

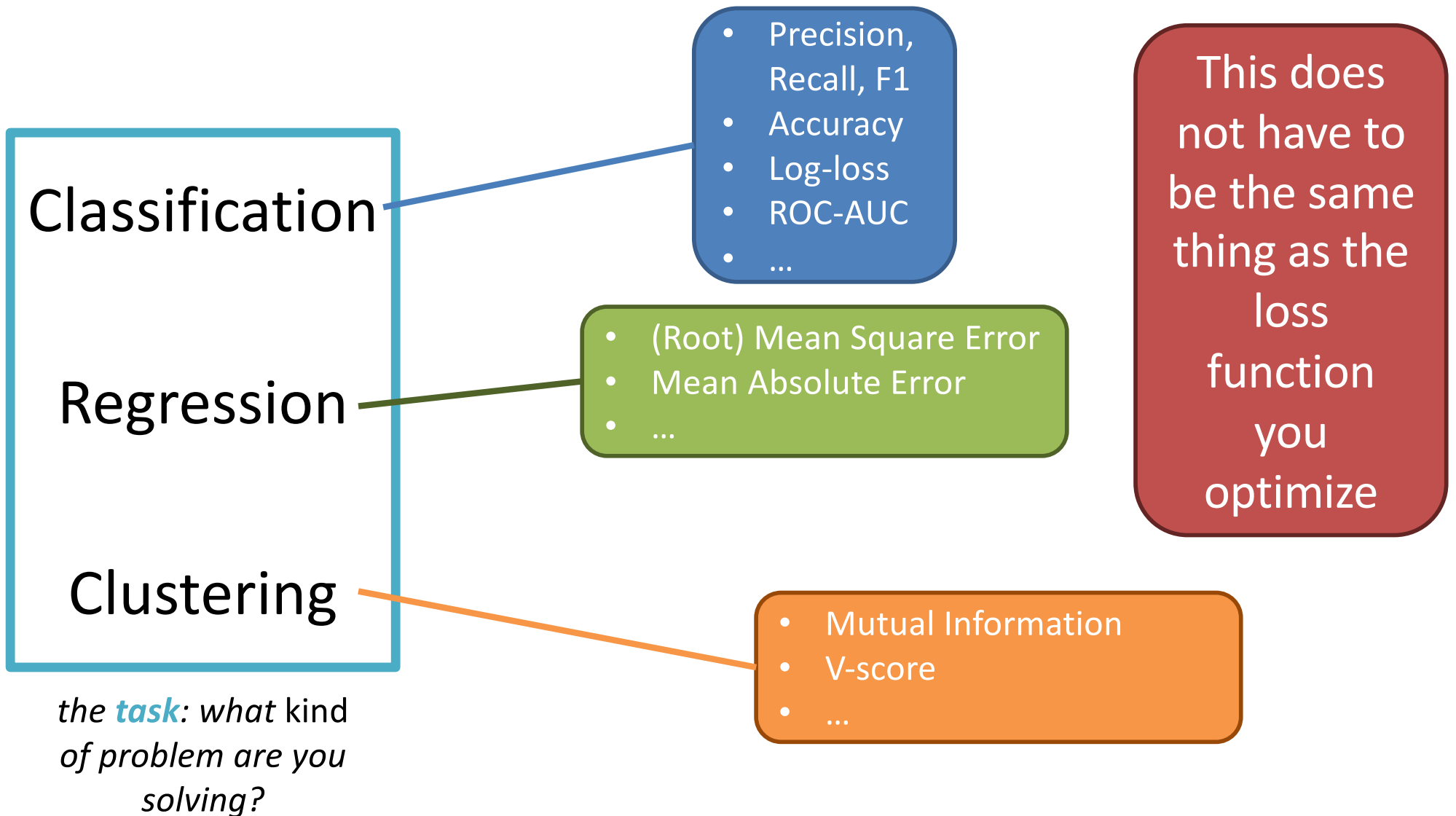
Maximum Entropy Models/ Logistic Regression

