



Machine Learning: The Probabilistic Perspective

Konstantin Tretyakov

<http://kt.era.ee>

**IFI Summer
School 2014**

STACC

Software Technology and
Applications Competence Center



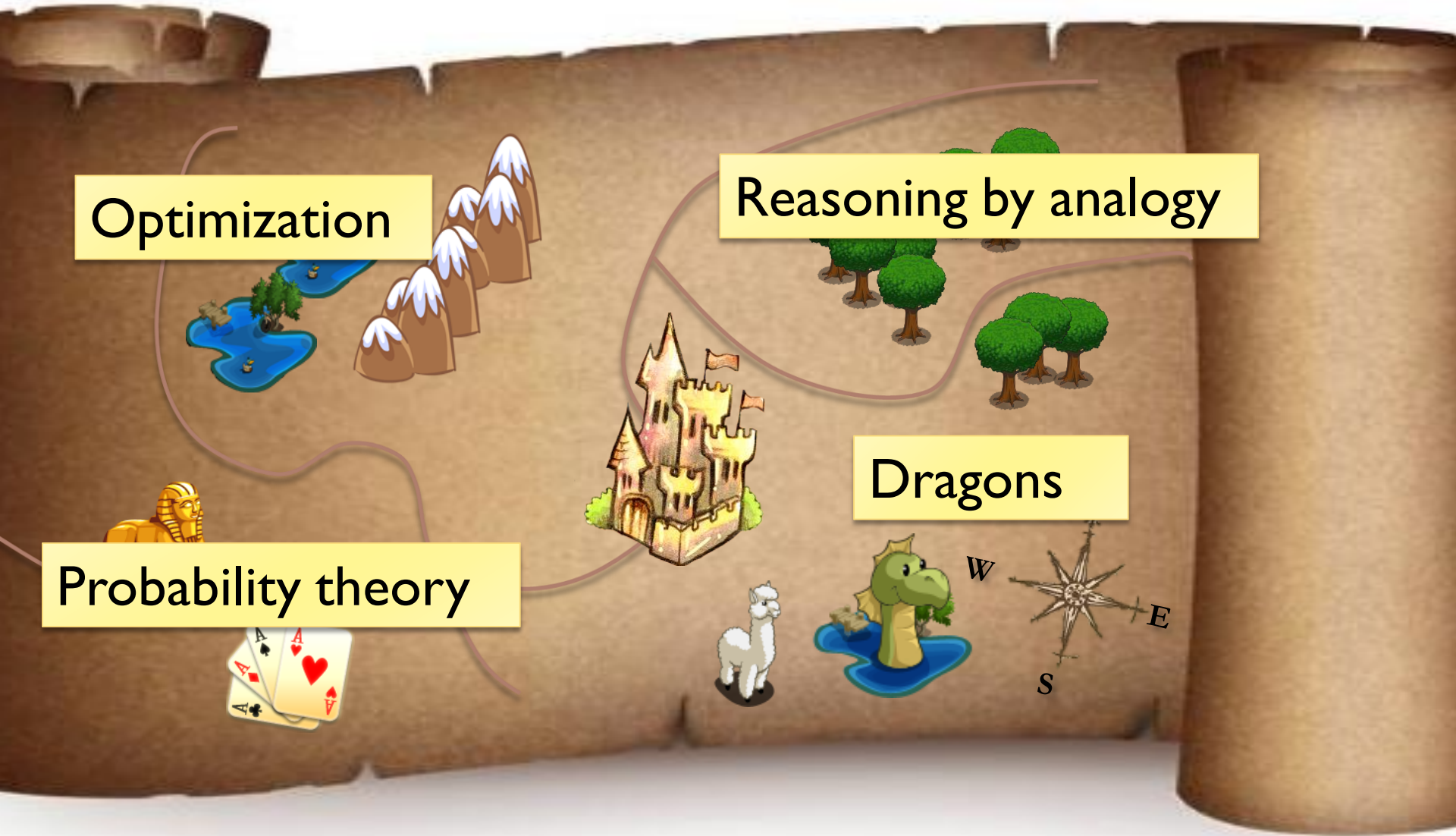
The Land of Machine Learning

Optimization

Reasoning by analogy

Probability theory

Dragons



So far...

- ▶ Machine learning is important and interesting
- ▶ The general concept:

Fitting models to data

So far...

- ▶ Machine learning is important and interesting
- ▶ The general concept:

Fitting models to data



Optimization



**Probability
Theory**

So far...

- ▶ Instance-based methods
- ▶ Tree learning methods
- ▶ The “soul” of machine learning:

$$\operatorname{argmin}_{\mathbf{w}} \operatorname{Error}(\text{Data}, \mathbf{w}) + \lambda \operatorname{Complexity}(\mathbf{w})$$

- ▶ Particular models:
 - ▶ OLS regression (ℓ_2 -loss, 0-penalty regression)
 - ▶ Ridge regression (ℓ_2 -loss, ℓ_2 -penalty regression)

Next



Why should the model,
tuned on the **training set**,
generalize to the test set?



The “No Free Lunch” Principle

Learning **purely from data** is, in general, impossible

X	Y	Output
0	0	False
0	1	True
1	0	True
1	1	?



The “No Free Lunch” Principle

Learning **purely from data** is, in general, impossible

- ▶ Is it good or bad?



- ▶ What should we do to enable learning?



The “No Free Lunch” Principle

Learning **purely from data** is, in general, impossible

▶ Is it good or bad?

▶ Good for cryptographers, bad for data miners

▶ What should we do to enable learning?

▶ Introduce **assumptions about data** (“inductive bias”):

1. **How does existing data relate to the future data?**

2. **What is the system we are learning?**

The “No Free Lunch” Principle

Learning **purely from data** is, in general, impossible

► Is it good or bad?

► Good for cryptographers, bad for data miners

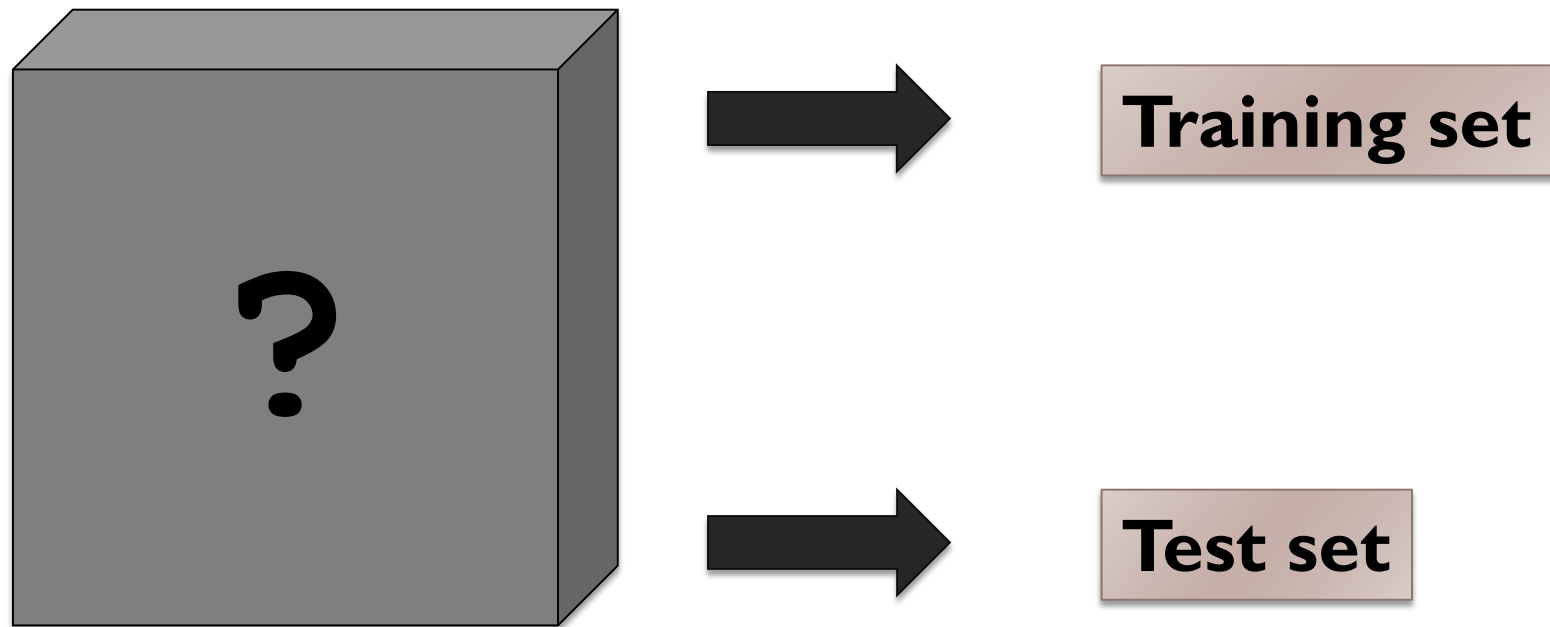
► What should we do to enable learning?

► Introduce **assumptions about data** (“inductive bias”).

1. **How does existing data relate to the future data?**

2. **What is the system we are learning?**

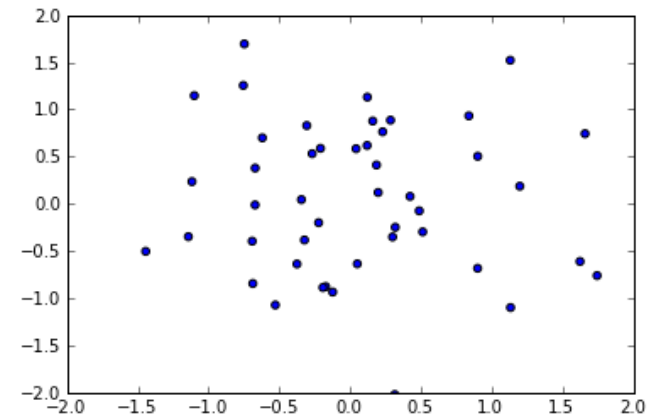
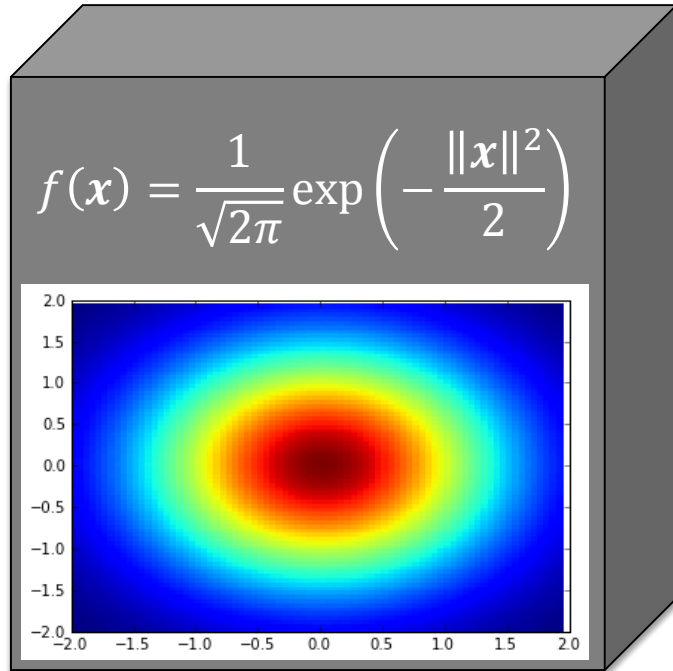
How does existing data relate to future data?



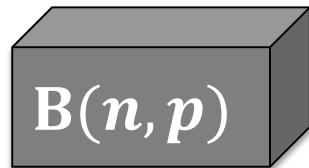
X	P(X)
heads	0.5
tails	0.5



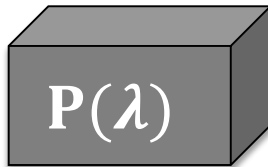
**heads,
heads,
tails,
heads,
tails,
...**

