

Machine Learning for NLP

Lecture 3: Optimization and machine learning



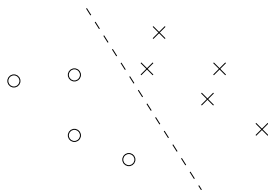
**UNIVERSITY OF
GOTHENBURG**

Richard Johansson

September 8, 2016

geometric view

- ▶ geometrically, a linear classifier can be interpreted as separating the vector space into two regions with a line (plane, hyperplane)



overview

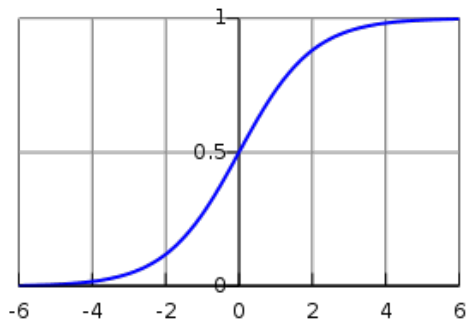
logistic regression

support vector machines

basic optimization

practical information

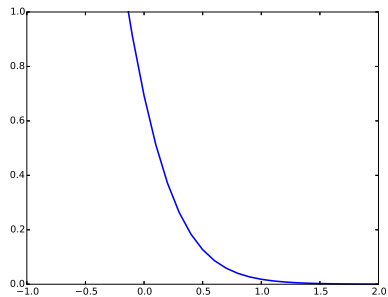
the logistic function



conversely

$$P(\text{negative output}|\mathbf{x}) = \frac{1}{1 + e^{\text{score}}}$$

plot of the log loss



recall: the fundamental tradeoff in machine learning



- ▶ **goodness of fit**: the learned classifier should be able to correctly classify the examples in the training data
- ▶ **regularization**: the classifier should be simple
- ▶ but so far in our LR description, we've just taken care of the first part!

combining the pieces

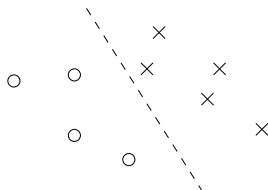
- ▶ we combine the loss and the regularizer:

$$\sum \text{Loss}(\mathbf{w}, \mathbf{x}_i, y_i) + \lambda \cdot |\mathbf{w}|^2$$

- ▶ in this formula, λ is a “tweaking” parameter that controls the tradeoff between loss and regularization
- ▶ note: in some formulations (including scikit-learn), there is a parameter C instead of the λ that is put before the loss

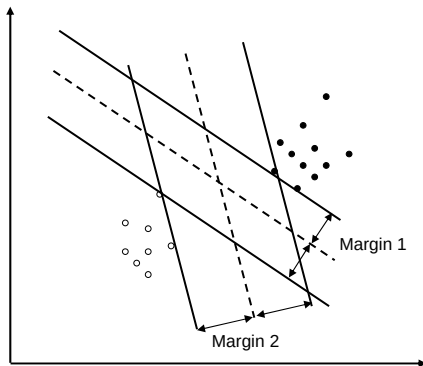
geometric view

- ▶ geometrically, a linear classifier can be interpreted as separating the vector space into two regions with a line (plane, hyperplane)



margin of separation

- ▶ the **margin** γ denotes how well $\mathbf{w}$ separates the classes:



large margins are good

- ▶ a result from statistical learning theory:

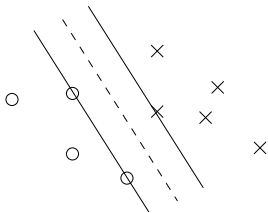
expected test error = training error + BigUglyFormula($\frac{1}{\gamma^2}$)

- ▶ larger margin $\rightarrow$ better generalization



support vector machines

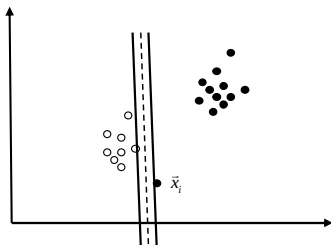
- ▶ **support vector machines** (SVMs) or support vector classifiers (SVC) are linear classifiers constructed by selecting the $\mathbf{w}$ that maximizes the margin



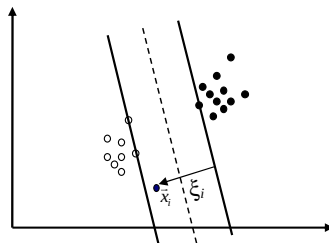
- ▶ note: the solution depends only on the borderline examples: the **support vectors**

soft-margin SVMs

- ▶ in some cases the dataset is inseparable, or nearly inseparable
- ▶ **soft-margin SVM**: allow some examples to be disregarded when maximizing the margin

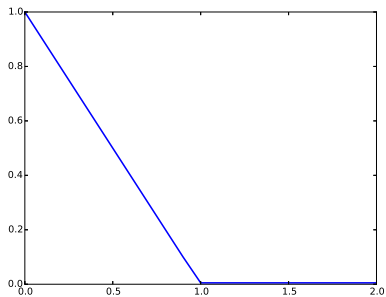


A) Hard Margin SVM

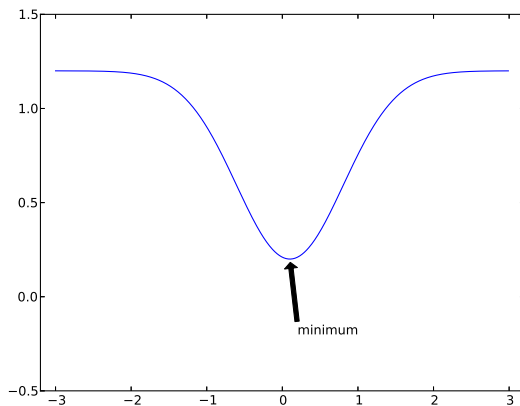


B) Soft Margin SVM

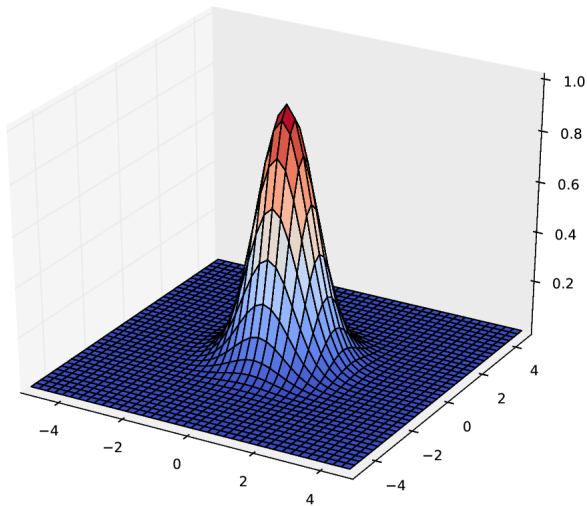
plot of the hinge loss



one-variable example

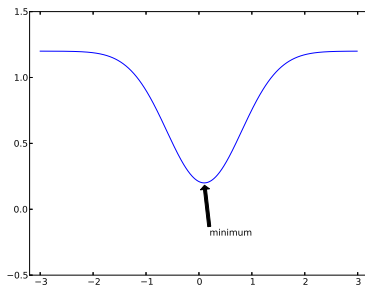


two-variable example



remember your highschool calculus...

