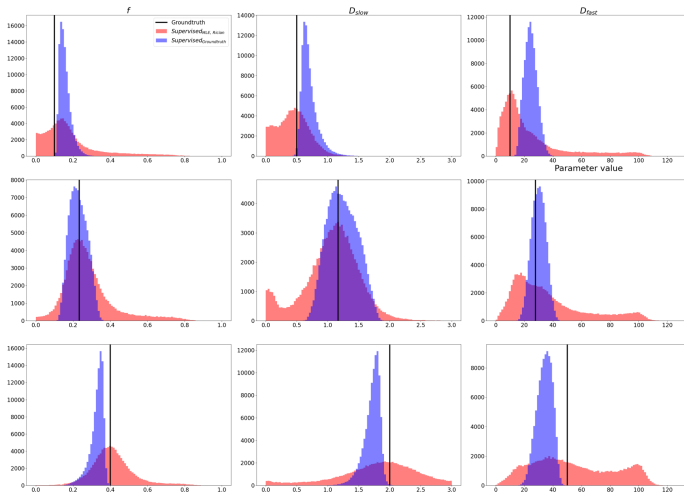


Choice of Training Label Matters: Deep Learning for Quantitative MRI Parameter Estimation

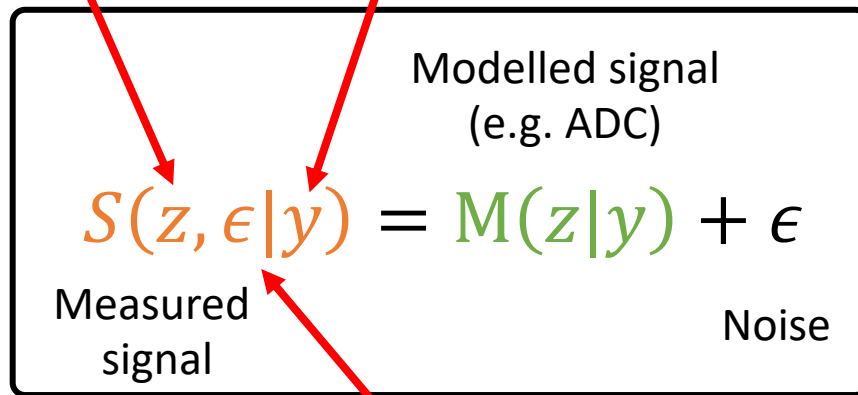
Sean Epstein

Centre for Medical Image Computing, UCL



Acquisition scheme
(e.g. b-value)

Tissue property
(e.g. diffusivity)



The diagram consists of a central black-bordered box containing the equation $S(z, \epsilon|y) = M(z|y) + \epsilon$. Above the box, two red arrows point downwards: one from the text 'Acquisition scheme (e.g. b-value)' pointing to the y in the equation, and another from 'Tissue property (e.g. diffusivity)' pointing to the z . Inside the box, the text 'Modelled signal (e.g. ADC)' is positioned above the $M(z|y)$ term. Below the equation, the text 'Measured signal' is aligned under $S(z, \epsilon|y)$ and 'Noise' is aligned under ϵ . A red arrow points from the text 'Noise instantiation' below the box to the ϵ term in the equation.

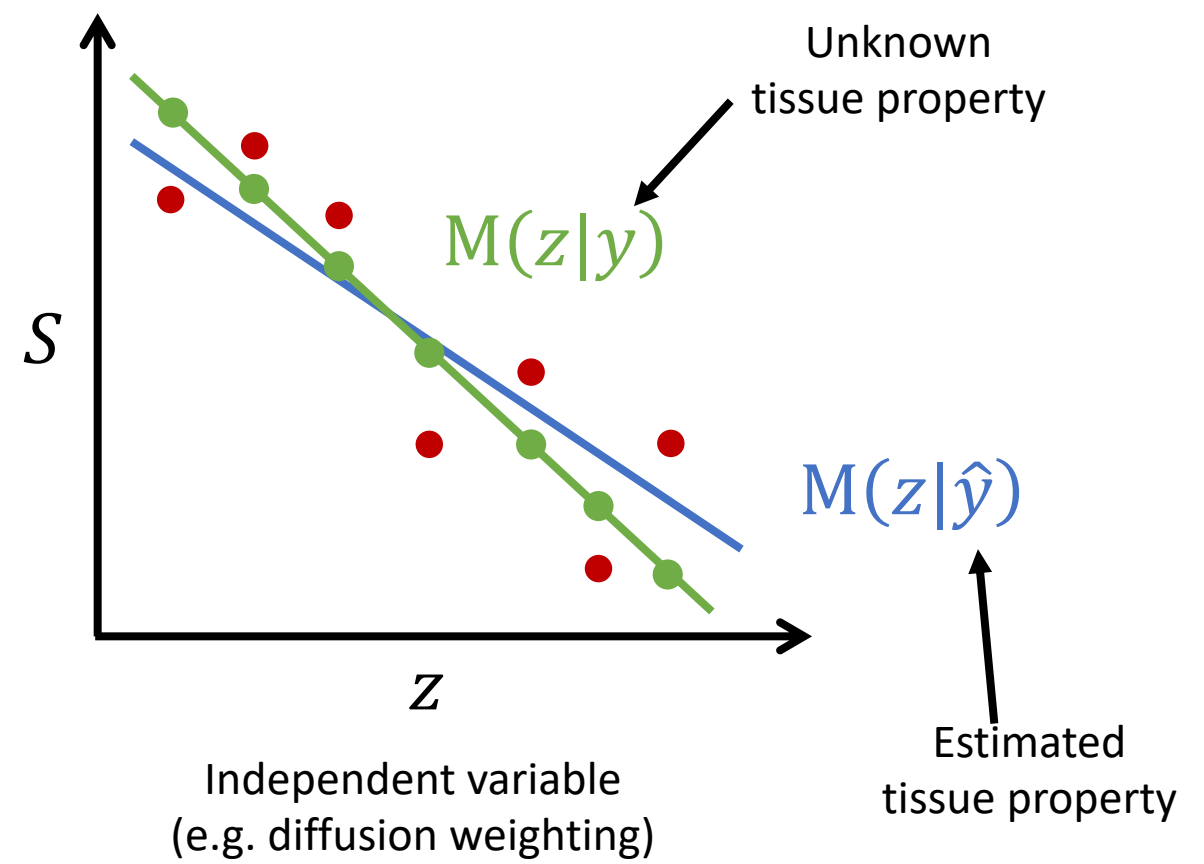
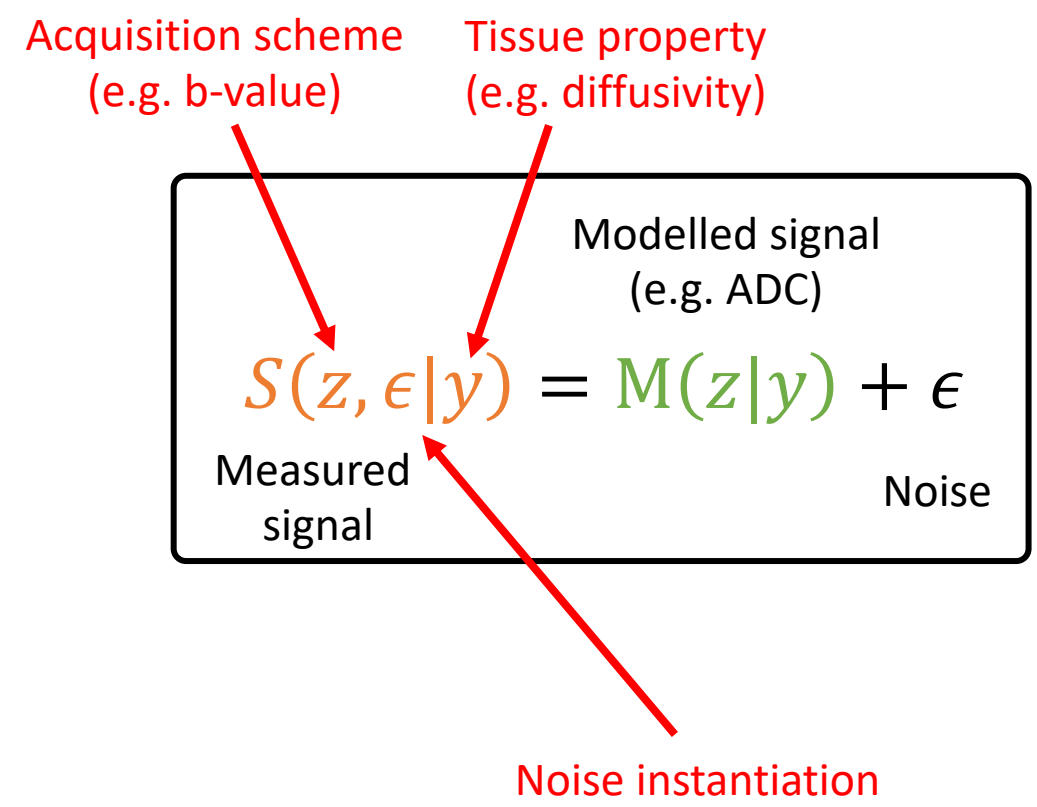
$$S(z, \epsilon|y) = M(z|y) + \epsilon$$

Modelled signal
(e.g. ADC)

Measured
signal

Noise

Noise instantiation



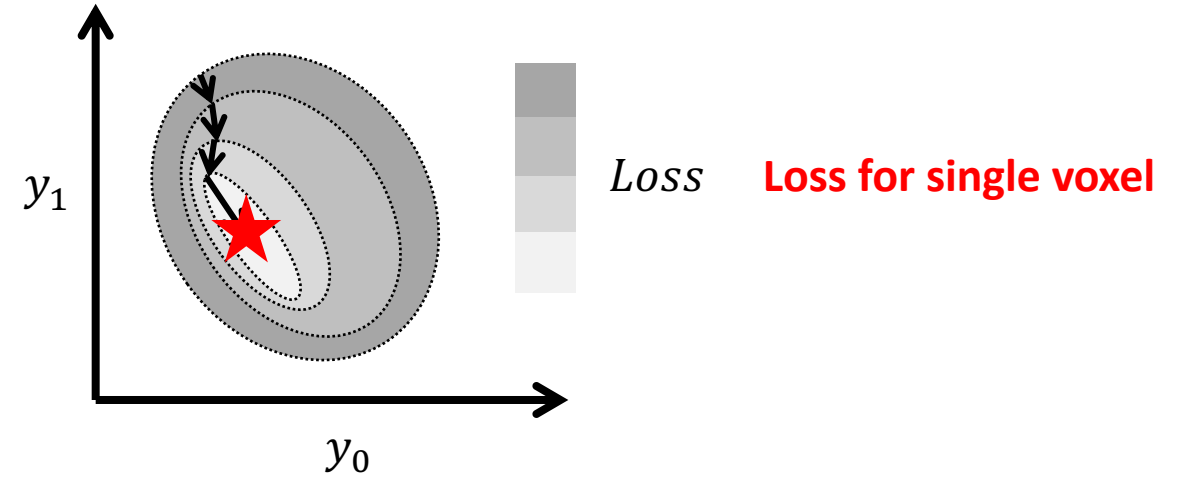
For a given noisy S , we want the corresponding y

Traditional (MLE): $\min_y \text{Loss}(y|S)$

Tackle the problem **directly**

Optimise over **variable of interest** y

Repeat optimisation for every signal



How do we solve this inverse problem?