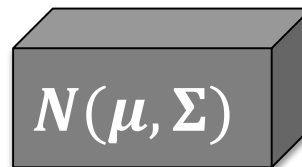


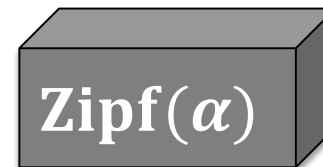
Probability theory

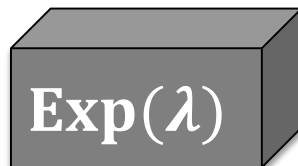

$$B(n, p)$$

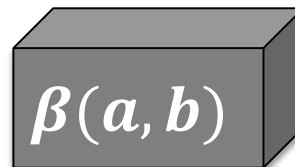

$$Hg(m, n, M, N)$$

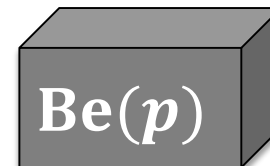

$$P(\lambda)$$

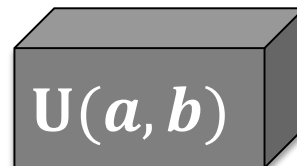

$$N(\mu, \Sigma)$$


$$\text{Zipf}(\alpha)$$


$$\text{Exp}(\lambda)$$


$$\beta(a, b)$$

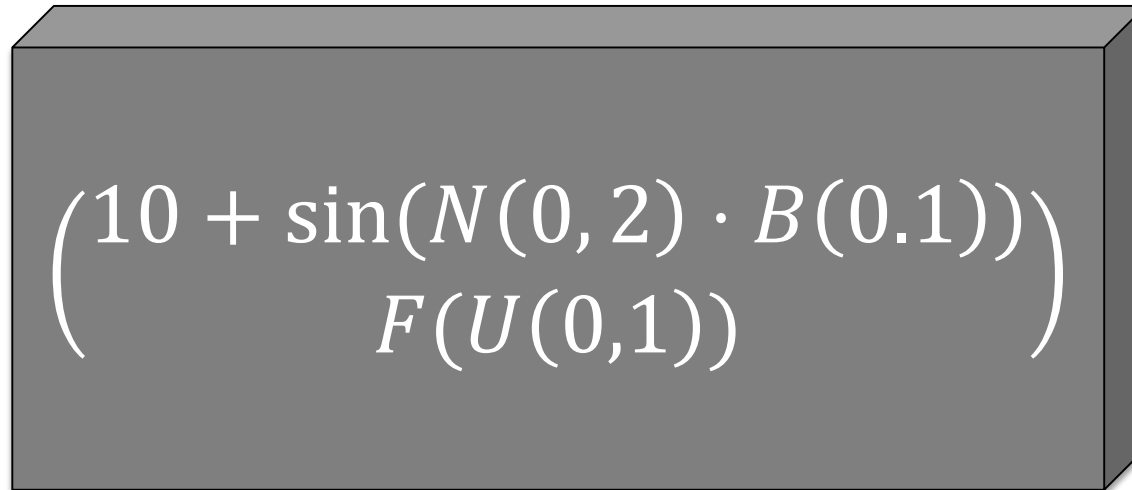

$$\text{Be}(p)$$


$$U(a, b)$$


$$\dots$$

V • T • E	Probability distributions	[hide]
	Discrete univariate with finite support	[hide]
	Benford • Bernoulli • Beta-binomial • binomial • categorical • hypergeometric • Poisson binomial • Rademacher • discrete uniform • Zipf • Zipf-Mandelbrot	
	Discrete univariate with infinite support	[hide]
	beta negative binomial • Boltzmann • Conway–Maxwell–Poisson • discrete phase-type • Delaporte • extended negative binomial • Gauss–Kuzmin • geometric • logarithmic • negative binomial • parabolic fractal • Poisson • Skellam • Yule–Simon • zeta	
	Continuous univariate supported on a bounded interval, e.g. [0,1]	[hide]
	Arcsine • ARGUS • Balding–Nichols • Bates • Beta • Beta rectangular • Irwin–Hall • Kumaraswamy • logit-normal • Noncentral beta • raised cosine • triangular • U-quadratic • uniform • Wigner semicircle	
	Continuous univariate supported on a semi-infinite interval, usually [0,∞)	[hide]
	Benini • Benktander 1st kind • Benktander 2nd kind • Beta prime • Bose–Einstein • Burr • chi-squared • chi • Coxian • Dagum • Davis • Erlang • exponential • F • Fermi–Dirac • folded normal • Fréchet • Gamma • generalized inverse Gaussian • half-logistic • half-normal • Hotelling's T-squared • hyper-exponential • hypoexponential • inverse chi-squared (scaled-inverse-chi-squared) • inverse Gaussian • inverse gamma • Kolmogorov • Lévy • log-Cauchy • log-Laplace • log-logistic • log-normal • Maxwell–Boltzmann • Maxwell speed • Mittag–Leffler • Nakagami • noncentral chi-squared • Pareto • phase-type • Rayleigh • relativistic Breit–Wigner • Rice • Rosin–Rammler • shifted Gompertz • truncated normal • type-2 Gumbel • Weibull • Wilks' lambda	
	Continuous univariate supported on the whole real line (−∞, ∞)	[hide]
	Cauchy • exponential power • Fisher's z • generalized normal • generalized hyperbolic • geometric stable • Gumbel • Holtzmark • hyperbolic secant • Landau • Laplace • Linnik • logistic • noncentral t • normal (Gaussian) • normal-inverse Gaussian • skew normal • slash • stable • Student's t • type-1 Gumbel • variance-gamma • Voigt	
	Continuous univariate with support whose type varies	[hide]
	generalized extreme value • generalized Pareto • Tukey lambda • q-Gaussian • q-exponential • shifted log-logistic	
	Mixed continuous-discrete univariate distributions	[hide]
	rectified Gaussian	
	Multivariate (joint)	[hide]
	<i>Discrete:</i> Ewens • multinomial • Dirichlet-multinomial • negative multinomial	
	<i>Continuous:</i> Dirichlet • Generalized Dirichlet • multivariate normal • Multivariate stable • multivariate Student • normal-scaled inverse gamma • normal-gamma	
	<i>Matrix-valued:</i> inverse matrix gamma • inverse-Wishart • matrix normal • matrix t • matrix gamma • normal-inverse-Wishart • normal-Wishart • Wishart	
	Directional	[hide]
	<i>Univariate (circular) directional:</i> Circular uniform • univariate von Mises • wrapped normal • wrapped Cauchy • wrapped exponential • wrapped Lévy	
	<i>Bivariate (spherical):</i> Kent • <i>Bivariate (toroidal):</i> bivariate von Mises	
	<i>Multivariate:</i> von Mises–Fisher • Bingham	
	Degenerate and singular	[hide]
	<i>Degenerate:</i> discrete degenerate • Dirac delta function	
	<i>Singular:</i> Cantor	
	Families	[hide]

Probability theory

A 3D rectangular box with a dark gray front face and lighter gray top and right side faces. The box is oriented horizontally. On the front face, a mathematical expression is written in white serif font.
$$\left(\begin{array}{c} 10 + \sin(N(0, 2) \cdot B(0.1)) \\ F(U(0,1)) \end{array} \right)$$

Probability theory

```
from numpy.random import beta, binomial,  
chisquare, dirichlet, exponential, f, gamma,  
geometric, gumbel, hypergeometric, ...
```

```
>>> numpy.random.seed(1)
```

```
>>> binomial(10, 0.2)
```

```
::: 2
```

Probability theory

```
from scipy.stats.distributions import beta,  
binom, chisquare, ...
```

```
>>> numpy.random.seed(1)
```

```
>>> X = binom(10, 0.2)
```

```
>>> X.rvs()
```

```
::: 2
```

```
>>> X.pmf(2), X.cdf(2), X.mean(), X.std(), ...
```

Everything is Probabilistic?

What is your height?



Everything is Probabilistic?

What is your height?

Is it a fixed number?



Everything is Probabilistic?

What is your height?

Is it a fixed number?

- ▶ Frequentist: **Yes, it is**, we just don't know it precisely.
- ▶ Bayesian: **No, it is not.**
It is a **distribution**.



Everything is Probabilistic?

What is your height?

Is it a fixed number?

- ▶ Frequentist: **Yes, it is**, we just don't know it precisely.
- ▶ Bayesian: **No, it is not**. It is a **distribution**.

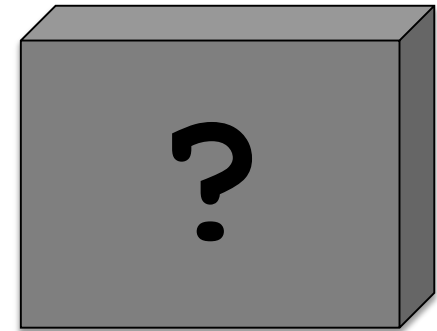
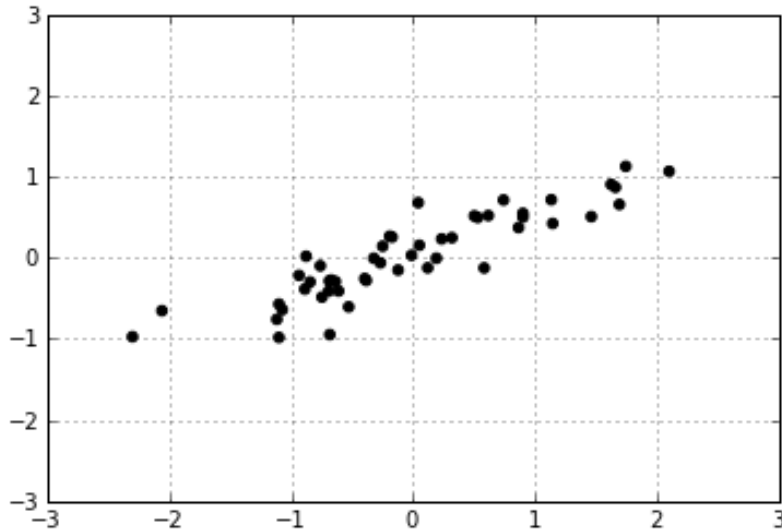
In any case, we need probabilistic reasoning.



Statistics & Decision Theory

► Statistics

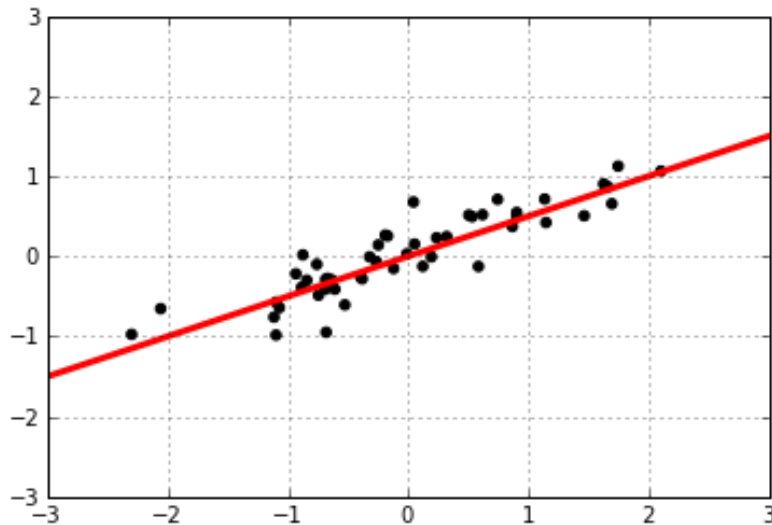
- How do we **infer** a probabilistic **model** based on data?



Statistics & Decision Theory

► Statistics

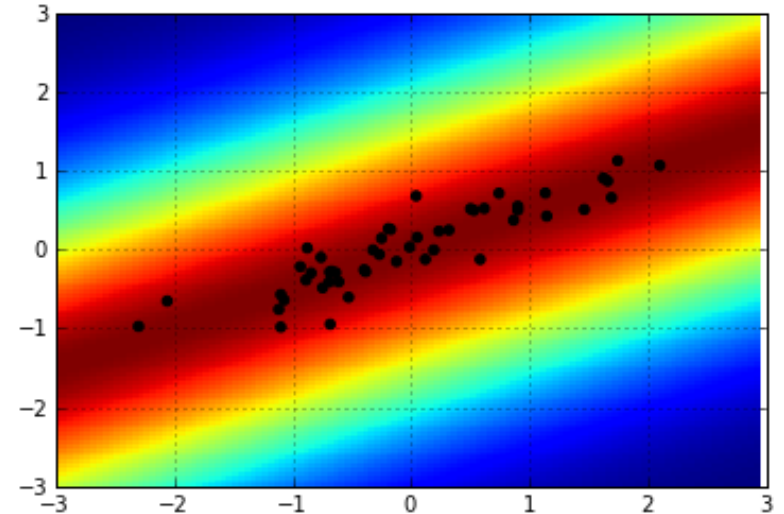
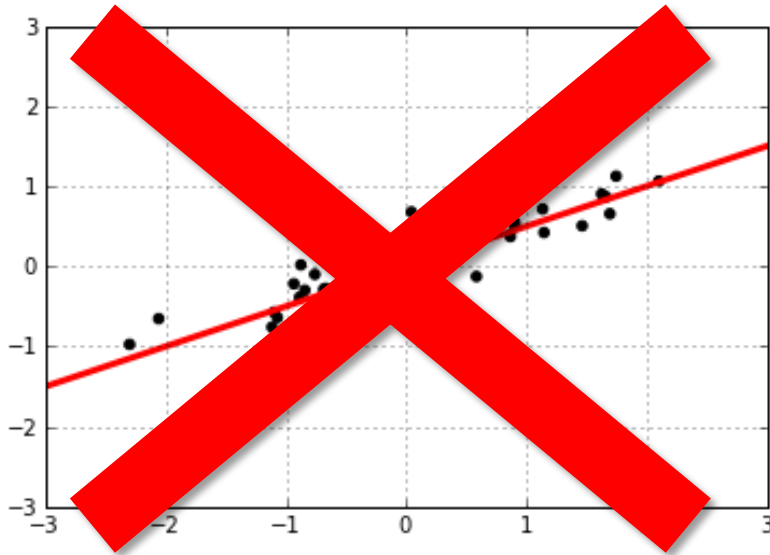
- How do we **infer** a probabilistic **model** based on data?



Statistics & Decision Theory

► Statistics

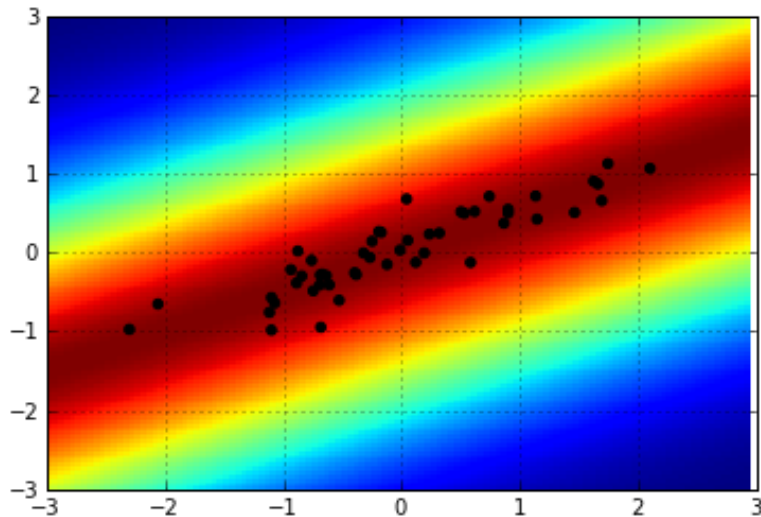
- How do we **infer** a probabilistic **model** based on data?



Statistics & Decision Theory

► Decision theory

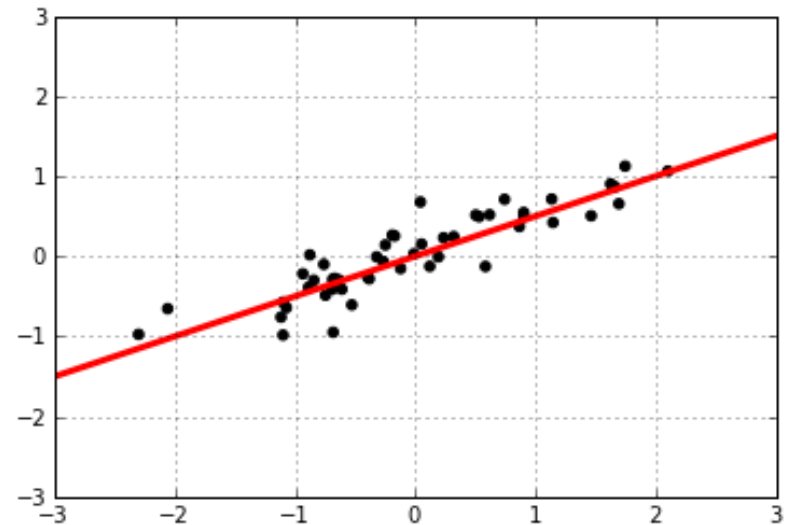
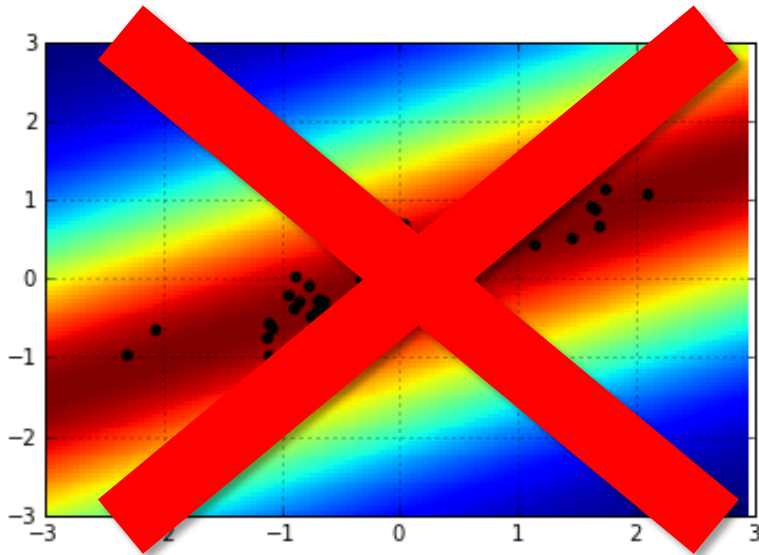
- How do we **use** a probabilistic model **to predict**?



