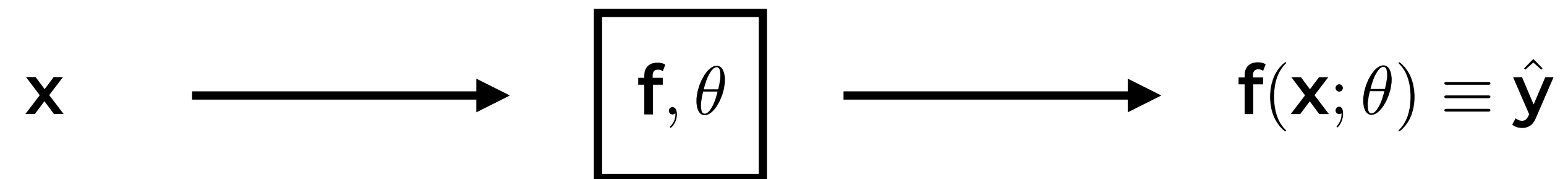
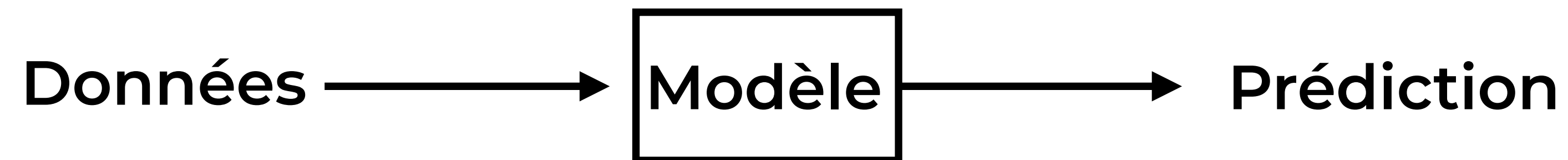


Apprentissage Automatique I MATH60629

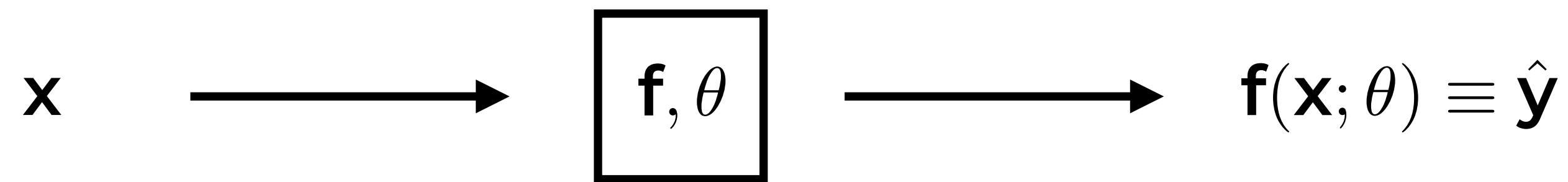
Sommaire jusqu'à maintenant

- **Bref survol de certains des concepts vus jusqu'à maintenant**
- **Emphase sur les concepts un peu plus avancés**

Apprentissage Supervisé



Fonction de perte (loss)



Différentes pertes pour selon le type de y

$y \in \mathcal{R}$

y catégorique e.g., {chat, chien, oiseau}

$y \in \{0, 1\}$

Régression

Classification

Classification binaire

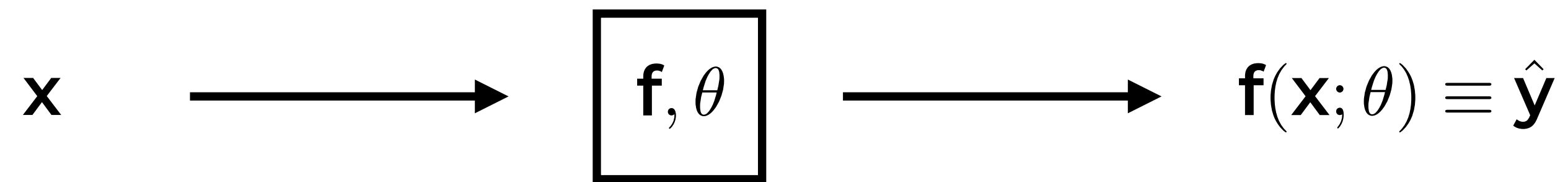
$(\hat{y} - y)^2$

taux de classification

AUC

Fonction de perte (loss)

Perte



Différentes pertes pour selon le type de y

$y \in \mathcal{R}$

y catégorique e.g., {chat, chien, oiseau}

$y \in \{0, 1\}$

Régression

Classification

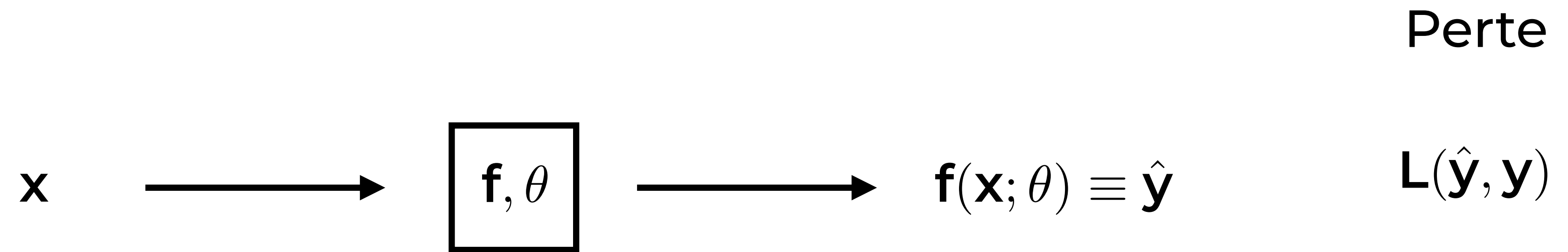
Classification binaire

$(\hat{y} - y)^2$

taux de classification

AUC

Fonction de perte (loss)