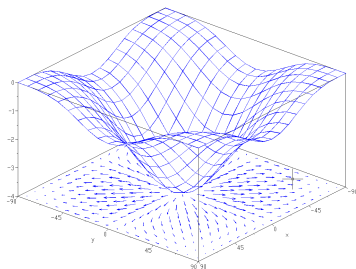
- ▶ in your early school days, you might have seen the **derivative** of a function
- ▶ intuition: the derivative measures the slope



- ▶ if a “nice” function has a maximum or minimum, then the derivative will be zero there

the gradient

- ▶ the multidimensional equivalent of the derivative is called the **gradient**
- ▶ if f is a function of n variables, then the gradient is an n -dimensional vector, often written $\nabla f(x)$
- ▶ intuition: the gradient points in the uphill direction



- ▶ again: the gradient is zero if we have an optimum

computing the gradient



gradient of $0.7 * \exp(-0.7*(x+1)^2 - 0.8*(y-1)^2)$



Examples Random

Assuming "gradient" is a function | Use as a [unit](#) instead

Input interpretation:

$$\text{grad}\left(0.7 e^{-0.7(1+x)^2 - 0.8(-1+y)^2}\right)$$

Result:

$$\text{grad}\left(0.7 e^{-0.7(x+1)^2 - 0.8(y-1)^2}\right)$$

Del operator form:

$$\nabla\left(0.7 e^{-0.7(1+x)^2 - 0.8(-1+y)^2}\right)$$

Result in 2D Cartesian coordinates:

$$\text{grad}\left(0.7 e^{-0.7(x+1)^2 - 0.8(y-1)^2}\right) = \left\{-0.98(x+1) e^{-0.7(x+1)^2 - 0.8(y-1)^2}, -1.12(y-1) e^{-0.7(x+1)^2 - 0.8(y-1)^2}\right\}$$

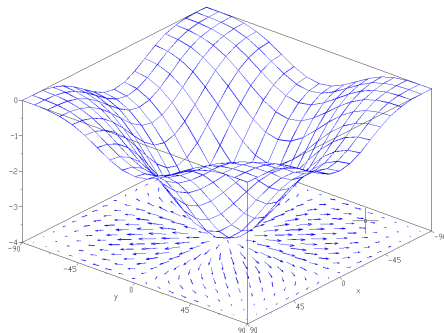
(x: first Cartesian coordinate | y: second Cartesian coordinate)

Download page

POWERED BY THE WOLFRAM LANGUAGE

gradient descent

- ▶ as we saw, the gradient points in the uphill direction:



- ▶ this intuition leads to a simple idea for finding the minimum:
 - ▶ take a small step in the direction opposite to the gradient
 - ▶ repeat until the gradient is close enough to zero
- ▶ this is called **gradient descent**

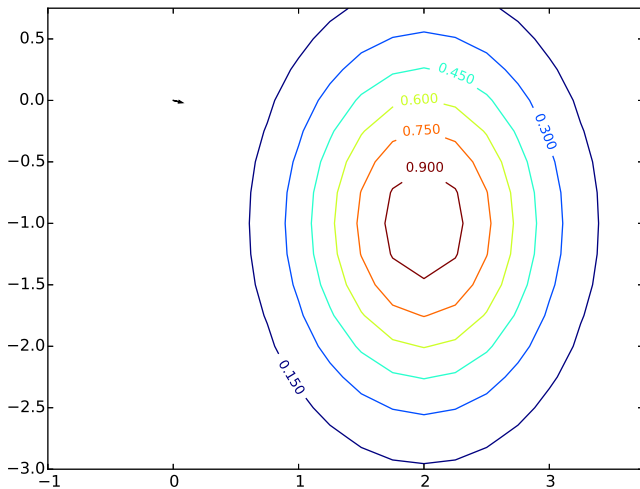
gradient descent, pseudocode

- ▶ the same thing again, in pseudocode:
 1. set x to some initial value, and select a suitable step size c
 2. compute the gradient $\nabla f(x)$
 3. if $\nabla f(x)$ is small enough, we are done
 4. otherwise, subtract $c \cdot \nabla f(x)$ from x and go back to step 2
- ▶ conversely, to find the maximum we can do **gradient ascent**: then we instead add $c \cdot \nabla f(x)$ to x

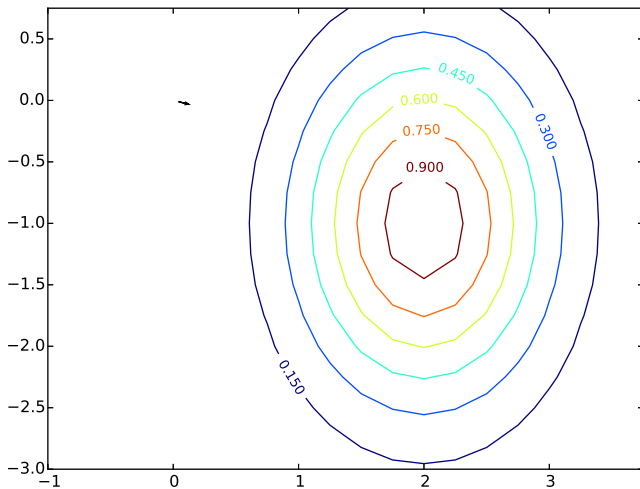
in Python

```
def gradient_ascent(x_init, y_init,
                    threshold = 0.001,
                    steplength = 0.01):
    x = x_init
    y = y_init
    done = False
    while not done:
        gxy = gradient_of_my_function(x, y)
        x += steplength * gxy[0]
        y += steplength * gxy[1]
        if numpy.linalg.norm(gxy) < threshold:
            done = True
    return (x, y)
```