Statistics & Decision Theory

► Decision theory

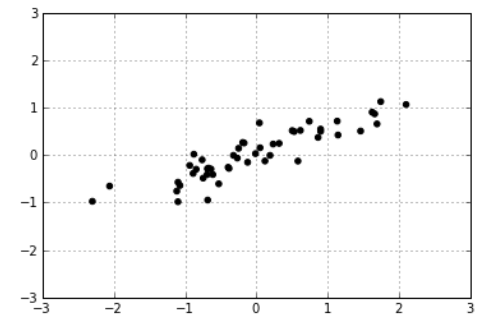
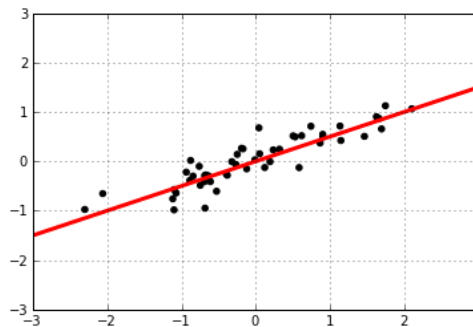
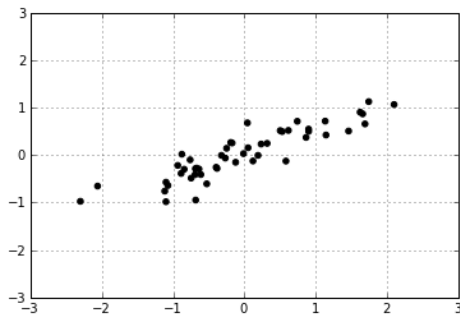
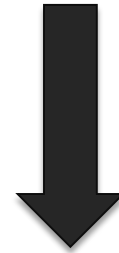
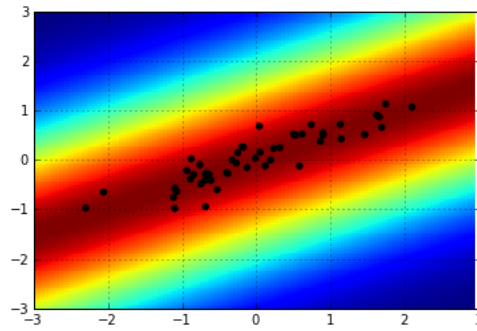
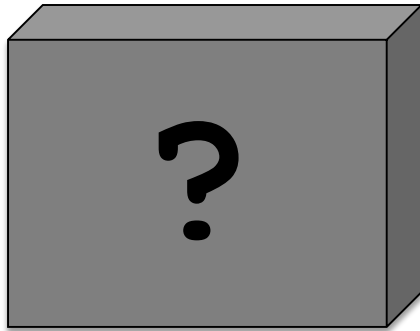
- How do we **use** a probabilistic model **to predict**?



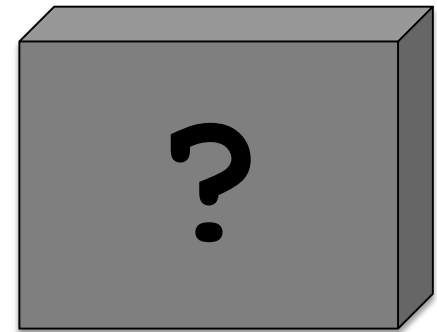
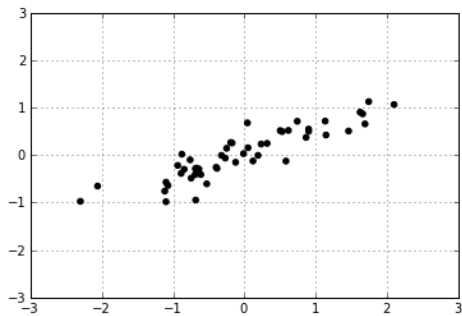
Quiz

- ▶ Model, trained on the training set might work well on the test set because:
 - ▶ Because we **assume** a single underlying mechanism.
 - ▶ Because we **use statistical inference** to infer the mechanism.
 - ▶ Because we **use decision theory** to produce optimal decisions.

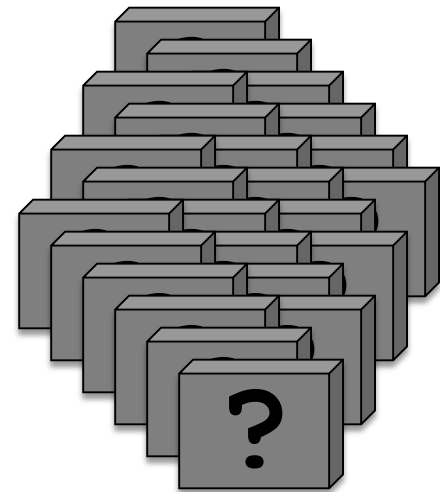
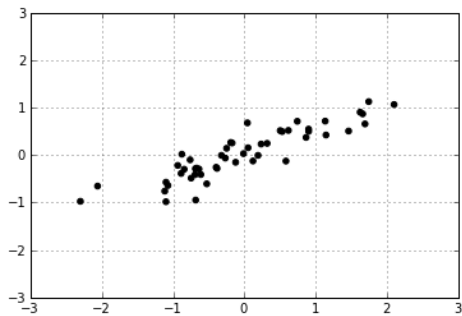
Quiz



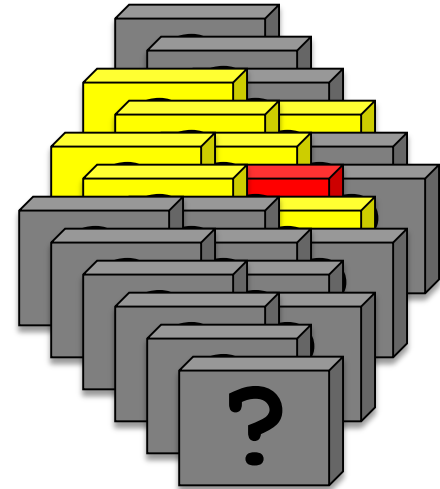
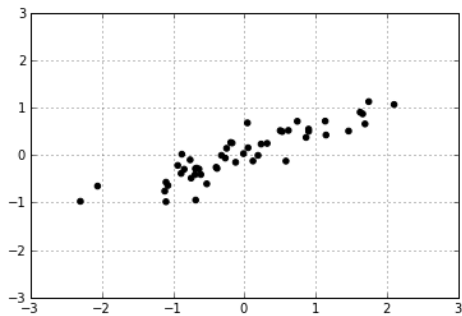
Statistics



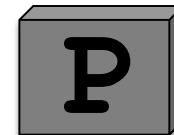
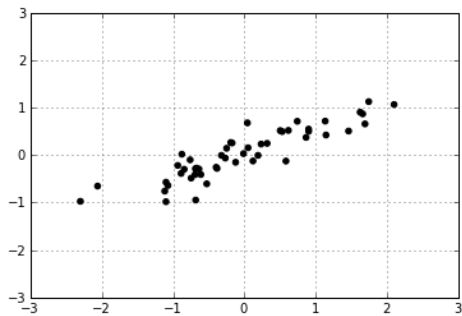
Space of candidate models



Statistics

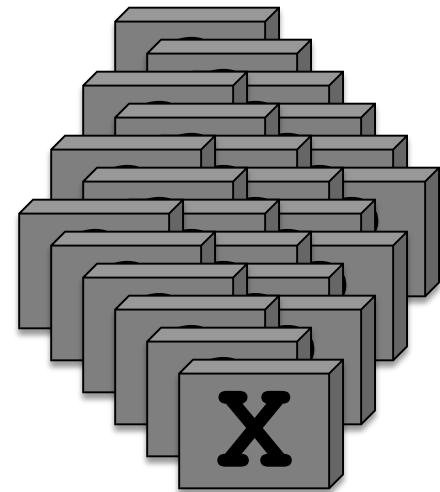
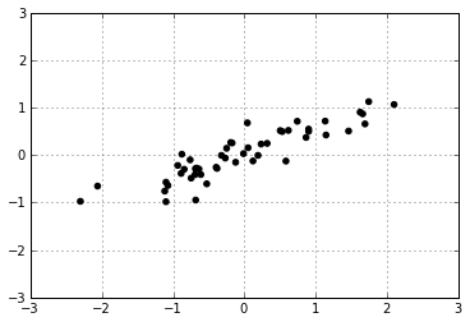


Hypothesis testing

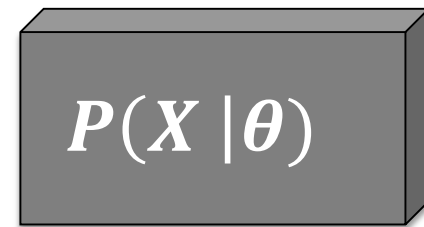
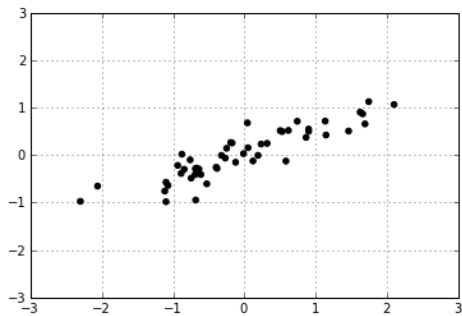


or not?

Model selection



Parameter inference



Parameter inference

Biased coin

$Be(p)$

X	$P(X)$
1	p
0	$1-p$



1,1,0,1,1

$$n_1 = 4$$

$$n_0 = 1$$

Maximum Likelihood Estimation

- ▶ **Data Likelihood:**

$$\Pr[\text{Data} \mid \text{Model}]$$

- ▶ **Example:**

- ▶ Model: $\text{Be}(0.5)$
- ▶ Data: 1,1,0,1,1
- ▶ Likelihood: ?

Maximum Likelihood Estimation

- ▶ Data Likelihood:

$$\Pr[\text{Data} \mid \text{Model}]$$

- ▶ Example:

- ▶ Model: $\text{Be}(0.5)$

- ▶ Data: 1,1,0,1,1

- ▶ Likelihood: $0.5 \cdot 0.5 \cdot 0.5 \cdot 0.5 \cdot 0.5 = 2^{-5}$

0.03125

Maximum Likelihood Estimation

- ▶ **Data Likelihood:**

$$\Pr[\text{Data} \mid \text{Model}]$$

- ▶ **Example:**

- ▶ Model: $\text{Be}(0.2)$
- ▶ Data: 1,1,0,1,1
- ▶ Likelihood: ?

Maximum Likelihood Estimation

- ▶ Data Likelihood:

$$\Pr[\text{Data} \mid \text{Model}]$$

- ▶ Example:

- ▶ Model: $\text{Be}(0.2)$

- ▶ Data: 1,1,0,1,1

- ▶ Likelihood: $0.2 \cdot 0.2 \cdot 0.8 \cdot 0.2 \cdot 0.2 = 0.2^4 \cdot 0.8$

0.00128

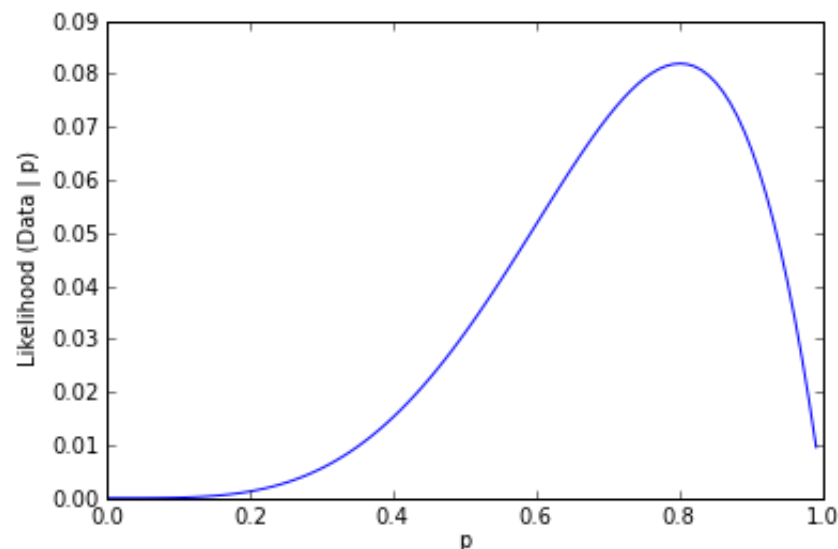
Maximum Likelihood Estimation

► Example:

► Model: $\text{Be}(p)$

► Data: 1,1,0,1,1

► Likelihood: $p \cdot p \cdot (1 - p) \cdot p \cdot p = p^{n_1} (1 - p)^{n_0}$



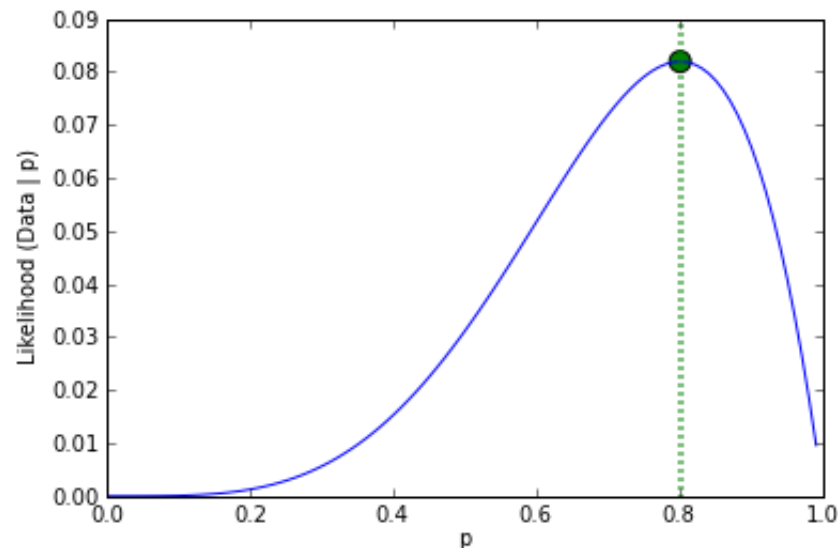
Maximum Likelihood Estimation

► Example:

► Model: $\text{Be}(p)$

► Data: 1,1,0,1,1

► Likelihood: $p \cdot p \cdot (1 - p) \cdot p \cdot p = p^{n_1} (1 - p)^{n_0}$



$$\hat{p} = \frac{n_1}{n_0 + n_1}$$

Maximum Likelihood Estimation

► Maximum Likelihood Estimation:

$$\operatorname{argmax}_{\text{Model}} \Pr(\text{Data} | \text{Model})$$

Problems of MLE

- ▶ You are on a trip in an exotic country and you meet a person who happens to be from Switzerland.
- ▶ Is he a member of the Swiss Parliament?

