

Comparison of Shannon, Renyi and Tsallis Entropy used in Decision Trees

Tomasz Maszczyk and Włodzisław Duch

{tmaszczyk,w duch}@is.umk.pl

Abstract

Shannon entropy used in standard top-down decision trees does not guarantee the best generalization. Split criteria based on generalized entropies offer different compromise between purity of nodes and overall information gain. Modified C4.5 decision trees based on Tsallis and Renyi entropies have been tested on several high-dimensional microarray datasets with interesting results. This approach may be used in any decision tree and information selection algorithm.

Introduction

Decision tree algorithms are still the foundation of most large data mining packages, offering easy and computationally efficient way to extract simple decision rules. They should always be used as a reference, with more complex classification models justified only if they give significant improvement. Trees are based on recursive partitioning of data and unlike most learning systems they use different sets of features in different parts of the feature space, automatically performing local feature selection. This is an important and unique property of general divide-and-conquer algorithms that has not been paid much attention. The hidden nodes in the first neural layer weight the inputs in a different way, calculating specific projections of the input data on a line defined by the weights. This is still non-local feature that captures information from a whole sector of the input space, not from the localized region. On the other hand localized radial basis functions capture only the local information around some reference points. Recursive partitioning in decision trees is capable of capturing local information in some dimensions and non-local in others. This is a desirable property that may be used in neural algorithms based on localized projected data.

The C4.5 algorithm to generate trees is still the basis of the most popular approach in this field. Tests for partitioning data in C4.5 decision trees are based on the concept of information entropy and applied to each feature individually. Such tests create two nodes that should on the one hand contain data that are as pure as possible (i.e. belong to a single class), and on the other hand increase overall separability of data. Tests that are based directly on indices measuring accuracy are optimal from Bayesian point of view, but are not so accurate as those based on information theory that may be evaluated with greater precision. Choice of the test is always a hidden compromise in how much weight is put on the purity of samples in one or both nodes and the total gain achieved by partitioning of data. It is therefore worthwhile to test other types of entropies that may be used as tests. Essentially the same reasoning may be used in applications of entropy-based indices in feature selection. For some data features that can help to distinguish rare cases are important, but standard approaches may rank them quite low.