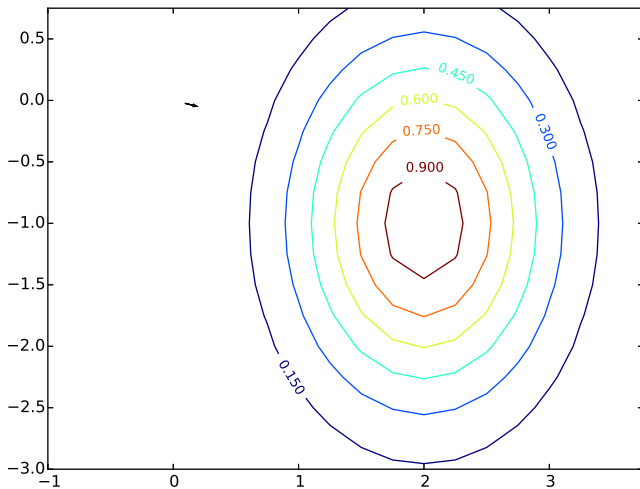

gradient ascent example



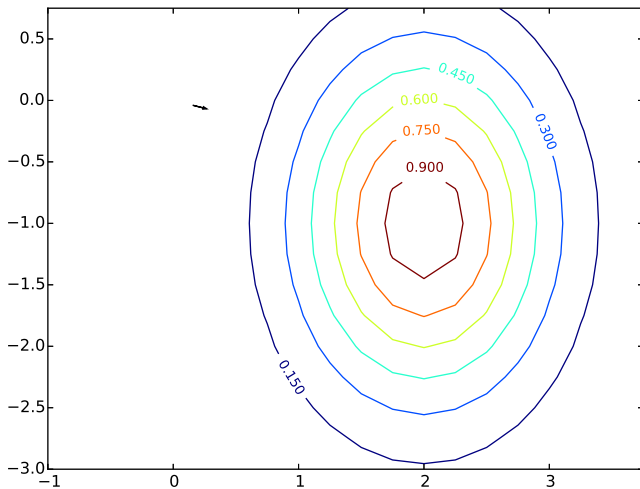
gradient ascent example



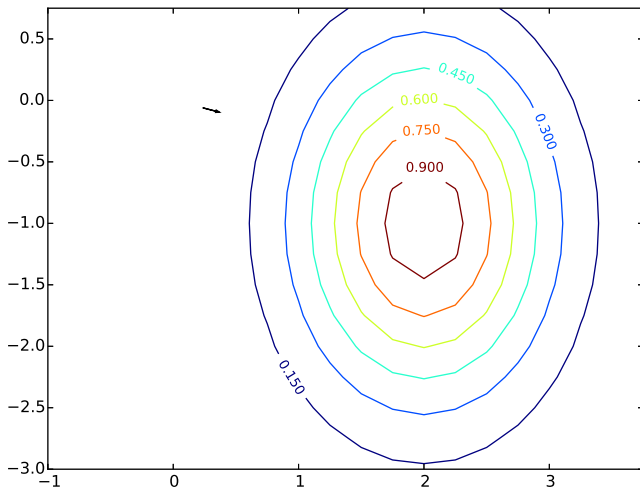
gradient ascent example



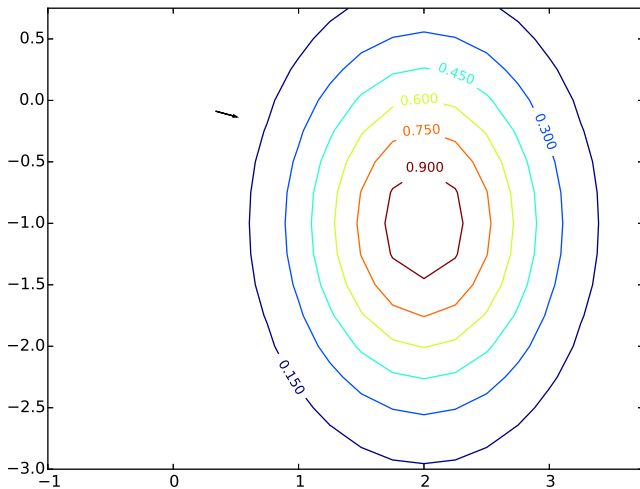
gradient ascent example



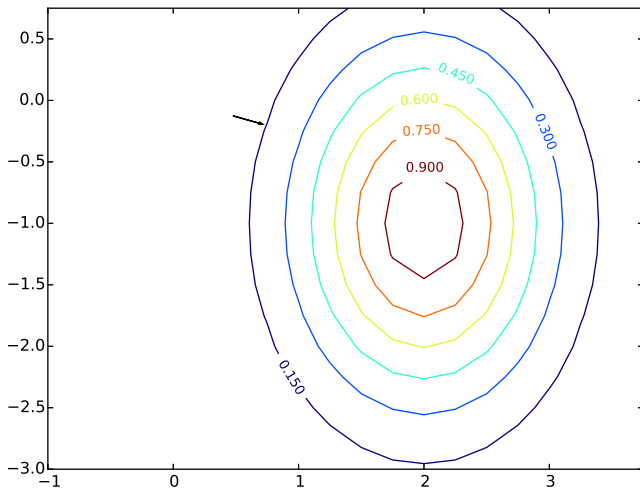
gradient ascent example



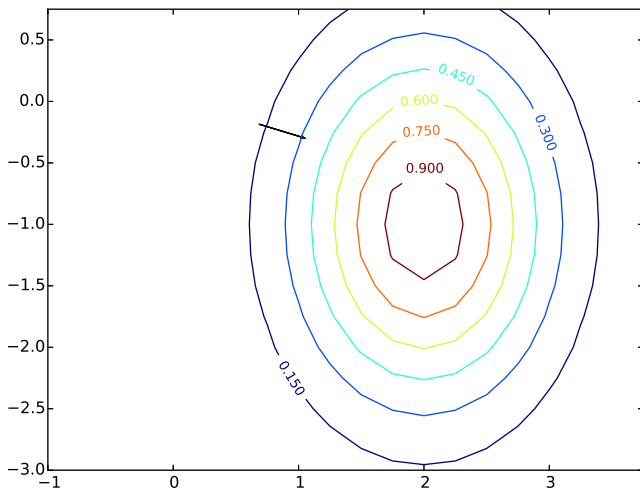
gradient ascent example



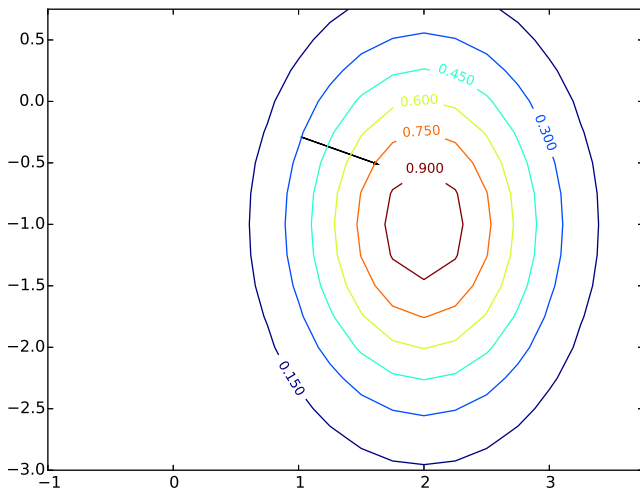
gradient ascent example