Différentes pertes pour selon le type de $\mathbf{y}$

$\mathbf{y} \in \mathcal{R}$

$\mathbf{y}$ catégorique e.g., {chat, chien, oiseau}

$\mathbf{y} \in \{0, 1\}$

Régression

Classification

Classification binaire

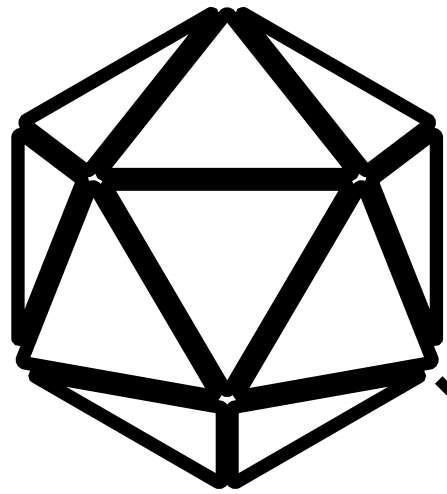
$(\hat{\mathbf{y}} - \mathbf{y})^2$

taux de classification

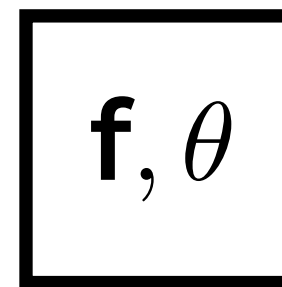
AUC

Processus d'apprentissage

Distribution
sur (x,y) :
 $P(x,y)$



$\mathbf{x}_{\text{train}}$

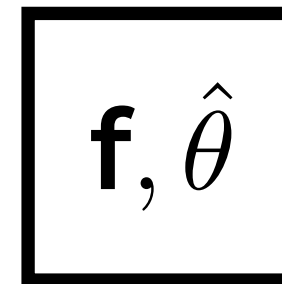


$\hat{\mathbf{y}}_{\text{train}}$

Perte

$$\mathbf{L}(\hat{\mathbf{y}}_{\text{train}}, \mathbf{y}_{\text{train}})$$

$\mathbf{x}_{\text{test}}$



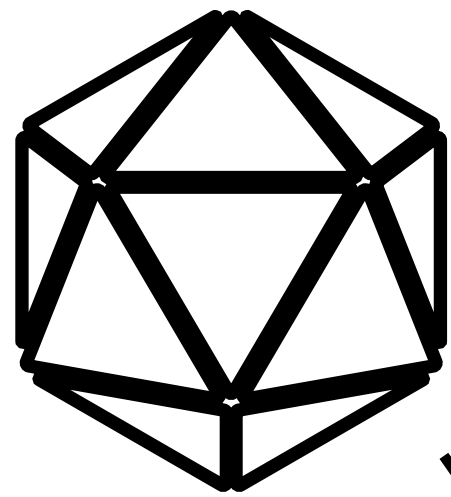
$\hat{\mathbf{y}}_{\text{test}}$

$$\mathbf{L}(\hat{\mathbf{y}}_{\text{test}}, \mathbf{y}_{\text{test}})$$

Processus d'apprentissage

En pratique

Distribution
sur (x,y) :
 $P(x,y)$



$\mathbf{x}_{\text{train}}$



$\mathbf{f}, \theta$



$\hat{\mathbf{y}}_{\text{train}}$

Perte

$\mathbf{L}(\hat{\mathbf{y}}_{\text{train}}, \mathbf{y}_{\text{train}})$

$\mathbf{x}_{\text{valid}}$



$\mathbf{f}, \hat{\theta}$



$\hat{\mathbf{y}}_{\text{valid}}$

$\mathbf{x}_{\text{test}}$



$\mathbf{f}, \hat{\theta}$



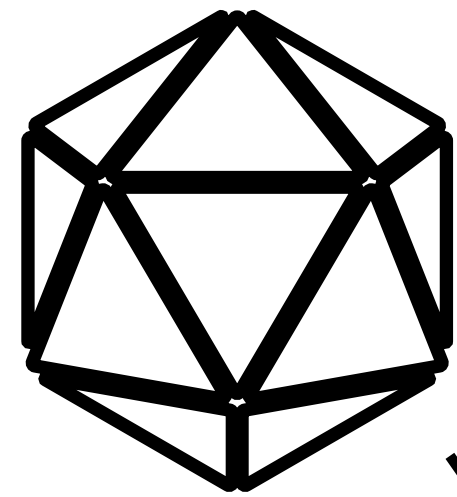
$\hat{\mathbf{y}}_{\text{test}}$

$\mathbf{L}(\hat{\mathbf{y}}_{\text{test}}, \mathbf{y}_{\text{test}})$

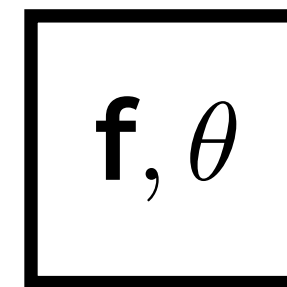
Processus d'apprentissage

En pratique

Distribution
sur (x,y) :
 $P(x,y)$



x_{train}

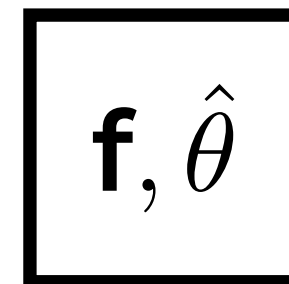


$\hat{y}_{\text{train}}$

Perte

$$L(\hat{y}_{\text{train}}, y_{\text{train}})$$

x_{valid}

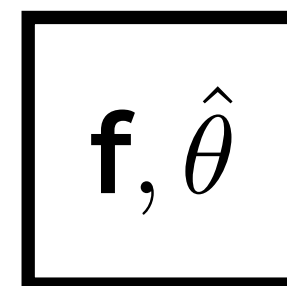


$\hat{y}_{\text{valid}}$

Pratique pour:

- Sélectionner les hyperparamètres
- Pour sélectionner le meilleur modèle

x_{test}



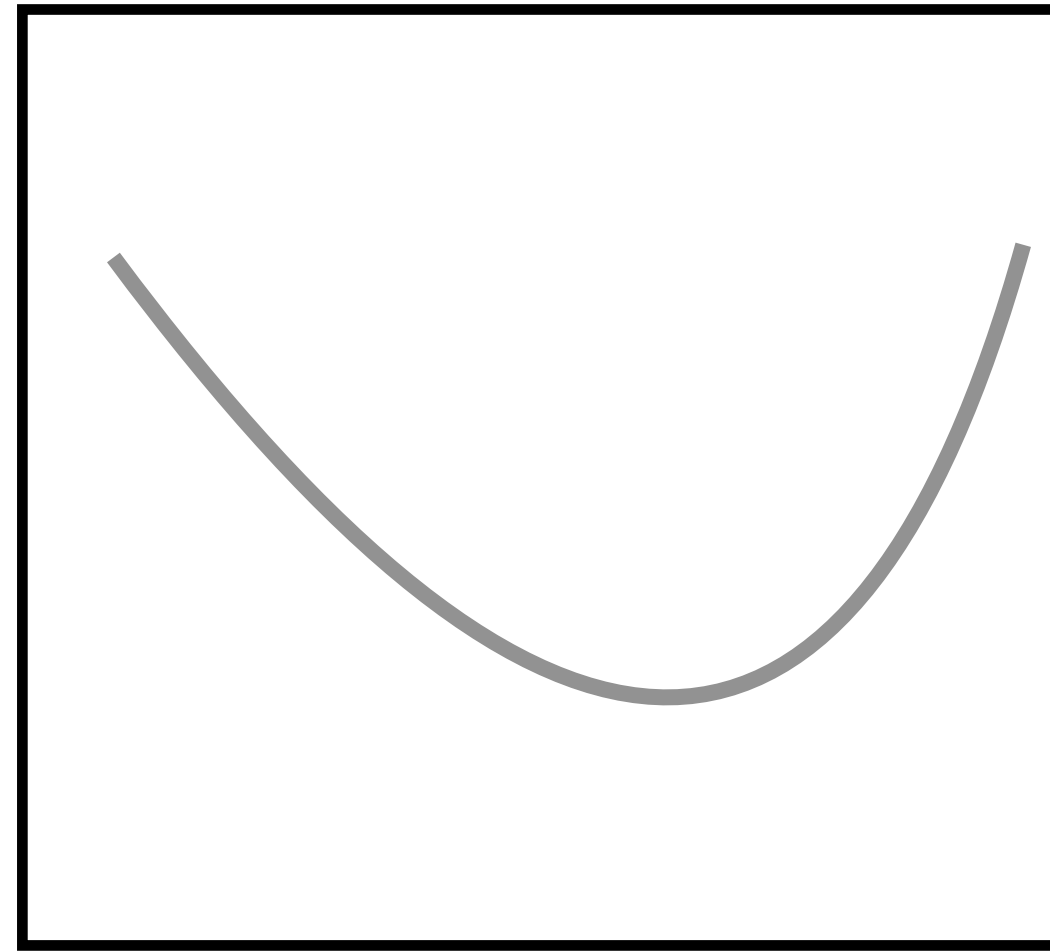
$\hat{y}_{\text{test}}$

$$L(\hat{y}_{\text{test}}, y_{\text{test}})$$

Apprentissage

- Apprendre: Changer les paramètres pour obtenir de meilleures prédictions

Loss



- En d'autres mots: changer les paramètres θ pour minimiser la perte
- Trouver les dérivées de la perte p.r. aux paramètres: $\frac{d \text{ Loss}}{d\theta}$

Différents modèles

- f : régression linéaire, θ a une solution analytique
- f : réseaux de neurones, θ n'a pas de solution analytique. On utilise la descente de gradient

- Given a training set: $\{(\mathbf{x}_{\text{train}}, \mathbf{y}_{\text{train}})\}$

- Initialize $\hat{\theta}$ randomly

for $t = 1, 2, \dots$ (epochs) **do**

for $i = 1, 2, \dots$ (datum) **do**

 - Obtain the prediction $\{\mathbf{f}(\mathbf{x}_i; \hat{\theta})\}$ (Forward propagation)

 - Compute the Loss: $\text{Loss}_{ti} := L(\mathbf{f}(\mathbf{x}_i; \hat{\theta}), \mathbf{y}_i)$

 - Find the derivative of the loss: $\frac{d \text{Loss}_{ti}}{d \hat{\theta}}$

 - Get new parameters: $\hat{\theta}' = \hat{\theta} - \alpha \frac{d \text{Loss}_{ti}}{d \hat{\theta}}$

 - If $\|\hat{\theta}' - \hat{\theta}\|_2^2 < \epsilon$ then stop

 - Update parameters: $\hat{\theta} = \hat{\theta}'$

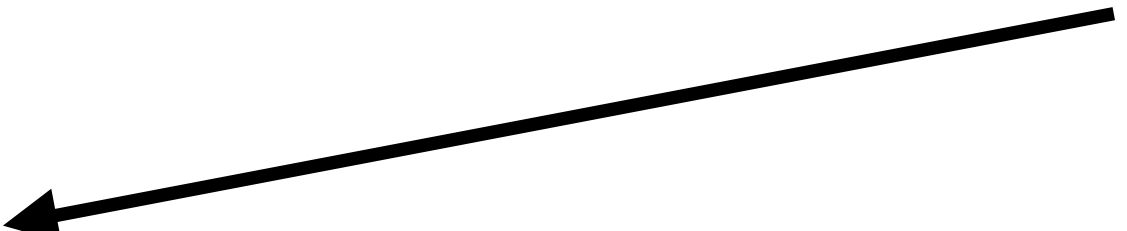
end for

end for

- Given a training set: $\{(\mathbf{x}_{\text{train}}, \mathbf{y}_{\text{train}})\}$

- Initialize $\hat{\theta}$ randomly

Stochastic Gradient Descent



for $t = 1, 2, \dots$ (epochs) **do**
 for $i = 1, 2, \dots$ (datum) **do**

- Obtain the prediction $\{f(\mathbf{x}_i; \hat{\theta})\}$ (Forward propagation)

- Compute the Loss: $\text{Loss}_{ti} := L(f(\mathbf{x}_i; \hat{\theta}), \mathbf{y}_i)$

- Find the derivative of the loss: $\frac{d \text{Loss}_{ti}}{d \hat{\theta}}$

- Get new parameters: $\hat{\theta}' = \hat{\theta} - \alpha \frac{d \text{Loss}_{ti}}{d \hat{\theta}}$

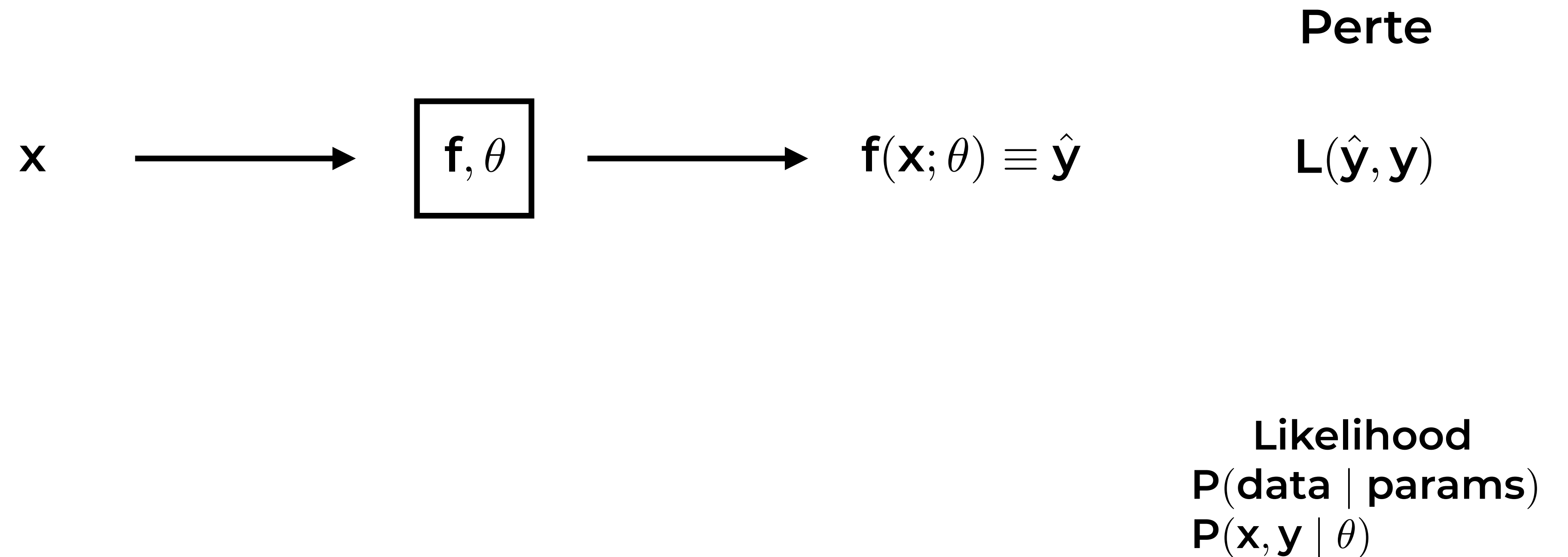
- If $\|\hat{\theta}' - \hat{\theta}\|_2^2 < \epsilon$ then stop