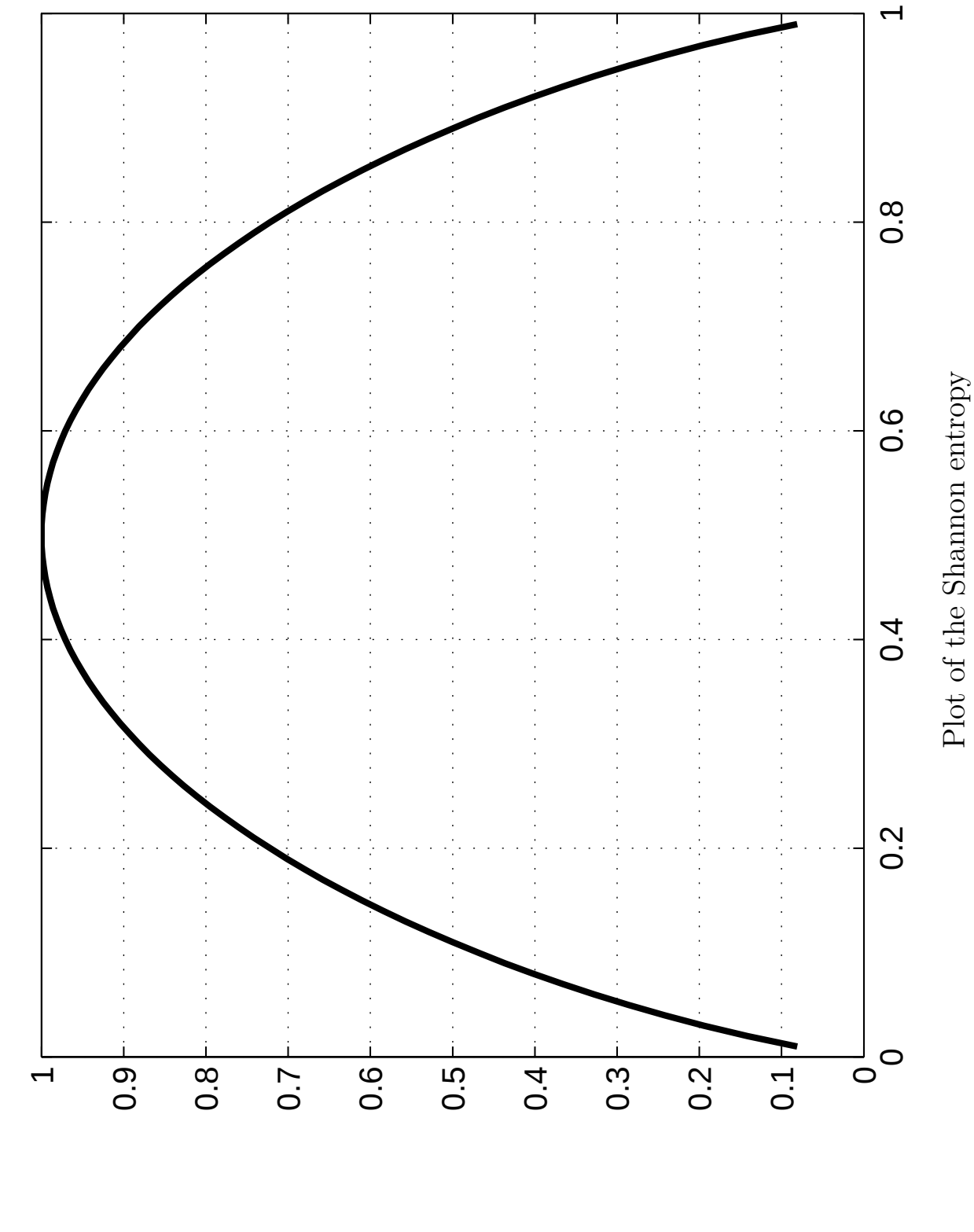
Theoretical framework

Entropy is the measure of disorder in physical systems, or an amount of information that may be gained by observations of disordered systems. Claude Shannon defined a formal measure of entropy, called Shannon entropy:

$$S = - \sum_{i=1}^n p_i \log_2 p_i$$

where p_i is the probability of occurrence of an event (feature value) x_i , being an element of the event (feature) X that can take values $\{x_1, \dots, x_n\}$.

The Shannon entropy is a decreasing function of a scattering of random variable, and is maximal when all the outcomes are equally likely.



Shannon entropy is may be used globally, for the whole data, or locally, to evaluate entropy of probability density distributions around some points. This notion of entropy can be generalized to provide additional information about the importance of specific events, for example outliers or rare events. Comparing entropy of two distributions, corresponding for example to two features, Shannon entropy assumes implicit certain tradeoff between contributions from the tails and the main mass of this distribution.

It should be worthwhile to control this tradeoff explicitly, as in many cases it may be important to distinguish weak signal overlapping with much stronger one. Entropy measures that depend on powers of probability, $\sum_{i=1}^n p(x_i)^\alpha$, provide such control. If α has large positive value this measure is more sensitive to events that occur often, while for large negative α it is more sensitive to the events which happen seldom. Constantino Tsallis and Alfred Renyi both proposed generalized entropies that for $\alpha = 1$ reduce to the Shannon entropy.

