

Correcting evaluation bias of relational classifiers with network cross validation

Jennifer Neville · Brian Gallagher · Tina Eliassi-Rad ·
Tao Wang

Received: 11 January 2010 / Revised: 18 September 2010 / Accepted: 13 December 2010
© The Author(s) 2010. This article is published with open access at Springerlink.com

Abstract Recently, a number of modeling techniques have been developed for data mining and machine learning in relational and network domains where the instances are not independent and identically distributed (i.i.d.). These methods specifically exploit the statistical dependencies among instances in order to improve classification accuracy. However, there has been little focus on how these same dependencies affect our ability to draw accurate conclusions about the performance of the models. More specifically, the complex link structure and attribute dependencies in relational data violate the assumptions of many conventional statistical tests and make it difficult to use these tests to assess the models in an unbiased manner. In this work, we examine the task of within-network classification and the question of whether two algorithms will learn models that will result in significantly different levels of performance. We show that the commonly used form of evaluation (paired *t*-test on overlapping network samples) can result in an unacceptable level of Type I error. Furthermore, we show that Type I error increases as (1) the correlation among instances increases and (2) the size of the evaluation set increases (i.e., the proportion of labeled nodes in the network decreases). We propose a method for network cross-validation that combined with paired *t*-tests produces more acceptable levels of Type I error while still providing reasonable levels of statistical power (i.e., 1–Type II error).

J. Neville (✉)

Departments of Computer Science and Statistics, Purdue University, West Lafayette, IN, USA
e-mail: neville@cs.purdue.edu

B. Gallagher

Lawrence Livermore National Laboratory, Livermore, CA, USA
e-mail: bgallagher@llnl.gov

T. Eliassi-Rad

Department of Computer Science, Rutgers University, Piscataway, NJ, USA
e-mail: eliassi@cs.rutgers.edu

T. Wang

Department of Computer Science, Purdue University, West Lafayette, IN, USA
e-mail: taowang@cs.purdue.edu

Report Documentation Page				Form Approved OMB No. 0704-0188	
Public reporting burden for the collection of information is estimated to average 1 hour per response, including the time for reviewing instructions, searching existing data sources, gathering and maintaining the data needed, and completing and reviewing the collection of information. Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing this burden, to Washington Headquarters Services, Directorate for Information Operations and Reports, 1215 Jefferson Davis Highway, Suite 1204, Arlington VA 22202-4302. Respondents should be aware that notwithstanding any other provision of law, no person shall be subject to a penalty for failing to comply with a collection of information if it does not display a currently valid OMB control number.					
1. REPORT DATE SEP 2010		2. REPORT TYPE		3. DATES COVERED 00-00-2010 to 00-00-2010	
4. TITLE AND SUBTITLE Correcting evaluation bias of relational classifiers with network cross validation				5a. CONTRACT NUMBER	
				5b. GRANT NUMBER	
				5c. PROGRAM ELEMENT NUMBER	
6. AUTHOR(S)				5d. PROJECT NUMBER	
				5e. TASK NUMBER	
				5f. WORK UNIT NUMBER	
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES) Purdue University, Computer Science Department, West Lafayette, IN, 47907				8. PERFORMING ORGANIZATION REPORT NUMBER	
9. SPONSORING/MONITORING AGENCY NAME(S) AND ADDRESS(ES)				10. SPONSOR/MONITOR'S ACRONYM(S)	
				11. SPONSOR/MONITOR'S REPORT NUMBER(S)	
12. DISTRIBUTION/AVAILABILITY STATEMENT Approved for public release; distribution unlimited					
13. SUPPLEMENTARY NOTES					
14. ABSTRACT Recently, a number of modeling techniques have been developed for data mining and machine learning in relational and network domains where the instances are not independent and identically distributed (i.i.d.). These methods specifically exploit the statistical dependencies among instances in order to improve classification accuracy. However, there has been little focus on how these same dependencies affect our ability to draw accurate conclusions about the performance of the models. More specifically, the complex link structure and attribute dependencies in relational data violate the assumptions of many conventional statistical tests and make it difficult to use these tests to assess the models in an unbiased manner. In this work, we examine the task of within-network classification and the question of whether two algorithms will learn models that will result in significantly different levels of performance. We show that the commonly used form of evaluation (paired t-test on overlapping network samples) can result in an unacceptable level of Type I error. Furthermore we show that Type I error increases as (1) the correlation among instances increases and (2) the size of the evaluation set increases (i.e., the proportion of labeled nodes in the network decreases). We propose a method for network cross-validation that combined with paired t-tests produces more acceptable levels of Type I error while still providing reasonable levels of statistical power (i.e., 1&#8722;Type II error).					
15. SUBJECT TERMS					
16. SECURITY CLASSIFICATION OF:			17. LIMITATION OF ABSTRACT Same as Report (SAR)	18. NUMBER OF PAGES 25	19a. NAME OF RESPONSIBLE PERSON
a. REPORT unclassified	b. ABSTRACT unclassified	c. THIS PAGE unclassified			