- Update parameters: $\hat{\theta} = \hat{\theta}'$

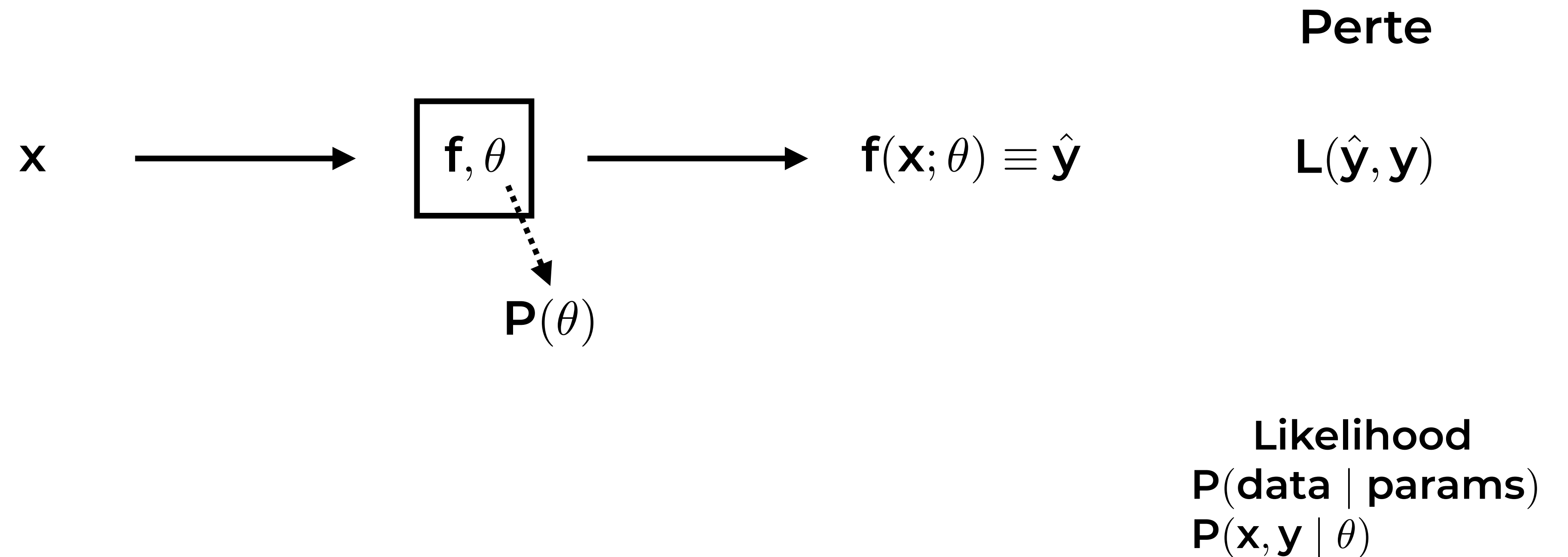
end for

end for

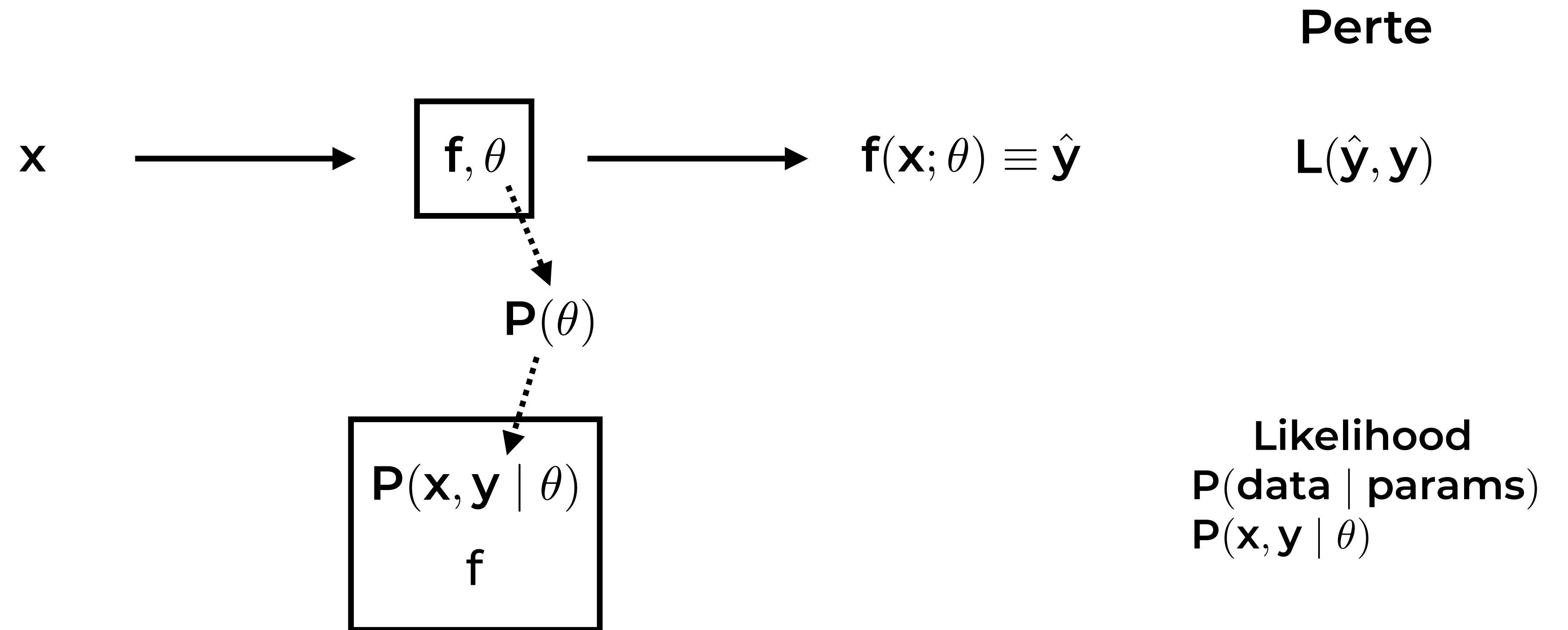
Modèles probabilistes



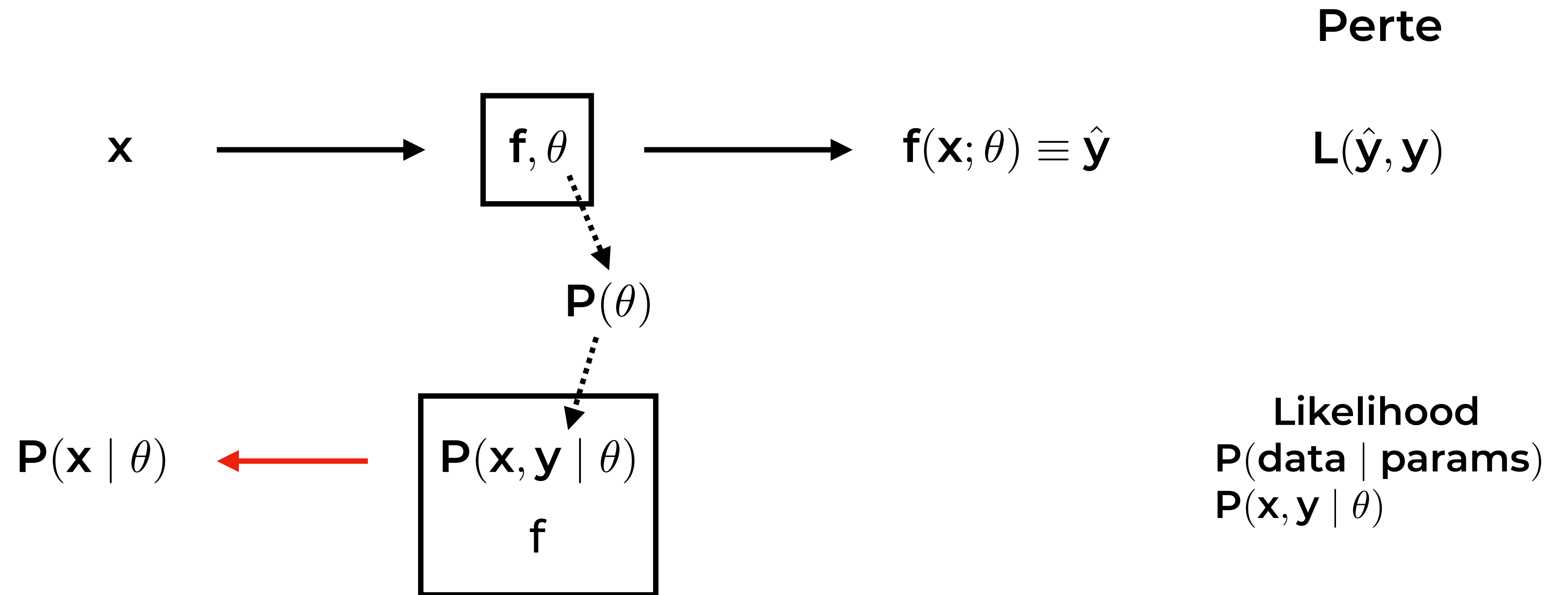
Modèles probabilistes