For a given noisy S , we want the corresponding y

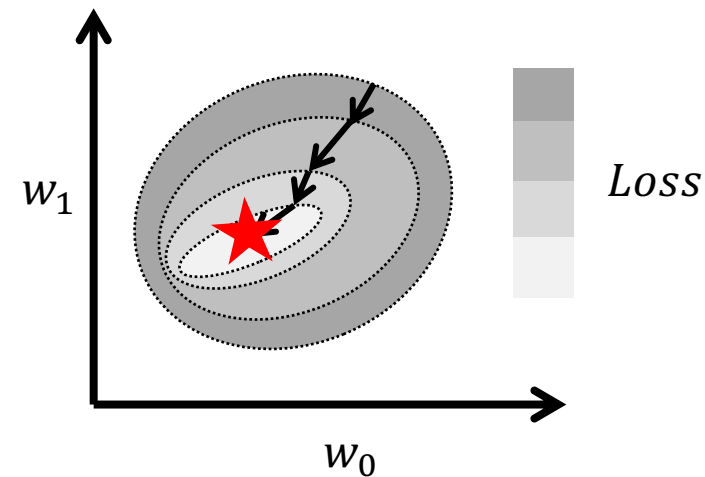
Machine learning (ML): $\min_w \sum_{i=1}^{N_{train}} Loss(w|S_i)$

Tackle the problem **indirectly**

Find **general mapping** from any S to corresponding y

Optimise over **latent variable** w

Perform optimisation **only once** (training)



**Loss averaged over
all training voxels**

For a given noisy S , we want the corresponding y

Traditional (MLE): $\min_y \text{Loss}(y|S)$

Tackle the problem **directly**

Optimise over **variable of interest** y

Repeat optimisation for every signal

Machine learning (ML): $\min_w \sum_{i=1}^{N_{train}} \text{Loss}(w|S_i)$

Tackle the problem **indirectly**

Find **general mapping** from any S to corresponding y

Optimise over **latent variable** w

Perform optimisation **only once** (training)

Limitations

- **Cost:** expensive, scales linearly with newly acquired data
- **Performance:** suffers from local minima; each optimisation solved in isolation

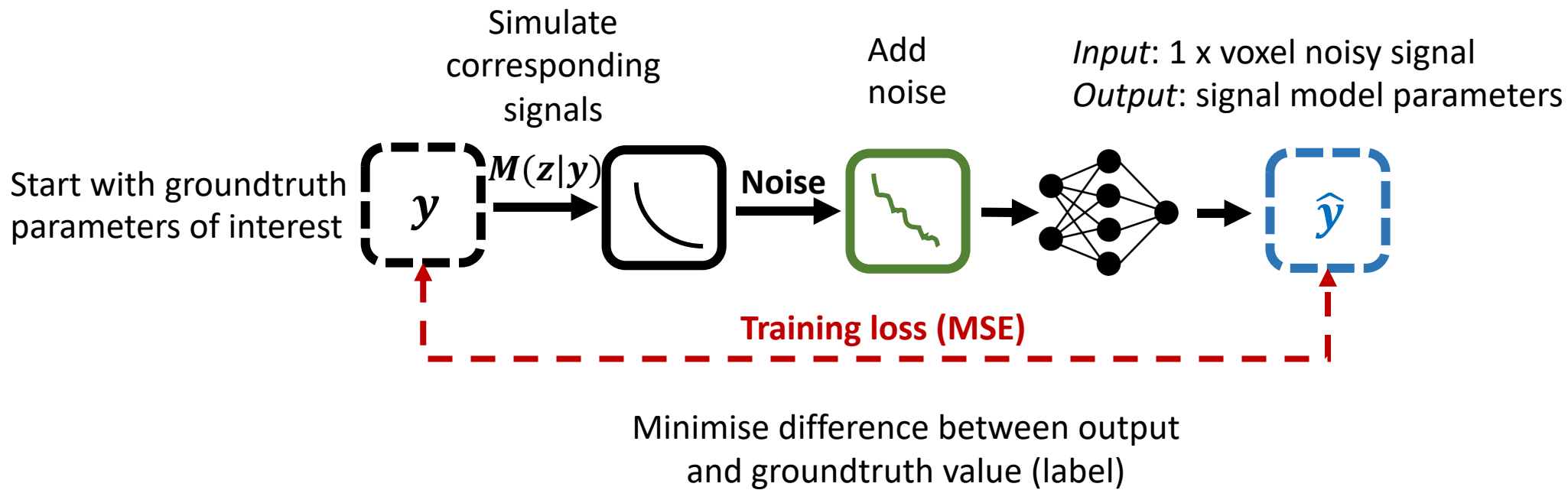
Solutions

- **Cost:** frontloaded (training), then negligible
- **Performance:** leverages patterns across training voxels

Supervised methods
e.g. Bertleff 2017, Gyori 2022

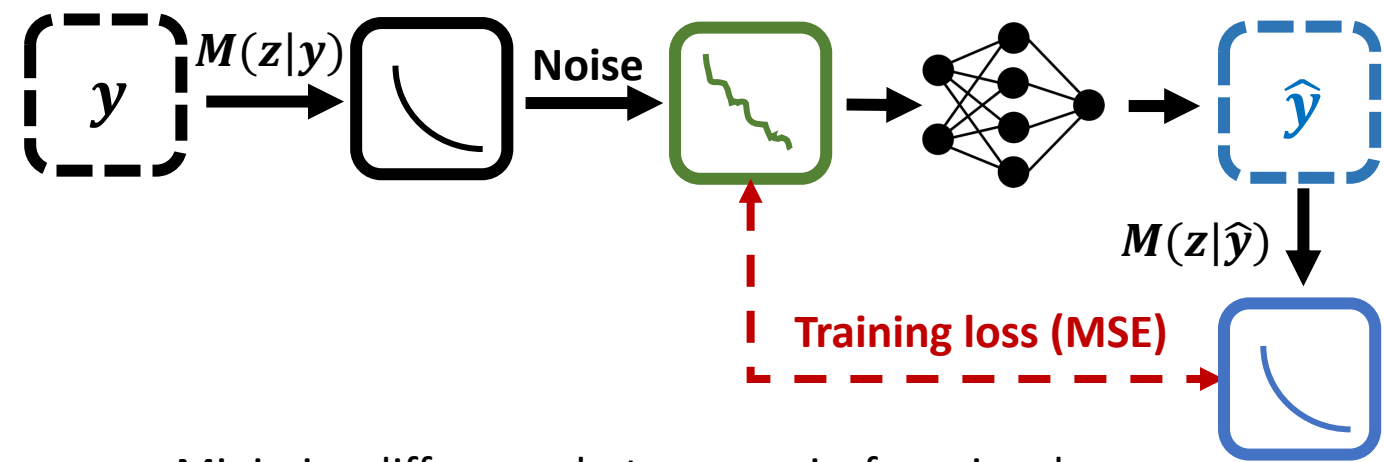
Synthetic training data

Self-supervised (“unsupervised”) methods
e.g. Barbieri 2019, Kaandorp 2021



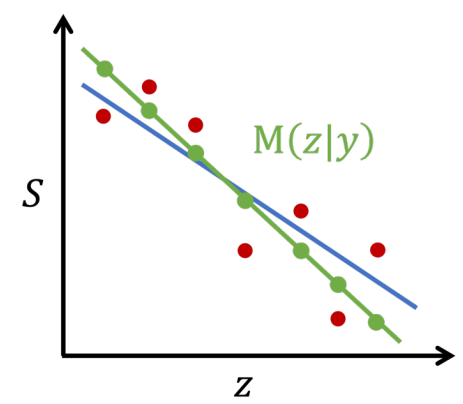
Supervised methods
e.g. Bertleff 2017, Gyori 2022

Self-supervised (“unsupervised”) methods
e.g. Barbieri 2019, Kaandorp 2021



Minimise difference between noisefree signal representation of output and noisy input

Replicating traditional MLE, regularised over a large training dataset



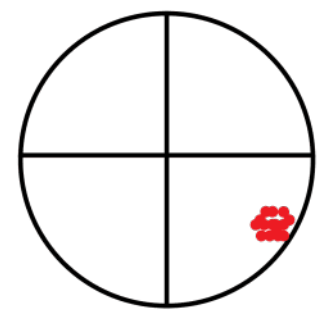
Supervised methods
e.g. Bertleff 2017, Gyori 2022



Low variance: consistent parameter estimates under noise repetition

High bias: estimates biased away from groundtruth

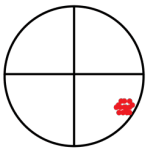
Low information content: bias depends on groundtruth, i.e. different groundtruths indistinguishable



e.g. Gyori 2022, Grussu 2021

Supervised methods

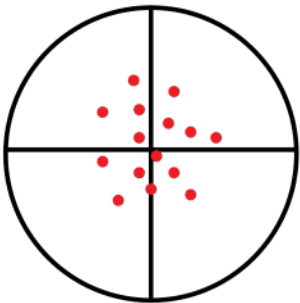
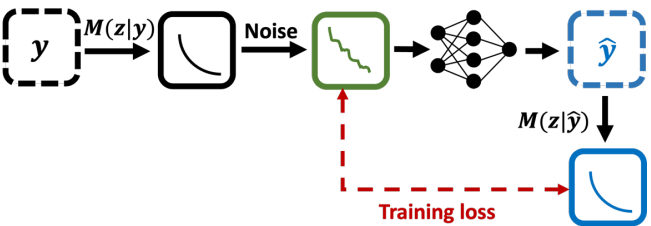
e.g. Bertleff 2017, Gyori 2022



e.g. Gyori 2022, Grussu 2021

Self-supervised methods

e.g. Barbieri 2019, Kaandorp 2021



e.g. Grussu 2021, Barbieri 2019

Lower bias: mean parameter estimates closer to groundtruth

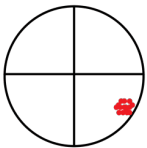
High variance: wide parameter estimation distribution under noise

Loss calculated in signal-space:

- requires differentiable loss formulation (i.e. MSE, Gaussian noise assumption; limits signal models);
- relative parameter loss weighting limited by acquisition protocol z