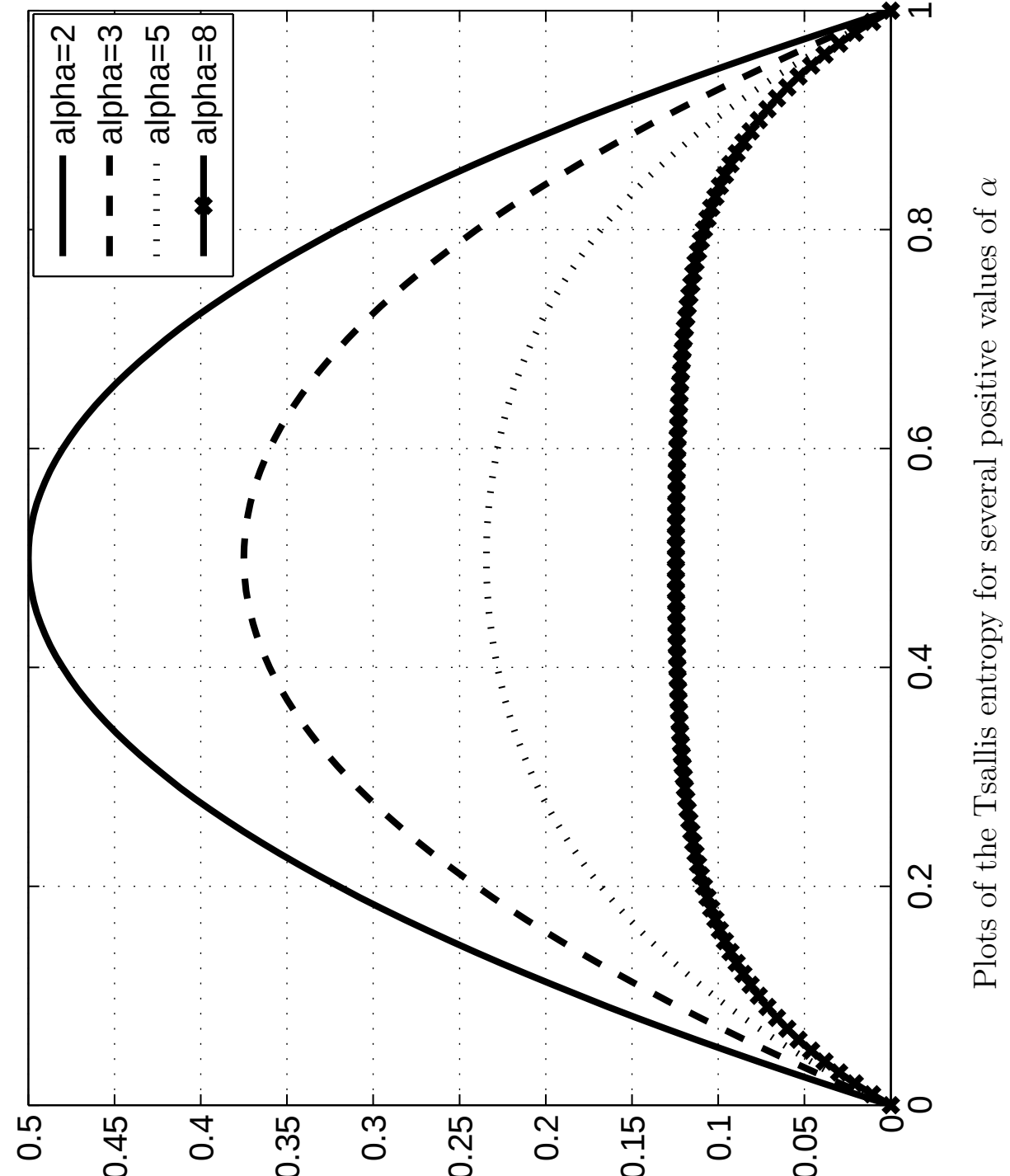