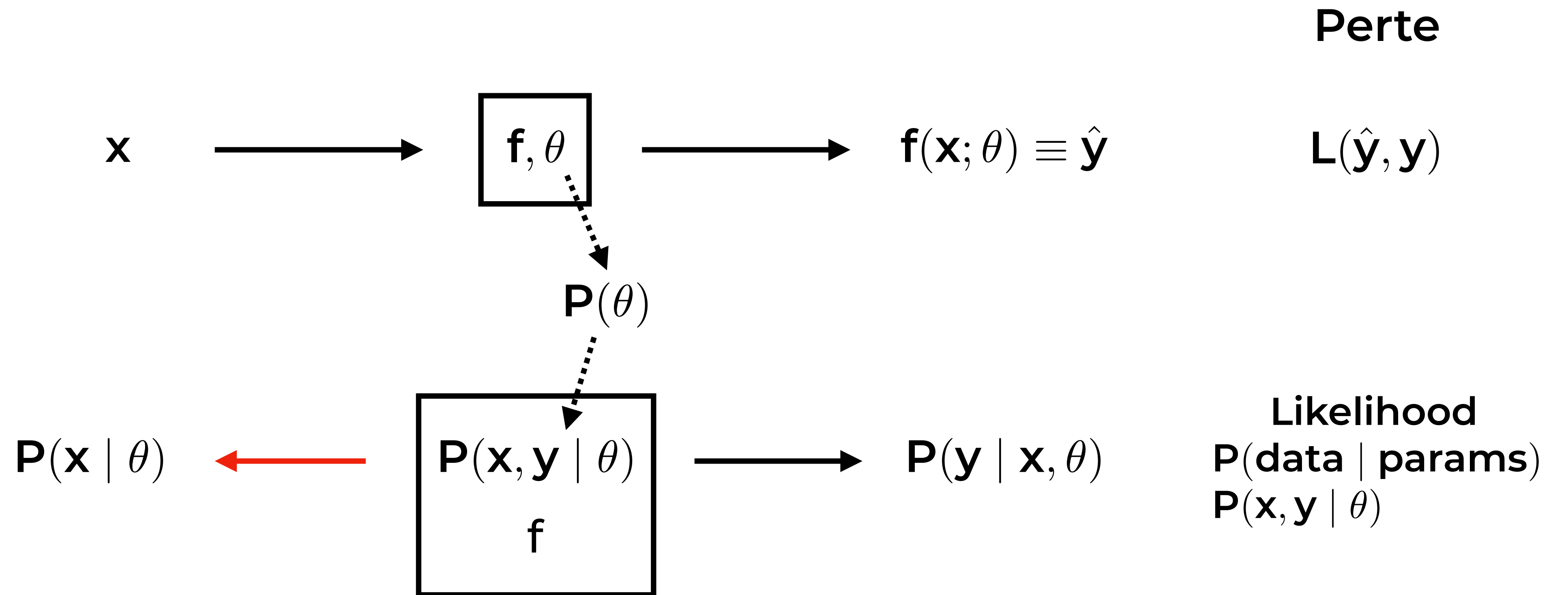
Modèles probabilistes



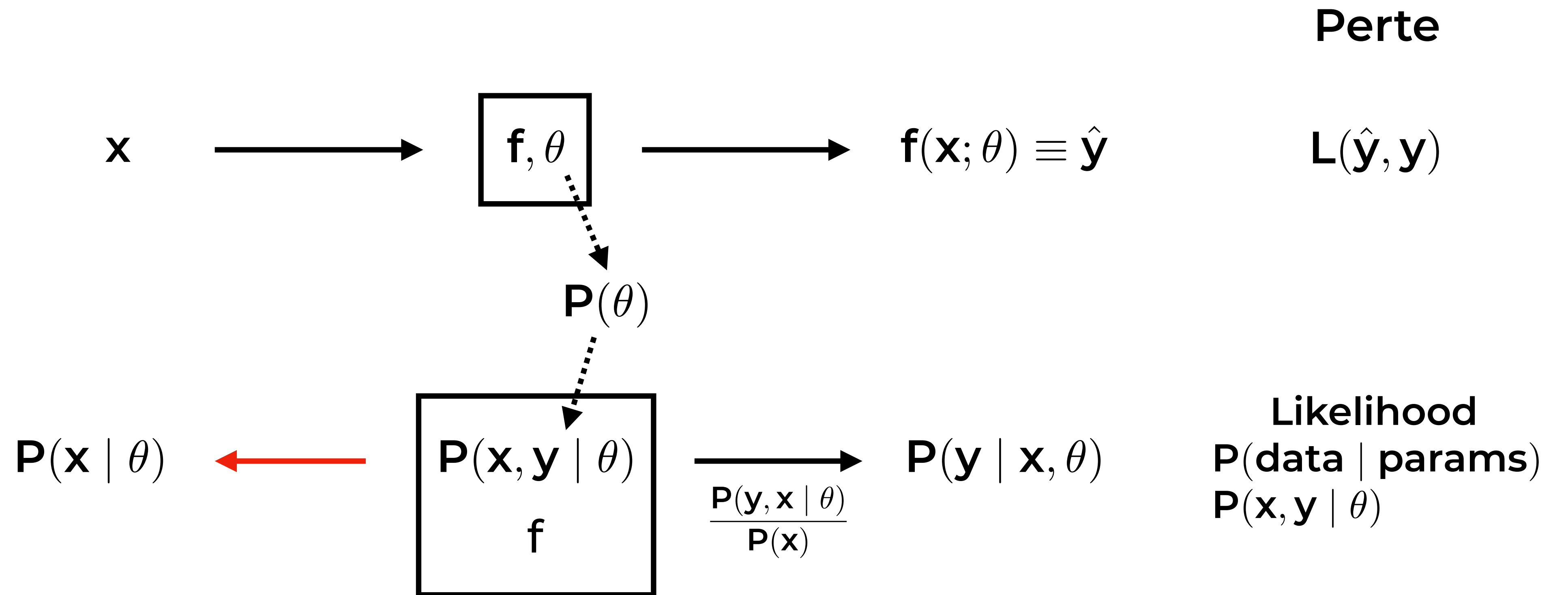
Modèles probabilistes



Modèles probabilistes



Modèles probabilistes



Exemple

Données: 952 1064 965 1037 871 1029 1138 (problème non-supervisé)