Supervised methods

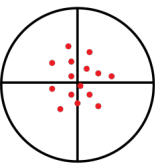
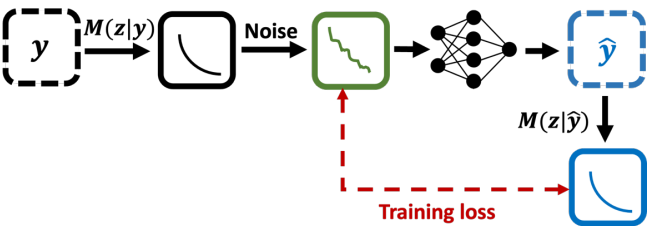
e.g. Bertleff 2017, Gyori 2022



e.g. Gyori 2022, Grussu 2021

Self-supervised methods

e.g. Barbieri 2019, Kaandorp 2021



e.g. Grussu 2021, Barbieri 2019

Spectrum: bias/variance trade-off

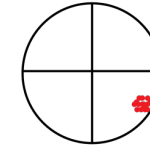
Self-supervised (low bias) has practical limitations

Would like to address these limitations and **move along this bias/variance spectrum**

Supervised training with a change of label

Supervised methods

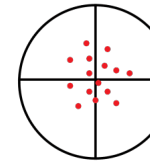
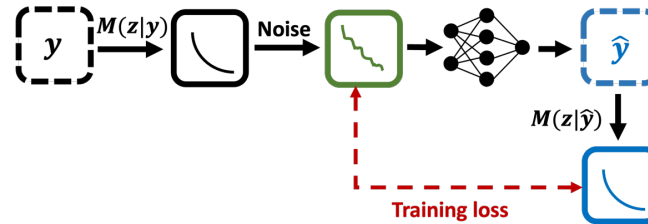
e.g. Bertleff 2017, Gyori 2022



e.g. Gyori 2022, Grussu 2021

Self-supervised methods

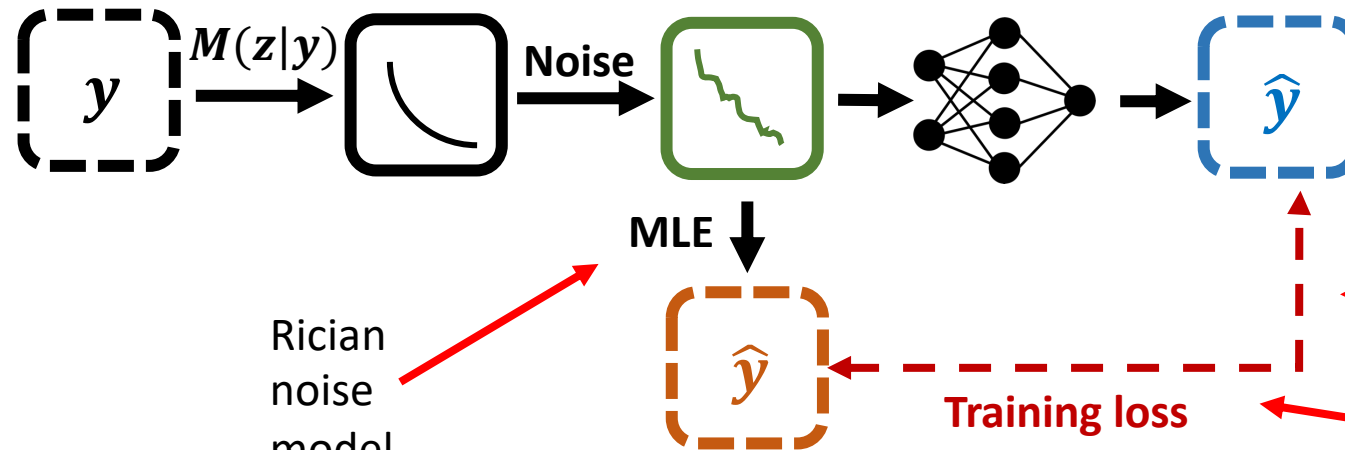
e.g. Barbieri 2019, Kaandorp 2021



e.g. Grussu 2021, Barbieri 2019

Supervised methods (MLE labels)

Epstein 2022

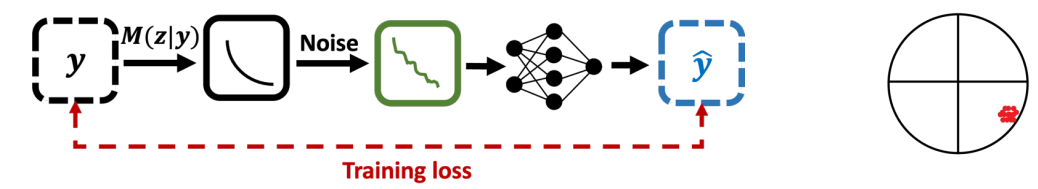


Loss computed in parameter space

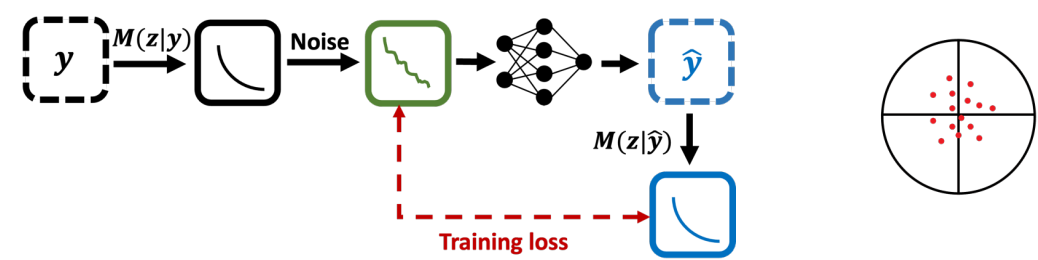
Can use any signal model

$$\text{Hybrid loss} = \alpha \cdot \text{Supervised}_{\text{MLE}} \text{ loss} + (1 - \alpha) \cdot \text{Supervised}_{\text{GT}} \text{ loss}$$

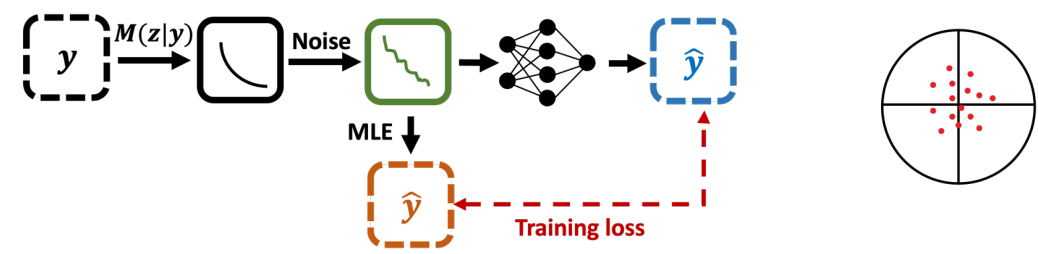
**Supervised methods
(groundtruth labels)**
e.g. Bertleff 2017, Gyori 2022



Self-supervised methods
e.g. Barbieri 2019, Kaandorp 2021



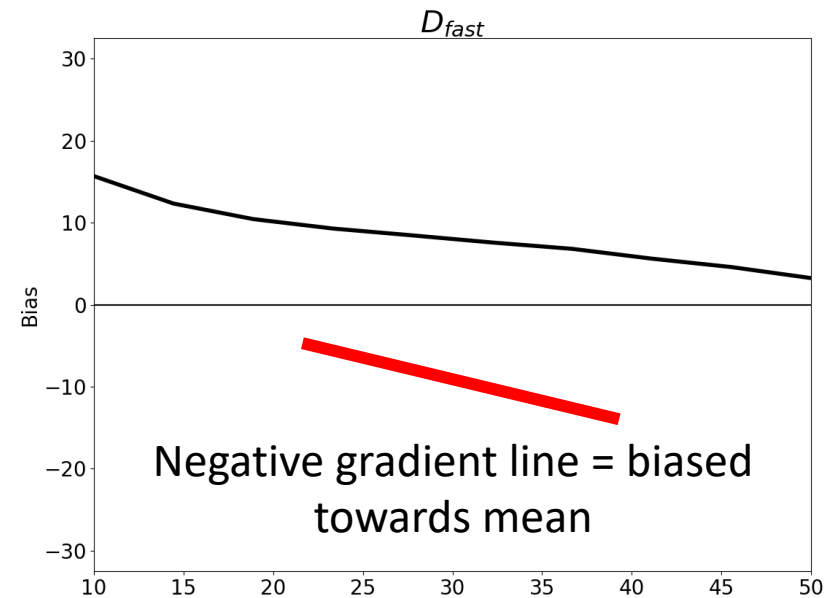
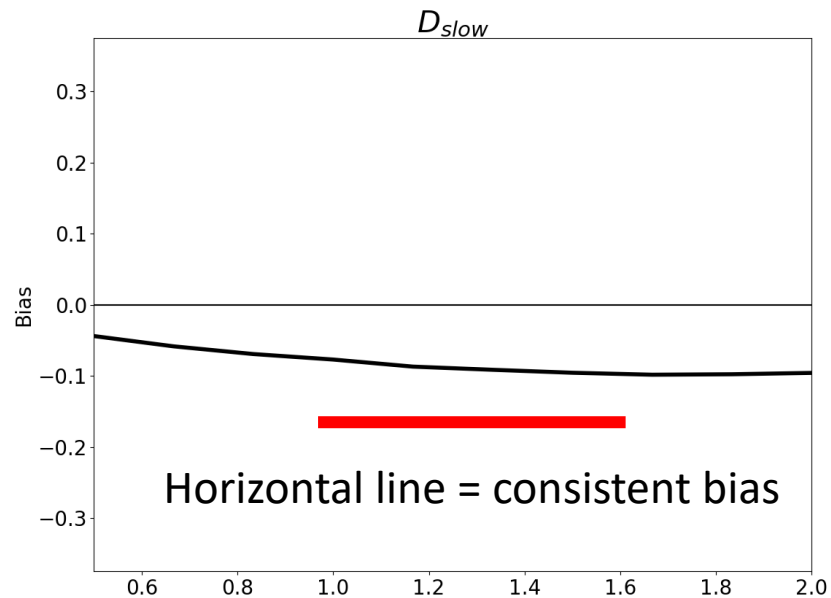
**Supervised methods
(MLE labels)**
e.g. Epstein 2022