CMSC 678

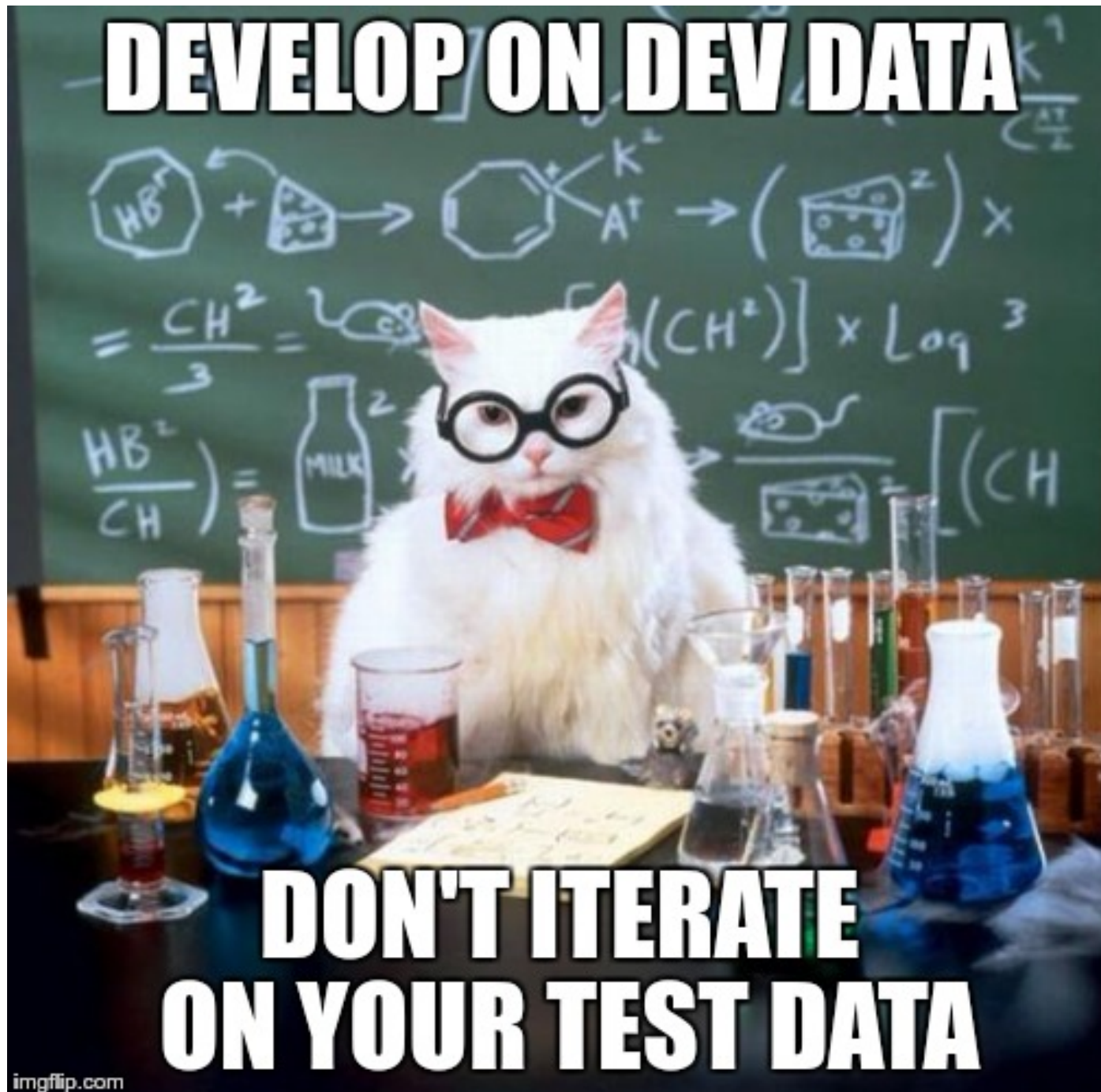
UMBC

Recap from last time...

Central Question: How Well Are We Doing?



Rule #1



We've only developed binary classifiers so far...

Option 1: Develop a multi-class version

Option 2: Build a one-vs-all (OvA) classifier

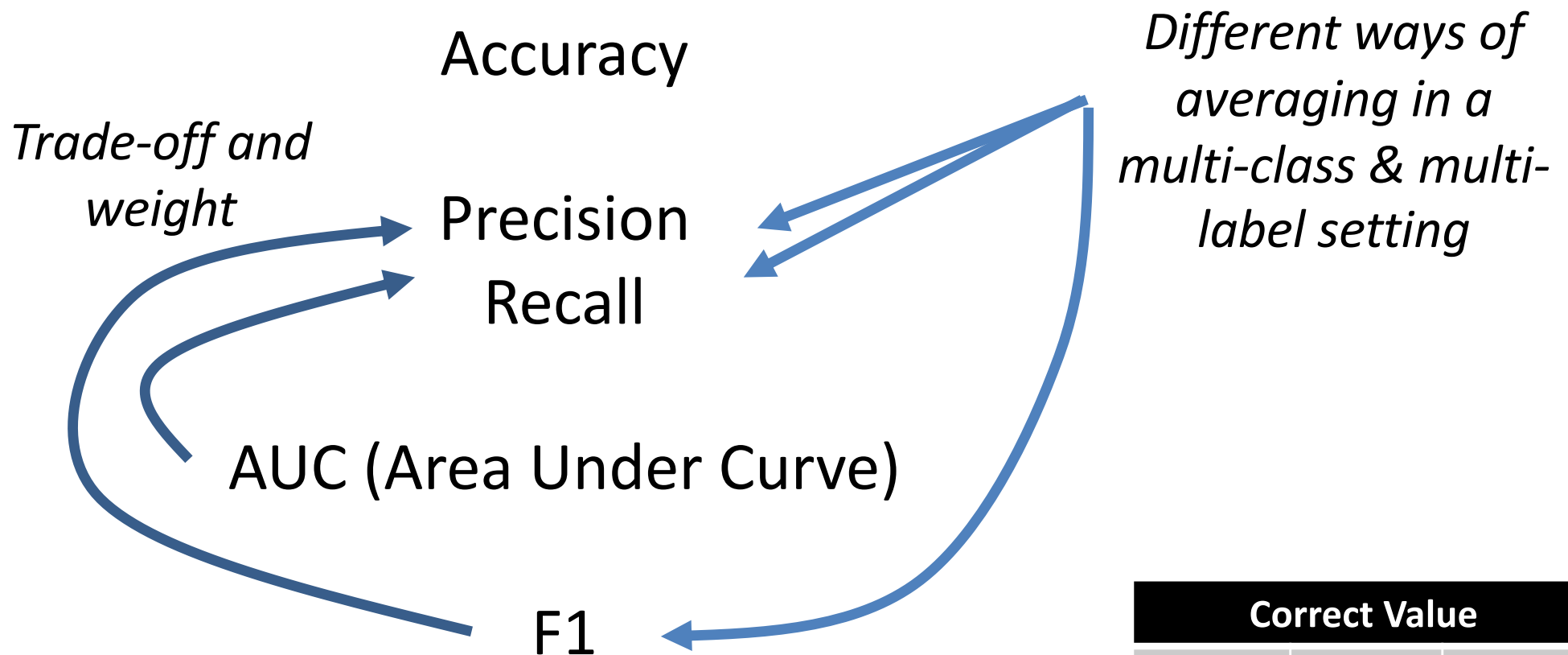
Option 3: Build an all-vs-all (AvA) classifier

(there can be others)

Which option you choose is problem-dependent:

1. Why might you want to use option 1 or options OvA/AvA?
2. What are the benefits of OvA vs. AvA?
3. What if you start with a balanced dataset, e.g., 100 instances per class?

