



Leave-two-out stability of ontology learning algorithm[☆]



Jianzhang Wu^{a,*}, Xiao Yu^b, Linli Zhu^c, Wei Gao^d

^aSchool of Computer Science and Engineer, Southeast University, Nanjing 210096, China

^bSchool of Continuing Education, Southeast University, Nanjing 210096, China

^cSchool of Computer Engineering, Jiangsu University of Technology, Changzhou, Jiangsu 213001, China

^dSchool of Information Science and Technology, Yunnan Normal University, Kunming 650500, China

ARTICLE INFO

Article history:

Available online 7 January 2016

Keywords:

Ontology

Similarity measure

Learning algorithm

Ontology loss function

Stability

ABSTRACT

Ontology is a semantic analysis and calculation model, which has been applied to many subjects. Ontology similarity calculation and ontology mapping are employed as machine learning approaches. The purpose of this paper is to study the leave-two-out stability of ontology learning algorithm. Several leave-two-out stabilities are defined in ontology learning setting and the relationship among these stabilities are presented. Furthermore, the results manifested reveal that leave-two-out stability is a sufficient and necessary condition for ontology learning algorithm.

© 2015 Elsevier Ltd. All rights reserved.

1. Introduction

It is in philosophy that the term “ontology” is first applied to describe the connection nature of things and the inherently hidden connections of their components. Ontology, being a model for storing and representing knowledge, has been widely applied in knowledge management, machine learning, information systems, image retrieval, information retrieval search extension, collaboration and intelligent information integration in information and computer science. Meanwhile, ontology is an effective concept semantic model and a powerful analysis tool. It has been used extensively in pharmacology, biology, medicine, geographic information system and social science in the past ten years (see Przydzial et al., [1], Koehler et al., [2], Ivanovic and Budimac [3], Hristoskova et al., [4], and Kabir [5]).

A simple graph can be used to express the structure of ontology. On that graph, each vertex represents a concept, object or element in ontology. Each (directed or undirected) edge refers to a relationship or hidden connection between two concepts (objects or elements). Let O be an ontology and G be a simple graph of O . The purpose of ontology engineer application is to get the similarity calculating function and then compute the similarities between ontology vertices. The inherent association between vertices in ontology graph can be illustrated by these similarities. Ontology mapping is to obtain the ontology similarity measuring function by measuring the similarity between vertices from different ontologies. Such a mapping connects different ontologies, through which a potential link between the objects or elements from different ontologies can be acquired. The ontology similarity function $Sim : V \times V \rightarrow \mathbb{R}^+ \cup \{0\}$ is a semi-positive score function mapping in which each pair of vertices maps to a non-negative real number.

An advanced usage of handling the ontology similarity computation is using ontology learning algorithm which gets an ontology function $f : V \rightarrow \mathbb{R}$. Using such an ontology function, the ontology graph is mapped into a line consisting of real numbers. After comparing the difference between their corresponding real numbers, the similarity

[☆] This work was supported in part by the National Natural Science Foundation of China (60903131), Science and Technology of Jiangsu Province (BE2011173), and Key Laboratory of Computer Network and Information Integration Founding in Southeast University.

* Corresponding author. Tel.: +86 4403877466372828.

E-mail address: jzwu@njnet.edu.cn (J. Wu).

between two concepts can be measured and in which the dimensionality reduction is the ore of the idea. If the ontology function is to be associated with ontology application, a vector is a very good choice as it expresses all the information of a vertex, for instance v . In a simpler representation, the notations are slightly confused and v is used to denote both the ontology vertex and its corresponding vector. The vector is mapped to a real number by ontology function $f: V \rightarrow \mathbb{R}$. The ontology function, which is a dimensionality reduction operator, maps vectors of multi-dimension into one dimensional ones.

All the related information of an arbitrary vertex in ontology graph G is expressed by a p dimensional vector, which includes its instance, structure, name, attribute, and semantic information of the concept which is corresponding to the vertex and that is contained in its vector. In order not to lose generality, it can be assumed that $v = \{v_1, \dots, v_p\}$ is a vector corresponding to a vertex v . Their notations are slightly confused and v is adopted to represent both the ontology vertex and its corresponding vector. In order to obtain an optimal ontology function $f: V \rightarrow \mathbb{R}$, ontology learning algorithms are used by the authors. Therefore, the value of $|f(v_i) - f(v_j)|$ is used to determine the similarity between two vertices v_i and v_j . Dimensionality reduction, i.e., using real number to represent p dimension vector is the core of such kind of ontology algorithm. In this way, we can regard an ontology function f as a dimensionality reduction operator $f: \mathbb{R}^p \rightarrow \mathbb{R}$.

There are many effective methods for getting efficient ontology similarity measure or ontology mapping algorithm. They have been studied in terms of ontology function. Moreover, the theoretical research of ontology algorithms has been contributed by several researchers. The uniform stability of multi-dividing ontology algorithm and the generalization bounds for stable multi-dividing ontology algorithms was put forth by Gao and Xu [6]. A gradient learning model for ontology similarity measuring and ontology mapping in multi-dividing setting was proposed by Gao and Zhu [7] cooperatively. In the setting, the sample error was determined in terms of the hypothesis space and the ontology dividing operator in which one can suppose that V is an instance space.