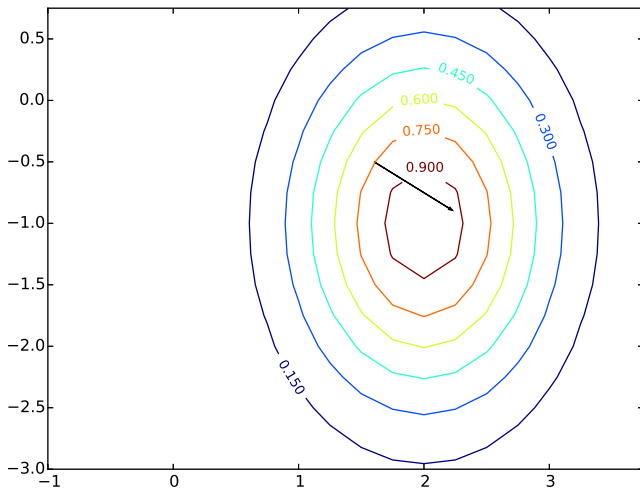
gradient ascent example



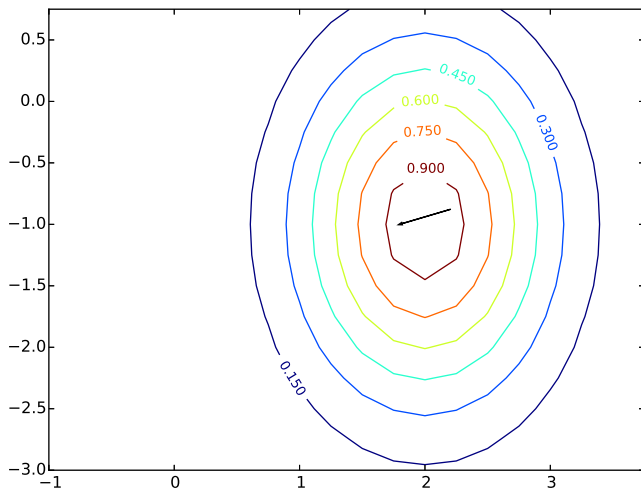
gradient ascent example



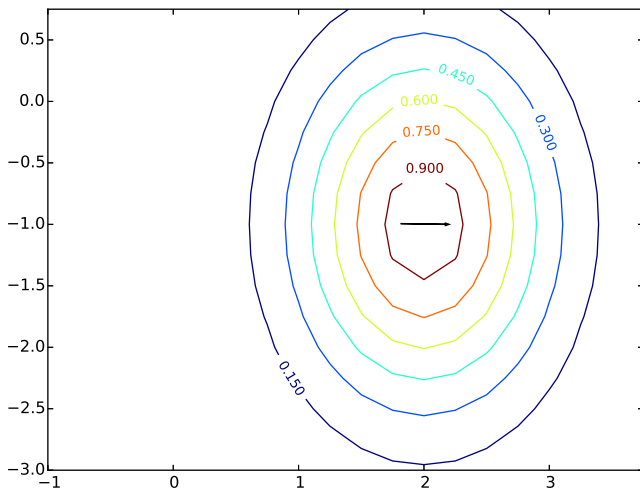
gradient ascent example



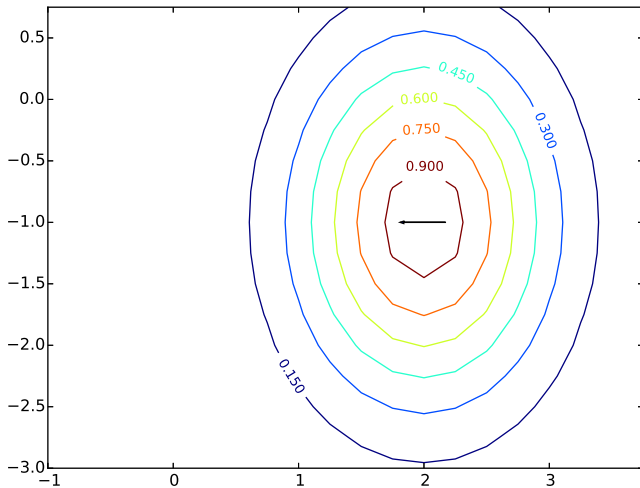
gradient ascent example



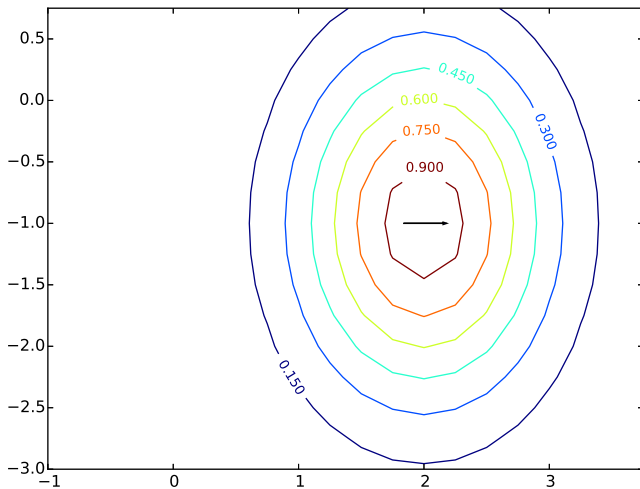
gradient ascent example



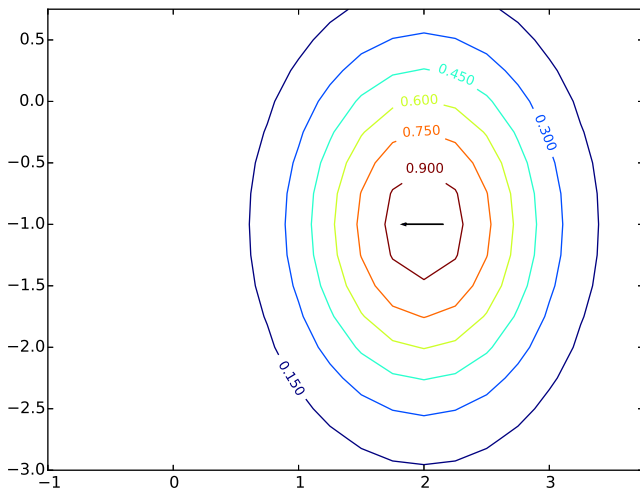
gradient ascent example



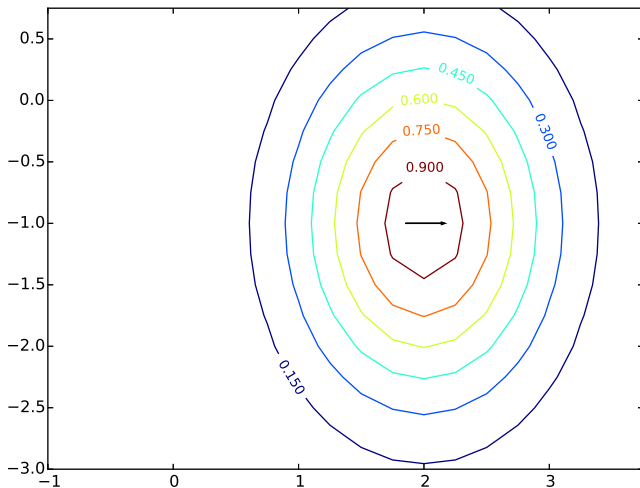
gradient ascent example



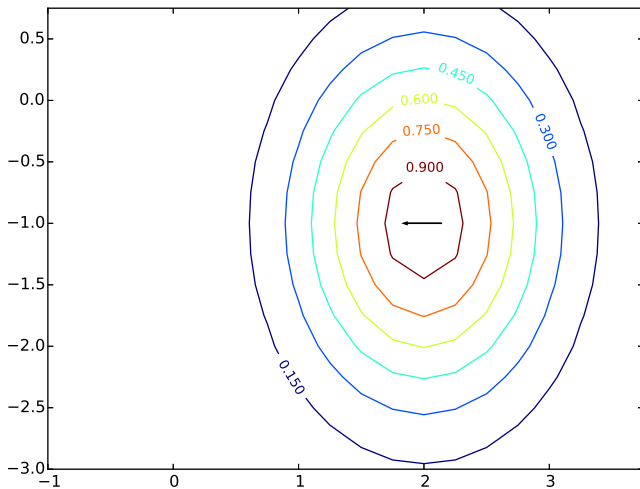