Some Classification Metrics



Confusion Matrix

		Correct Value		
				
Guesse d Value		#	#	#
		#	#	#
		#	#	#

Outline

Log-Linear (Maximum Entropy) Models

Basic Modeling

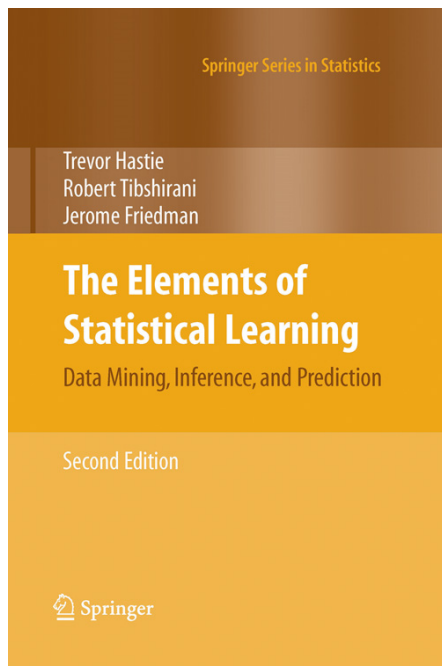
Connections to other techniques (“... *by any other name...*”)

Objective to optimize

Regularization

Maximum Entropy (Log-linear) Models

$$p(y | x) \propto \exp(\theta^T f(x, y))$$

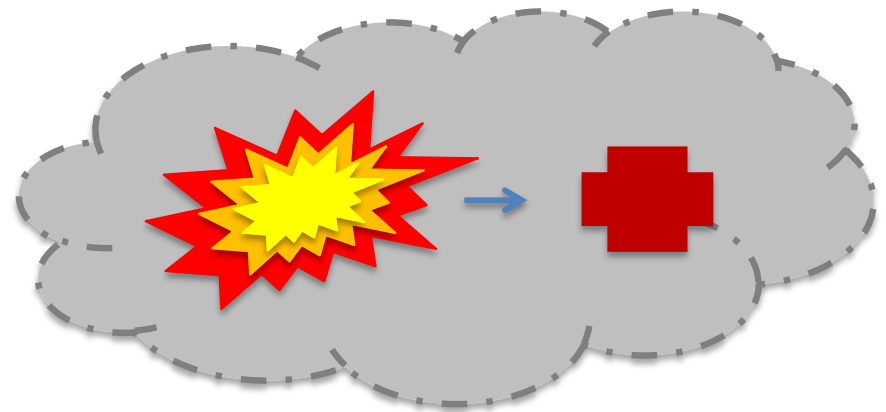


“model the posterior probabilities of the K classes via linear functions in θ , while at the same time ensuring that they sum to one and remain in $[0, 1]$ ” ~ *Ch 4.4*

Document Classification

Three people have been fatally shot, and five people, including a mayor, were seriously wounded as a result of a Shining Path attack today against a community in Junin department, central Peruvian mountain region.

ATTACK



Observed document

Label

Q: What *features* of this document could indicate an **ATTACK**?

Document Classification



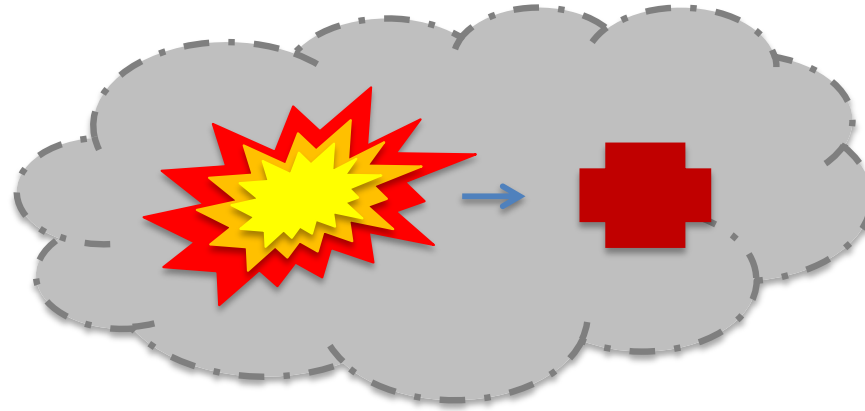
Three people have been fatally shot, and five people, including a mayor, were seriously wounded as a result of a Shining Path attack today against a community in Junin department, central Peruvian mountain region.

attack

ATTACK

- # killed:
- Type:
- Perp:

Document Classification



Three people have been fatally shot, and five people, including a mayor, were seriously wounded as a result of a Shining Path attack today against a community in **Junin** department, central Peruvian mountain region.

ATTACK

there could be many relevant clues

Features

The “clues” that help our system make its decision

Apply a vector of features

$f(\text{document}, y) = (f_1(\text{document}, y), \dots, f_K(\text{document}, y))$
to a given document  and possible label y

$f_{\text{fatally shot, ATTACK}}(\text{document}, \text{ATTACK})$
 $f_{\text{seriously wounded, ATTACK}}(\text{document}, \text{ATTACK})$
 $f_{\text{Shining Path, ATTACK}}(\text{document}, \text{ATTACK})$
 $f_{\text{happy cat, ATTACK}}(\text{document}, \text{ATTACK})$
...

Features

The “clues” that help our system make its decision

Apply a vector of features

$f(\text{document}, y) = (f_1(\text{document}, y), \dots, f_K(\text{document}, y))$
to a given document  and possible label y

Each feature function f_k can take any real value:

binary

count-based

likelihood

$f_{\text{fatally shot, ATTACK}}(\text{document}, \text{ATTACK})$
 $f_{\text{seriously wounded, ATTACK}}(\text{document}, \text{ATTACK})$
 $f_{\text{Shining Path, ATTACK}}(\text{document}, \text{ATTACK})$
 $f_{\text{happy cat, ATTACK}}(\text{document}, \text{ATTACK})$
...

Features

The “clues” that help our system make its decision

Apply a vector of features

$f(\text{document}, y) = (f_1(\text{document}, y), \dots, f_K(\text{document}, y))$ to a given document  and possible label y

Each feature function f_k can take any real value:

binary

count-based

likelihood

Features that don't “fire” don't apply to the pair

$$f_k(\text{document}, y) = 0$$

$f_{\text{fatally shot, ATTACK}}(\text{document}, \text{ATTACK})$

$f_{\text{seriously wounded, ATTACK}}(\text{document}, \text{ATTACK})$

$f_{\text{Shining Path, ATTACK}}(\text{document}, \text{ATTACK})$

$f_{\text{happy cat, ATTACK}}(\text{document}, \text{ATTACK})$

...

Features:

Score and Combine Our Possibilities

define for
each key
phrase/
clue...

$$\theta_{\text{fatally shot}, \text{ATTACK}}(\text{document}, \text{ATTACK})$$

$$\theta_{\text{seriously wounded}, \text{ATTACK}}(\text{document}, \text{ATTACK})$$

$$\theta_{\text{Shining Path}, \text{ATTACK}}(\text{document}, \text{ATTACK})$$

$$\theta_{\text{happy cat}, \text{ATTACK}}(\text{document}, \text{ATTACK})$$

...