Problems of MLE

- ▶ Data: “X is from Switzerland”
- ▶ Models:
 - ▶ “X is a member of Swiss Parliament”,
 - ▶ “X is not a member of Swiss Parliament”

Problems of MLE

- ▶ Data: “X is from Switzerland”
- ▶ Models:
 - ▶ “X is a member of Swiss Parliament”,
 - ▶ “X is not a member of Swiss Parliament”
- ▶ Likelihoods:
 - ▶ $P(\text{X is from Switzerland} \mid \text{X is a member of SP}) =$
 - ▶ $P(\text{X is from Switzerland} \mid \text{X is **not** a member of SP}) =$

Problems of MLE

- ▶ Data: “X is from Switzerland”
- ▶ Models:
 - ▶ “X is a member of Swiss Parliament”,
 - ▶ “X is not a member of Swiss Parliament”
- ▶ Likelihoods:
 - ▶ $P(\text{X is from Switzerland} \mid \text{X is a member of SP}) = 1$
 - ▶ $P(\text{X is from Switzerland} \mid \text{X is **not** a member of SP}) = \frac{8}{7000}$

Problems of MLE

► Data: “X is from Switzerland”

► Models:

► “X is a member of Swiss Parliament”,

► **MLE treats all candidate models as equal and can thus overfit**

► $P(X \text{ is from Switzerland} \mid X \text{ is a member of SP}) = 1$

► $P(X \text{ is from Switzerland} \mid X \text{ is **not** a member of SP}) = \frac{8}{7000}$

Maximum A-posteriori Estimation

- ▶ Maximum Likelihood Estimate (MLE):

$$\operatorname{argmax}_{\text{Model}} \Pr(\text{Data} | \text{Model})$$

- ▶ Maximum A-posteriori Estimate (MAP):

$$\operatorname{argmax}_{\text{Model}} \Pr(\text{Model} | \text{Data})$$

MAP Estimation

$$\operatorname{argmax}_{\text{Model}} \Pr(\text{Model}|\text{Data})$$

MAP Estimation

$$\operatorname{argmax}_{\text{Model}} \Pr(\text{Model}|\text{Data})$$

$$\operatorname{argmax}_{\text{Model}} \frac{\Pr(\text{Model}, \text{Data})}{\Pr(\text{Data})}$$

$$\operatorname{argmax}_{\text{Model}} \Pr(\text{Model}, \text{Data})$$

MAP Estimation

$$\operatorname{argmax}_{\text{Model}} \Pr(\text{Model}|\text{Data})$$

$$\operatorname{argmax}_{\text{Model}} \frac{\Pr(\text{Model}, \text{Data})}{\Pr(\text{Data})}$$

$$\operatorname{argmax}_{\text{Model}} \Pr(\text{Model}, \text{Data})$$

$$\operatorname{argmax}_{\text{Model}} \Pr(\text{Data} | \text{Model}) \cdot \Pr(\text{Model})$$

MAP Estimation

$$\operatorname{argmax}_{\text{Model}} \boxed{\Pr(\text{Model}|\text{Data})}$$

$$\operatorname{argmax}_{\text{Model}} \frac{\Pr(\text{Model}, \text{Model posterior})}{\Pr(\text{Data})}$$

$$\operatorname{argmax}_{\text{Model}} \Pr(\text{Model}, \text{Data})$$

$$\operatorname{argmax}_{\text{Model}} \boxed{\Pr(\text{Data} | \text{Model})} \boxed{\Pr(\text{Model})}$$

Likelihood

Model prior

Summary

- ▶ Maximum Likelihood Estimate (MLE):

$$\operatorname{argmax}_{\text{Model}} \Pr(\text{Data} \mid \text{Model})$$

- ▶ Maximum A-posteriori Estimate (MAP):

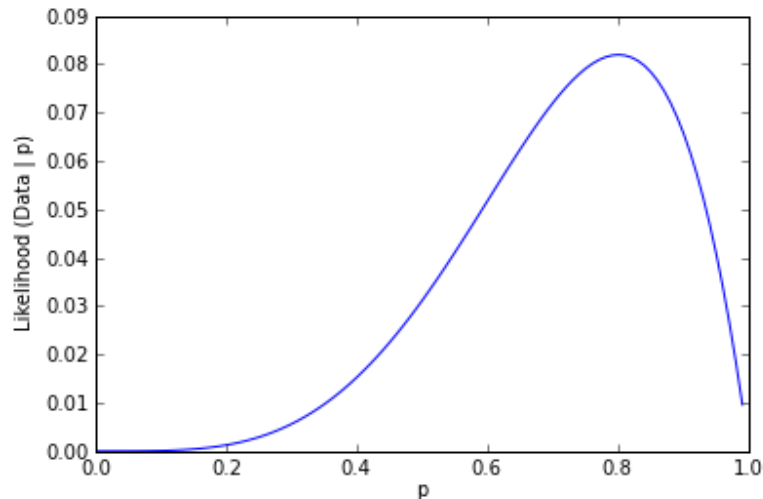
$$\operatorname{argmax}_{\text{Model}} \Pr(\text{Data} \mid \text{Model}) \Pr(\text{Model})$$

MAP Estimation

► Model: $\text{Be}(p)$

Data: 1,1,0,1,1

Likelihood: $p^4(1 - p)$



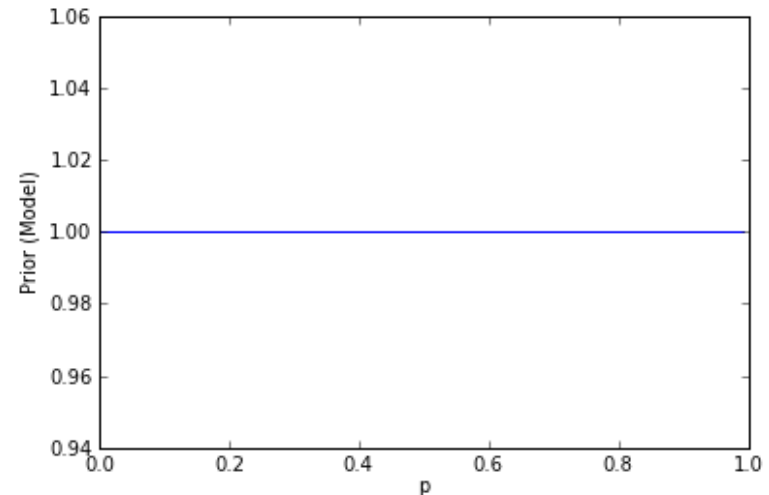
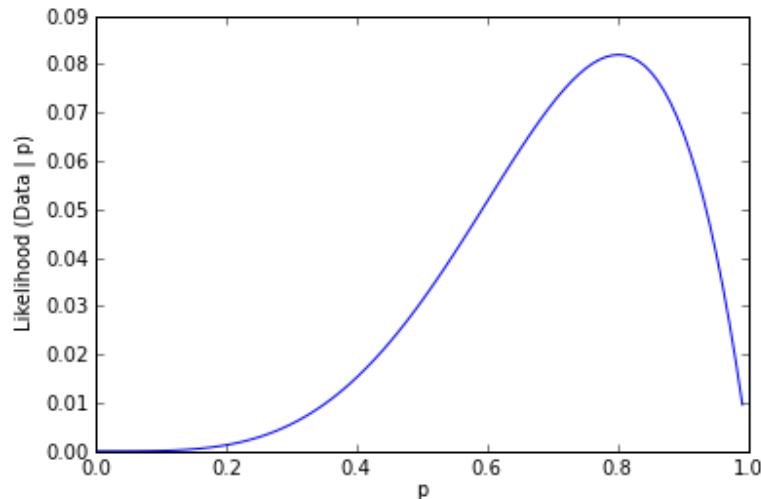
MAP Estimation

► Model: $\text{Be}(p)$

Data: 1,1,0,1,1

Likelihood: $p^4(1 - p)$

Prior: $U(0,1)$



$$\hat{p}_{MAP} = \hat{p}_{MLE} = \frac{n_1}{n_0 + n_1}$$

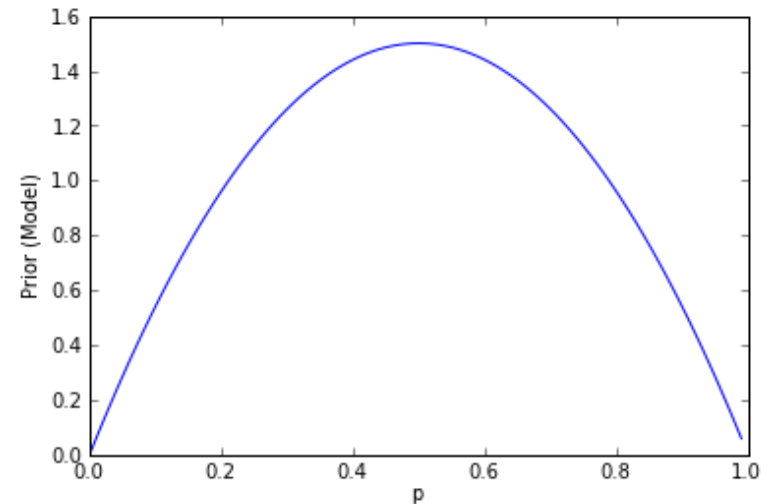
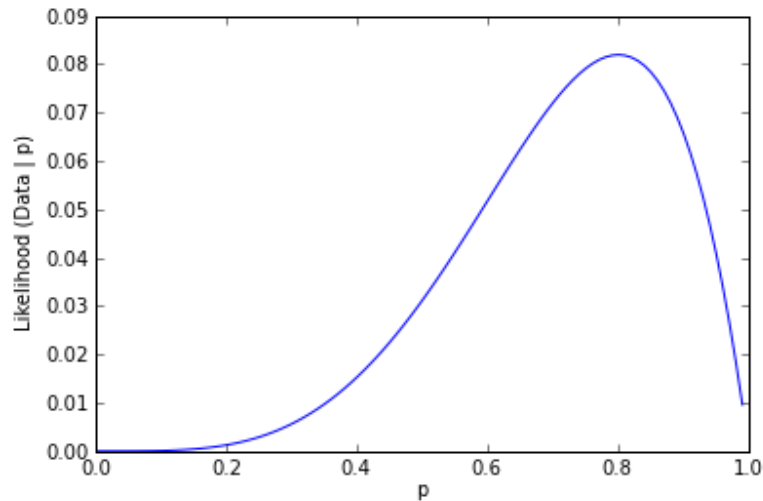
MAP Estimation

► Model: $\text{Be}(p)$

Data: 1,1,0,1,1

Likelihood: $p^4(1 - p)$

Prior: $\text{Beta}(2, 2)$

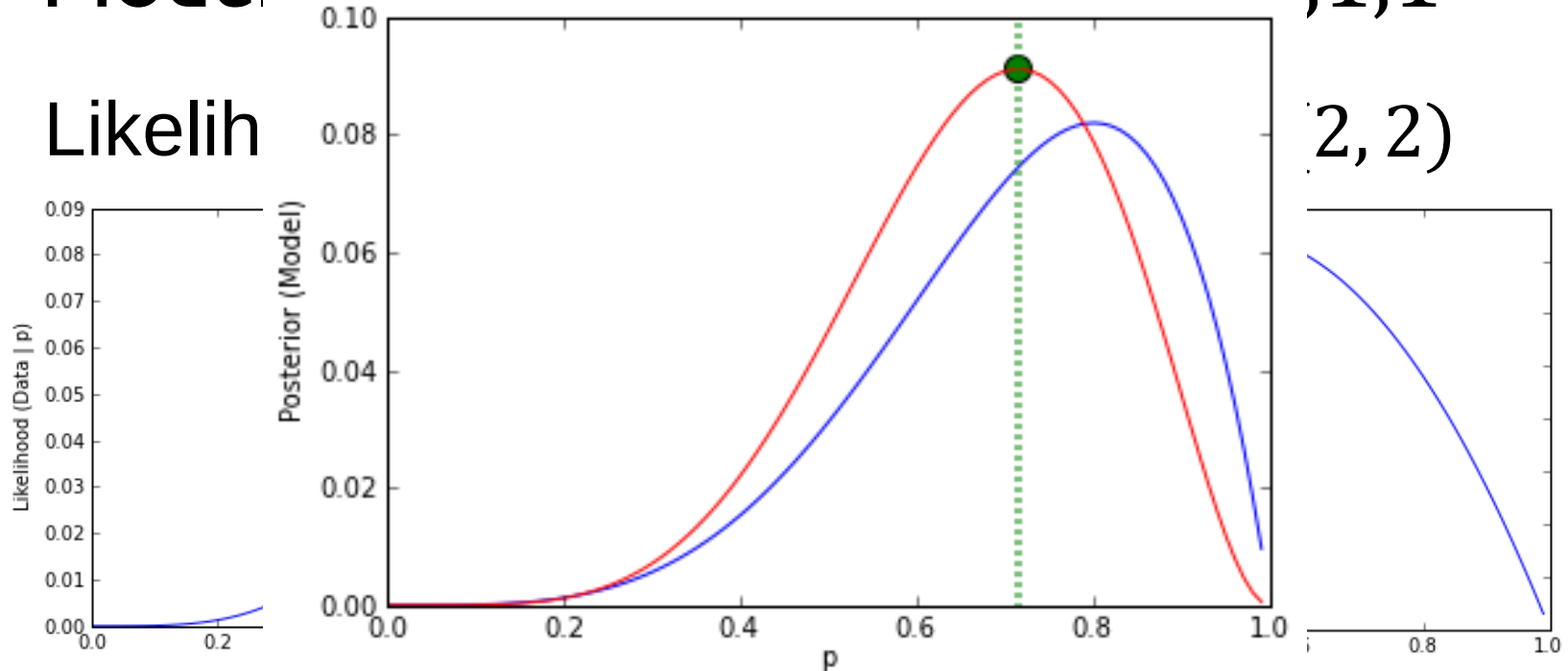


MAP Estimation

► Model: $B_\theta(p)$

Data: 1 1 0, 1, 1

(2, 2)



$$\hat{p}_{MAP} = \frac{n_1 + 1}{n_0 + n_1 + 2}$$

MAP Estimation

$$\operatorname{argmax}_{\text{Model}} \Pr(\text{Data} \mid \text{Model}) \cdot \Pr(\text{Model})$$

MAP Estimation

$$\operatorname{argmax}_{\text{Model}} \Pr(\text{Data} \mid \text{Model}) \cdot \Pr(\text{Model})$$

$$\operatorname{argmax}_{\text{Model}} \log (\Pr(\text{Data} \mid \text{Model}) \cdot \Pr(\text{Model}))$$

MAP Estimation

$$\operatorname{argmax}_{\text{Model}} \Pr(\text{Data} \mid \text{Model}) \cdot \Pr(\text{Model})$$

$$\operatorname{argmax}_{\text{Model}} \log (\Pr(\text{Data} \mid \text{Model}) \cdot \Pr(\text{Model}))$$

$$\operatorname{argmax}_{\text{Model}} \log \Pr(\text{Data} \mid \text{Model}) + \log \Pr(\text{Model})$$

MAP Estimation

$$\operatorname{argmax}_{\text{Model}} \Pr(\text{Data} \mid \text{Model}) \cdot \Pr(\text{Model})$$

$$\operatorname{argmax}_{\text{Model}} \log (\Pr(\text{Data} \mid \text{Model}) \cdot \Pr(\text{Model}))$$

$$\operatorname{argmax}_{\text{Model}} \log \Pr(\text{Data} \mid \text{Model}) + \log \Pr(\text{Model})$$

MAP Estimation

$$\operatorname{argmax}_{\text{Model}} \Pr(\text{Data} \mid \text{Model}) \cdot \Pr(\text{Model})$$

$$\operatorname{argmax}_{\text{Model}} \log (\Pr(\text{Data} \mid \text{Model}) \cdot \Pr(\text{Model}))$$