In this article, we research the influences of ontology learning algorithm when two sample vertices are deleted from ontology sample set. In next section, we describe the detailed notations, definitions and setting of ontology learning problem. And then, the main conclusions are drawn in Section 3.

2. The notations, definitions and setting of ontology learning problem

Suppose that V is a compact domain in Euclidean space and Y is a set of labels. Let $\mu(v, y)$ be an unknown probability distribution on $Z = V \times Y$ and $S = \{(v_1, y_1), \dots, (v_n, y_n)\} = (v_i, y_i)_{i=1}^n = (z_i)_{i=1}^n$ be an ontology sample set consisting of n samples drawn i.i.d. from the probability distribution on Z_n . The aim of ontology learning is to predict an ontology function $f_S: V \rightarrow \mathbb{R}$ using the empirical ontology data S which evaluates at a new ontology vertex v to predict its corresponding value of y .

Let $L: \mathbb{R}^V \times V \times \mathbb{R} \rightarrow \mathbb{R}$ be the ontology loss function and $L(f, z)$ be the value of punishment for fixed ontology function f and $z = (v, y)$. Throughout our paper, we always assume that the loss function L is square ontology loss $L(f, z) = (f(v) - y)^2$ and there exists M which satisfies $0 \leq L(f, z) \leq M$ for any $f \in \mathcal{F}$ (here $\mathcal{F}$ is a hypothesis space in ontology setting) and $z \in Z$. Denote $l(z) = L(f, z)$ for convenience. Thus $l(z): V \times Y \rightarrow \mathbb{R}$ and we set $\mathcal{L} = \{l(f): f \in \mathcal{F}\}$ as the space of ontology loss function.

The ontology expected error for fixed ontology function f , ontology loss function L and a probability distribution μ is defined by: $R(f) = \mathbb{E}_Z L(f, z)$. When L is square ontology loss, we have

$$\begin{aligned} R(f) &= \mathbb{E}_Z L(f, z) = \int_{V,Y} (f(v) - y)^2 d\mu(v, y) \\ &= \mathbb{E}_\mu |f(v) - y|^2. \end{aligned}$$

However, $R(f)$ can't be calculated directly since μ is unknown. In reality, we compute the ontology empirical error instead which is presented as $\hat{R}_S(f) = \frac{1}{n} \sum_{i=1}^n L(f, z_i)$. In addition, in our square ontology loss setting, it is equal to $\hat{R}_S(f) = \frac{1}{n} \sum_{i=1}^n (f(v_i) - y_i)^2 = \mathbb{E}_{\mu_n} (f(v) - y)^2$, where μ_n is the ontology empirical supported on $\{v_1, \dots, v_n\}$ which means $\mu_n = \sum_{i=1}^n \delta_{v_i}/n$ and δ_{v_i} is the vertex evaluation functional on v_i .

In what follows, set $S^{i,j}$ as the ontology training set from S by deleting two vertices v_i and v_j ($1 \leq i < j \leq n$). For our ontology learning setting, the functions f_S and $f_{S^{i,j}}$ are the minimizers of $\hat{R}_S(f)$ and $\hat{R}_{S^{i,j}}(f)$, respectively. The notations $\mathbb{E}_S$ and $\mathbb{P}_S$ are used to express the expectation and the probability on the ontology training set S which is drawn i.i.d, according to probability distribution on Z_n .

An ontology algorithm is called symmetric if the optimal ontology function can't be changed when the elements in training set S are re-arranged. Given an ontology training set S and an ontology function space $\mathcal{F}$, the almost ontology learning algorithm is defined as a symmetric procedure which chooses an ontology function $f_S^{\mathcal{E}}$ that minimizes the ontology empirical risk over all ontology functions $f \in \mathcal{F}$, we infer

$$\hat{R}_S(f_S^{\mathcal{E}}) \leq \inf_{f \in \mathcal{F}} \hat{R}_S(f) + \mathcal{E}^E, \quad (1)$$

where $\mathcal{E}^E > 0$.

An ontology learning map is called (universally, weakly) consistent if for any positive number $\varepsilon_c > 0$, we have

$$\lim_{n \rightarrow \infty} \sup_{\mu} \mathbb{P}\{R(f_S) > \inf_{f \in \mathcal{F}} R(f) + \varepsilon_c\} = 0.$$

The consistency is universal implies that the inequality above is established with respect to the set of any measure on Z , whereas weak consistency means only convergence in probability and strong consistency requires almost sure convergence.

Let $\mathcal{F}$ be any class of functions. $\mathcal{F}$ is a (weak) uniform Glivenko–Cantelli (in short, uGC) class if for any positive number ε ,

$$\lim_{n \rightarrow \infty} \sup_{\mu} \mathbb{P}\{\sup_{f \in \mathcal{F}} |\mathbb{E}_{\mu_n} f - \mathbb{E}_{\mu} f| > \varepsilon\} = 0.$$

When comes to the ontology loss functions l , the definition implies that for all distributions μ there exist ε_n and $\delta_{\varepsilon_n, n}$

Download English Version:

<https://daneshyari.com/en/article/1895407>

Download Persian Version:

<https://daneshyari.com/article/1895407>

[Daneshyari.com](https://daneshyari.com)