Remember: each
 $\theta_{w, l}(\text{document}, y)$ is actually
computed as
 $\theta_{w, l} * f_{w, l}(\text{document}, y)$

Features:

Score and Combine Our Possibilities

... and for each label

define for
each key
phrase/
clue...

$$\begin{array}{ll} \theta_{\text{fatally shot}, \text{ATTACK}}(\text{document}, \text{ATTACK}) & \theta_{\text{fatally shot}, \text{TECH}}(\text{document}, \text{ATTACK}) \\ \theta_{\text{seriously wounded}, \text{ATTACK}}(\text{document}, \text{ATTACK}) & \theta_{\text{seriously wounded}, \text{TECH}}(\text{document}, \text{ATTACK}) \\ \theta_{\text{Shining Path}, \text{ATTACK}}(\text{document}, \text{ATTACK}) & \theta_{\text{Shining Path}, \text{TECH}}(\text{document}, \text{ATTACK}) \\ \theta_{\text{happy cat}, \text{ATTACK}}(\text{document}, \text{ATTACK}) & \theta_{\text{happy cat}, \text{TECH}}(\text{document}, \text{ATTACK}) \\ \dots & \dots \end{array}$$

Remember: each
 $\theta_{w, l}(\text{document}, y)$ is actually
computed as
 $\theta_{w, l} * f_{w, l}(\text{document}, y)$

Features:

Score and Combine Our Possibilities

... and for each label

*define for
each key
phrase/
clue...*

$\theta_{\text{fatally shot, ATTACK}}(\text{document}, \text{ATTACK})$	$\theta_{\text{fatally shot, TECH}}(\text{document}, \text{ATTACK})$
$\theta_{\text{seriously wounded, ATTACK}}(\text{document}, \text{ATTACK})$	$\theta_{\text{seriously wounded, TECH}}(\text{document}, \text{ATTACK})$
$\theta_{\text{Shining Path, ATTACK}}(\text{document}, \text{ATTACK})$	$\theta_{\text{Shining Path, TECH}}(\text{document}, \text{ATTACK})$
$\theta_{\text{happy cat, ATTACK}}(\text{document}, \text{ATTACK})$	$\theta_{\text{happy cat, TECH}}(\text{document}, \text{ATTACK})$
...	...

*Remember: each
 $\theta_{w, l}(\text{document}, y)$ is actually
computed as
 $\theta_{w, l} * f_{w, l}(\text{document}, y)$*

Not all of these will be relevant

Features:

Score and Combine Our Possibilities

... and for each label

define for
each key
phrase/
clue...

$\theta_{\text{fatally shot, ATTACK}}(\text{document}, \text{ATTACK})$

$\theta_{\text{fatally shot, TECH}}(\text{document}, \text{ATTACK})$

$\theta_{\text{seriously wounded, ATTACK}}(\text{document}, \text{ATTACK})$

$\theta_{\text{seriously wounded, TECH}}(\text{document}, \text{ATTACK})$

$\theta_{\text{Shining Path, ATTACK}}(\text{document}, \text{ATTACK})$

$\theta_{\text{Shining Path, TECH}}(\text{document}, \text{ATTACK})$

$\theta_{\text{happy cat, ATTACK}}(\text{document}, \text{ATTACK})$

$\theta_{\text{happy cat, TECH}}(\text{document}, \text{ATTACK})$

...

...

Remember: each
 $\theta_{w, l}(\text{document}, y)$ is actually
computed as
 $\theta_{w, l} * f_{w, l}(\text{document}, y)$

Each of these scored features describes how “good” a
particular phrase is for a given document type if the
provided document document document has a proposed type

Score and Combine Our Possibilities

*Shortcut notation: focus only
on the features that “fire”*

θ_1 (fatally shot, ATTACK)

θ_2 (seriously wounded, ATTACK)

θ_3 (Shining Path, ATTACK)

...



Weight each of these: score
how “important” each feature
(clue) is

Q: How many features
are there?

A: As many as you want
there to be (but be
careful of
underfitting/overfitting)

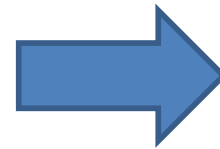
Score and Combine Our Possibilities

θ_1 (fatally shot, ATTACK)

θ_2 (seriously wounded, ATTACK)

θ_3 (Shining Path, ATTACK)

COMBINE



posterior
probability of
ATTACK

...

Weight each of these: score
how “important” each feature
(clue) is

Scoring Our Possibilities

score(

Three people have been fatally shot, and five people, including a mayor, were seriously wounded as a result of a Shining Path attack today against a community in Junin department, central Peruvian mountain region .

, ATTACK) =

$\theta_1(\text{fatally shot, ATTACK})$ +

$\theta_2(\text{seriously wounded, ATTACK})$ +

$\theta_3(\text{Shining Path, ATTACK})$ +

...

our linear regression model

Maxent Modeling

$$p(\text{ATTACK} \mid \text{Three people have been fatally shot, and five people, including a mayor, were seriously wounded as a result of a Shining Path attack today against a community in Junin department, central Peruvian mountain region .}) \propto$$

$$\text{SNAP}(\text{score}(\text{Three people have been fatally shot, and five people, including a mayor, were seriously wounded as a result of a Shining Path attack today against a community in Junin department, central Peruvian mountain region .}, \text{ATTACK}))$$

What function...