Modèle: $P(\mathbf{x} \mid \theta) := \mathcal{N}(\boldsymbol{\mu}, \mathbf{I})$

Exemple

Données: 952 1064 965 1037 871 1029 1138 (problème non-supervisé)

Modèle: $P(\mathbf{x} \mid \theta) := \mathcal{N}(\boldsymbol{\mu}, \mathbf{I})$

Exemple

Données: 952 1064 965 1037 871 1029 1138 (problème non-supervisé)

Modèle: $P(\mathbf{x} \mid \theta) := \mathcal{N}(\boldsymbol{\mu}, \mathbf{I})$

Likelihood d'une donnée:

Exemple

Données: 952 1064 965 1037 871 1029 1138 (problème non-supervisé)

Modèle: $P(\mathbf{x} \mid \theta) := \mathcal{N}(\boldsymbol{\mu}, 1)$

Likelihood d'une donnée:

$$\text{Likelihood}(\mathbf{x} \mid \boldsymbol{\mu}, 1) = \frac{1}{\sqrt{2\pi}} \exp - \frac{(\mathbf{x} - \boldsymbol{\mu})^2}{2}$$

Exemple

Données: 952 1064 965 1037 871 1029 1138 (problème non-supervisé)

Modèle: $P(\mathbf{x} \mid \theta) := \mathcal{N}(\boldsymbol{\mu}, 1)$

Likelihood d'une donnée:

$$\text{Likelihood}(\mathbf{x} \mid \boldsymbol{\mu}, 1) = \frac{1}{\sqrt{2\pi}} \exp - \frac{(\mathbf{x} - \boldsymbol{\mu})^2}{2}$$

Log-Likelihood

$$= \log \frac{1}{\sqrt{2\pi}} \exp - \frac{(\mathbf{x} - \boldsymbol{\mu})^2}{2}$$

$$= \log 1 - \frac{1}{2} \log 2\pi - \frac{(\mathbf{x} - \boldsymbol{\mu})^2}{2}$$

Exemple

Données: 952 1064 965 1037 871 1029 1138 (problème non-supervisé)

Modèle: $P(\mathbf{x} \mid \theta) := \mathcal{N}(\boldsymbol{\mu}, 1)$

Likelihood d'une donnée:

$$\text{Likelihood}(\mathbf{x} \mid \boldsymbol{\mu}, 1) = \frac{1}{\sqrt{2\pi}} \exp - \frac{(\mathbf{x} - \boldsymbol{\mu})^2}{2}$$

Log-Likelihood

$$= \log \frac{1}{\sqrt{2\pi}} \exp - \frac{(\mathbf{x} - \boldsymbol{\mu})^2}{2}$$

$$= \log 1 - \frac{1}{2} \log 2\pi - \frac{(\mathbf{x} - \boldsymbol{\mu})^2}{2}$$

Quelle valeur de $\boldsymbol{\mu}$ la maximise?

Exemple

Données: 952 1064 965 1037 871 1029 1138 (problème non-supervisé)

Modèle: $P(\mathbf{x} \mid \theta) := \mathcal{N}(\boldsymbol{\mu}, 1)$

Likelihood d'une donnée:

$$\text{Likelihood}(\mathbf{x} \mid \boldsymbol{\mu}, 1) = \frac{1}{\sqrt{2\pi}} \exp - \frac{(\mathbf{x} - \boldsymbol{\mu})^2}{2}$$

Log-Likelihood

$$\begin{aligned} &= \log \frac{1}{\sqrt{2\pi}} \exp - \frac{(\mathbf{x} - \boldsymbol{\mu})^2}{2} \\ &= \log 1 - \frac{1}{2} \log 2\pi - \frac{(\mathbf{x} - \boldsymbol{\mu})^2}{2} \end{aligned}$$

Quelle valeur de $\boldsymbol{\mu}$ la maximise?

$$\frac{d \text{ Log-Likelihood}}{d \boldsymbol{\mu}}$$

$$= \frac{d \frac{(\mathbf{x} - \boldsymbol{\mu})^2}{2}}{d \boldsymbol{\mu}}$$

$$= (\mathbf{x} - \boldsymbol{\mu})$$

set to 0

$$\boldsymbol{\mu} = \mathbf{x}$$

Exemple

Données: 952 1064 965 1037 871 1029 1138 (problème non-supervisé)

Modèle: $P(\mathbf{x} \mid \theta) := \mathcal{N}(\boldsymbol{\mu}, 1)$

Likelihood d'une donnée:

$$\text{Likelihood}(\mathbf{x} \mid \boldsymbol{\mu}, 1) = \frac{1}{\sqrt{2\pi}} \exp - \frac{(\mathbf{x} - \boldsymbol{\mu})^2}{2}$$

Log-Likelihood

$$= \log \frac{1}{\sqrt{2\pi}} \exp - \frac{(\mathbf{x} - \boldsymbol{\mu})^2}{2}$$

$$= \log 1 - \frac{1}{2} \log 2\pi - \frac{(\mathbf{x} - \boldsymbol{\mu})^2}{2}$$

Quelle valeur de $\boldsymbol{\mu}$ la maximise?

$$\frac{d \text{ Log-Likelihood}}{d \boldsymbol{\mu}}$$

$$= \frac{d \frac{(\mathbf{x} - \boldsymbol{\mu})^2}{2}}{d \boldsymbol{\mu}}$$

$$= (\mathbf{x} - \boldsymbol{\mu})$$

set to 0

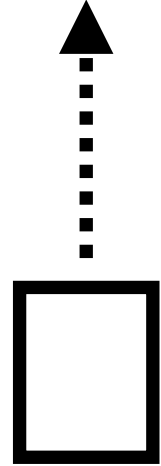
$$\boldsymbol{\mu} = \mathbf{x} = 952$$

Données: (x,y)
Naive Bayes

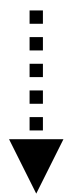
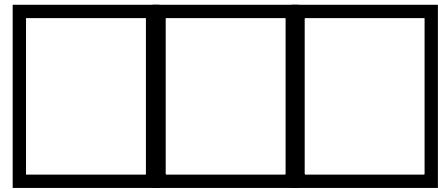
Modele : $P(\mathbf{x}, \mathbf{y} \mid \theta) = P(\mathbf{x} \mid \mathbf{y}, \theta)P(\mathbf{y} \mid \theta)$