gradient ascent example



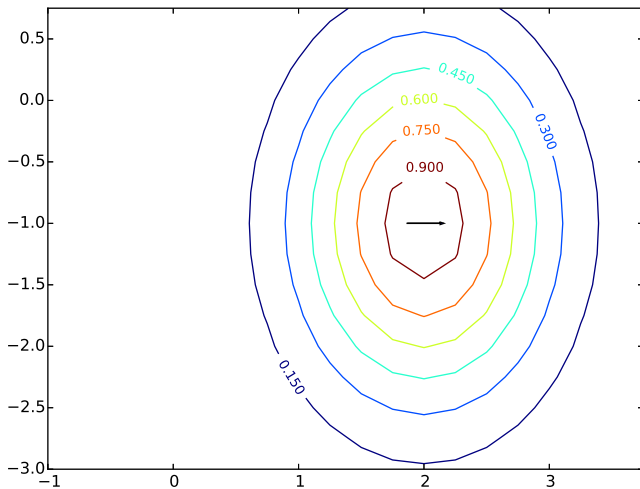
gradient ascent example



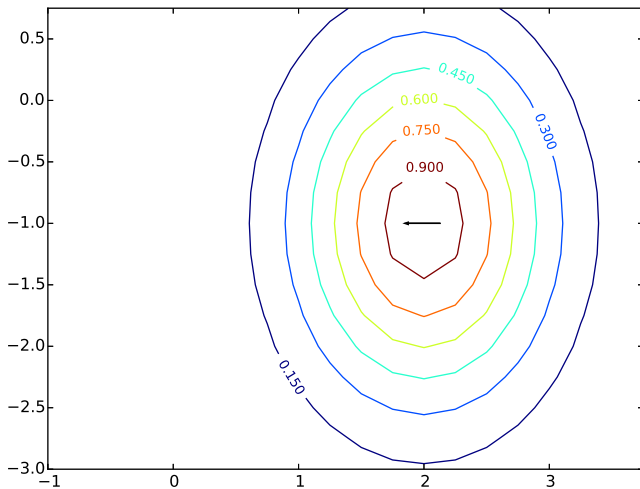
gradient ascent example



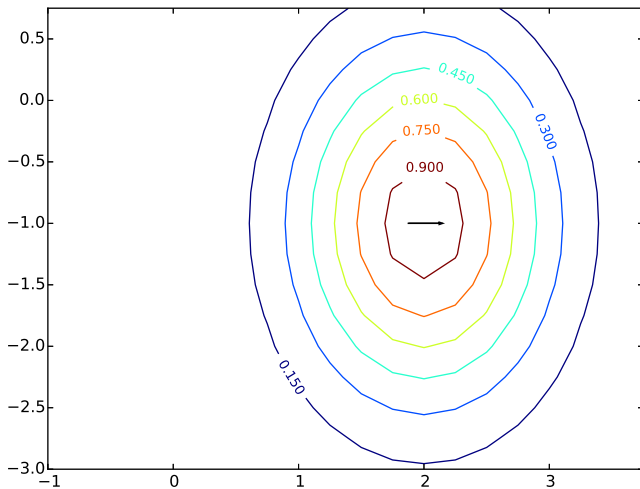
gradient ascent example



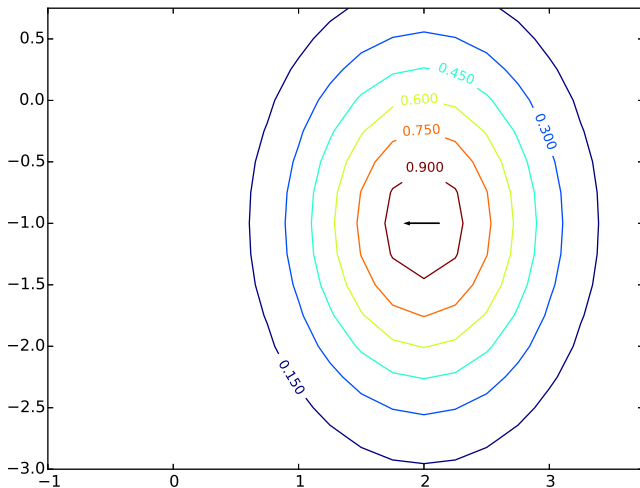
gradient ascent example



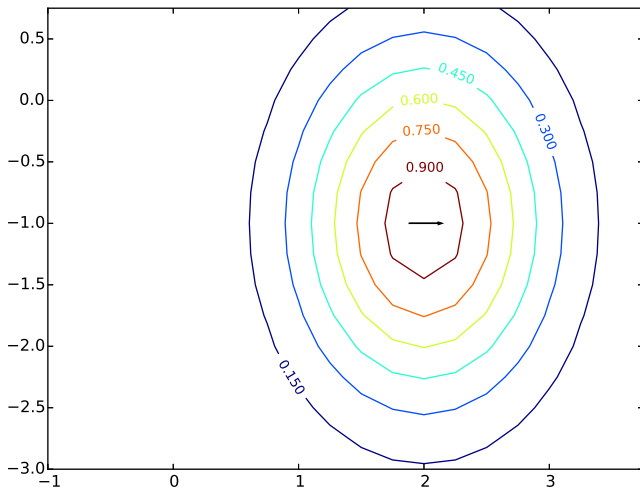
gradient ascent example



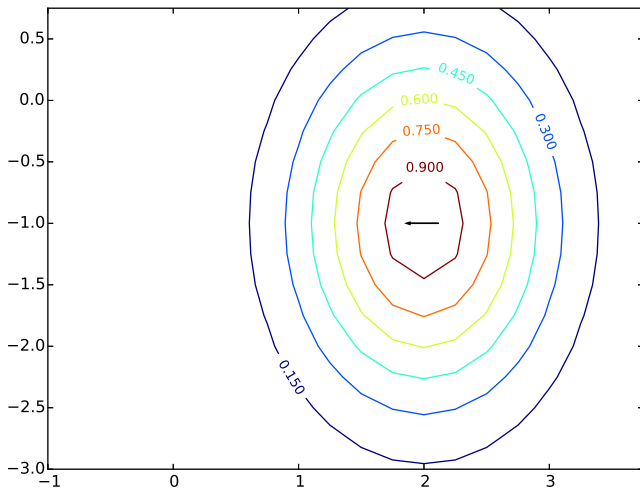
gradient ascent example



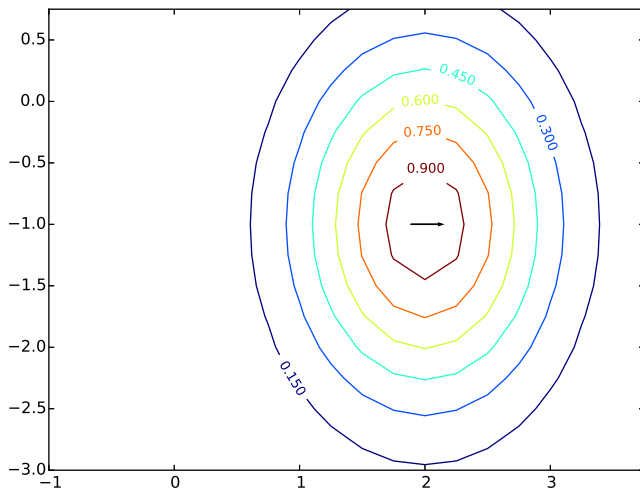