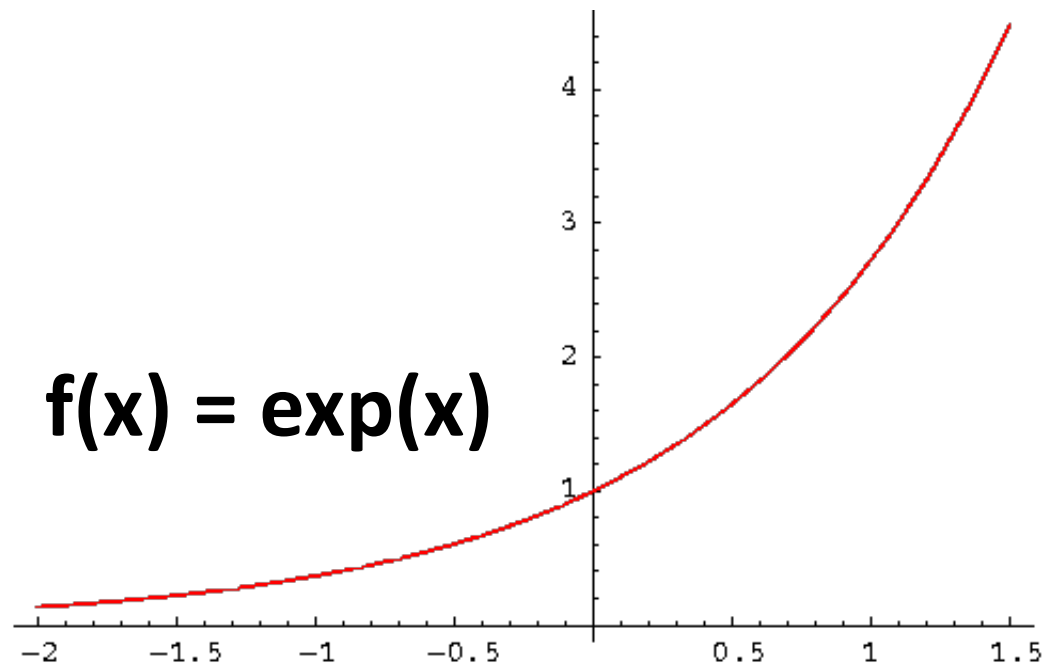
operates on any real number?

is never less than 0?

What function...

operates on any real number?

is never less than 0?



Maxent Modeling

$$p(\text{ATTACK} \mid \text{Three people have been fatally shot, and five people, including a mayor, were seriously wounded as a result of a Shining Path attack today against a community in Junin department, central Peruvian mountain region .}) \propto$$

$$\exp(\text{score}(\text{Three people have been fatally shot, and five people, including a mayor, were seriously wounded as a result of a Shining Path attack today against a community in Junin department, central Peruvian mountain region .}, \text{ATTACK}))$$

Maxent Modeling

$$p(\text{ATTACK} \mid \text{Three people have been fatally shot, and five people, including a mayor, were seriously wounded as a result of a Shining Path attack today against a community in Junin department, central Peruvian mountain region .}) \propto$$

$$\exp(\theta_1(\text{fatally shot, ATTACK}) + \theta_2(\text{seriously wounded, ATTACK}) + \theta_3(\text{Shining Path, ATTACK}) + \dots)$$

*this is assuming binary features, but
they don't have to be*

Maxent Modeling

$$p(\text{ATTACK} \mid \text{Three people have been fatally shot, and five people, including a mayor, were seriously wounded as a result of a Shining Path attack today against a community in Junin department, central Peruvian mountain region .}) \propto$$

$$\exp(\text{weight}_1 * f_1(\text{fatally shot, ATTACK}) + \text{weight}_2 * f_2(\text{seriously wounded, ATTACK}) + \text{weight}_3 * f_3(\text{Shining Path, ATTACK}) + \dots)$$

Maxent Modeling

Three people have been fatally shot, and five people, including a mayor, were seriously wounded as a result of a Shining Path attack today against a community in Junin department, central Peruvian mountain region .

Q: How do we define Z?

$$\frac{1}{Z} \exp \left(\begin{aligned} & \text{weight}_1 * f_1(\text{fatally shot, ATTACK}) \\ & \text{weight}_2 * f_2(\text{seriously wounded, ATTACK}) \\ & \text{weight}_3 * f_3(\text{Shining Path, ATTACK}) \\ & \dots \end{aligned} \right)$$

Normalization for Classification

$$Z =$$

$$\sum_{\text{label } y} \exp(\begin{array}{l} \text{weight}_1 * f_1(\text{fatally shot}, y) \\ \text{weight}_2 * f_2(\text{seriously wounded}, y) \\ \text{weight}_3 * f_3(\text{Shining Path}, y) \\ \dots \end{array} + + +)$$

Q: What if none of our features apply?

Guiding Principle for **Maximum Entropy** Models

*“[The log-linear estimate] is the least biased estimate possible on the given information; i.e., it is **maximally noncommittal with regard to missing information.**”*

Edwin T. Jaynes, 1957



$$\exp(\theta \cdot f) \rightarrow \exp(\theta \cdot 0) = 1$$

<https://www.csee.umbc.edu/courses/graduate/678/spring19/loglin-tutorial>

Lesson 1: Basic Feature Design

Welcome! This interactive visualization will help you understand the popular technique of log-linear modeling.

Try it out! The sliders below control the parameters ("weights") of a log-linear model. When you increase the `circle` weight, which filled shapes get bigger? Which ones get smaller?

One game is to try to match all 4 shapes to the gray outlines. You will need to use both sliders. A shape will turn `gray` if it matches well. It turns `red` if it is too small, `blue` if it is too big. Note: You may like to zoom in with your browser.

Log Likelihood Score:
Current log: -10.000000 / 45.176

Data & Model Options

Change the data

Regularization

☒ None ☐ L_1 ☐ L_2

Hints

☒ Show gradient

Type Counts: Observed and Expected

	circle	square
filled	30	15
outline	15	30

$N = 60$

Feature Weights

circle: square:

Authors: Frank Fellers and Jason Eisner.
Find a bug? Want a feature? Think something can be improved? Please file an [issue report](#).
Last change: [2019-03-15](#)

The source code is released [here](#) under the [GNU Affero 3.0 license](#).
The lesson text and materials are licensed under a [Creative Commons Attribution-ShareAlike 3.0 Unported License](#).
[CC BY-SA](#)

To refer to this work, cite the paper that describes it [Fellers & Eisner, 2019](#).

Ingredients for classification

Inject *your* knowledge into a learning system

Feature representation

*Training data:
labeled examples*

Model

Ingredients for classification

Inject *your* knowledge into a learning system

Problem specific

Difficult to learn from bad
ones

Feature representation

*Training data:
labeled examples*

Model

What features would you extract to...

distinguish a **picture** of **me** from a picture of **someone else**?

determine whether a **sentence** is **grammatical** or **not**?

distinguish **cancerous cells** from **normal** cells?