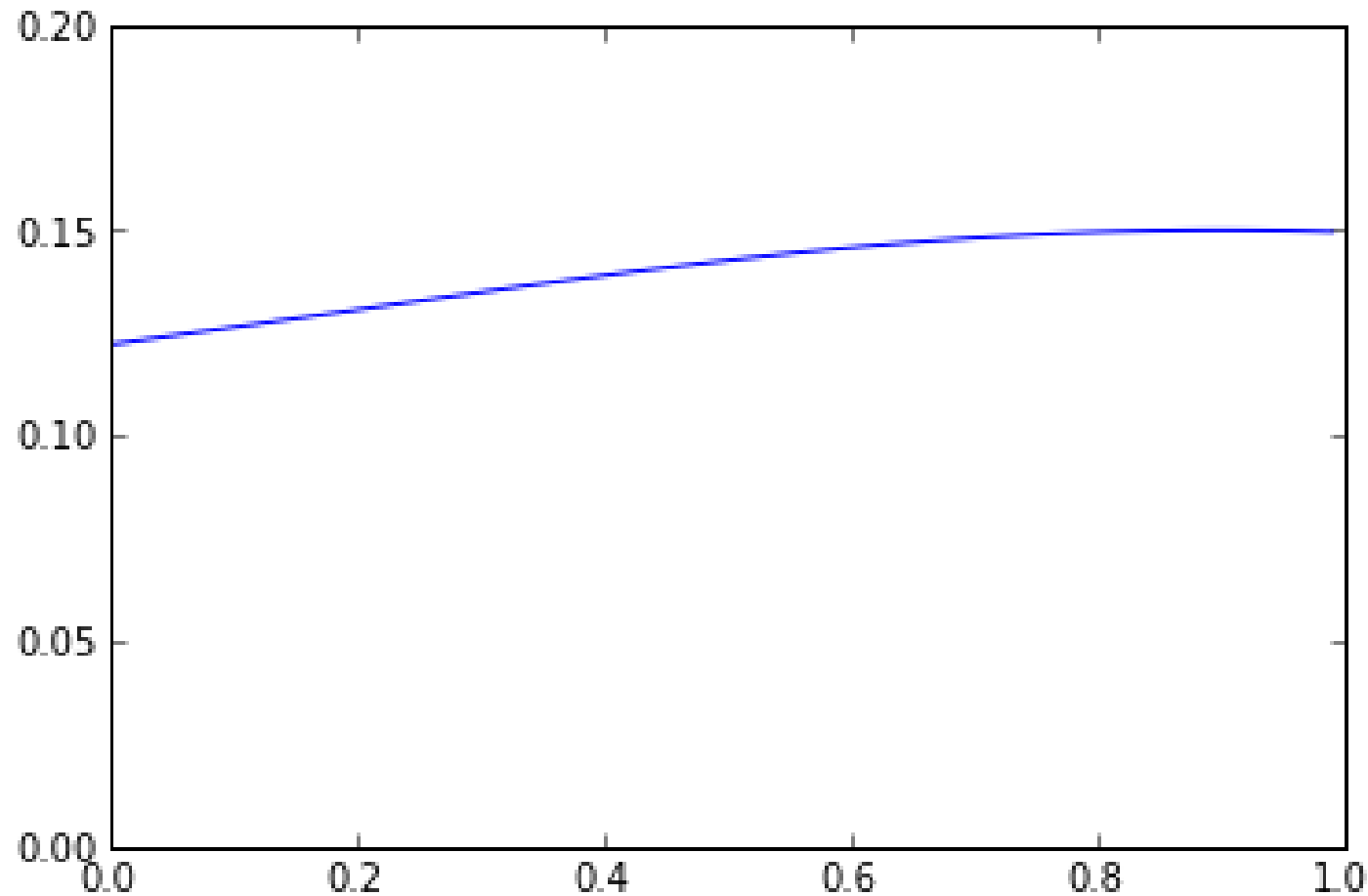
$$\operatorname{argmax}_{\text{Model}} \log \Pr(\text{Data} \mid \text{Model}) + \log \Pr(\text{Model})$$

$$\operatorname{argmin}_{\mathbf{w}} \text{Error}(\text{Data}, \mathbf{w}) + \text{Complexity}(\mathbf{w})$$

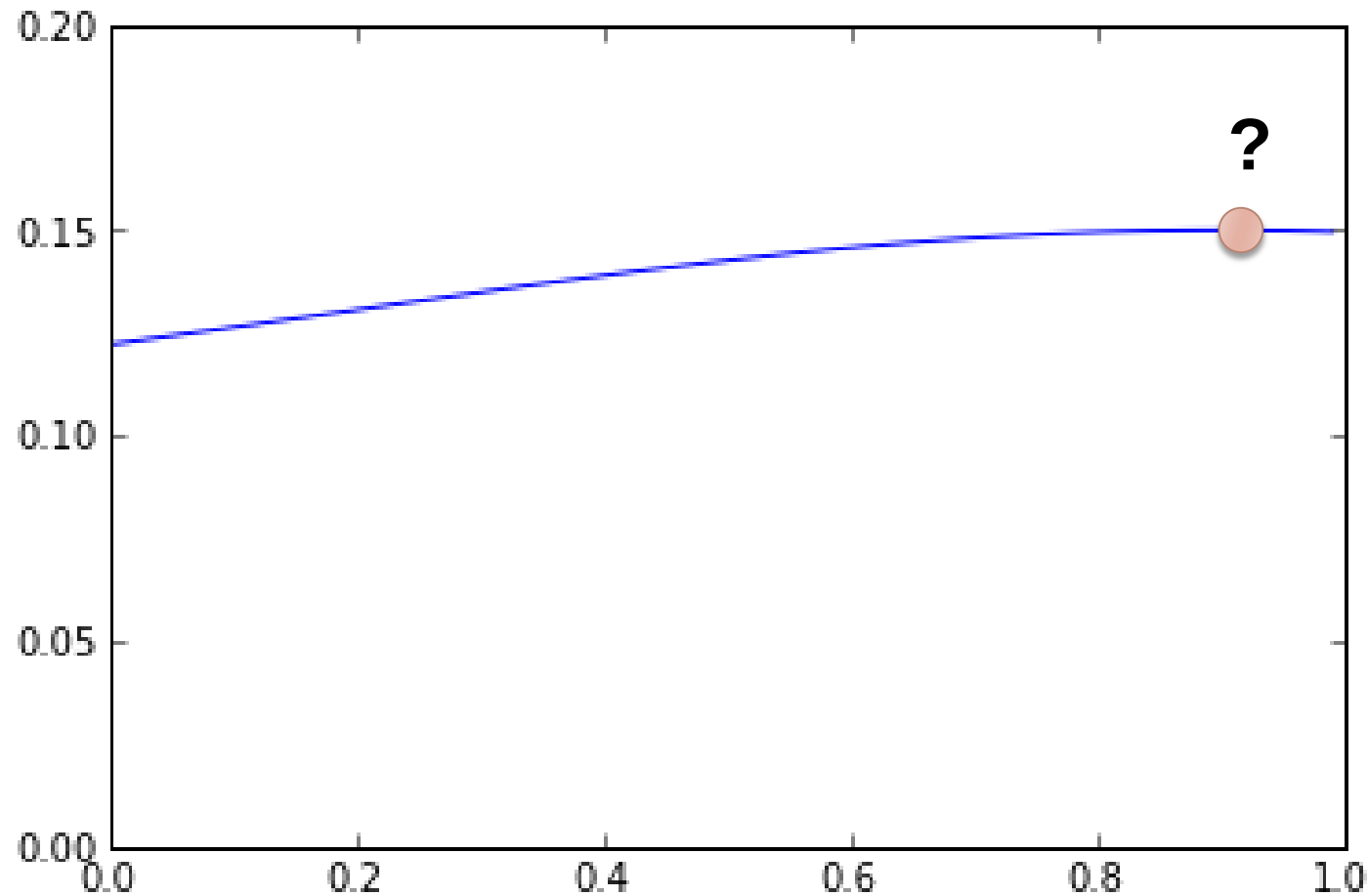
Problems of MAP estimation



Problems of MAP estimation

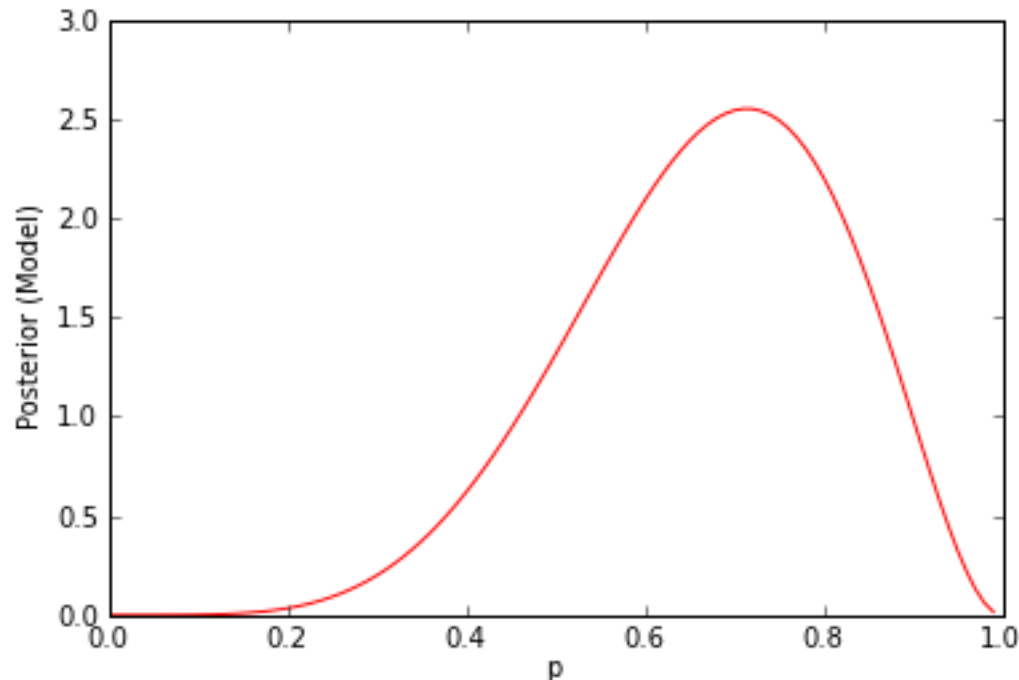


Problems of MAP estimation



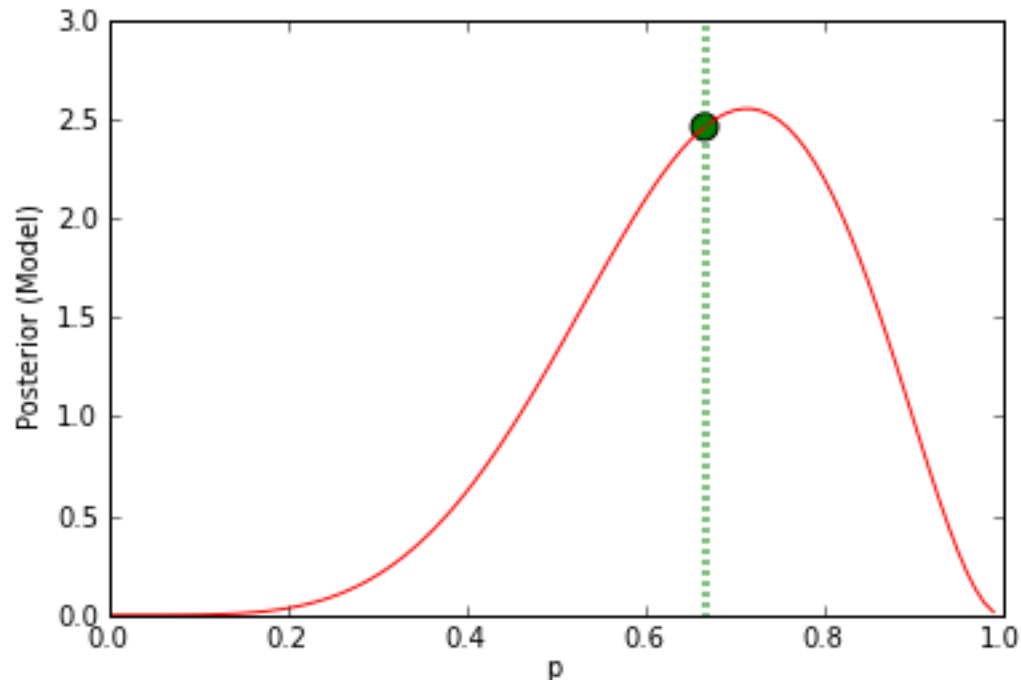
Bayesian estimation

- Pick the model with minimal *expected risk*
 $E(\text{Model} | \text{Data})$



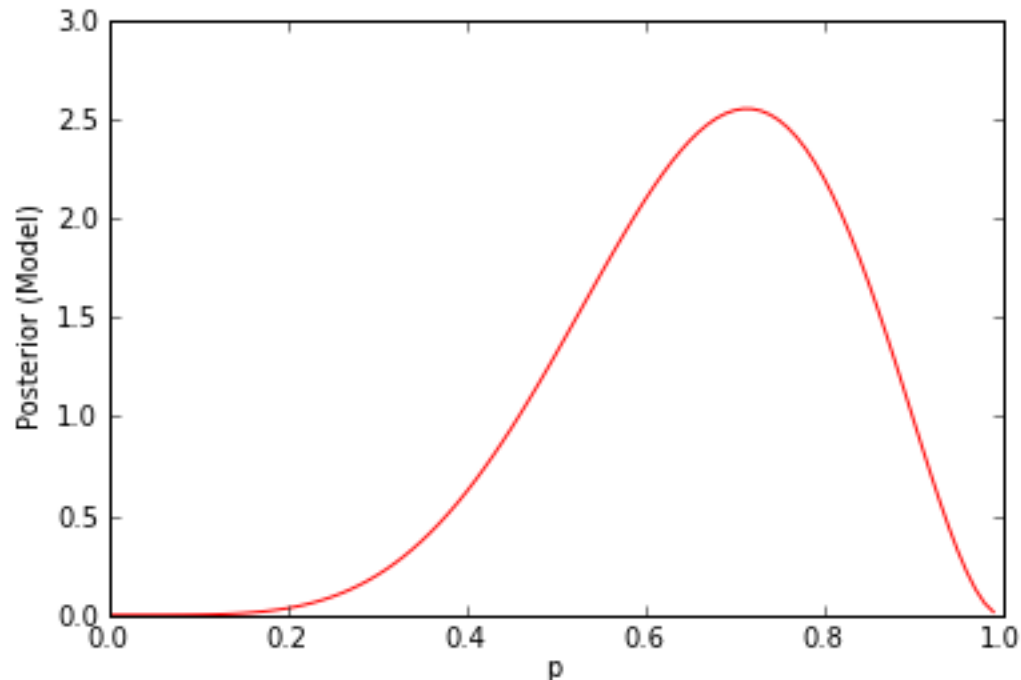
Bayesian estimation

- Pick the model with minimal *expected risk*
 $E(\text{Model} | \text{Data})$



Bayesian estimation +

- Use the full posterior distribution
 $\Pr(\text{Model} | \text{Data})$



Quiz

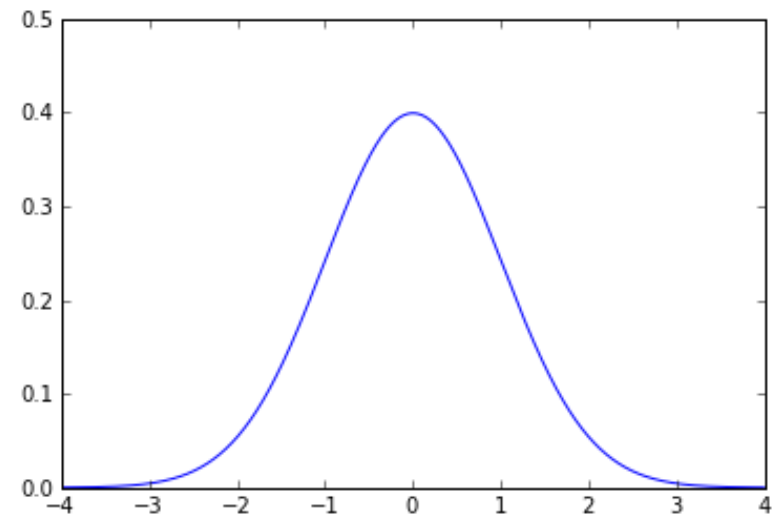
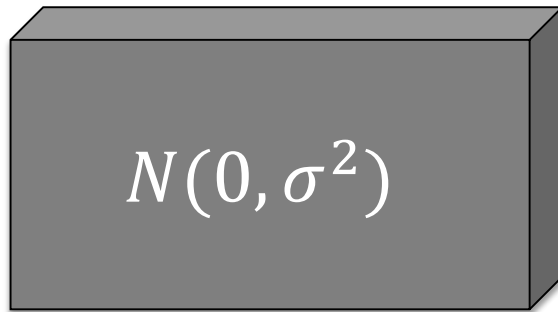


- ▶ Three major model inference methods are:



Linear Regression (again)