gradient ascent example



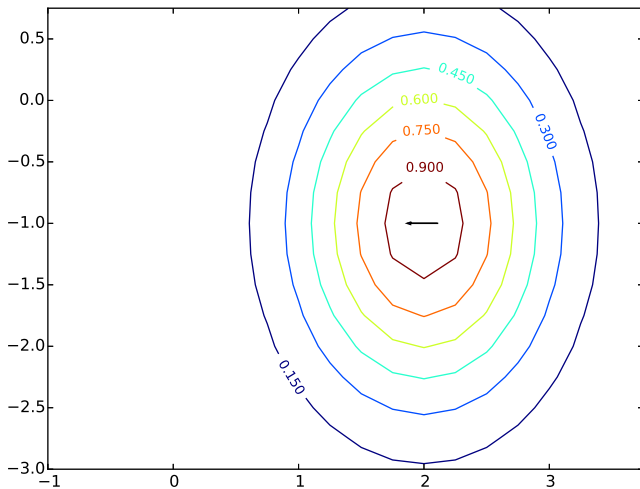
gradient ascent example



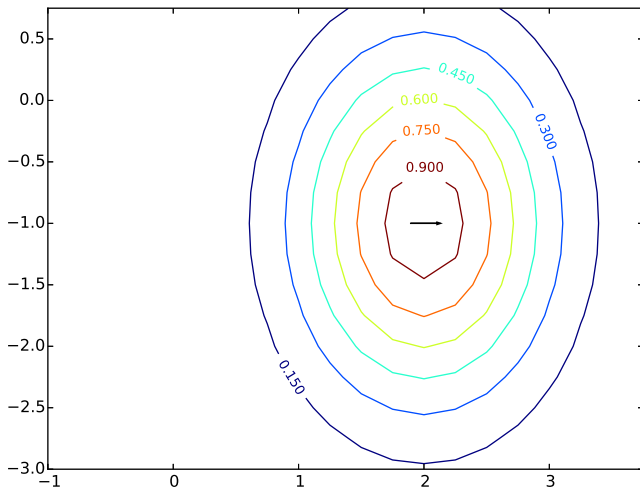
gradient ascent example



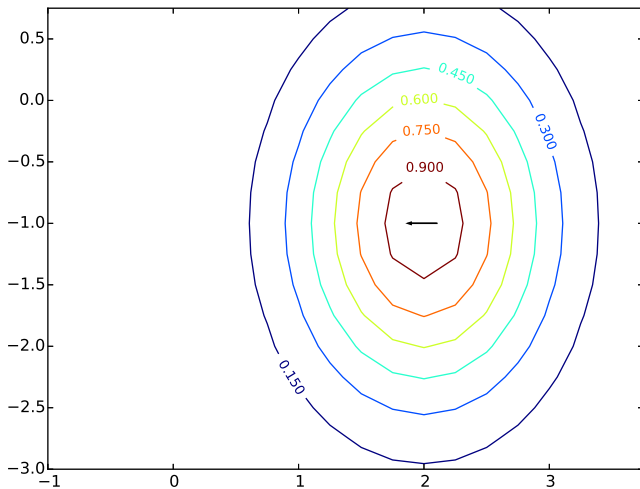
gradient ascent example



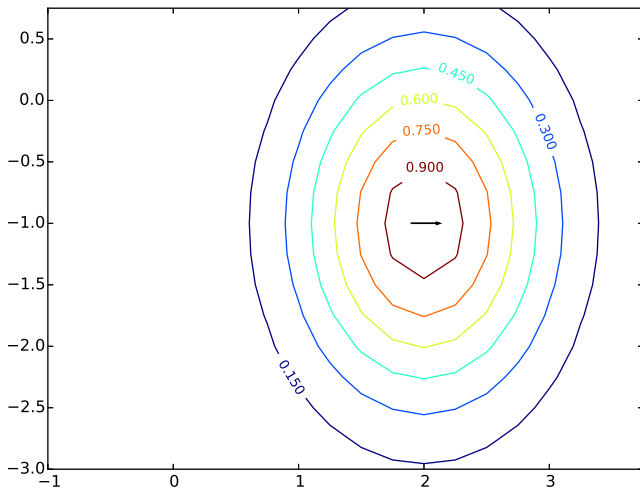
gradient ascent example



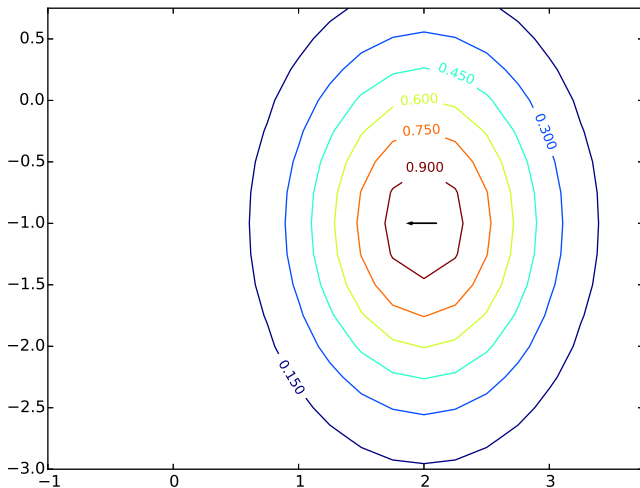
gradient ascent example



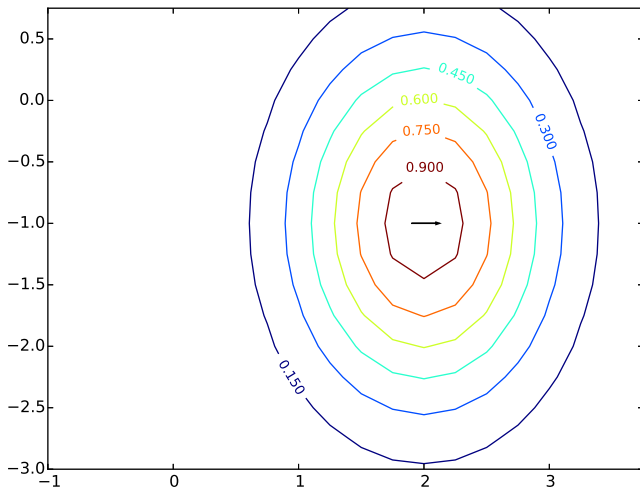
gradient ascent example



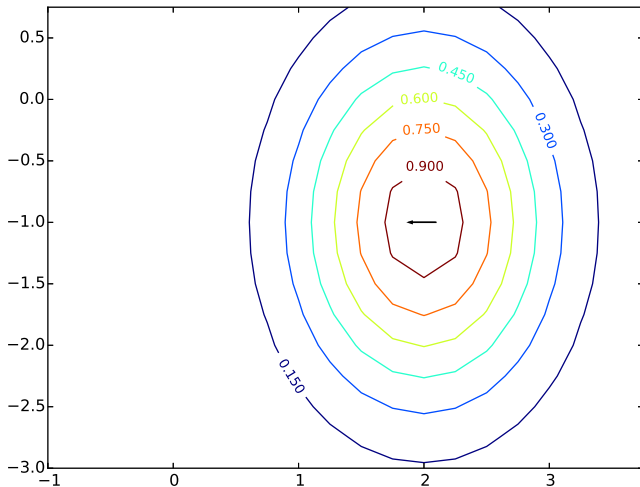
gradient ascent example