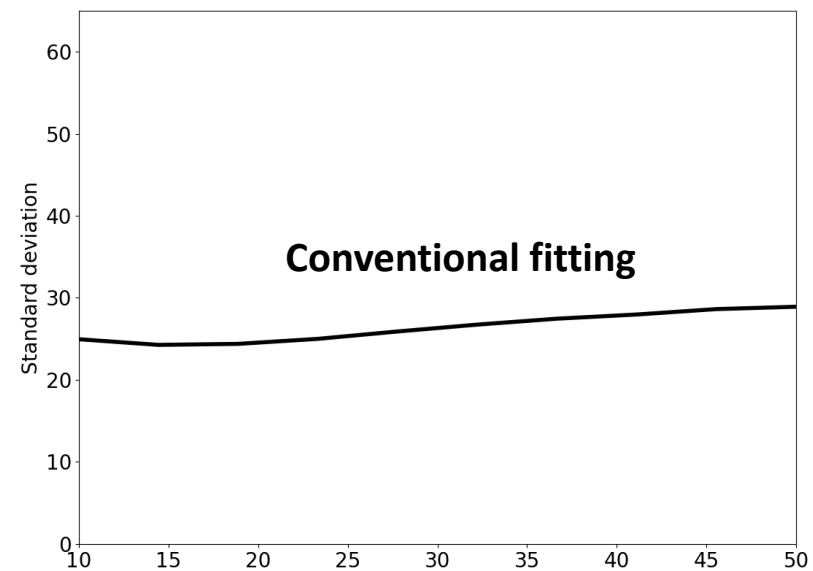
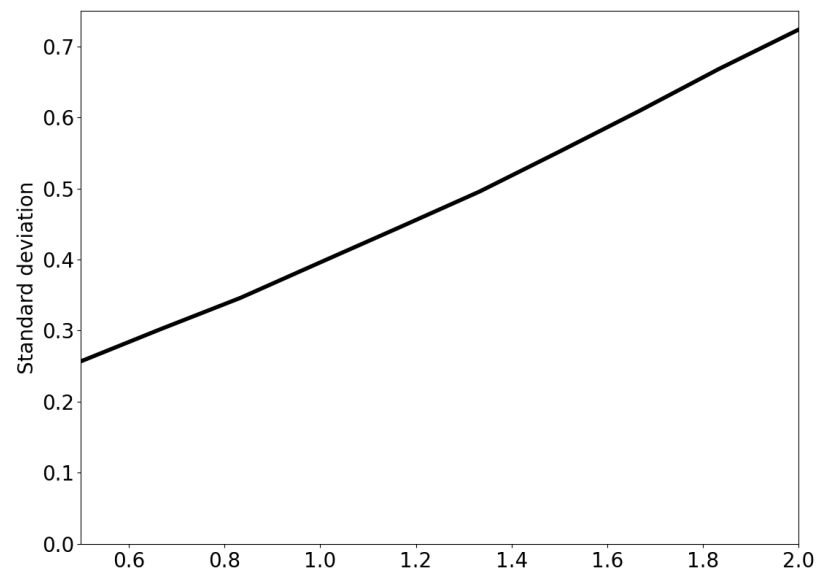
y_1 : easy to fit

y_2 : difficult to fit

BIAS



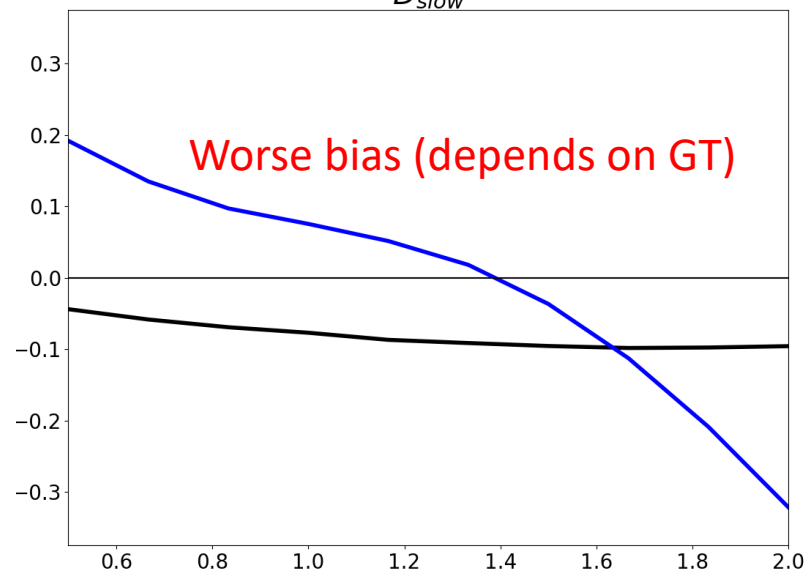
VARIANCE



y_1 : easy to fit

D_{slow}

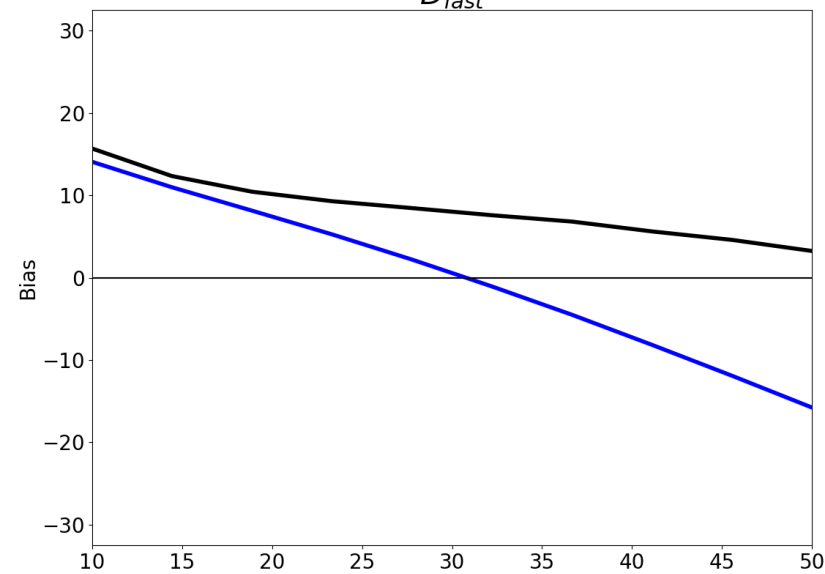
BIAS



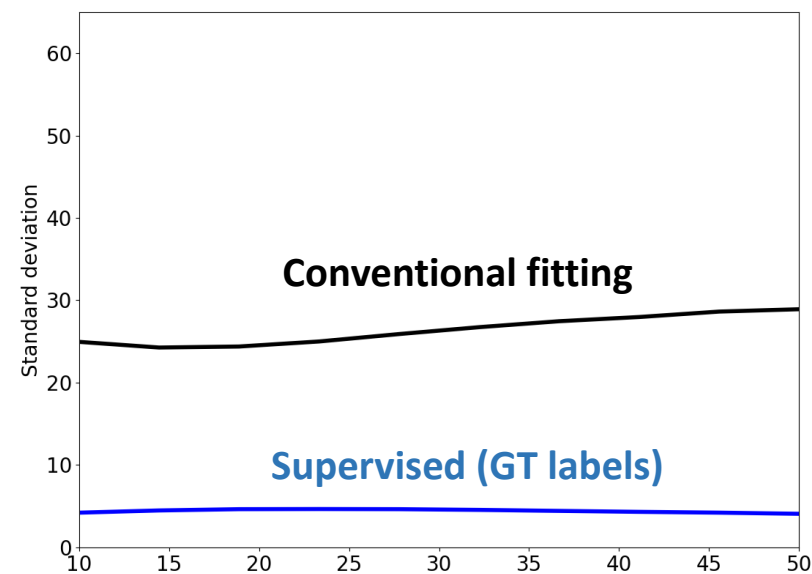
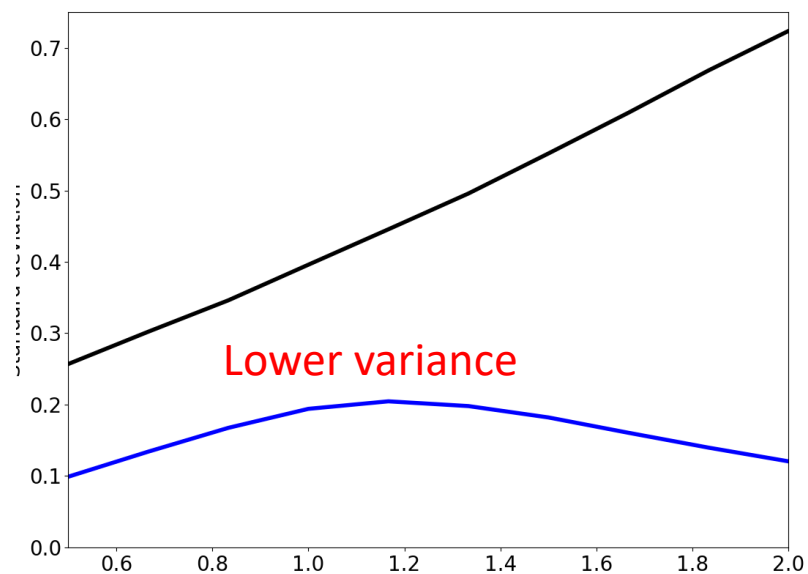
y_2 : difficult to fit