Outline

Log-Linear (Maximum Entropy) Models

Basic Modeling

Connections to other techniques (“... by
any other name...”)

Objective to optimize

Regularization

Connections to Other Techniques

Log-Linear Models

Connections to Other Techniques

Log-Linear Models

as statistical
regression

(Multinomial) logistic regression

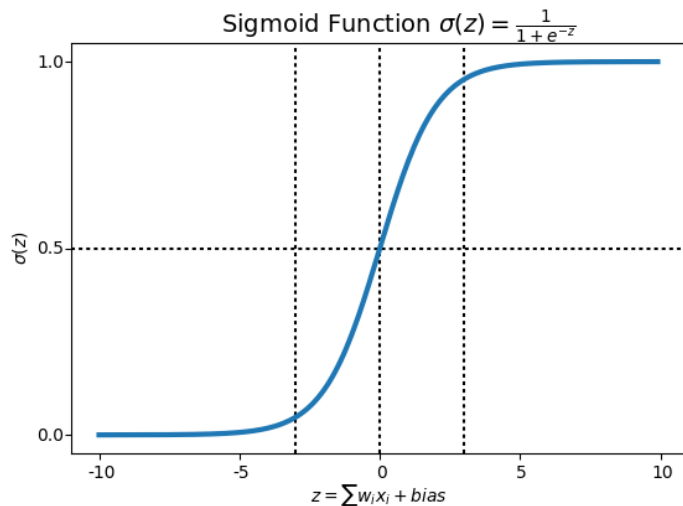
Softmax regression

“Solution” 1: A Simple Probabilistic (Linear*) Classifier

decision rule:

$$\hat{y}_i = \begin{cases} 0, & \sigma(\mathbf{w}^T \mathbf{x}_i + b) < .5 \\ 1, & \sigma(\mathbf{w}^T \mathbf{x}_i + b) \geq .5 \end{cases}$$

*turn responses
into probabilities*



Remember from
“Linear
regression”

minimize posterior 0-1 loss:

$$\min_{\mathbf{w}} \sum_i \mathbb{E}_{\hat{y}_i} [1[y_i p(\hat{y}_i = 1|x_i) < 0]] =$$

$$\max_{\mathbf{w}} \sum_i p(\hat{y}_i = y_i | x_i)$$

*why MAP
classifiers are
reasonable*

Connections to Other Techniques

Log-Linear Models

as statistical regression (Multinomial) logistic regression

Softmax regression

based in information theory Maximum Entropy models (MaxEnt)

Connections to Other Techniques

Log-Linear Models

as statistical regression (Multinomial) logistic regression

Softmax regression

based in information theory Maximum Entropy models (MaxEnt)

a form of Generalized Linear Models

Generalized Linear Models

$$y = \sum_k \theta_k x_k + b$$

response *linear* wrt parameters*

the *response* can be a general (transformed) version of another *response*

**affine is okay*

Generalized Linear Models

$$y = \sum_k \theta_k x_k + b$$

response *linear* wrt parameters*

the *response* can be a general (transformed) version of another *response*

logistic regression

$$\frac{\log p(x = i)}{\log p(x = K)} = \sum_k \theta_k f(x_k, i) + b$$

**affine is okay*

Connections to Other Techniques

Log-Linear Models

as statistical regression (Multinomial) logistic regression

Softmax regression

based in information theory Maximum Entropy models (MaxEnt)

a form of Generalized Linear Models

viewed as Discriminative Naïve Bayes

Connections to Other Techniques

Log-Linear Models

as statistical regression (Multinomial) logistic regression

Softmax regression

based in information theory Maximum Entropy models (MaxEnt)

a form of Generalized Linear Models

viewed as Discriminative Naïve Bayes

to be cool today :) Very shallow (sigmoidal) neural nets

Outline

Log-Linear (Maximum Entropy) Models

Basic Modeling

Connections to other techniques (“... *by any other name...*”)

Objective to optimize

Regularization

Version 1: Minimize Cross Entropy Loss

loss uses y (random variable), or model's probabilities $\ell^{\text{xent}}(\vec{y}^*, p(y))$

$$\ell^{\text{xent}}(\vec{y}^*, y) = - \sum_k \vec{y}^*[k] \log p(y = k)$$

index of "1"
indicates
correct value

$$\begin{pmatrix} 0 \\ 0 \\ \dots \\ 1 \\ \dots \\ 0 \end{pmatrix}$$

one-hot
vector

minimize xent loss $\rightarrow$
maximize log-likelihood (A2, Q2)

objective is convex

Version 2: Maximize (Full/Log) Likelihood

$$\prod_i p_{\theta}(y_i|x_i) \propto \prod_i \exp(\theta^T f(x_i, y_i))$$

These values can have very
small magnitude → underflow

Differentiating this
product could be a pain

Version 2: Maximize Log-Likelihood

Wide range of (negative) numbers

Sums are more stable

$$\log \prod_i p_{\theta}(y_i|x_i) = \sum_i \log p_{\theta}(y_i|x_i)$$

Version 2: Maximize Log-Likelihood

Wide range of (negative) numbers

Sums are more stable

$$\log \prod_i p_{\theta}(y_i|x_i) = \sum_i \log p_{\theta}(y_i|x_i)$$
$$= \sum_i \underbrace{\theta^T f(x_i, y_i) - \log Z(x_i)}$$

Differentiating this
becomes nicer (even
though Z depends on θ)

Log-Likelihood Gradient

Each component k is the difference
between:

Log-Likelihood Gradient

Each component k is the difference
between:

the total value of feature f_k in the
training data

$$\sum_i f_k(x_i, y_i)$$

Log-Likelihood Gradient

Each component k is the difference between:

the total value of feature f_k in the training data

$$\sum_i f_k(x_i, y_i)$$

and

the total value the current model p_θ *thinks* it computes for feature f_k

$$\sum_i \mathbb{E}_p[f(x_i, y')]$$

“Moment Matching”

A1 Q4, Eq-1 (what were the feature functions)?