Keywords Relational learning · Collective classification · Statistical tests · Methodology

1 Introduction

The seminal work of Dietterich [5] focused on enumerating the types of statistical questions that analysts could ask of the models and algorithms that they develop and/or learn. His work outlined a taxonomy of questions that differentiates between algorithm and model performance, and whether the goal is to estimate accuracy or to choose between models/algorithms. Within this taxonomy, Dietterich formulated the question that is most central to data mining and machine learning research: *Given two learning algorithms A and B and a dataset of size S from a domain D, which algorithm will produce more accurate classifiers when trained on other datasets of size S drawn from D?* This question explicitly formulates the notion of generalization and provides a means to test the notion statistically.

Within this framework, Dietterich analyzed the characteristics of five statistical tests that can be used to assess generalization performance and showed that two of the tests in widespread use (at that time) had a high probability of Type I error (i.e., the tests will likely lead to an erroneous conclusion of algorithm difference when there is none). Overall, Dietterich's work showed that the overlap in training/test sets combined with *imbalanced* samples can lead to higher Type I errors due to biased estimates of mean performance difference between two algorithms. Therefore, a methodology that reduces the overlap between the training and test sets leads to lower Type I errors. Based on this analysis, Dietterich developed a novel 5×2 cross-validation test, which has lower Type I error than the standard cross-validation test but slightly reduced statistical power (i.e., higher Type II error).

However, Dietterich's work only considered conventional machine learning algorithms that assume i.i.d. data. In this work, we consider the task of comparing algorithm performance in the context of *within-network* relational learning. In contrast to *across-network* relational learning,¹ in which a model is learned on one network and then applied to another disconnected network, *within-network* learning aims to generalize within a single relational data graph. In within-network learning, models are learned on a partially labeled network and then applied to predict the class labels in the remainder of the network (i.e., the unlabeled portion). In many real world applications, relational learning tasks fall naturally into the within-network classification setting. For example, in the task of research paper classification, new papers to be classified usually have citation links to papers in the past whose topics are known. Similarly, in fraud detection, brokers whose fraud status is yet to be determined might associate with other brokers who have already been identified as fraudulent or not.

Within-network relational learning tasks have two characteristics that can complicate the application of conventional statistical tests for comparing generalization performance. First, the instances in the network are not independent. Indeed, relational learning algorithms are specifically trying to exploit the dependencies among instances to improve prediction accuracy. The dependencies among instances, however, tend to result in correlated errors among the instances. These correlated errors can increase the imbalance between network samples, and this can lead to increased Type I errors. Second, the size of the training and test sets are dependent, and thus as the proportion of (labeled) training data decreases, the size of the test set increases. This results from the fact that the models are learned/applied to a partially labeled network with varying levels of labeled instances, and the full set of unlabeled

¹ We use the more general term *relational learning* to refer to both within-network and across-network learning.