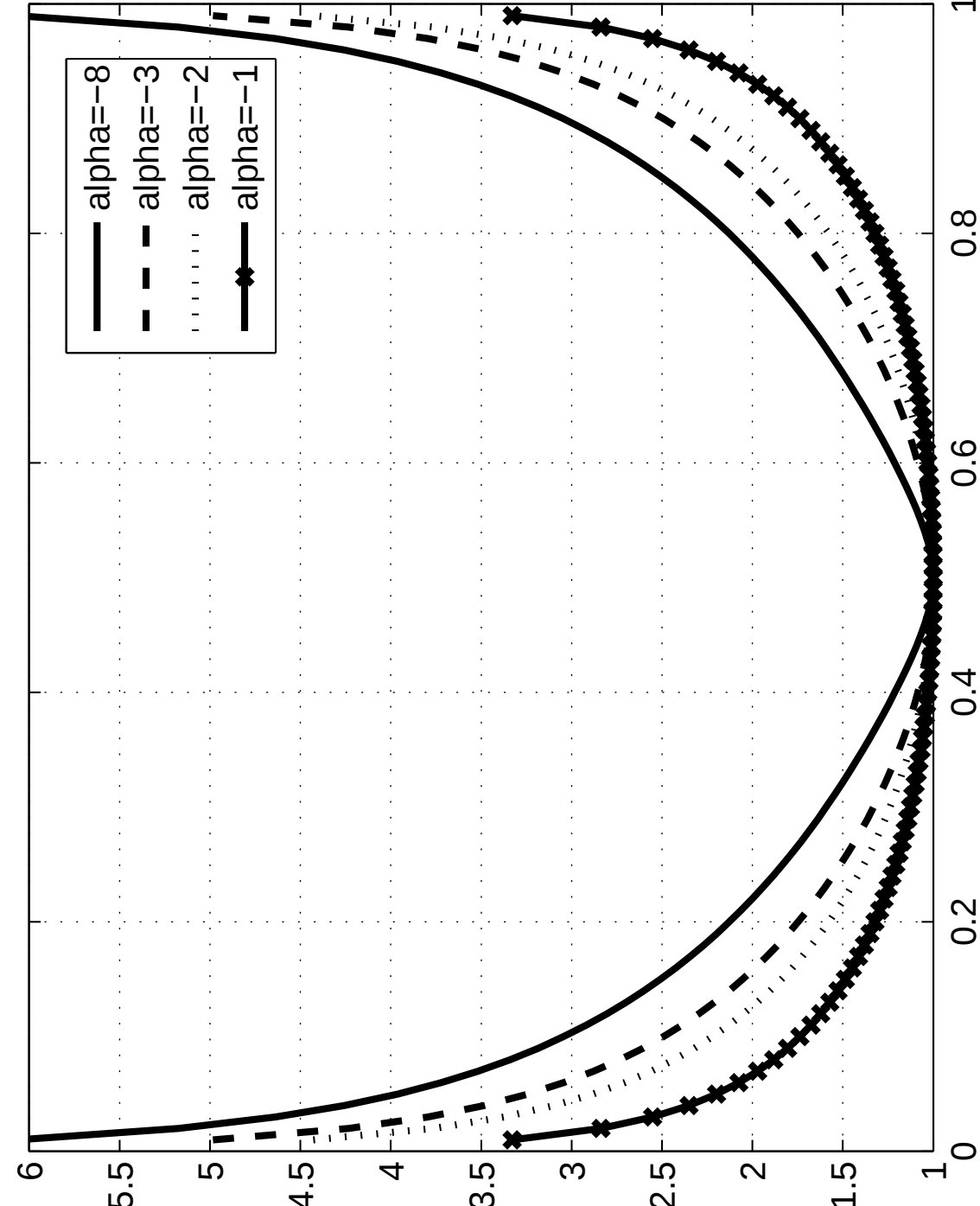
The Renyi entropy is defined as:

$$I_\alpha = \frac{1}{1-\alpha} \log \left(\sum_{i=1}^n p_i^\alpha \right)$$

Tsallis defined his entropy as:

$$S_\alpha = \frac{1}{\alpha-1} \left(1 - \sum_{i=1}^n p_i^\alpha \right)$$

They has similar properties as the Shannon entropy but contains additional parameter α which can be used to make it more or less sensitive to the shape of probability distributions.



The modification of the standard C4.5 algorithm has been done by simply replacing the Shannon measure with one of the two other entropies, as the goal here was to evaluate their influence on the properties of decision trees. This means that the final split criterion is based on the gain ratio: a test on attribute A that partitions the data D in two branches with D_i and D_j data, with a set of classes *omega* has gain value:

$$G(\omega, A|D) = H(\omega|D) - \frac{|D_i|}{|D|} H(\omega|D_i) - \frac{|D_j|}{|D|} H(\omega|D_j)$$

where $|D|$ is the number of elements in the D set and $H(\omega|S)$ is one of the 3 entropies considered here: Shannon, Renyi or Tsallis. Parameter α has clearly an influence on what type of splits are going to be created, with preference for negative values of α given to rare events or longer tails of probability distribution.

Results

To evaluate the usefulness of Renyi and Tsallis entropy measures in decision trees the classical C4.5 algorithm has been modified and applied first to artificial data to verify its usefulness. The results were encouraging, therefore experiments on three data sets of gene expression profiles were carried out. Such data are characterized by large number of features and very small number of samples. In such situations several features may by pure statistical chance seem to be quite informative and allow for good generalization. Therefore it is deceivably simple to reach high accuracy on such data, although it is very difficult to classify these data reliably. Only the simplest models may avoid overfitting, therefore decision trees providing simple rules may have an advantage over other models.

Entropy	Class	-0.1	0.1	0.3	0.5	0.7	0.9
Renyi	1	77.3±4.1	75.4±2.1	77.7±5.3	77.8±4.7	79.1±2.6	78.8±4.4
	Tsallis	64.6±0.2	77.3±3.7	75.4±4.0	74.4±4.3	71.3±5.4	73.0±3.4
	Class	1.1	1.5	2.0	3.0	4.0	5.0
	Renyi	82.1±4.2	82.9±2.5	84.0±3.9	79.4±3.0	80.8±3.1	78.9±2.2
Tsallis	1	74.9±1.8	71.1±4.0	70.2±3.9	73.9±4.4	72.8±3.6	71.1±4.4
	Shannon	81.2±3.7					

Accuracy on Colon cancer data set; Shannon and $\alpha = 1$ results are identical.

Entropy	Class	-0.1	0.1	0.3	0.5	0.7	0.9
Renyi	1	59.7±4.7	58.7±6.8	69.8±6.4	63.2±7.3	66.0±6.5	65.8±6.5
	2	87.2±4.1	84.7±2.4	87.2±2.8	85.8±4.0	86.2±3.9	85.8±4.6
Tsallis	1	0.0±0.0	58.2±5.1	59.8±9.3	59.8±5.2	50.7±10.2	55.7±7.7
	2	100±0.0	87.6±4.8	83.9±4.3	82.8±4.4	82.6±3.9	82.8±3.4
	Class	1.1	1.5	2.0	3.0	4.0	5.0
	Renyi	1	70.0±7.2	67.3±5.4	69.2±6.6	58.5±2.9	61.0±3.8
Tsallis	1	88.5±4.5	91.5±2.1	92.1±2.9	90.7±4.3	91.6±3.7	90.1±3.6
	2	58.5±4.9	58.3±8.2	53.2±7.6	67.2±9.2	60.5±7.4	60.0±10.4
Shannon	1	84.2±2.2	78.9±4.6	80.0±3.7	77.7±6.0	79.7±5.3	77.3±3.4
	2	69.5±4.2					

Conclusions

First experiments with modified split criterion of the C4.5 decision trees were presented here. Theoretical results and tests on artificial data show that this can be useful particularly for datasets with one or more small classes. Additional parameter α that may be easily adapted using crossvalidation tests gives possibility to tune the tree to the discrimination of classes of different size. This algorithm makes it more attractive than standard approach based on Shannon entropy that does not allow for exploration of the tradeoff between the probability of different classes and the overall information gain. This opens a way to applications in many types of decision trees, encouraging also modification of split criteria that are not based on entropy to account for the same tradeoff. An explicit formula for such tradeoff may be defined, aimed at separation of nodes of high purity at the expense of lower overall information gain. This as well as applications of non-standard entropies to the text classification data, where small classes are very important, remain to be explored. Bearing in mind the results and the simplicity of the approach presented here the use of non-standard entropies may be highly recommended.