<https://www.csee.umbc.edu/courses/graduate/678/spring19/loglin-tutorial>

Lesson 6: Gradient Optimization

Welcome! This interactive visualization will help you understand the popular technique of log-linear modeling.

Try it out! The sliders below control the parameters ("weights") of a log-linear model. When you increase the `circle` weight, which filled shapes get bigger? Which ones get smaller?

One game is to try to match all 4 shapes to the gray outlines. You will need to use both sliders. A shape will turn `gray` if it matches well. It turns `red` if it is too small, `blue` if it is too big. Note: You may like to zoom in with your browser.

Log Likelihood Score:
Current LL: -45.176

Data & Model Options

Change the data

Regularization

☒ None ☐ L_1 ☐ L_2

Hints

☒ Show gradient

Type Counts: Observed and Expected

	circle	square
filled	30	15
outline	15	30

$N = 60$

Feature Weights

circle:

square:

Authors: Frank Fellers and Jason Eisner.
Find a bug? Want a feature? Think something can be improved? Please file an [issue report](#).
Last change: 2019-03-15 14:00:00

The source code is released [here](#) under the [GNU Affero 3.0 license](#).
The lesson text and materials are licensed under a [Creative Commons Attribution-ShareAlike 3.0 Unported license](#).
[CC BY-SA](#)

To refer to this work, cite the paper that describes it [Fellers & Eisner, 2019](#).

Log-Likelihood Gradient Derivation

$$\nabla_{\theta} F(\theta) = \nabla_{\theta} \sum_i [\theta^T f(x_i, y_i) - \log Z(x_i)]$$

Log-Likelihood Gradient Derivation

$$\nabla_{\theta} F(\theta) = \nabla_{\theta} \sum_i [\theta^T f(x_i, y_i) - \log Z(x_i)]$$

$$= \nabla_{\theta} \sum_i f(x_i, y_i) -$$

$$Z(x_i) = \sum_{y'} \exp(\theta \cdot f(x_i, y'))$$


Log-Likelihood Gradient Derivation

$$\begin{aligned}\nabla_{\theta} F(\theta) &= \nabla_{\theta} \sum_i [\theta^T f(x_i, y_i) - \log Z(x_i)] \\ &= \nabla_{\theta} \sum_i f(x_i, y_i) - \sum_i \sum_{y'} \underbrace{\frac{\exp(\theta^T f(x_i, y'))}{Z(x_i)}}_{\text{scalar } p(y' | x_i)} \underbrace{f(x_i, y')}_{\text{vector of functions}}\end{aligned}$$

use the (calculus) chain rule

$$\frac{\partial}{\partial \theta} \log g(h(\theta)) = \left(\frac{\partial g}{\partial h(\theta)} \right) \left(\frac{\partial h}{\partial \theta} \right)$$

scalar $p(y' | x_i)$

vector of functions

Log-Likelihood Gradient Derivation

$$\begin{aligned}\nabla_{\theta} F(\theta) &= \nabla_{\theta} \sum_i [\theta^T f(x_i, y_i) - \log Z(x_i)] \\ &= \nabla_{\theta} \sum_i f(x_i, y_i) - \sum_i \sum_{y'} \frac{\exp(\theta^T f(x_i, y'))}{Z(x_i)} f(x_i, y')\end{aligned}$$

Do we want these to *fully* match?

What does it mean if they do?

What if we have missing values in our data?

Outline

Log-Linear (Maximum Entropy) Models

Basic Modeling

Connections to other techniques (“... *by any other name...*”)

Objective to optimize

Regularization

Weight regularization $R(\mathbf{w})$

Nice if $R(\mathbf{w})$ is **convex**

Small weights regularization

$$R^{(\text{norm})}(\mathbf{w}) = \sqrt{\sum_d w_d^2}$$

$$R^{(\text{sqr})}(\mathbf{w}) = \sum_d w_d^2$$

Sparsity regularization

$$R^{(\text{count})}(\mathbf{w}) = \sum_d \mathbf{1}[|w_d| > 0] \quad \text{not convex}$$

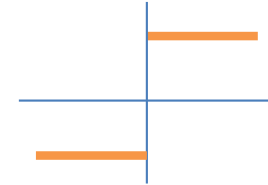
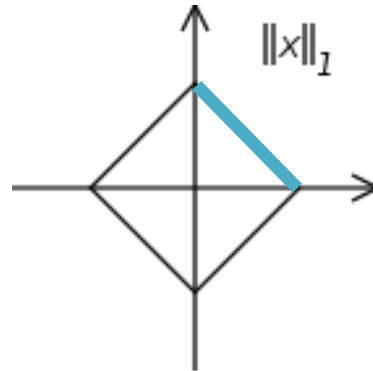
Family of “**p-norm**” regularization

$$R^{(\text{p-norm})}(\mathbf{w}) = \left(\sum_d |w_d|^p \right)^{1/p}$$

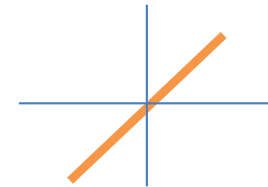
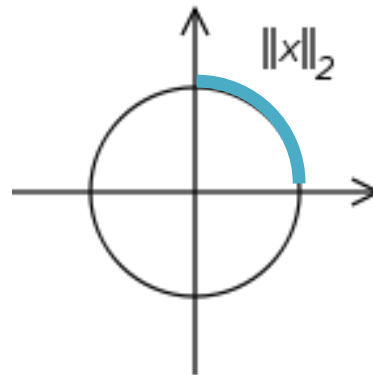
convex: $p \geq 1$
not convex: $0 \leq p < 1$