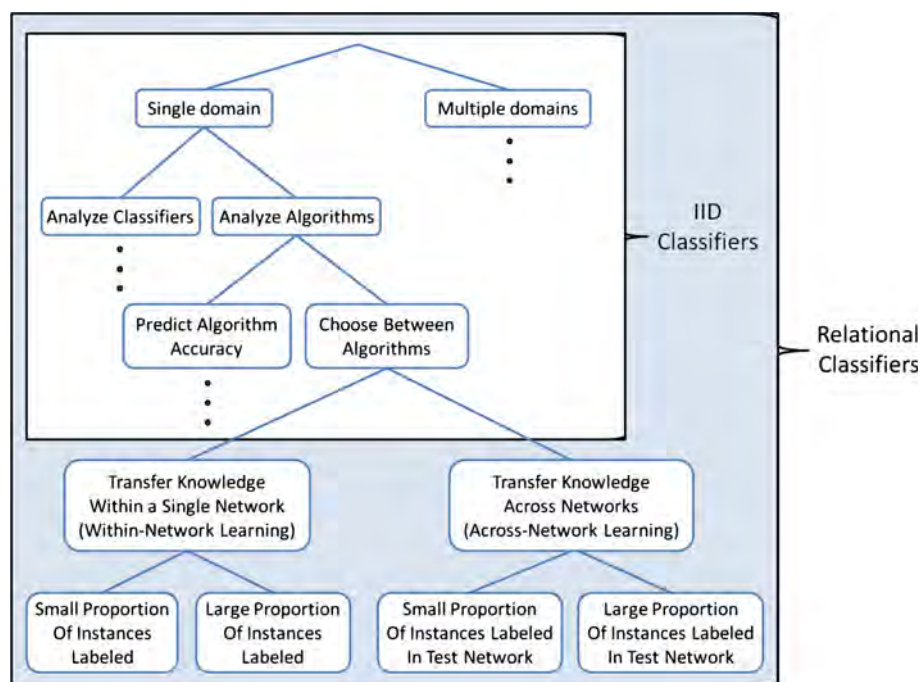


Fig. 1 A taxonomy of statistical questions in statistical relational machine learning. The white box contains Dietterich's original taxonomy of questions for i.i.d. classifiers. All questions in the hierarchy are applicable to relational classifiers

instances are typically used for evaluation. As the size of the unlabeled set increases, the dependencies between samples increase, and this can also lead to increased Type I errors.

Figure 1 outlines a taxonomy of statistical questions for relational classifiers that builds on the taxonomy proposed by Dietterich [5]. Our taxonomy includes two additional dimensions of variation: (1) transfer of knowledge within a single network (i.e., within-network classification) versus across multiple networks (i.e., across-network classification) and (2) the proportion of known labels available in the test network (small vs. large). While the focus of this study is within-network classification, our findings will also apply to across-network classification if the test networks are partially labeled. The distinction between a small versus large proportion of available labels is important because with a sufficiently small label set (i.e., $\leq 50\%$), standard cross-validation is no longer applicable as test sets will start to overlap—leading to an increase in Type I error. Our study considers a range of labeling proportions, spanning the “small” and “large” label-proportion categories.

More specifically, in this paper, we consider the following question: *Given two learning algorithms A and B and a partially labeled network from domain D , and with S_L labeled instances and S_U unlabeled instances ($S = S_L + S_U$), which algorithm will produce more accurate classifiers when trained on other partially labeled networks of size S drawn from D ?* We consider statistical tests that can be used to evaluate this question and discuss the characteristics of relational data that can complicate accurate evaluation. We show theoretically how these network characteristics can lead to increased Type I error, then we investigate the nature of this bias empirically. Our empirical investigation compares the performance of a number of common statistical tests, using both simulated and real classifiers, and both

synthetic and real datasets. The experimental methodology and results for each combination can be found in the following sections:

- Theoretical Analysis: *Sect. 4*
- Simulated Classifiers/Synthetic Data: *Sect. 5*
- Real Classifiers/Synthetic Data: *Sect. 6*
- Real Classifiers/Real Data: *Sect. 7*

Our findings indicate that a commonly used method of statistical assessment—paired *t*-tests on repeated samples of randomly selected network samples (labeled training set and unlabeled test set)—results in unacceptably high levels of Type I error. We propose a method for *network cross-validation* that combined with unpaired *t*-tests produces low levels of Type I error at the expense of reduced statistical power. Combining network cross-validation with paired *t*-tests is a good compromise, resulting in both acceptable levels of Type I error and reasonable levels of statistical power.

This article builds on preliminary work [24]. The most significant additions for this article include:

- Development of a taxonomy of statistical questions for comparing network classifiers.
- Development of a theoretical explanation of how relational data characteristics lead to type I error in standard statistical tests.

Additional contributions of the overall work include:

- Formulation of important statistical questions for comparing network classifiers.
- Discussion and demonstration of the challenges of relational data for using conventional statistical tests to compare classifiers.
- A proposed solution, *network cross-validation*, which addresses these challenges.
- Empirical evaluation of statistical test characteristics (Type I error/power) on real-world and synthetic data.

2 Comparing classifiers in network domains

Statistical tests for comparing classifiers generally consist of two parts: (1) The *resampling procedure* dictates how the available data are partitioned into training and test sets for estimation of classifier performance (i.e., *how many times is the classifier trained and tested?*, *which data are used to train the classifier?*, *which data are used to test the classifier?*) and (2) The *significance test* takes the classification results from the resampling trials and makes a determination as to whether observed differences reflect a true difference in classifier performance or whether it is likely to have occurred by chance alone.