D_{fast}

Bias



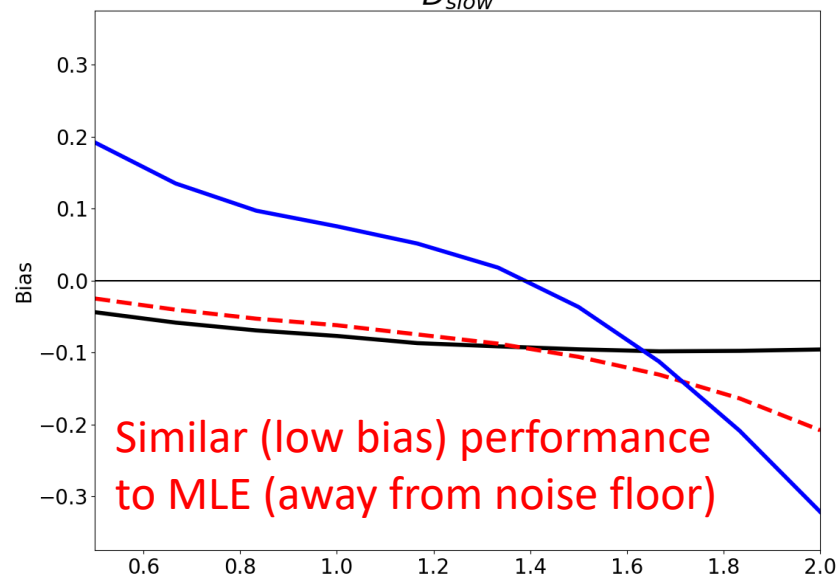
VARIANCE



y_1 : easy to fit

D_{slow}

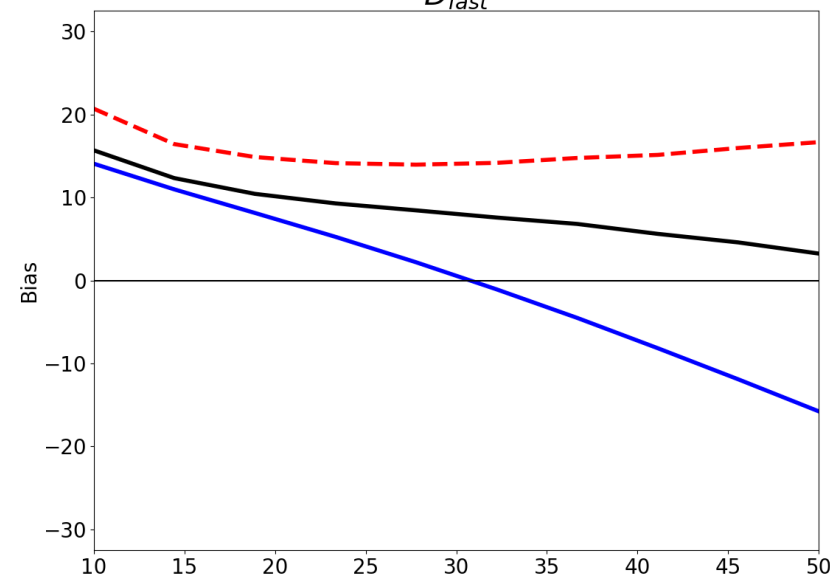
BIAS



y_2 : difficult to fit

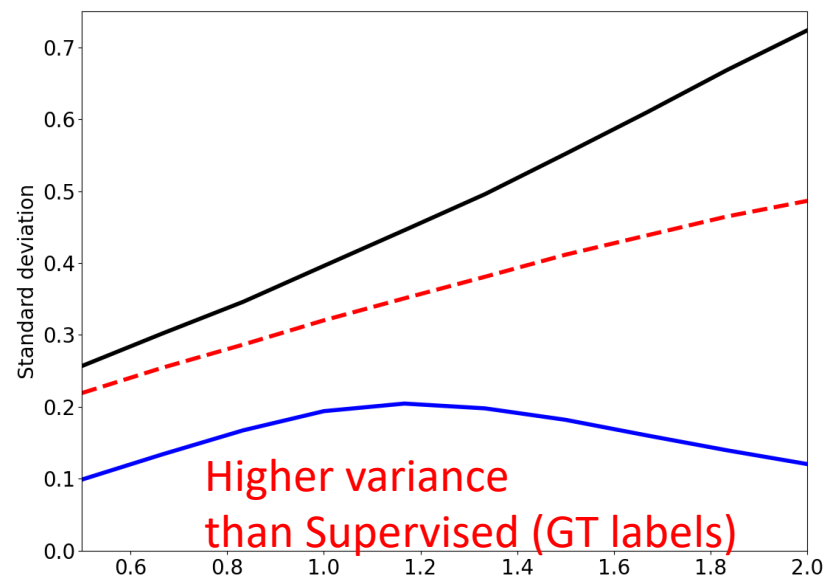
D_{fast}

BIAS

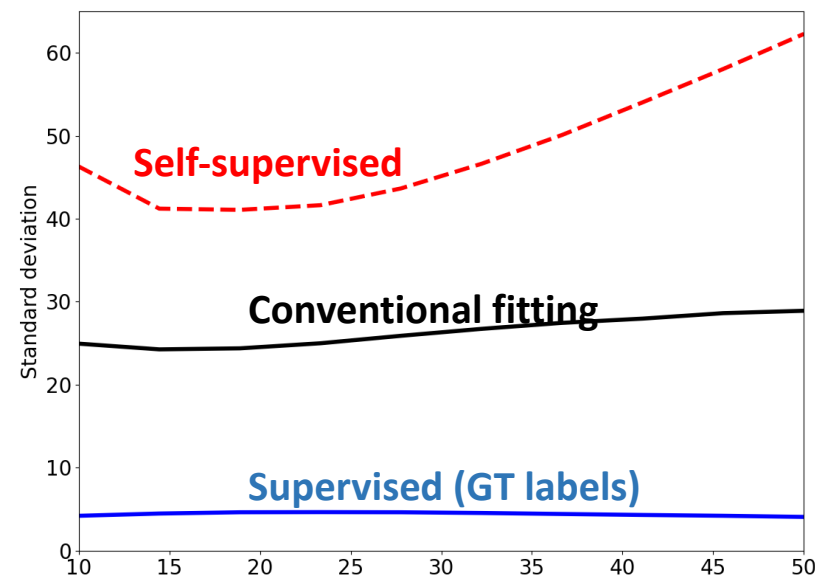


Worse y_2 performance: underweight in signal space

VARIANCE



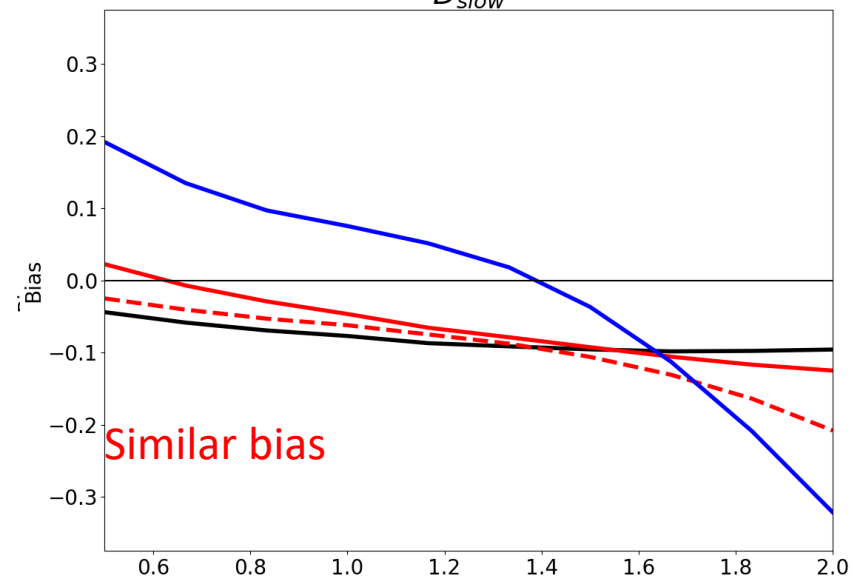
Higher variance than Supervised (GT labels)



y_1 : easy to fit

D_{slow}

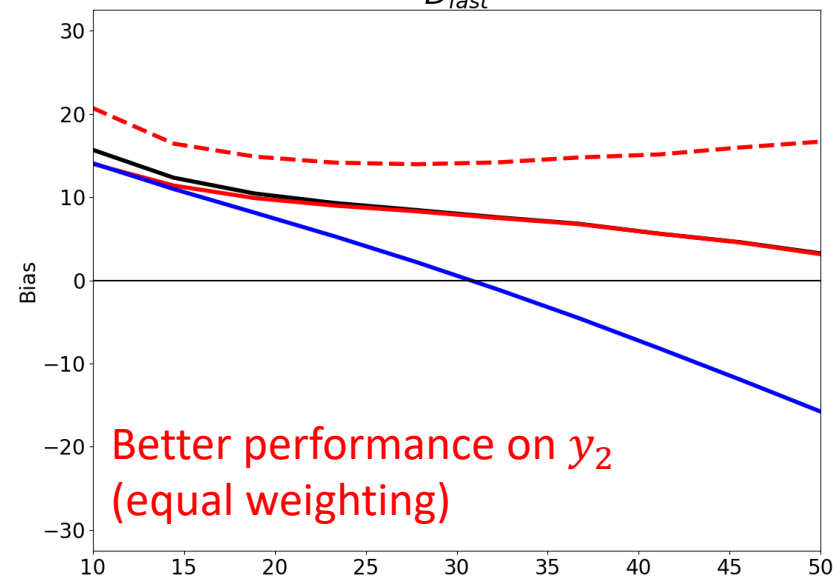
BIAS



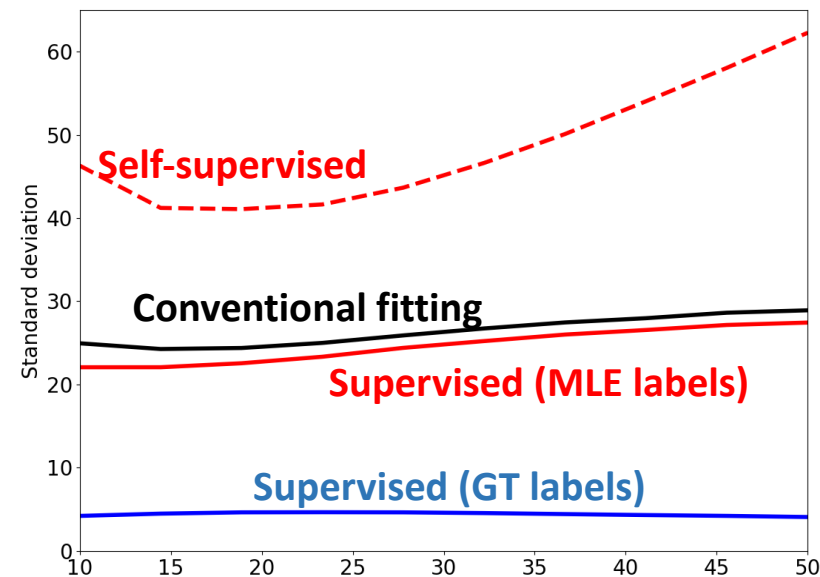
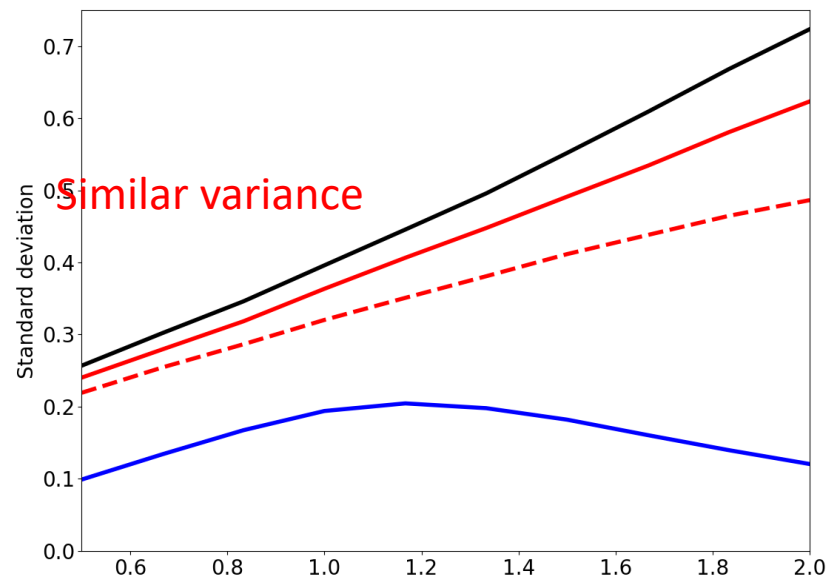
y_2 : difficult to fit