Données: (x,y)
Naive Bayes

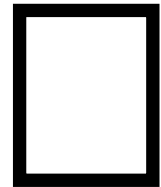
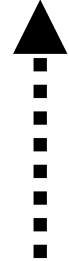
Modele : $P(\mathbf{x}, \mathbf{y} \mid \theta) = P(\mathbf{x} \mid \mathbf{y}, \theta) P(\mathbf{y} \mid \theta)$



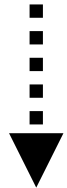
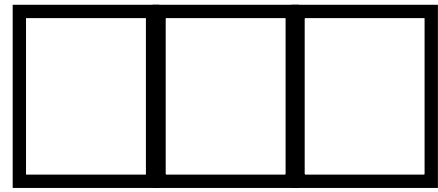
Données: (x,y)
Naive Bayes



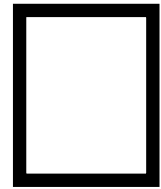
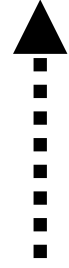
Modele : $P(\mathbf{x}, \mathbf{y} \mid \theta) = P(\mathbf{x} \mid \mathbf{y}, \theta) P(\mathbf{y} \mid \theta)$



Données: (x,y)
Naive Bayes

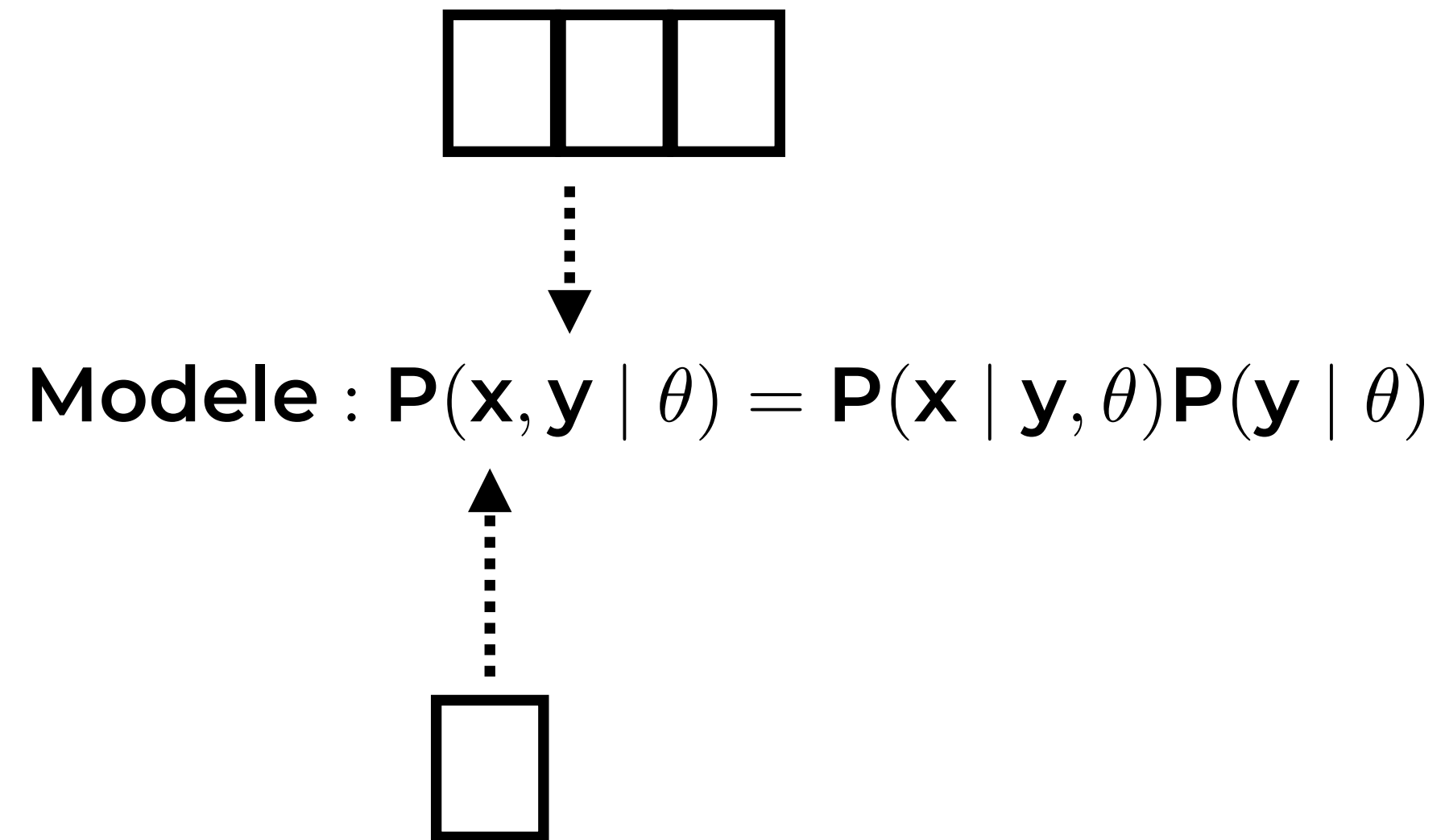


Modele : $P(x, y \mid \theta) = P(x \mid y, \theta) P(y \mid \theta)$

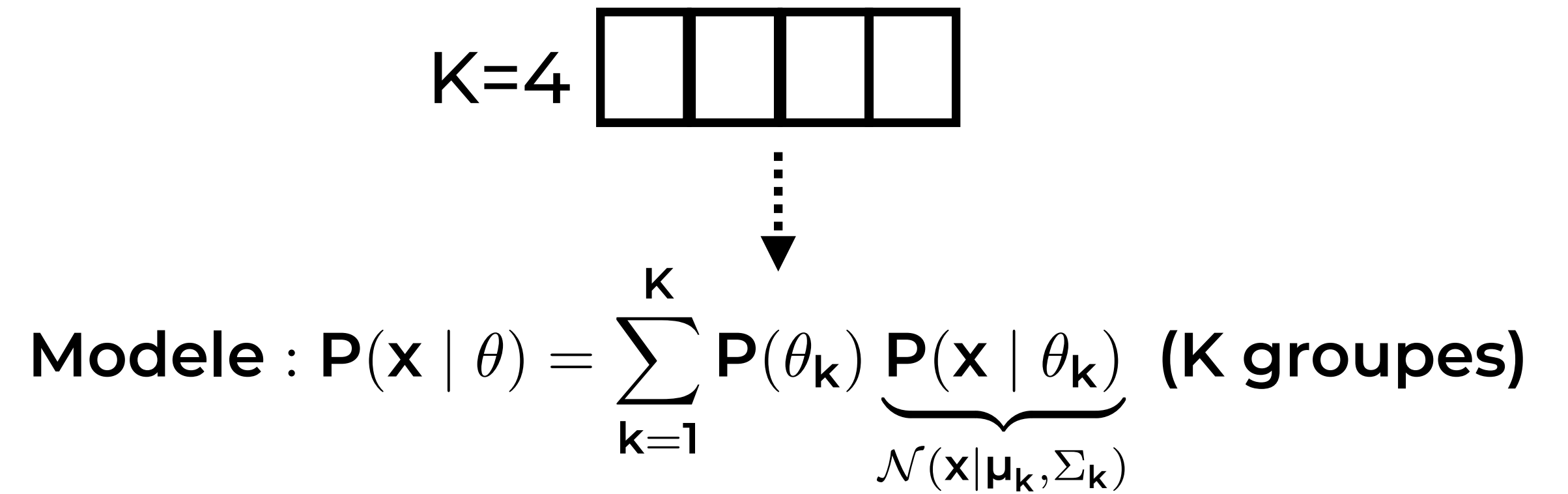


Données: x
Gaussian Mixture Models

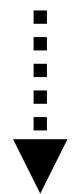
Données: (x,y) Naive Bayes



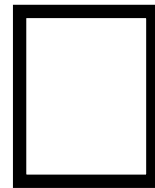
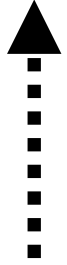
Données: x Gaussian Mixture Models