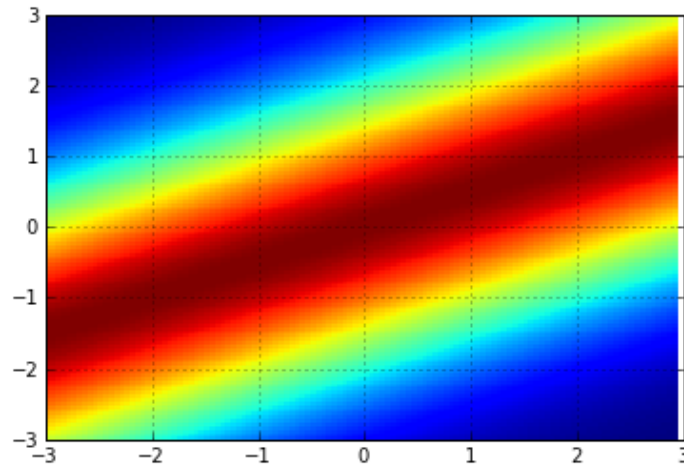
Normal distribution



$$f(x) = \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left(-\frac{1}{2} \frac{x^2}{\sigma^2}\right)$$

Linear Regression (again)

x $\rightarrow$ $Y = w^T x + N(0, \sigma^2)$



$$Y = \mathbf{w}^T \mathbf{x} + N(0, \sigma^2)$$

$$Y = \mathbf{w}^T \mathbf{x} + N(0, \sigma^2)$$

$$e = Y - \mathbf{w}^T \mathbf{x} \sim N(0, \sigma^2)$$

$$Y = \mathbf{w}^T \mathbf{x} + N(0, \sigma^2)$$

$$e = Y - \mathbf{w}^T \mathbf{x} \sim N(0, \sigma^2)$$

$$\Pr((x, y) | \mathbf{w}, \sigma^2) = \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left(-\frac{1}{2} \frac{e^2}{\sigma^2}\right)$$

$$\Pr((x_1, y_1), (x_2, y_2), \dots, (x_n, y_n) | \mathbf{w}, \sigma^2) \\ = \prod_i \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left(-\frac{1}{2} \frac{e_i^2}{\sigma^2}\right)$$

$$\Pr((x_1, y_1), (x_2, y_2), \dots, (x_n, y_n) | \mathbf{w}, \sigma^2) \\ \propto \prod_i \exp\left(-\frac{1}{2} \frac{e_i^2}{\sigma^2}\right)$$

$$\log \Pr((x_1, y_1), (x_2, y_2), \dots, (x_n, y_n) | \mathbf{w}, \sigma^2) \\ \propto \log \prod_i \exp \left(-\frac{1}{2} \frac{e_i^2}{\sigma^2} \right)$$

$$\begin{aligned} & \log \Pr((x_1, y_1), (x_2, y_2), \dots, (x_n, y_n) | \mathbf{w}, \sigma^2) \\ & \propto \log \prod_i \exp \left(-\frac{1}{2} \frac{e_i^2}{\sigma^2} \right) \\ & = \sum_i \left(-\frac{1}{2} \frac{e_i^2}{\sigma^2} \right) \end{aligned}$$

$$\begin{aligned} & \log \Pr((x_1, y_1), (x_2, y_2), \dots, (x_n, y_n) | \mathbf{w}, \sigma^2) \\ & \propto \log \prod_i \exp \left(-\frac{1}{2} \frac{e_i^2}{\sigma^2} \right) \\ & = \sum_i \left(-\frac{1}{2} \frac{e_i^2}{\sigma^2} \right) \\ & = -\frac{1}{2\sigma^2} \sum_i e_i^2 \end{aligned}$$

MLE and OLS

$$\operatorname{argmax}_{\mathbf{w}} \operatorname{LogLikelihood}(\text{Data}, \mathbf{w}) = \operatorname{argmin}_{\mathbf{w}} \sum_i e_i^2$$

Linear Regression + MAP

$$\Pr(\mathbf{w}|\text{Data}) \propto \Pr(\text{Data}|\mathbf{w}) \cdot \Pr(\mathbf{w})$$

Linear Regression + MAP

$$\log \Pr(\mathbf{w}|\text{Data}) \propto \log \Pr(\text{Data}|\mathbf{w}) + \log \Pr(\mathbf{w})$$

Linear Regression + MAP

$$\log \Pr(\mathbf{w}|\text{Data}) \propto \log \Pr(\text{Data}|\mathbf{w}) + \log \Pr(\mathbf{w})$$

$$-\sum_i e_i^2$$

Linear Regression + MAP

$$\log \Pr(\mathbf{w}|\text{Data}) \propto \log \Pr(\text{Data}|\mathbf{w}) + \log \Pr(\mathbf{w})$$

$$-\sum_i e_i^2$$

Let the prior on w_j be Gaussian:

$$\Pr(w_j) \propto \exp\left(-\frac{w_j^2}{2\alpha^2}\right)$$

Linear Regression + MAP

$$\log \Pr(\mathbf{w}|\text{Data}) \propto \log \Pr(\text{Data}|\mathbf{w}) + \log \Pr(\mathbf{w})$$

$$-\sum_i e_i^2$$

Let the prior on w_j be Gaussian:

$$\Pr(\mathbf{w}) \propto \prod_j \exp\left(-\frac{w_j^2}{2\alpha^2}\right)$$

Linear Regression + MAP

$$\log \Pr(\mathbf{w}|\text{Data}) \propto \log \Pr(\text{Data}|\mathbf{w}) + \log \Pr(\mathbf{w})$$

$$-\sum_i e_i^2$$

Let the prior on w_j be Gaussian:

$$\log \Pr(\mathbf{w}) \propto \sum_j -\frac{w_j^2}{2\alpha^2}$$

Linear Regression + MAP

$$\log \Pr(\mathbf{w}|\text{Data}) \propto \log \Pr(\text{Data}|\mathbf{w}) + \log \Pr(\mathbf{w})$$

$$-\sum_i e_i^2$$

Let the prior on w_j be Gaussian:

$$\log \Pr(\mathbf{w}) \propto -\sum_j w_j^2$$

Linear Regression + MAP

$$\log \Pr(\mathbf{w}|\text{Data}) \propto \log \Pr(\text{Data}|\mathbf{w}) + \log \Pr(\mathbf{w})$$

$$-\sum_i e_i^2 \quad -\sum_j w_j^2$$

Let the prior on w_j be Gaussian:

$$\log \Pr(\mathbf{w}) \propto -\sum_j w_j^2$$

Linear Regression + MAP

$$\log \Pr(\mathbf{w}|\text{Data}) \propto \log \Pr(\text{Data}|\mathbf{w}) + \log \Pr(\mathbf{w})$$

$$-\sum_i e_i^2$$

$$-\sum_j w_j^2$$

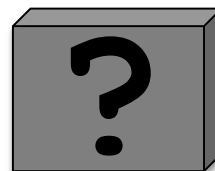
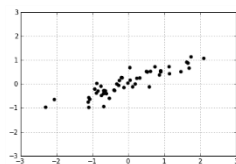
Loss $\Leftrightarrow$ Error distribution

Penalty $\Leftrightarrow$ Model prior

Summary



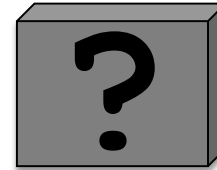
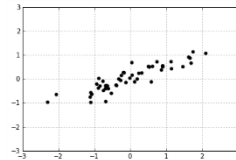
Summary



MLE

MAP

Summary



MLE

MAP

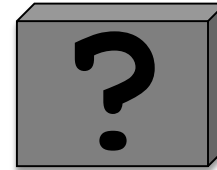
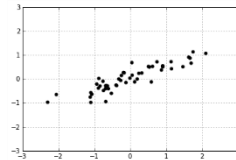
**Loss $\Leftrightarrow$ Error
distribution**

**Loss $\Leftrightarrow$ Error
distribution**

+

Penalty $\Leftrightarrow$ Model prior

Summary



MLE

MAP

**Loss $\Leftrightarrow$ Error
distribution**

**Loss $\Leftrightarrow$ Error
distribution**

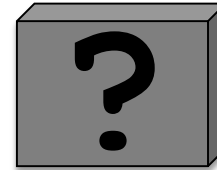
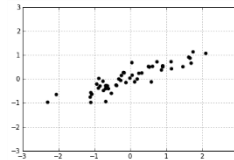
+