D_{fast}

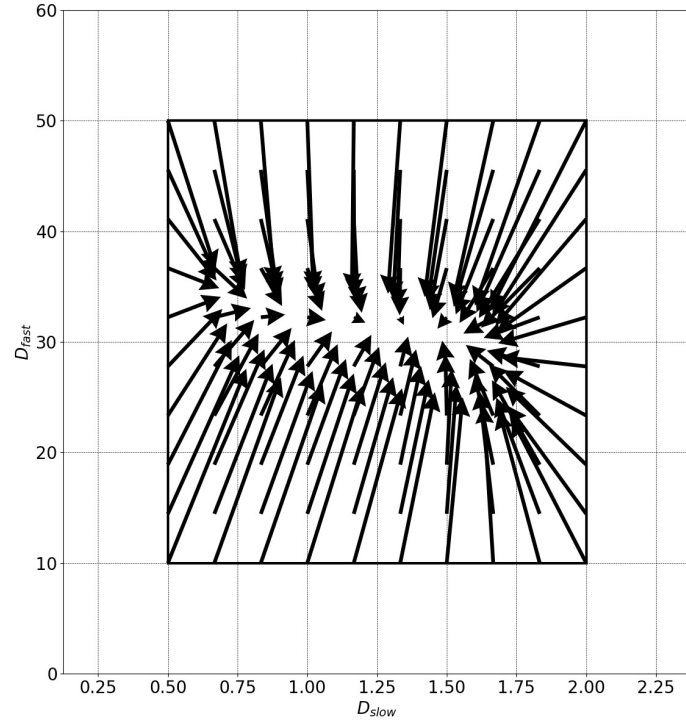
Bias



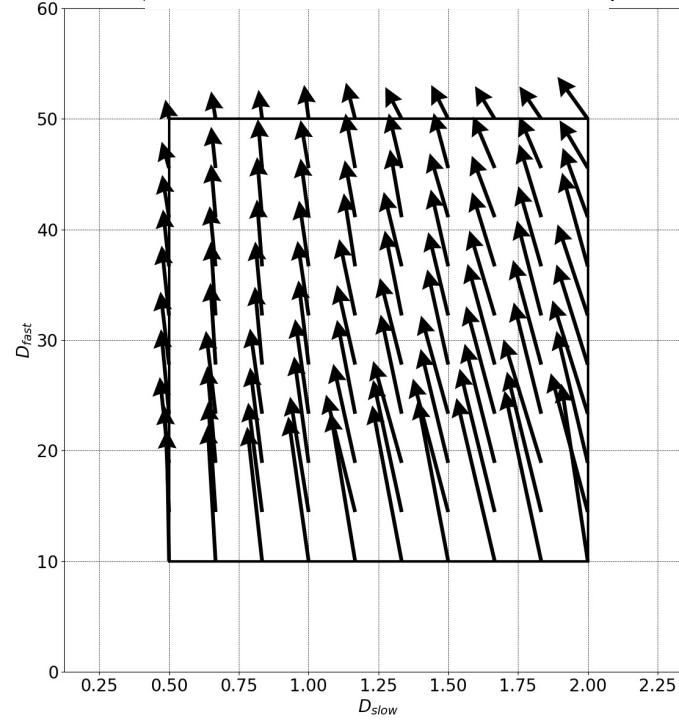
VARIANCE



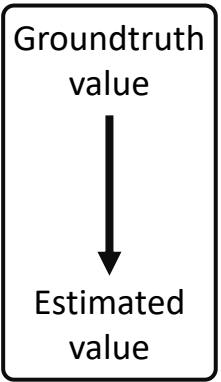
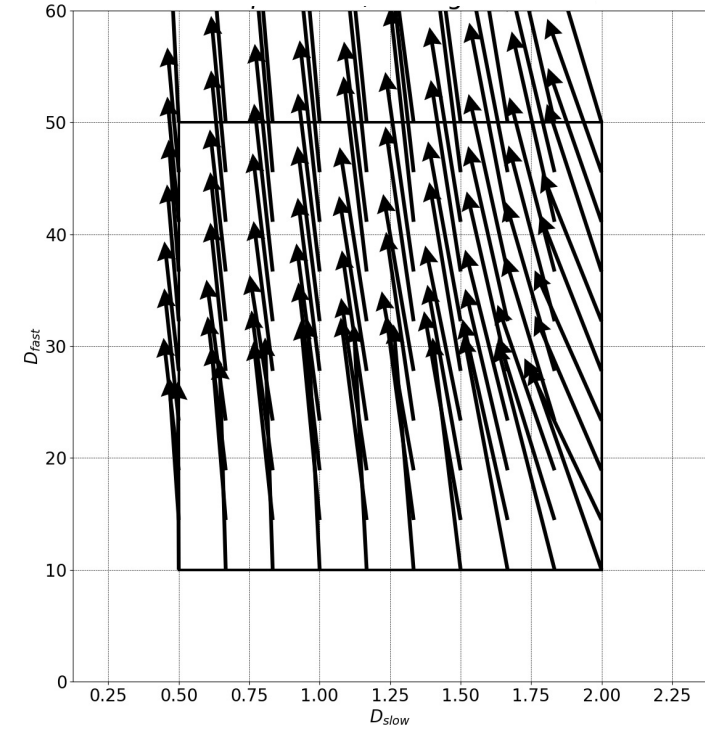
Supervised (GT labels)



Supervised (MLE labels)

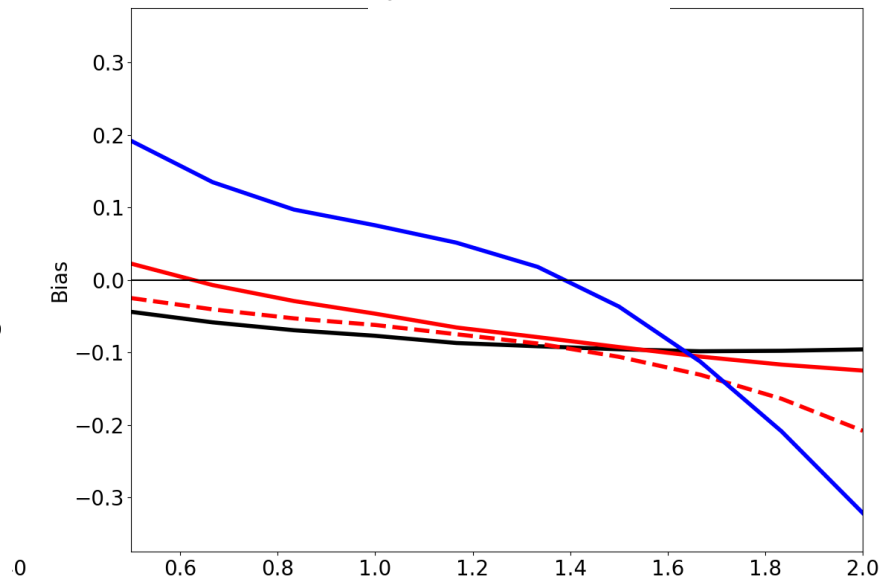


Self-supervised

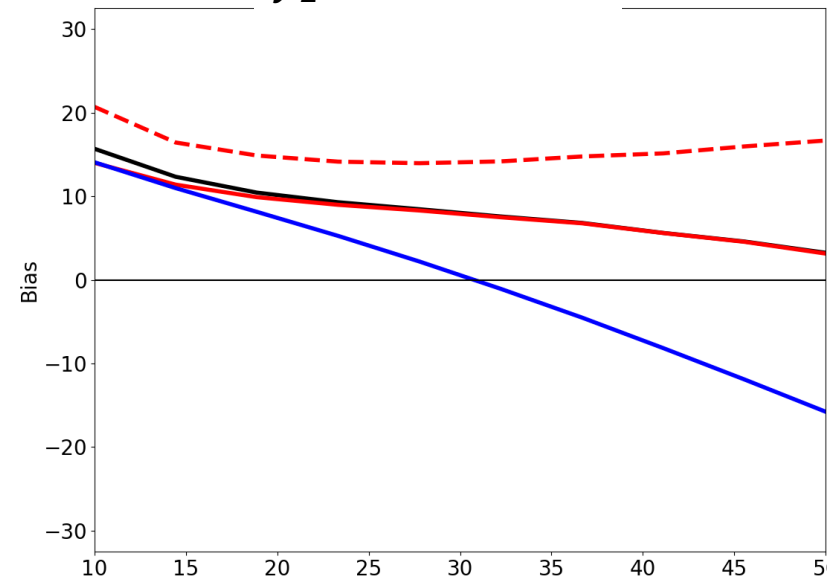


see Gyori 2022 for more examples

y_1 : easy to fit

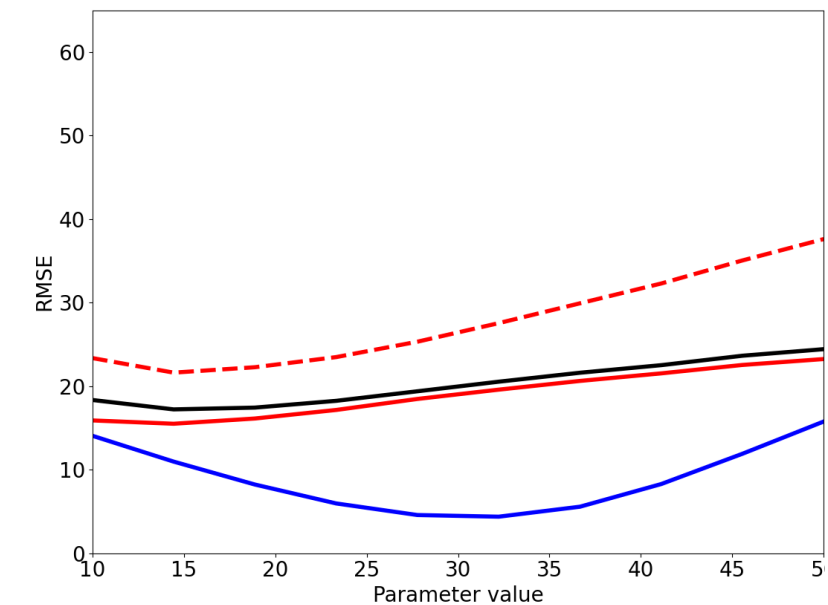
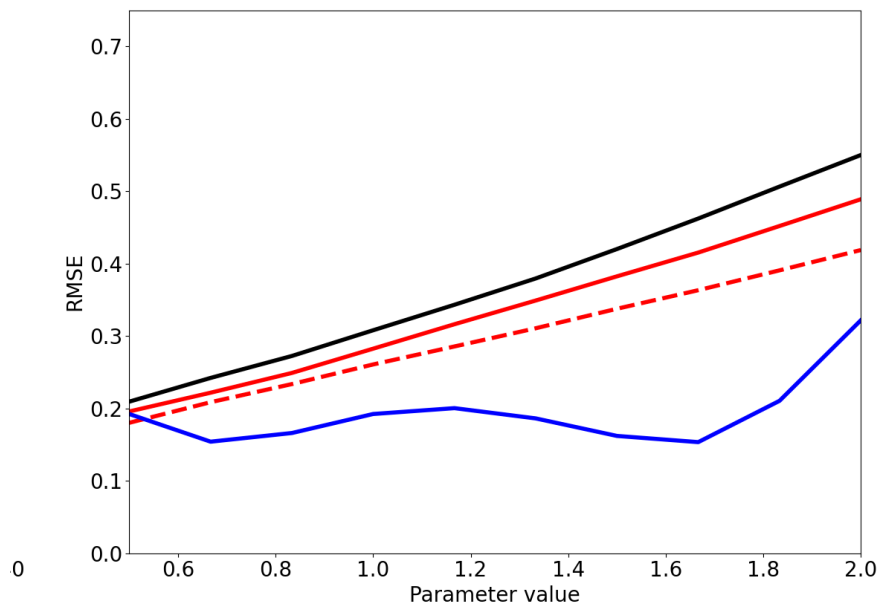