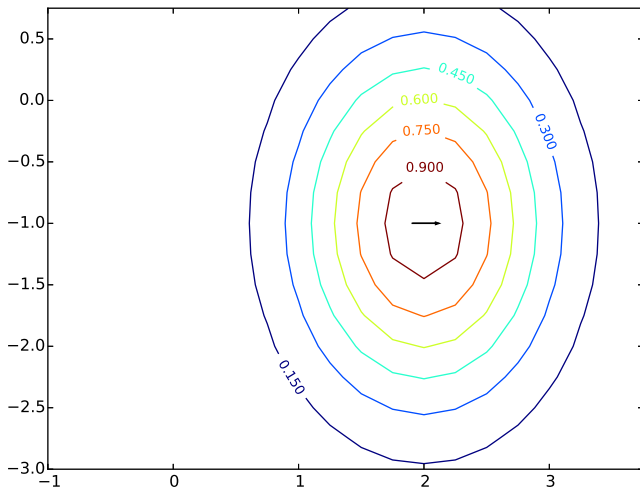
gradient ascent example



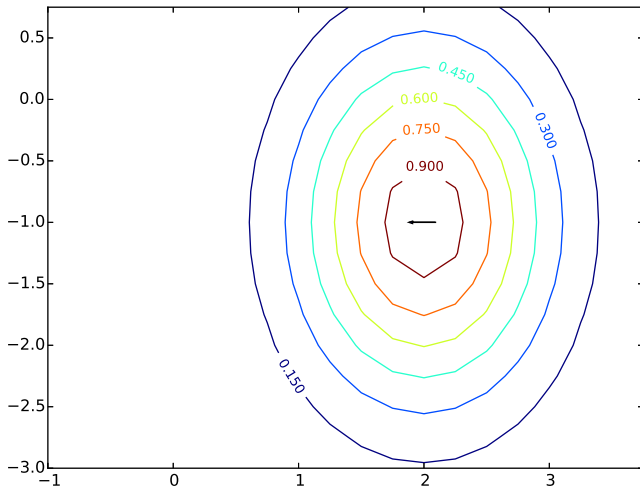
gradient ascent example



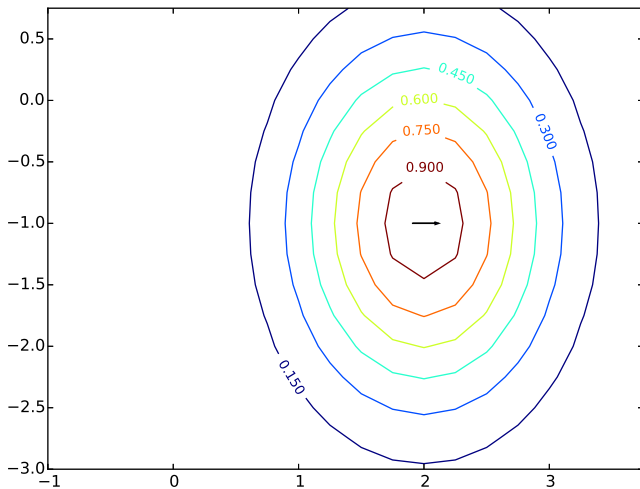
gradient ascent example



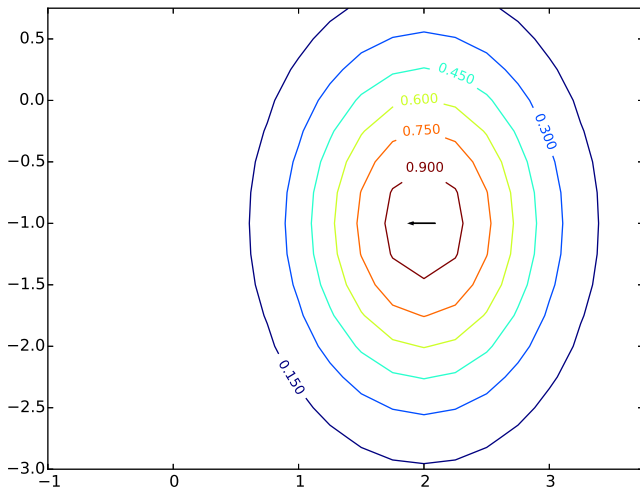
gradient ascent example



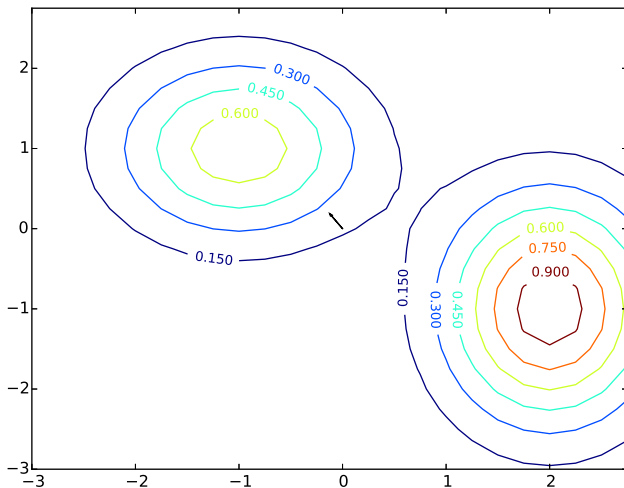
gradient ascent example



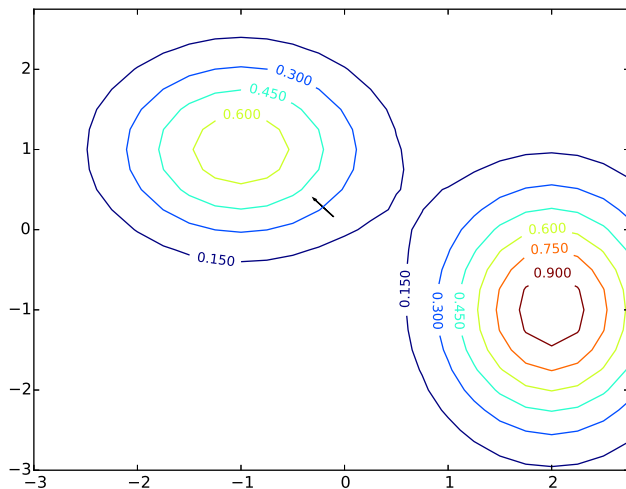
gradient ascent example



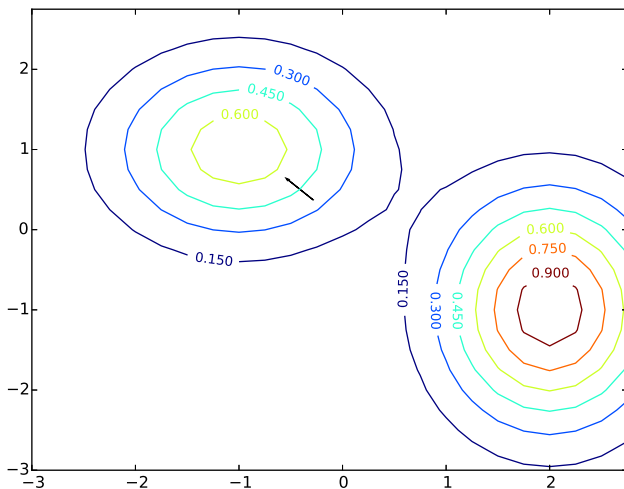
gradient ascent example (2)