Penalty $\Leftrightarrow$ Model prior

OLS Regression

Ridge Regression

Summary



MLE

MAP

**Loss $\Leftrightarrow$ Error
distribution**

**Loss $\Leftrightarrow$ Error
distribution**

+

Penalty $\Leftrightarrow$ Model prior

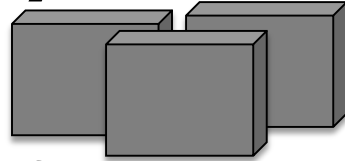
OLS Regression

Ridge Regression

ℓ_2 -loss/penalty $\Leftrightarrow$ Normal distribution

Summary

- ▶ **Probability** for modeling



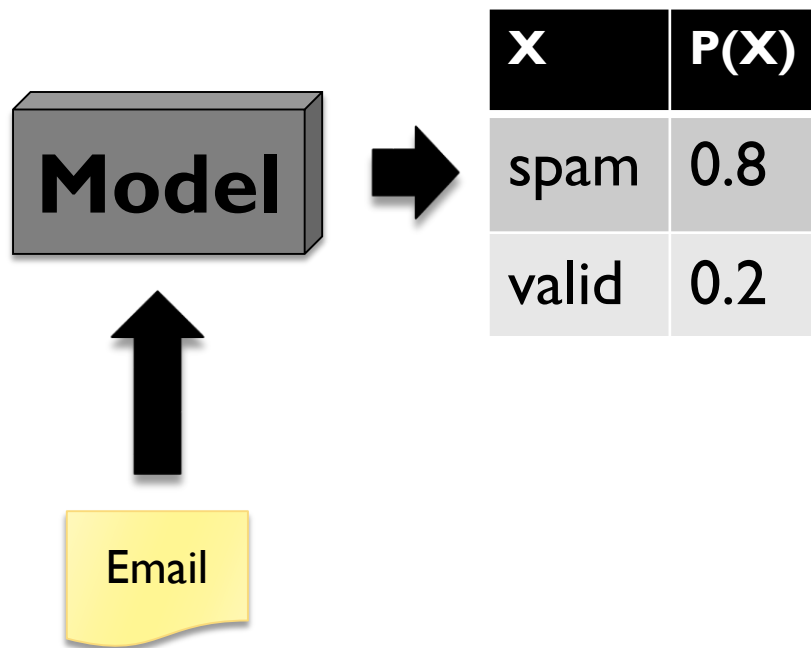
- ▶ **Statistics** for estimation

MLE

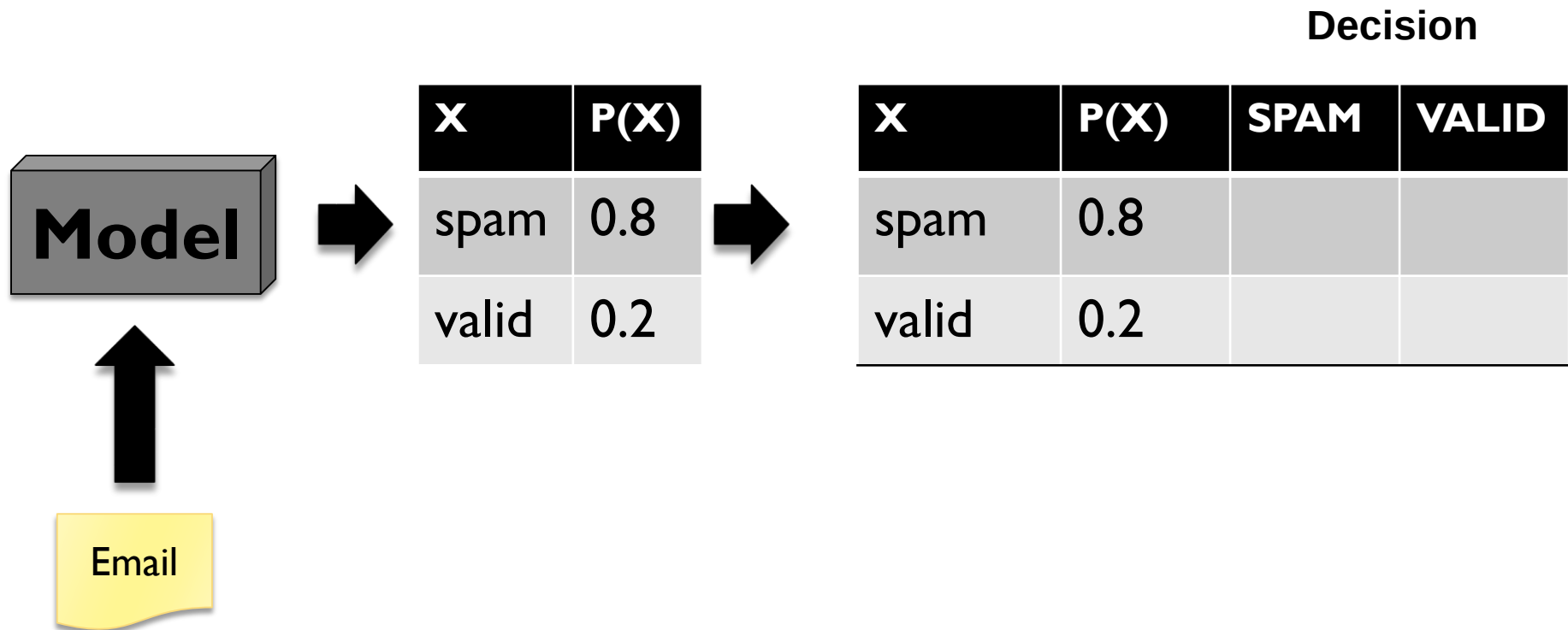
MAP

- ▶ **Decision theory** for prediction

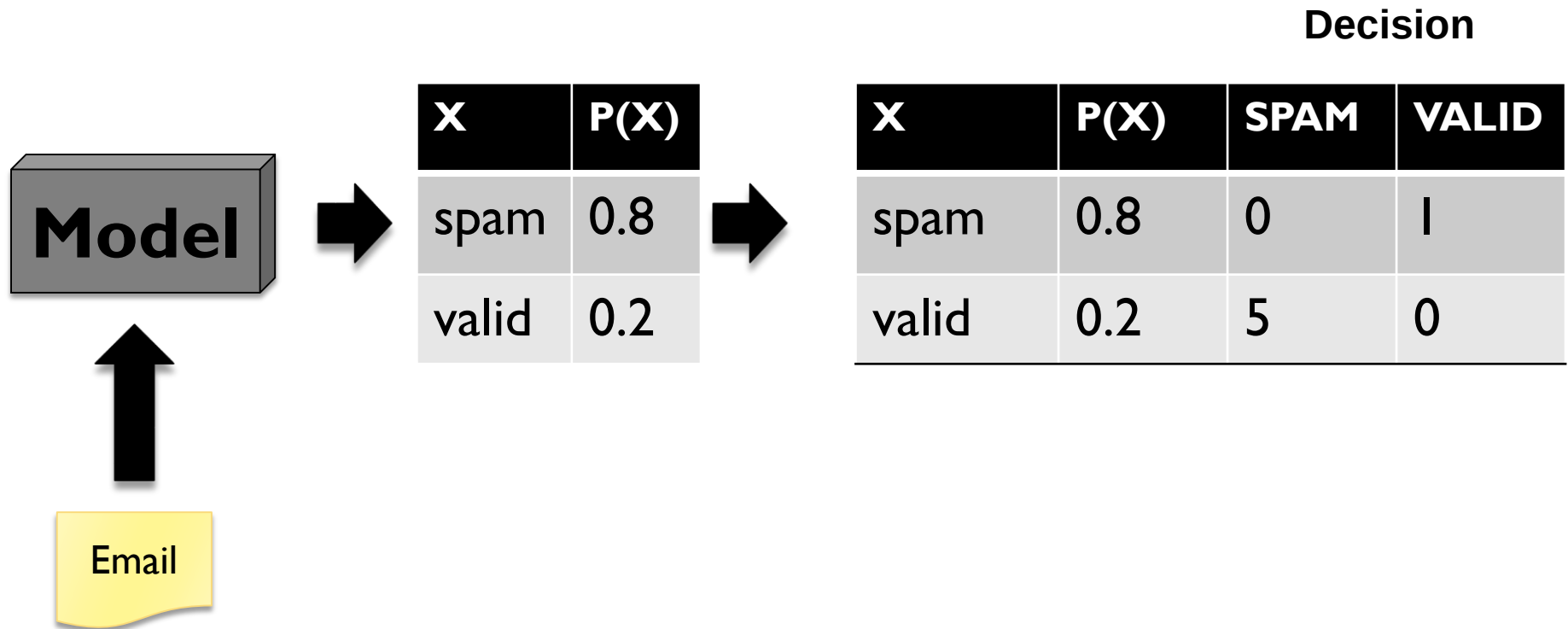
Decision Theory



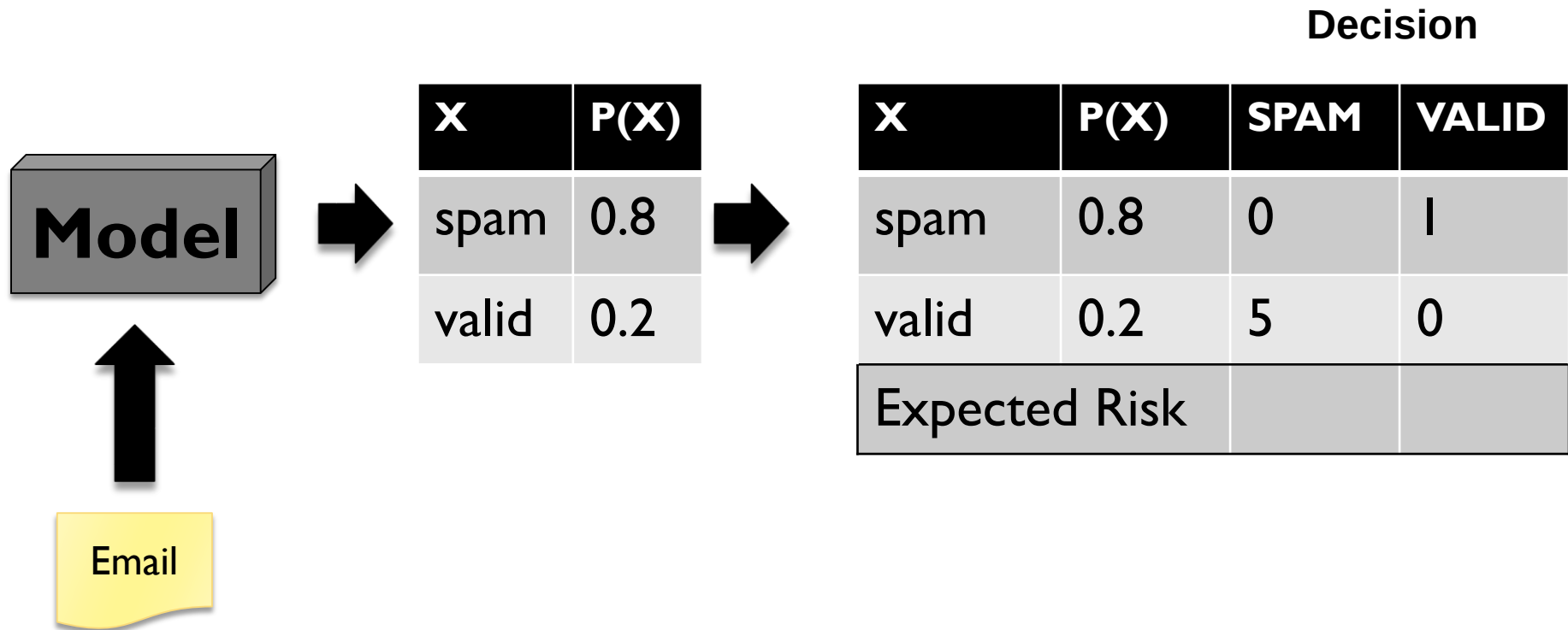
Decision Theory



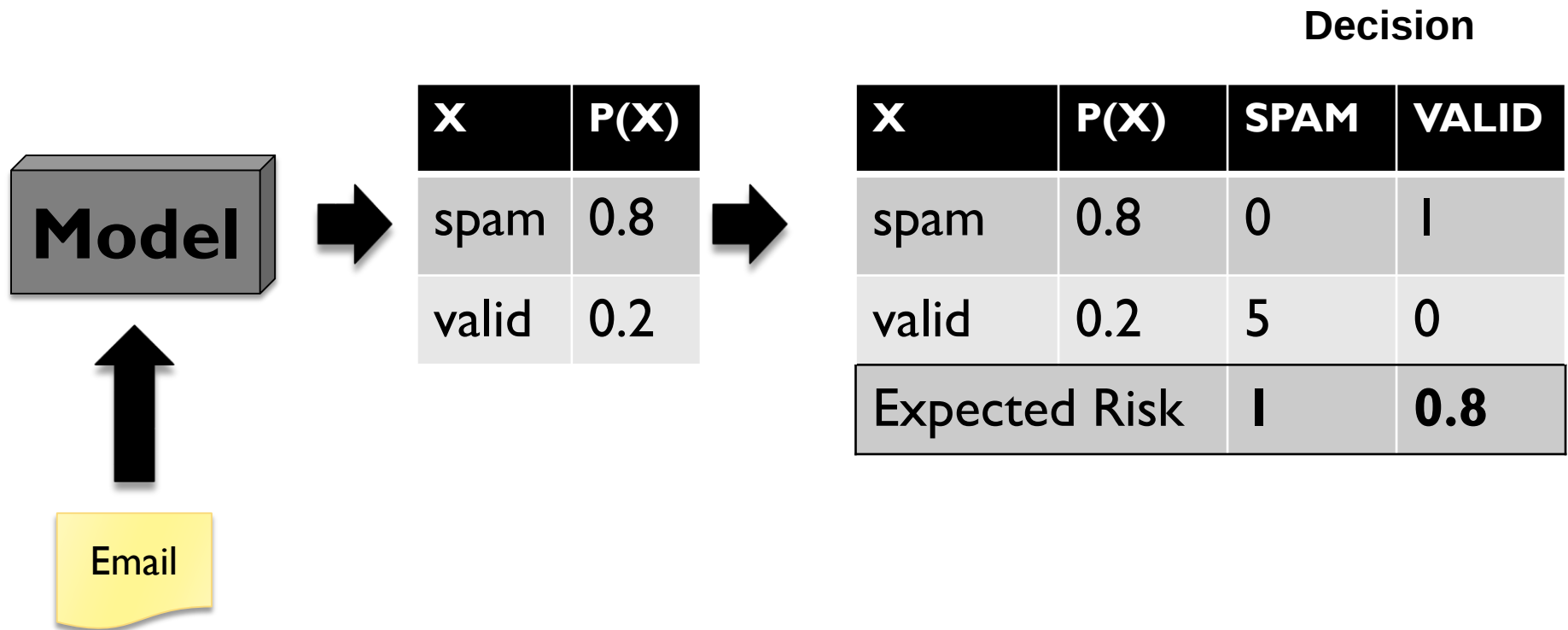
Decision Theory



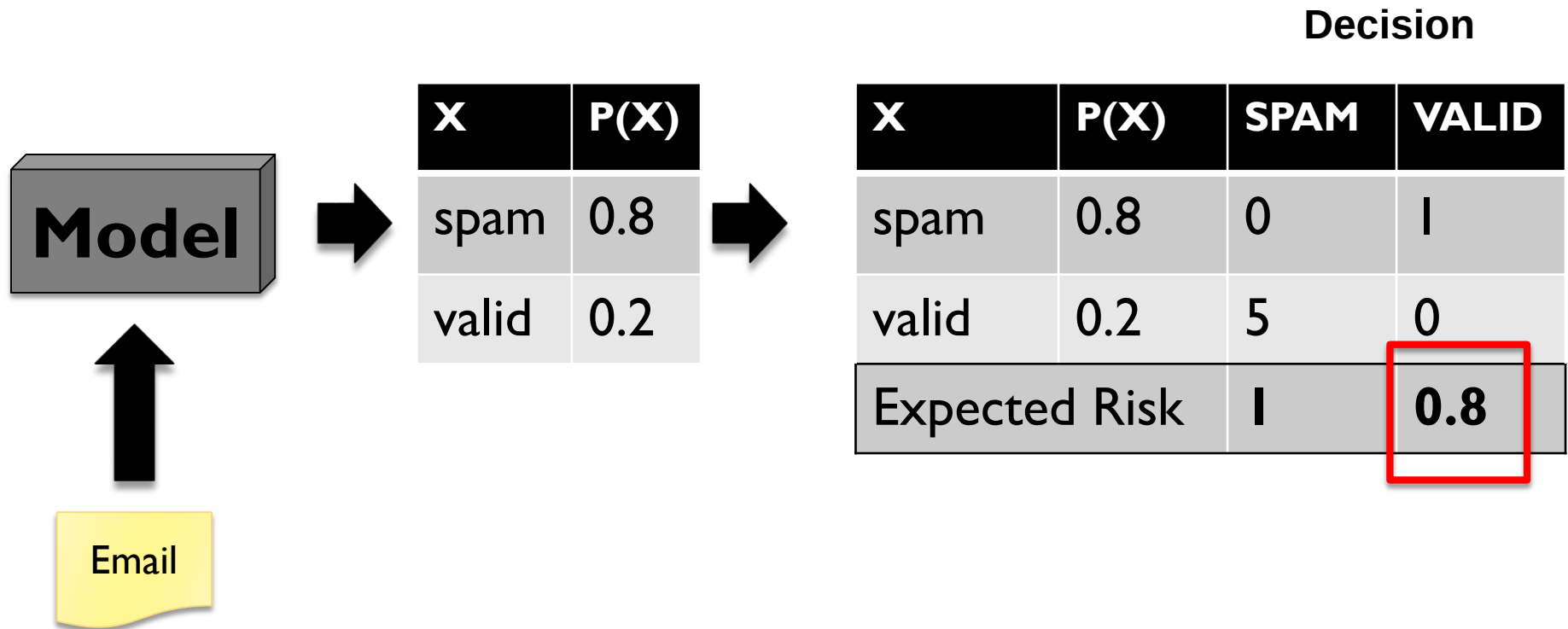
Decision Theory



Decision Theory



Decision Theory



Expected Risk (Supervised Learning)

$$R(\hat{y}|x) = \sum_y \Pr(y|x) \ell(\hat{y}, y)$$

Expected Risk (Supervised Learning)

$$R(\hat{v} | \mathcal{X}) = \sum_y \Pr(y | \mathcal{X}) \ell(\hat{v}, y)$$

Diagram illustrating the Expected Risk formula with visual components:

- SPAM, VALID** (Yellow box) and **Email** (Yellow box) represent the input data $\mathcal{X}$.
- Model** (Grey box) represents the function $\ell(\hat{v}, y)$.
- spam, valid** (Blue box) represents the output $\hat{v}$.
- The summation $\sum_y$ is over the possible outputs y .
- The probability $\Pr(y | \mathcal{X})$ is shown as a matrix:

0	1
5	0

Expected Risk (Supervised Learning)

$$R(\hat{y}|x) = \sum_y \Pr(y|x) \ell(\hat{y}, y)$$

Diagram illustrating the components of the Expected Risk formula:

- SPAM, VALID** (Yellow box) and **Email** (Yellow box) represent the input x .
- Model** (Gray box) represents the function $\hat{y}$.
- spam, valid** (Blue box) represents the output $\hat{y}$.
- 0 1** and **5 0** (Brown box) represent the loss function $\ell(\hat{y}, y)$.

$$R(\hat{y}|x) = \int_y \ell(\hat{y}, y) dF(y|x)$$

Bayesian Classifier

► Optimal classifier:

For a given x and a conditional probabilistic model $\Pr(y|x)$
predict $\hat{y}$, that has the **smallest expected risk**.

Bayesian Classifier

► Optimal classifier:

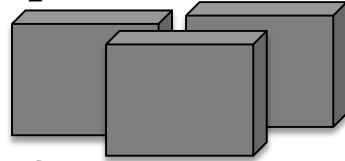
For a given x and a conditional probabilistic model $\Pr(y|x)$
predict $\hat{y}$, that has the **smallest expected risk**.

For *symmetric risk* ℓ , this corresponds to picking the option with the highest probability.

0	1
1	0

Summary

- ▶ **Probability** for modeling



- ▶ **Statistics** for estimation

MLE

MAP

- ▶ **Decision theory** for prediction

Bayesian Classifier



The Tennis Dataset

Day	Outlook	Temp	Humidity	Wind	PlayTennis
D1	Sunny	Hot	High	Weak	No
D2	Sunny	Hot	High	Strong	No
D3	Overcast	Hot	High	Weak	Yes
D4	Rain	Mild	High	Weak	Yes
D5	Rain	Cool	Normal	Weak	Yes
D6	Rain	Cool	Normal	Strong	No
D7	Overcast	Cool	Normal	Strong	Yes
D8	Sunny	Mild	High	Weak	No
D9	Sunny	Cool	Normal	Weak	Yes
D10	Rain	Mild	Normal	Weak	Yes
D11	Sunny	Mild	Normal	Strong	Yes
D12	Overcast	Mild	High	Strong	Yes
D13	Overcast	Hot	Normal	Weak	Yes
D14	Rain	Mild	High	Strong	No

Shall we play tennis today?

<i>PlayTennis</i>
<i>No</i>
<i>No</i>
<i>Yes</i>
<i>Yes</i>
<i>Yes</i>
<i>No</i>
<i>Yes</i>
<i>No</i>
<i>Yes</i>
<i>Yes</i>
<i>Yes</i>
<i>Yes</i>
<i>Yes</i>
<i>No</i>

Shall we play tennis today?

<i>PlayTennis</i>
<i>No</i>
<i>No</i>
<i>Yes</i>
<i>Yes</i>
<i>Yes</i>
<i>No</i>
<i>Yes</i>
<i>No</i>
<i>Yes</i>
<i>Yes</i>
<i>Yes</i>
<i>Yes</i>
<i>Yes</i>
<i>No</i>

Estimate a
probabilistic model
and predict:

$$\Pr(\text{Yes}) = 9/14 = 0.64$$

$$\Pr(\text{No}) = 5/14 = 0.36$$

➔ Yes

It's windy today. Tennis, anyone?

Wind	<i>PlayTennis</i>
<i>Weak</i>	<i>No</i>
<i>Strong</i>	<i>No</i>
<i>Weak</i>	<i>Yes</i>
<i>Weak</i>	<i>Yes</i>
<i>Weak</i>	<i>Yes</i>
<i>Strong</i>	<i>No</i>
<i>Strong</i>	<i>Yes</i>
<i>Weak</i>	<i>No</i>
<i>Weak</i>	<i>Yes</i>
<i>Weak</i>	<i>Yes</i>
<i>Strong</i>	<i>Yes</i>
<i>Strong</i>	<i>Yes</i>
<i>Weak</i>	<i>Yes</i>
<i>Strong</i>	<i>No</i>

It's windy today. Tennis, anyone?