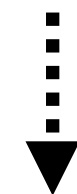
Données: (x,y) Naive Bayes



Modele : $P(\mathbf{x}, \mathbf{y} \mid \theta) = P(\mathbf{x} \mid \mathbf{y}, \theta) P(\mathbf{y} \mid \theta)$



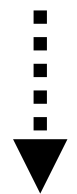
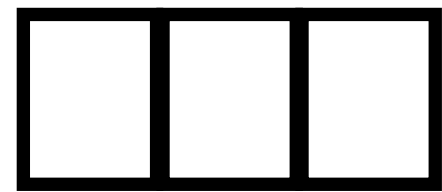
Données: x Gaussian Mixture Models



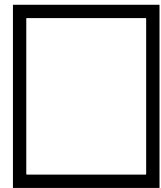
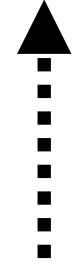
Modele : $P(\mathbf{x} \mid \theta) = \sum_{k=1}^K P(\theta_k) \underbrace{P(\mathbf{x} \mid \theta_k)}_{\mathcal{N}(\mathbf{x} \mid \boldsymbol{\mu}_k, \Sigma_k)} \text{ (K groupes)}$

Max. likelihood (MLE) : $\hat{\theta}_{\text{MLE}} = \arg \max_{\theta} P(\mathbf{x} \mid \theta)$

Données: (x,y)
Naive Bayes



Modele : $P(\mathbf{x}, \mathbf{y} \mid \theta) = P(\mathbf{x} \mid \mathbf{y}, \theta) P(\mathbf{y} \mid \theta)$



Données: x
Gaussian Mixture Models



Modele : $P(\mathbf{x} \mid \theta) = \sum_{k=1}^K P(\theta_k) \underbrace{P(\mathbf{x} \mid \theta_k)}_{\mathcal{N}(\mathbf{x} \mid \boldsymbol{\mu}_k, \Sigma_k)}$ (K groupes)

Max. likelihood (MLE) : $\hat{\theta}_{\text{MLE}} = \arg \max_{\theta} P(\mathbf{x} \mid \theta)$

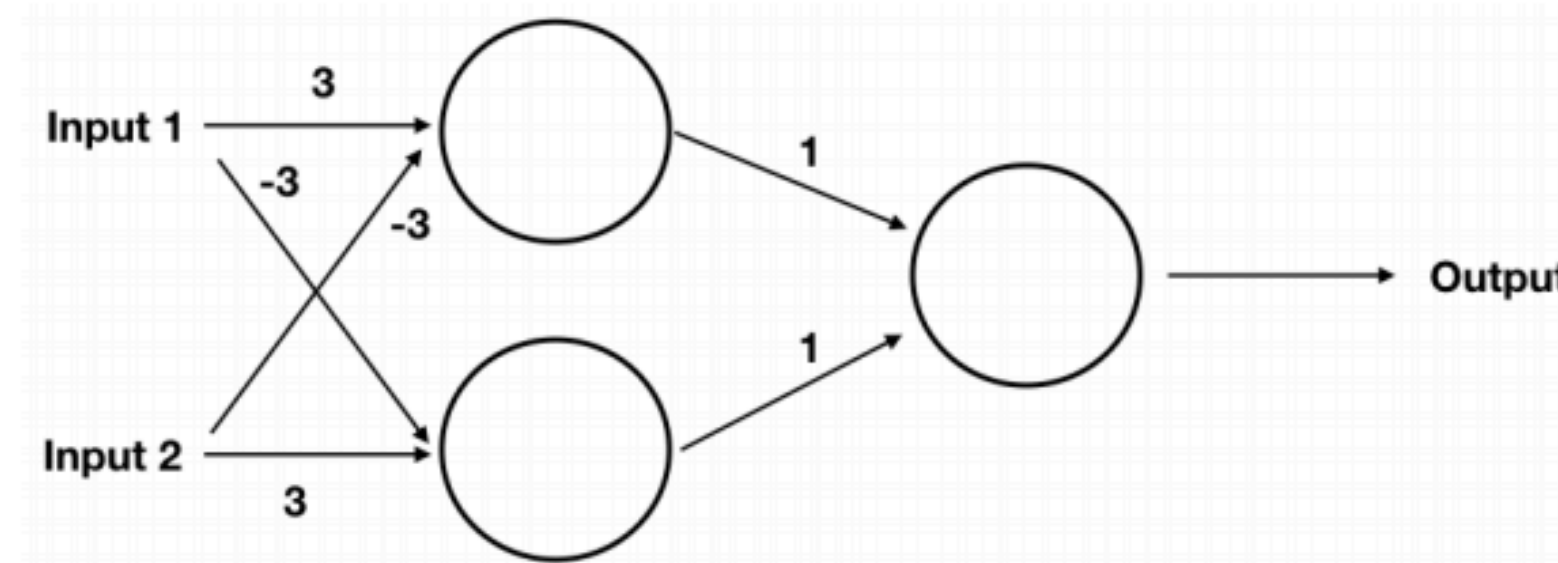
Max. a posteriori (MAP) : $\hat{\theta}_{\text{MAP}} = \arg \max_{\theta} P(\theta \mid \mathbf{x}, \gamma) = \arg \max_{\theta} \frac{P(\mathbf{x} \mid \theta) P(\theta \mid \gamma)}{P(\mathbf{x})}$

MLPs / RNNs / CNNs

- MLPs: les couches sont complètement connectées à la prochaine couche
- RNNs: entrées à chaque couche
 - Application standard: modélisation des séries temporelles
- CNNs: remplacent les multiplications de matrices par des convolutions (sparse connections, weight sharing) + et du pooling
 - Application standard: reconnaissance d'objets dans des images

MLPs

- (e) (4 points) Consider the neural network below. We have estimated its parameters (shown next to their corresponding arrows).



The activation function of each unit in the network is a simple thresholding function:

$$\text{threshold}(x) = \begin{cases} 0 & \text{if } x \leq 0, \\ 1 & \text{if } x > 0. \end{cases} \quad (1)$$

For each of these four sets of inputs write down the network's output (i.e., its prediction) in the "Output" column of the table below.

Input 1	Input 2	Output
0	0	
1	1	
0	1	
1	0	