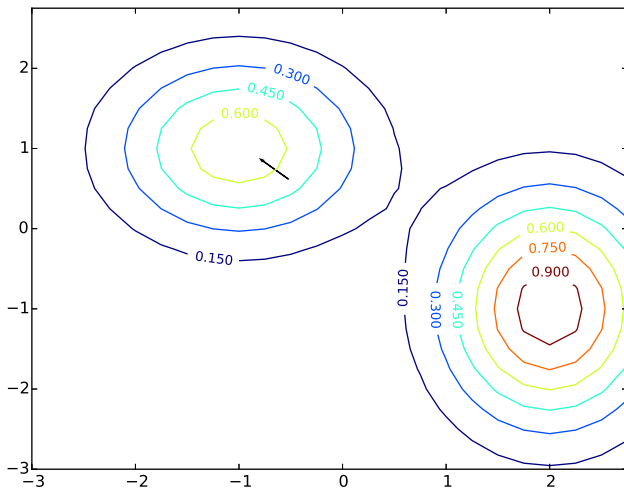
gradient ascent example (2)



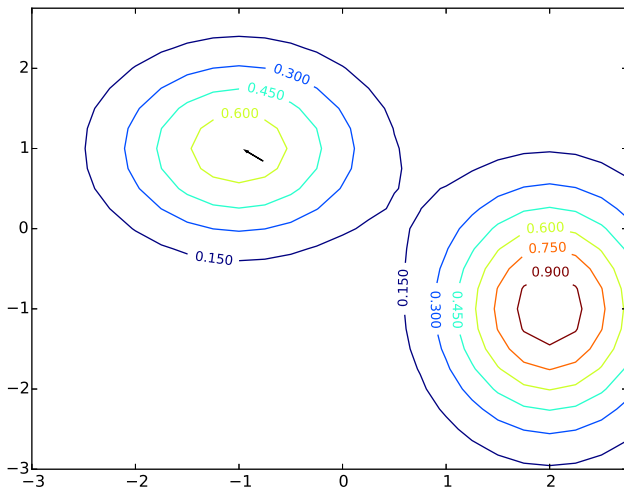
gradient ascent example (2)



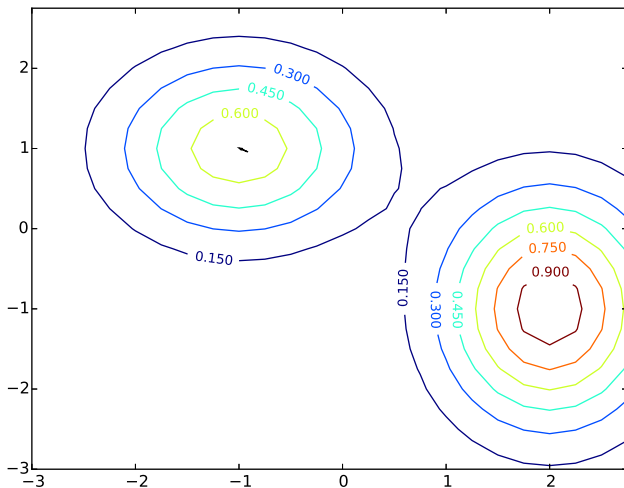
gradient ascent example (2)



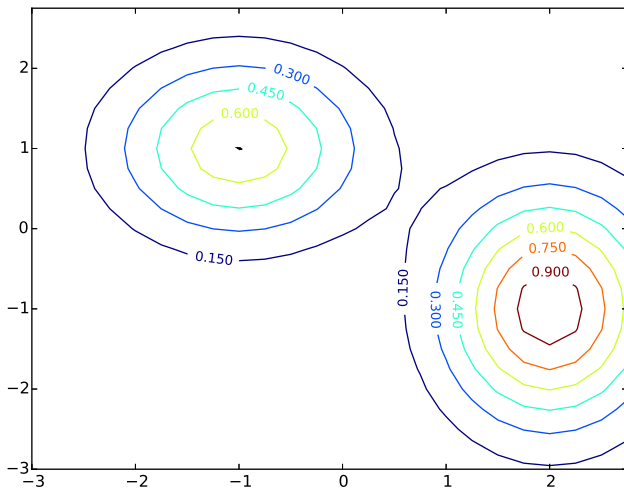
gradient ascent example (2)



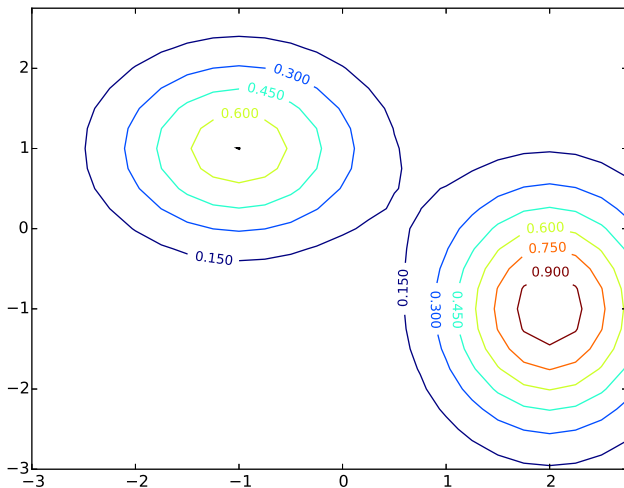
gradient ascent example (2)



gradient ascent example (2)

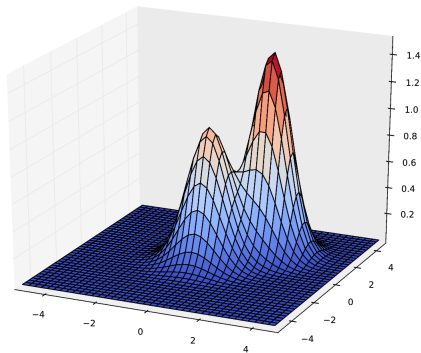


gradient ascent example (2)



local and global maxima/minima

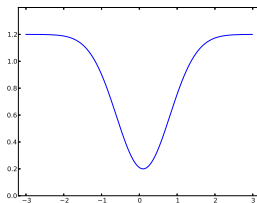
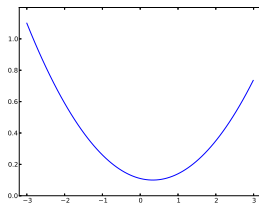
- ▶ some functions have **local maxima or minima**



- ▶ these functions are harder to optimize because the local (but not global) optima also have a gradient of 0

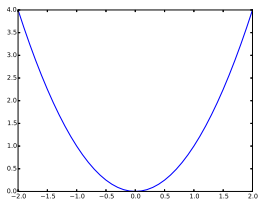
convex and concave functions

- ▶ a function is **convex** if it always curves downwards
- ▶ equivalently, if we draw a line between two points of the surface, the surface is always below the line

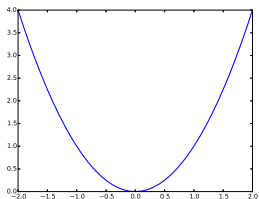
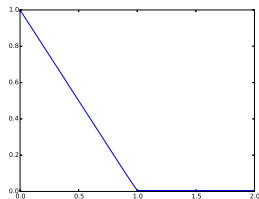


- ▶ the point of this: if we find a local optimum (gradient is 0) of a convex function, this is **guaranteed to be the minimum**
- ▶ conversely, a function is **concave** if it always curves upwards

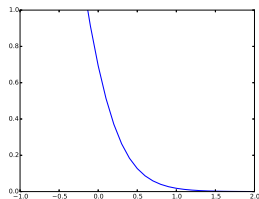
SVM and LR have convex objective functions



+



+



overview

logistic regression

support vector machines

basic optimization

practical information