Wind	PlayTennis
Weak	No
Strong	No
Weak	Yes
Weak	Yes
Weak	Yes
Strong	No
Strong	Yes
Weak	No
Weak	Yes
Weak	Yes
Strong	Yes
Strong	Yes
Weak	Yes
Strong	No

$$\Pr(\text{Yes} \mid \text{Weak}) = 6/8$$

$$\Pr(\text{No} \mid \text{Weak}) = 2/8$$

$$\Pr(\text{Yes} \mid \text{Strong}) = 3/6$$

$$\Pr(\text{No} \mid \text{Strong}) = 3/6$$

More attributes

Humidity	Wind	PlayTennis
High	Weak	No
High	Strong	No
High	Weak	Yes
High	Weak	Yes
Normal	Weak	Yes
Normal	Strong	No
Normal	Strong	Yes
High	Weak	No
Normal	Weak	Yes
Normal	Weak	Yes
Normal	Strong	Yes
High	Strong	Yes
Normal	Weak	Yes
High	Strong	No

$$\Pr(\text{Yes} \mid \text{High, Weak}) = 2/4$$

$$\Pr(\text{No} \mid \text{High, Weak}) = 2/4$$

$$\Pr(\text{Yes} \mid \text{High, Strong}) = 1/3$$

$$\Pr(\text{No} \mid \text{High, Strong}) = 2/3$$

...

The Bayesian Classifier

In general:

1. **Estimate from data:**

$$\Pr(\text{Class} | x_1, x_2, x_3, \dots)$$

2. **For a given instance** $(x_1, x_2, x_3, \dots)$
predict class whose conditional
probability is greater:

$$\Pr(C_1 | x_1, x_2, x_3, \dots) > \Pr(C_2 | x_1, x_2, x_3, \dots)$$

➔ predict C_1

Problem

► We need exponential amount of data

Humidity	Wind	PlayTennis
High	Weak	No
High	Strong	No
High	Weak	Yes
High	Weak	Yes
Normal	Weak	Yes
Normal	Strong	No
Normal	Strong	Yes
High	Weak	No
Normal	Weak	Yes
Normal	Weak	Yes
Normal	Strong	Yes
High	Strong	Yes
Normal	Weak	Yes
High	Strong	No

$$\Pr(\text{Yes} \mid \text{High, Weak}) = 2/4$$

$$\Pr(\text{No} \mid \text{High, Weak}) = 2/4$$

$$\Pr(\text{Yes} \mid \text{High, Strong}) = 1/3$$

$$\Pr(\text{No} \mid \text{High, Strong}) = 2/3$$

...

Naïve Bayes Classifier

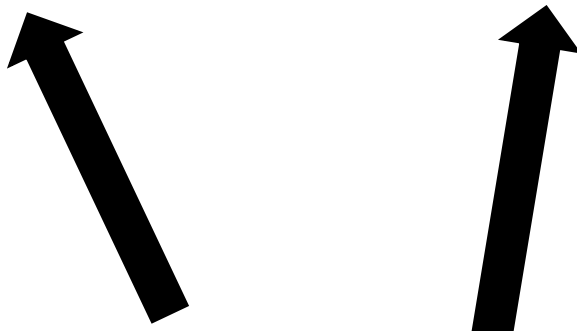
To scale beyond 2-3 attributes, use a hack:

Assume that attributes are independent within each class:

$$\begin{aligned} \Pr(x_1, x_2, x_3 \mid \text{Class}) \\ = \Pr(x_1 \mid \text{Class}) \Pr(x_2 \mid \text{Class}) \Pr(x_3 \mid \text{Class}) \dots \end{aligned}$$

Naïve Bayes Classifier

1. $\Pr(C_1 | \mathbf{x}) > \Pr(C_2 | \mathbf{x})$
 $\rightarrow$ predict C_1


$$P(A|B) = \frac{P(B|A) P(A)}{P(B)}$$

Naïve Bayes Classifier

1. $\Pr(C_1 | \mathbf{x}) > \Pr(C_2 | \mathbf{x})$
 $\rightarrow$ predict C_1

2. $\frac{\Pr(C_1) \Pr(\mathbf{x} | C_1)}{\Pr(\mathbf{x})} > \frac{\Pr(C_2) \Pr(\mathbf{x} | C_2)}{\Pr(\mathbf{x})}$
 $\rightarrow$ predict C_1

Naïve Bayes Classifier

1. $\Pr(C_1 | \mathbf{x}) > \Pr(C_2 | \mathbf{x})$

→ predict C_1

2. $\frac{\Pr(C_1) \Pr(\mathbf{x} | C_1)}{\Pr(\mathbf{x})} > \frac{\Pr(C_2) \Pr(\mathbf{x} | C_2)}{\Pr(\mathbf{x})}$

→ predict C_1

3. $\Pr(C_1) \Pr(\mathbf{x} | C_1) > \Pr(C_2) \Pr(\mathbf{x} | C_2)$

→ predict C_1

Naïve Bayes Classifier

1. $\Pr(C_1 | \mathbf{x}) > \Pr(C_2 | \mathbf{x})$

→ predict C_1

2. $\frac{\Pr(C_1) \Pr(\mathbf{x} | C_1)}{\Pr(\mathbf{x})} > \frac{\Pr(C_2) \Pr(\mathbf{x} | C_2)}{\Pr(\mathbf{x})}$

→ predict C_1

3. $\Pr(C_1) \Pr(\mathbf{x} | C_1) > \Pr(C_2) \Pr(\mathbf{x} | C_2)$

→ predict C_1

4. $\Pr(C_1) \cdot \Pr(x_1 | C_1) \Pr(x_2 | C_1) \dots \Pr(x_m | C_1) >$
 $\Pr(C_2) \cdot \Pr(x_1 | C_2) \Pr(x_2 | C_2) \dots \Pr(x_m | C_2)$

→ predict C_1

Naïve Bayes Classifier

1. $\Pr(C_1 | \mathbf{x}) > \Pr(C_2 | \mathbf{x})$
 $\rightarrow$ predict C_1

2. $\frac{\Pr(C_1) \Pr(\mathbf{x} | C_1)}{\Pr(\mathbf{x})} > \frac{\Pr(C_2) \Pr(\mathbf{x} | C_2)}{\Pr(\mathbf{x})}$
 $\rightarrow$ predict C_1

3. $\Pr(C_1) \Pr(\mathbf{x} | C_1) > \Pr(C_2) \Pr(\mathbf{x} | C_2)$
 $\rightarrow$ predict C_1

4. $\frac{\Pr(C_1) \cdot \Pr(x_1 | C_1) \Pr(x_2 | C_1) \dots \Pr(x_m | C_1)}{\Pr(C_2) \cdot \Pr(x_1 | C_2) \Pr(x_2 | C_2) \dots \Pr(x_m | C_2)} >$
 $\rightarrow$ predict C_1

Naïve Bayes Classifier

- ▶ Works for both discrete and continuous attributes.

- ▶ The goods:
 - ▶ Easy to implement, efficient
 - ▶ Won't overfit, interpretable
 - ▶ Works better than you would expect (e.g. spam filtering)

- ▶ The bads
 - ▶ “Naïve”, linear
 - ▶ Usually won't work well for too many classes
 - ▶ Not a good probability estimator

Naïve Bayes Classifier

```
from sklearn.naive_bayes import  
    BernoulliNB,  
    MultinomialNB,  
    GaussianNB
```

Quiz

► MLE:

$\operatorname{argmax}_{\text{Model}} \underline{\hspace{2cm}}$

► MAP:

$\operatorname{argmax}_{\text{Model}} \underline{\hspace{2cm}}$

► Gaussian distribution:

$$f(x) = \text{const} \times \exp(\underline{\hspace{2cm}})$$

Quiz

- ▶ Bayesian classifier has optimal _____

- ▶ Naïve Bayesian classifier assumption:
 - ▶ $\Pr(C|x_1, x_2) = \Pr(C|x_1) \Pr(C|x_2)$
 - ▶ $\Pr(x_1, x_2|C) = \Pr(x_1|C) \Pr(x_2|C)$
 - ▶ $\Pr(C_1, C_2|x) = \Pr(C_1|x) \Pr(C_2|x)$
 - ▶ $\Pr(x|C_1, C_2) = \Pr(x|C_1) \Pr(x|C_2)$

-
- ▶ All machine learning methods we have considered so far rely on MLE or MAP
 - ▶ Yes
 - ▶ No

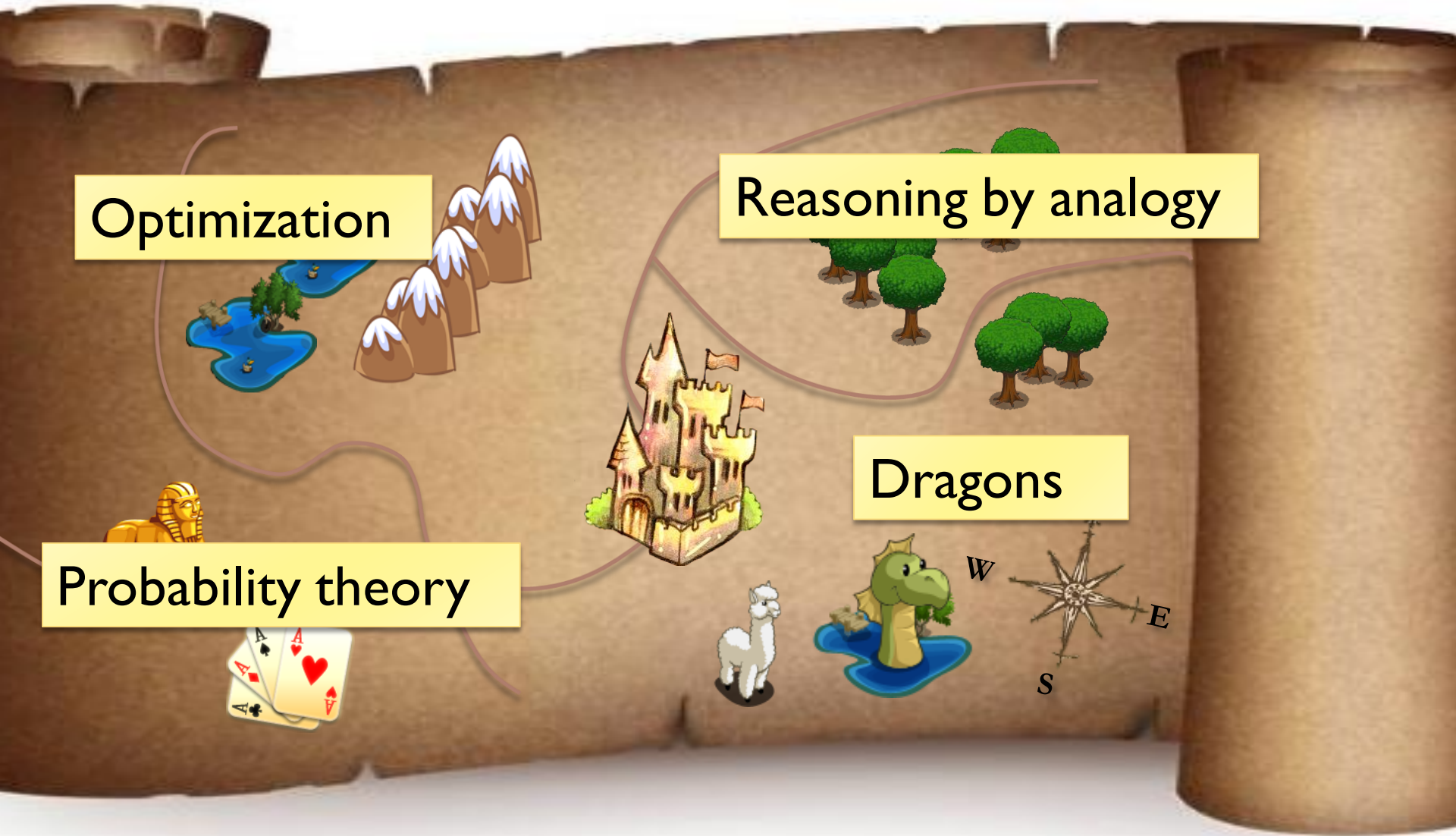
The Land of Machine Learning

Optimization

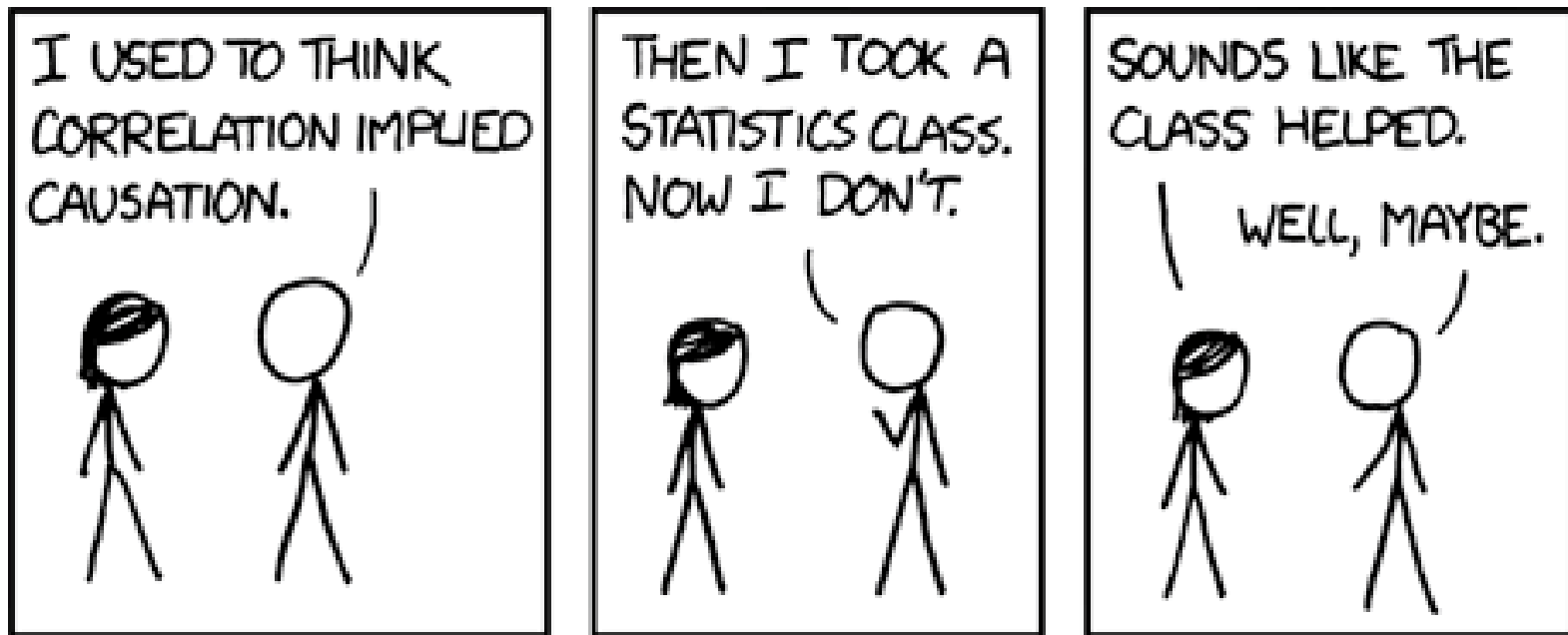
Reasoning by analogy

Probability theory

Dragons



Questions?



<http://xkcd.com/552/>