2.1 Resampling procedures

Given a fully labeled network of size S , we consider three resampling procedures to generate training (labeled set S_L) and test (unlabeled set S_U) sets to evaluate within-network classification algorithms: *simple random resampling* (RRS), *equal-instance random resampling* (ERS), and *network cross-validation* (NCV). The first two methods have been used extensively in past work on relational learning algorithms (see Sect. 3 for more detail). The third method is a new approach, based on the incremental cross-validation procedure outlined in Cohen [2] for generating learning curves, which we will show is more robust to Type I error.

Table 1 Simple random resampling procedure

input: *network*, *propLabeled*, *k*
 S = total number of instances in *network*
 $F = \emptyset$
for *fold* 1 to *k*
 testSet = uniform random sample of $((1 - \text{propLabeled}) * S)$ nodes from *network*
 trainSet = *network* – *testSet*
 $F = F \cup \langle \text{trainSet}, \text{testSet} \rangle$
end for
output: *F*

Table 2 Equal-instance resampling procedure

input: *network*, *propLabeled*, *k*
 S = total number of instances in *network*
// Split data into overlapping folds so that each instance occurs in the same number of folds.
testSetSize = $(1 - \text{propLabeled}) * S$
 $\text{numCopies} = \frac{k * \text{testSetSize}}{S}$
pool = sorted list with *numCopies* of each instance in *network*
testSet[*i*] = {}, for *i* = 1 to *k*
for *instance* $\in$ *pool*
 add *instance* to smallest *testSet*[*i*] that does not already contain it
end for
// create training/test splits
 $F = \emptyset$
for *testSet* 1 to *k*
 trainSet = *network* – *testSet*
 $F = F \cup \langle \text{trainSet}, \text{testSet} \rangle$
end for
output: *F*

Tables 1 and 2 outline the procedures for RRS and ERS, respectively. Both methods involve repeated random draws from the sample population to generate the training/test splits and therefore produce overlapping test sets. However, ERS ensures that each instance in the original sample occurs in exactly the same number of test sets in the collection of resamples.

Table 3 outlines the NCV procedure that eliminates overlap between test sets altogether. The procedure samples for *k* disjoint test sets that will be used for evaluation. Then for each test set fold, the remaining folds are merged together, and the training set of size S_L is randomly sampled from the merged set. When the training set size is less than the size of the merged folds (i.e., $S_L < (k - 1) \frac{S}{k}$), this will leave a set of unlabeled nodes that are neither in the test set nor the training set. Since these unlabeled instances will likely be connected to nodes in the test set, we will run collective inference over the full set of unlabeled nodes (the *inference* set), and then only evaluate model performance on the nodes assigned to the test set.

Table 3 Network cross-validation procedure

given: *network*, *propLabeled*, *k*
 S = total number of instances in *network*
 $F = \emptyset$
Split data into k disjoint folds
for *fold* 1 to k
 current *fold* becomes *testSet*
 remaining folds are merged and become *trainPool*
 trainSet = $\bar{\text{uniform}}$ random sample of $(\text{propLabeled} * S)$ nodes from *trainPool*
 inferenceSet = *network* – *trainSet*
 $F = F \cup \langle \text{trainSet}, \text{testSet}, \text{inferenceSet} \rangle$
end for
output: F

NCV addresses a limitation of standard cross-validation for within-network classification tasks. Namely, standard CV forces us to label $k - 1$ of every k instances. So, if $k = 10$, we are forced to experiment with 90% labeled data. The NCV approach accommodates a lower proportion of labeled instances because it samples a smaller labeled set from the $k - 1$ non-test folds. Since NCV applies collective inference to the full unlabeled portion of the network, it can fully exploit the network structure to improve model predictions. However, since NCV only *evaluates* the model on the disjoint test set instances, it will not suffer the same problems experienced by resampling due to overlapping test sets.

2.2 Significance tests

Once a sampling procedure is chosen to create training/test splits within a network, the algorithms are learned on each training set, and then the models are applied for collective inference over the associated unlabeled portion of the network. The predictions on the test set instances are evaluated to generate an estimate of algorithm performance (e.g., accuracy, AUC, squared loss). The training/test splits result in a set of performance measurements for each algorithm, and a significance test is then used to determine whether the observed performance differences are *significantly* different than what would be expected if the performance measures were drawn from the same underlying distribution (i.e., the algorithms perform equivalently).