y_2 : difficult to fit



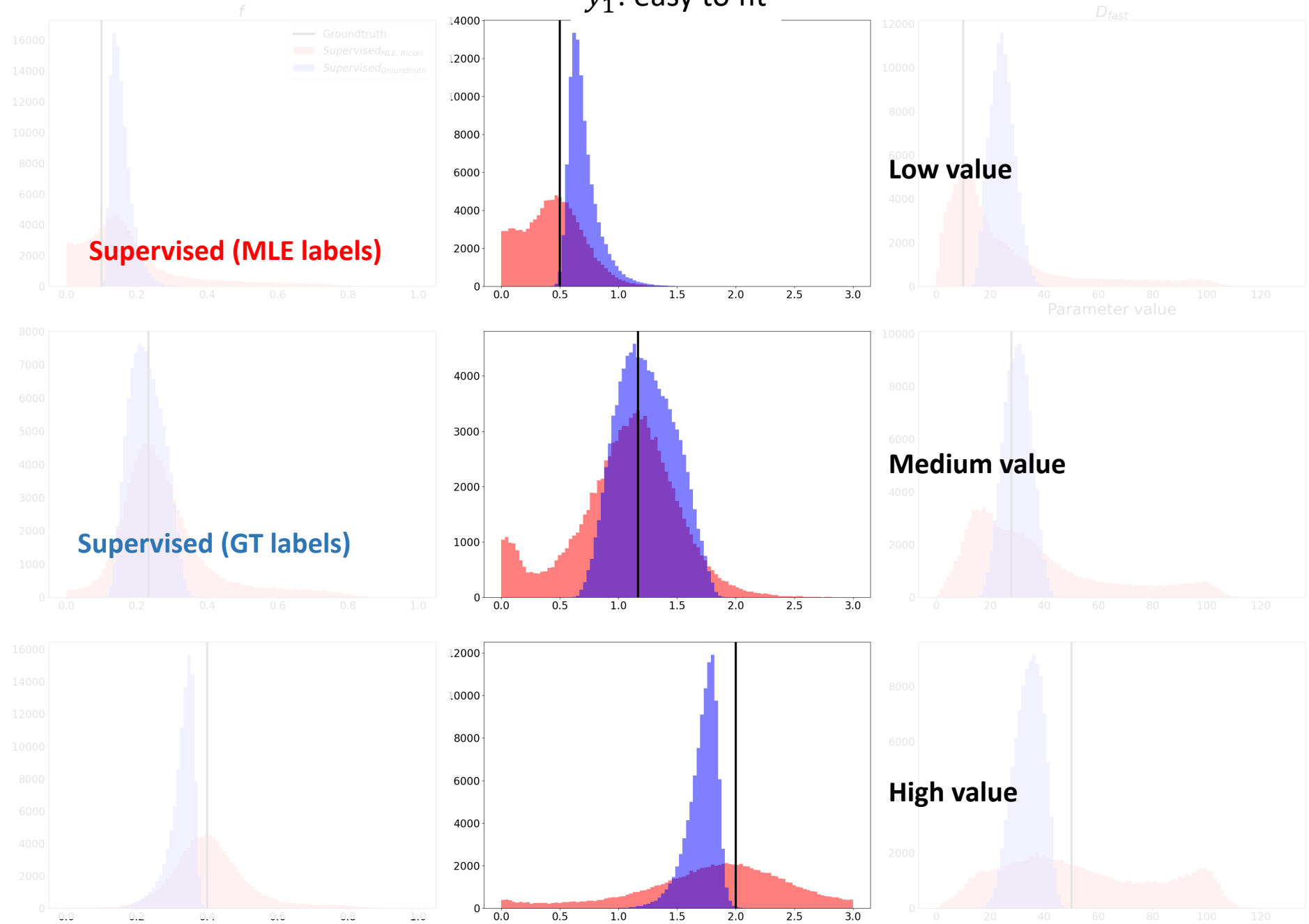
Bias worse

RMSE

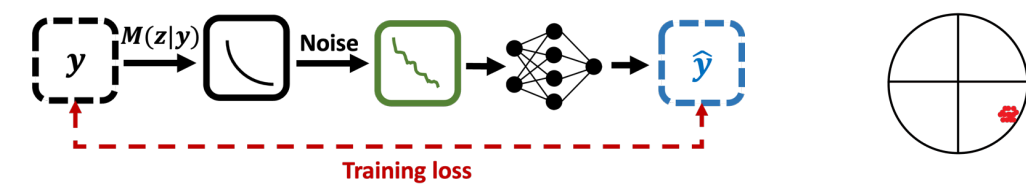


RMSE better

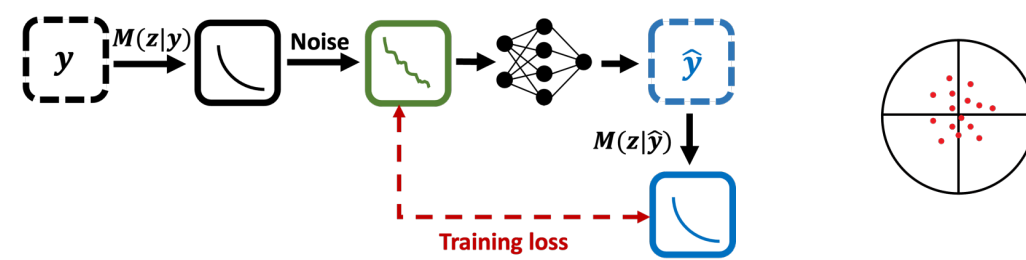
y_1 : easy to fit



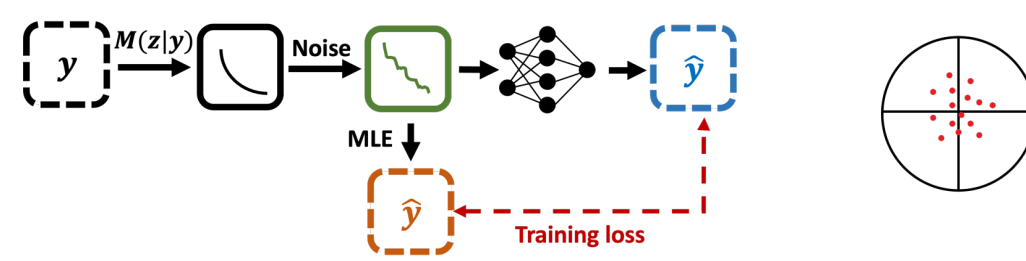
**Supervised methods
(groundtruth labels)**
e.g. Bertleff 2017, Gyori 2022



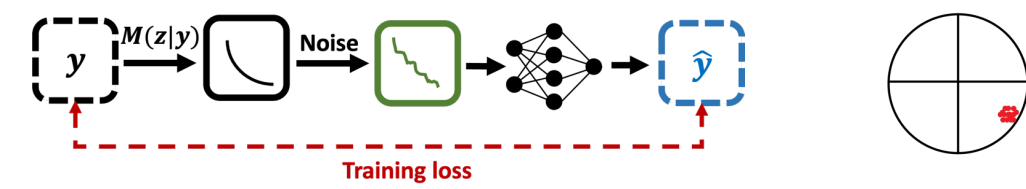
Self-supervised methods
e.g. Barbieri 2019, Kaandorp 2021



**Supervised methods
(MLE labels)**
e.g. Epstein 2022



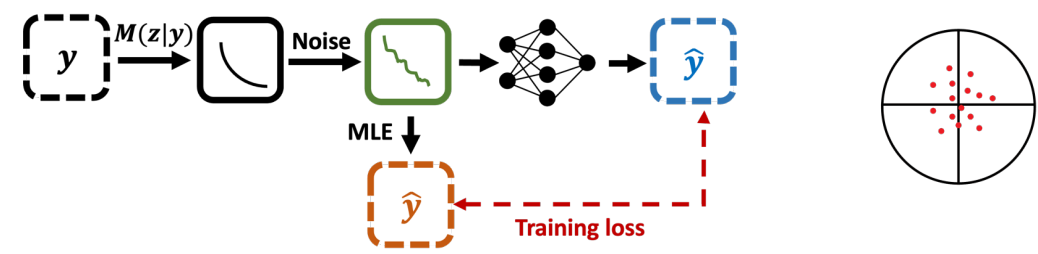
**Supervised methods
(groundtruth labels)**
e.g. Bertleff 2017, Gyori 2022



Hybrid loss function during training:

$$\text{Hybrid loss} = \alpha \cdot \text{Supervised}_{\text{MLE}} \text{ loss} + (1 - \alpha) \cdot \text{Supervised}_{\text{GT}} \text{ loss}$$