Contours of p-norms

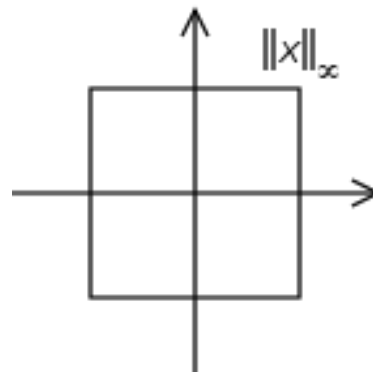
$$\|x\|_1 = \sum_{i=1}^n |x_i|$$



$$\|x\|_2 = \sqrt{\sum_{i=1}^n |x_i|^2}$$



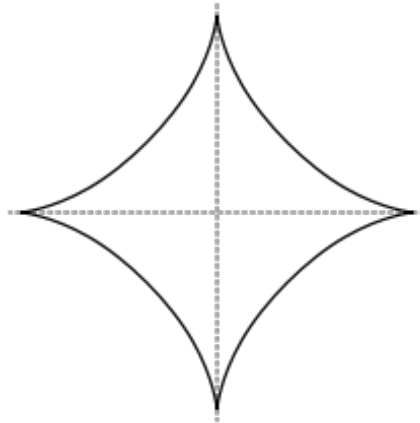
$$\|x\|_\infty = \max_{i=1, \dots, n} |x_i|$$



examine **shape**
(**slope**) of **surfaces** to
determine effect on
the regularized
parameters

Contours of p-norms

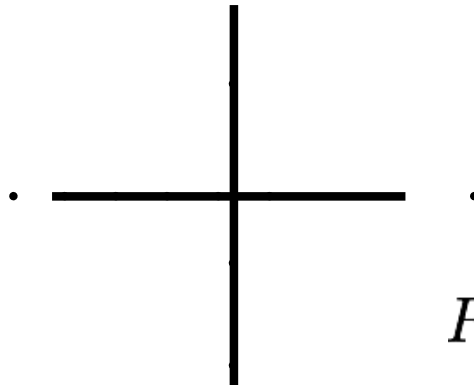
$$p = \frac{2}{3}$$



examine **shape**
(**slope**) of **surfaces** to
determine effect on
the regularized
parameters

Counting non-zeros:

$$p = 0$$



$$R^{(\text{count})}(\mathbf{w}) = \sum_d \mathbf{1}[|w_d| > 0]$$

A Simple Regularized Linear Classifier

decision rule: $\hat{y}_i = \begin{cases} 0, & \mathbf{w}^T \mathbf{x}_i < 0 \\ 1, & \mathbf{w}^T \mathbf{x}_i \geq 0 \end{cases}$

loss function: $\ell = 1[y_i \mathbf{w}^T \mathbf{x}_i < 0]$

$$\min_{\mathbf{w}} \sum_n \mathbf{1}[y_n \mathbf{w}^T \mathbf{x}_n < 0] + \lambda R(\mathbf{w})$$

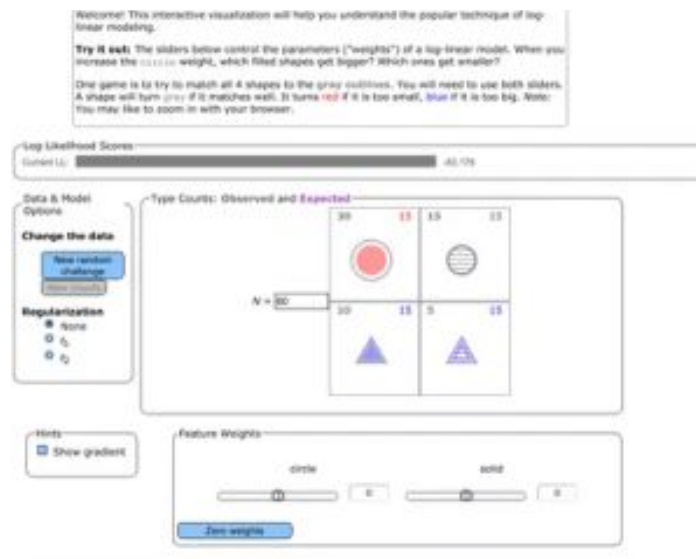
hyperparameter $\downarrow$

$\uparrow$ fewest mistakes on training

$\nwarrow$ regularize toward a simpler model

<https://www.csee.umbc.edu/courses/graduate/678/spring19/loglin-tutorial>

Lesson 8: Regularization



Authors: Frank Fahrenberg and Jason Eisner.
Find a bug? Want a feature? Think something can be improved? Please file out an [issue report](#).
Last change:

The source code is released [here](#) under the [GNU Affero 3.0 license](#).
The lesson text and materials are licensed under a [Creative Commons Attribution-ShareAlike 3.0 Unported License](#).
[CC BY-SA](#)

To refer to this work, cite the paper that describes it [Fahrenberg & Eisner, 2013](#).

Understanding Conditioning

$$p(y | x) \propto \exp(\theta \cdot f(x))$$

Is this a good posterior classifier? (no)

<https://www.csee.umbc.edu/courses/graduate/678/spring19/loglin-tutorial>

Lesson 11: Global vs. Conditional Modeling

Welcome! This interactive visualization will help you understand the popular technique of log-linear modeling.

Try it out! The sliders below control the parameters ("weights") of a log-linear model. When you increase the `circle` weight, which filled shapes get bigger? Which ones get smaller?

One game is to try to match all 4 shapes to the gray outlines. You will need to use both sliders. A shape will turn `gray` if it matches well. It turns `red` if it is too small, `blue` if it is too big. Note: You may like to zoom in with your browser.

Log Likelihood Score:
Current LL: -45.176

Data & Model Options

Change the data

Regularization

☒ None ☐ L_1 ☐ L_2

Hints

☒ Show gradient

Type Counts: Observed and Expected

30	15	15	15
			
20	15	5	15
			

$N = 60$

Feature Weights

circle: solid:

Authors: Frank Fahrenberg and Jason Eisner.
Find a bug? Want a feature? Think something can be improved? Please file out an [issue report](#).
Last change:

The source code is released [here](#) under the [GNU Affero 3.0 license](#).
The lesson text and materials are licensed under a [Creative Commons Attribution-ShareAlike 3.0 Unported License](#).
[CC BY-SA](#)

To refer to this work, cite the paper that describes it [Fahrenberg & Eisner, 2017](#).

Connections to Other Techniques

Log-Linear Models

as statistical
regression

(Multinomial) logistic regression

Softmax regression

based in
information theory

Maximum Entropy models (MaxEnt)

a form of

Generalized Linear Models

viewed as

Discriminative Naïve Bayes

to be cool
today :)

Very shallow (sigmoidal) neural nets

Outline

Log-Linear (Maximum Entropy) Models

Basic Modeling

Connections to other techniques (“... *by any other name...*”)

Objective to optimize

Regularization