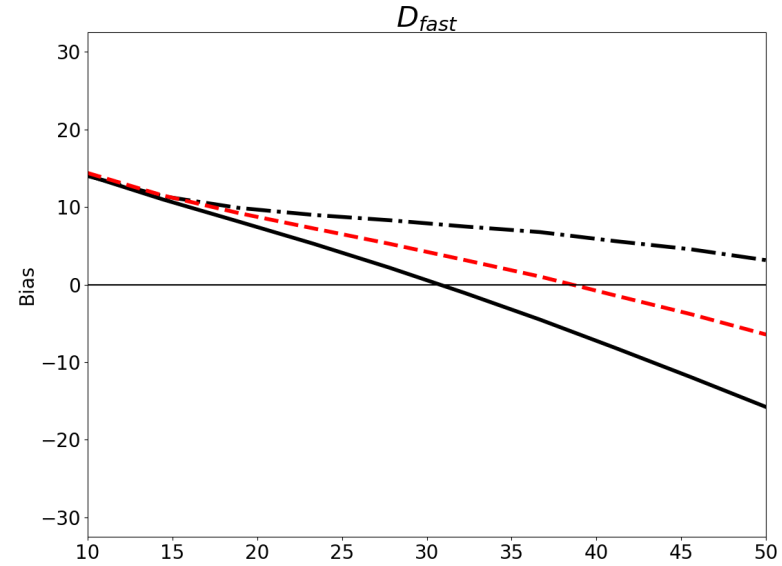
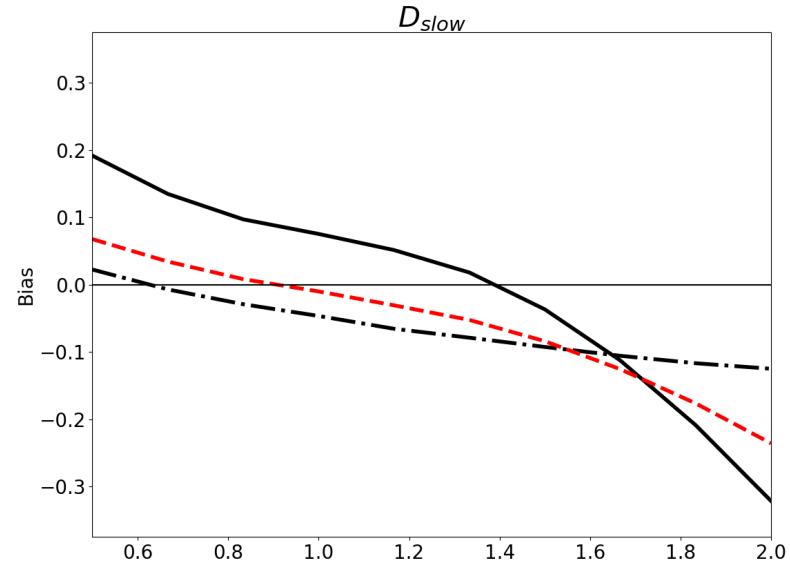
**Supervised methods
(MLE labels)**
e.g. Epstein 2022



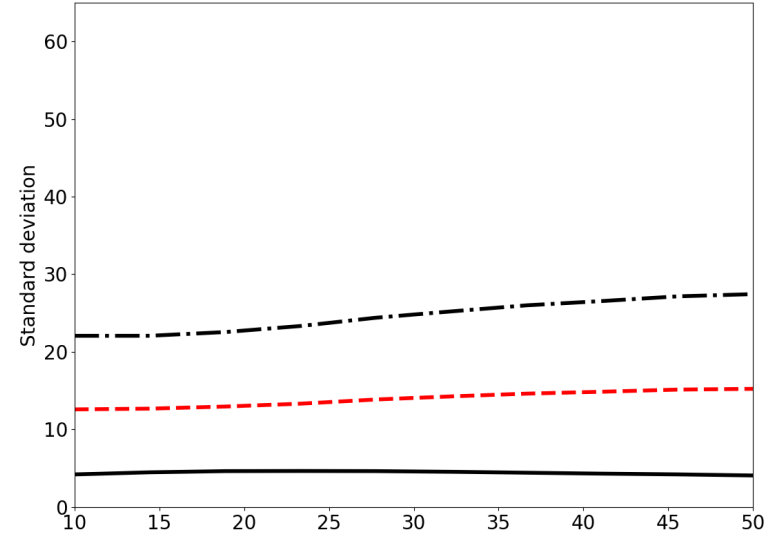
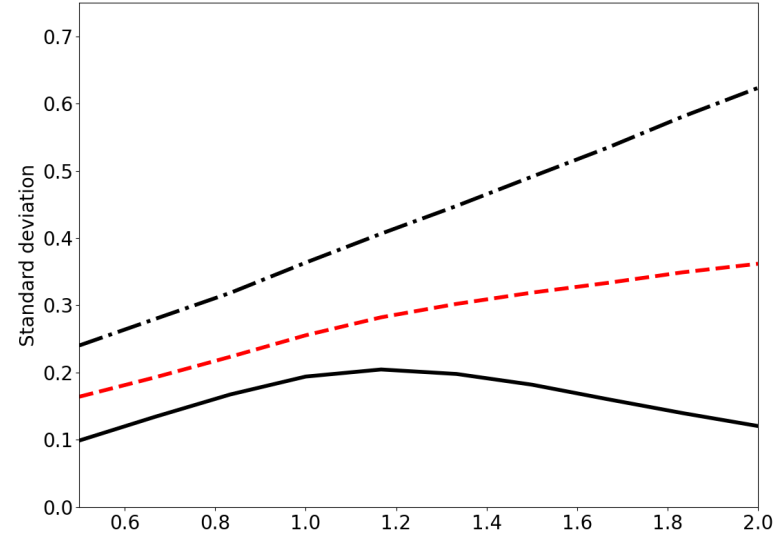
y_1 : easy to fit

y_2 : difficult to fit

BIAS

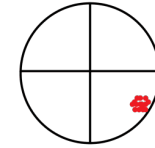


VARIANCE



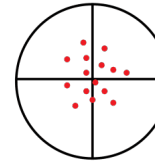
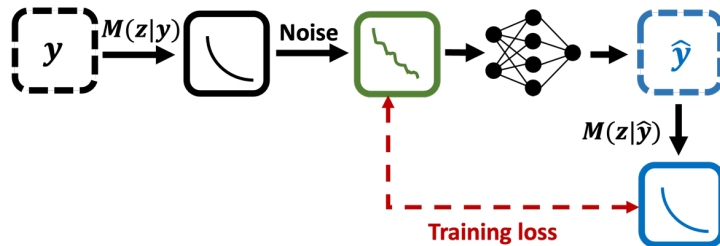
Supervised methods (groundtruth labels)

e.g. Bertleff 2017, Gyori 2022



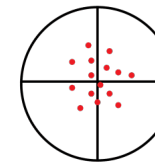
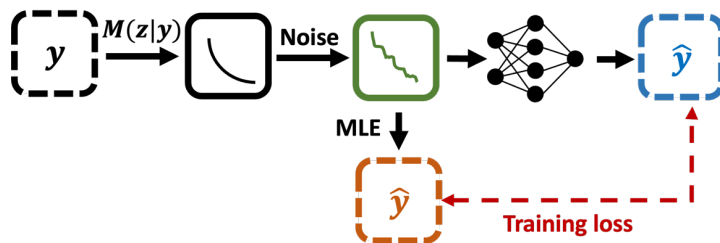
Self-supervised methods

e.g. Barbieri 2019, Kaandorp 2021



Supervised methods (MLE labels)

Epstein 2022



- **Bias/variance tradeoff**
- **Look beyond RMSE:** misleading summary metric
- Supervised training has **practical advantages** over self-supervised
- **Don't always use GT labels** – even if you have access to them
- **Can adjust network performance by tailoring contribution of different labels**



Tim Bray



Margaret Hall-Craggs



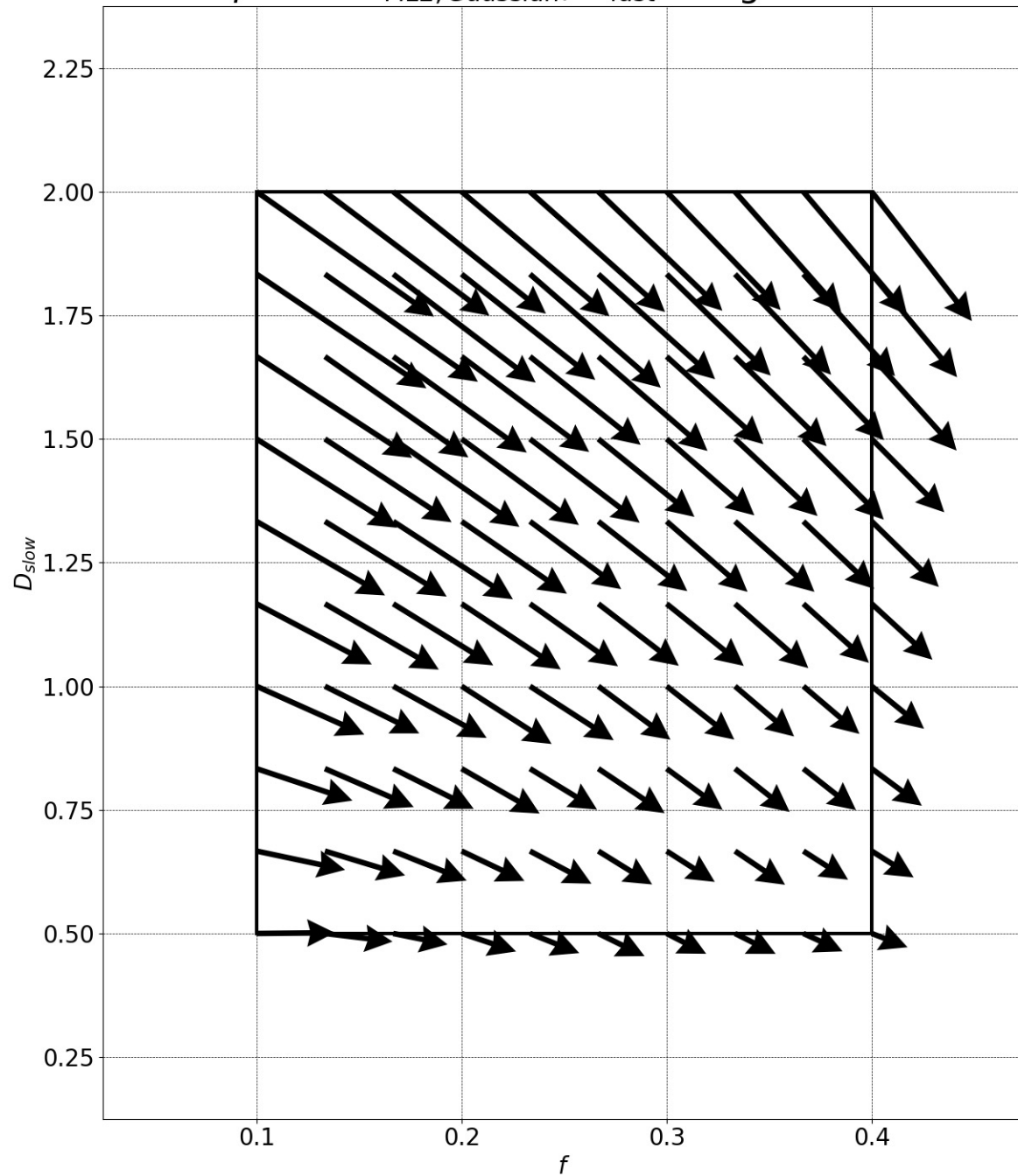
Gary Zhang

arXiv:2205.05587

Choice of training label matters: how to best use deep learning for quantitative MRI parameter estimation



Supervised_{MLE, Gaussian}, D_{fast} marginalisation



Supervised_{MLE, Rician}, D_{fast} marginalisation