In this work, we investigated the following three statistical tests: (1) paired t -test, (2) unpaired t -test, and (3) Wilcoxon signed rank test. Both the paired t -test and Wilcoxon signed rank test assume independence of the paired differences between classifiers. We observed no substantive differences due to the use of the Wilcoxon test versus the t -test. Therefore, we focus on the more commonly used t -test for the remainder of this paper.

3 Survey of previous methodology

Over the past 10 years, there has been a great deal of work on classifiers for relational domains. Here, we survey the methodological design of 23 research papers most relevant to our work [1, 7–12, 16–21, 25–29, 31, 34, 35, 38, 40]. Relevant papers compare the performance of two or more classifiers on a relational classification task. Based on our survey, there are two common evaluation methodologies which emerge:

Independent set size The salient feature of the independent set size methodology is that there is no dependency between training and test set sizes. The task may be either within-network or across-network classification (i.e., there may or may not be relations between instances in the labeled and unlabeled set). For across-network classification, classifiers are generally trained on a fully labeled training network and then evaluated on a disjoint (partially labeled) test network. In this case, the proportion of labeled data available in the training and test networks may be varied independently. This means that there is no dependence between training and test set sizes. For within-network classification, there is a single network, and classifiers are trained on the labeled instances and evaluated on the unlabeled instances. So, the only way to achieve independent set sizes in within-network classification is to *fix* the proportion of labeled data available in the network. All of the papers in our survey that employ independent set sizes (15 out of the 23) use some form of random resampling or cross-validation for evaluation.

Dependent set size This methodology applies to within-network classification only, where training and test sizes are dependent. Any change in the labeling proportion over the network will affect the number of instances available for both training and testing. All of the papers in our survey that employ dependent set sizes (9 out of the 23) use some sort of random resampling for evaluation (i.e., test sets overlap). Note that standard cross-validation is not possible here since it assumes a fixed proportion of labeled data.

Both of the above methodologies generally address the statistical question that we consider in our work: *Given two learning algorithms A and B and a partially labeled network from domain D , and with S_L labeled instances and S_U unlabeled instances ($S = S_L + S_U$), which algorithm will produce more accurate classifiers when trained on other partially labeled networks of size S drawn from D ?* However, the independent set size methodology makes a rather strong assumption about the value of S_L . In particular, most studies that employ independent set sizes use 10-fold cross-validation, which means that their results only generalize to graphs where 90% of the data are labeled to begin with. This is limiting because: (1) most interesting real-world problems have far less than 90% of data labeled to begin with and (2) many algorithms that perform well at 90% labeled will perform poorly at sparser labelings (e.g., 10%). The dependent set size methodology is more powerful since it generalizes over different values of S_L , so we focus on this version in our work.

Table 4 provides a summary of related work along a number of methodological dimensions. Note that counts in each section do not necessarily sum to 23 since papers may fit in more than one category. The majority of the studies in our survey (13/23) make use of resampling procedures that produce overlap between test sets. This includes all resampling procedures except cross-validation and temporal sampling. *Temporal sampling* involves training on past instances (e.g., previously published papers with known topics) and evaluating on present instances (e.g., a newly submitted paper with unknown topic). *Controlled random sampling* procedures attempt to control or account for the amount of overlap between test sets (e.g., as in the *equal-instance* resampling procedure described in Sect. 2.1).

In our survey, within-network classification tasks are more common than across-network tasks (13 papers vs. 8). However, in a substantial number of cases, it is unclear exactly how the experiments are set up. For example, authors will often say something like: we split the network into training and test sets. It is not clear from this description whether the links are preserved between instances in the training set and instances in the test set. In other cases, authors are explicit regarding whether such links are retained or removed.

Table 4 Experimental characteristics of 23 surveyed papers

Resampling procedure		Systematic variation of % labeled	
Cross validation	14	No	13
Simple random	8	Yes	10
Controlled random	3	Within-network	8
Snowball sampling	2	Across network	2
Temporal resampling	2		
Statistical test		Number of resampling folds	
t-test	10	10	14
StDev/Var/StErr	6	<10	7
None	6	>10	2
Wilcoxon signed rank	2	Unspecified	2
Within versus across network classification		Performance measure	
Within-network	13	Accuracy	14
Across network	8	AUC	10
Unspecified	6	Precision/Recall/F1	1

Note that section counts do not necessary sum to 23 since papers can fit in >1 category

The majority of studies in our survey (13/23) do not vary the proportion of labeled data available. Of the studies that do vary the proportion of labeled data, most (8/10) are within-network studies (i.e., *dependent set size*).

The most common number of resampling folds used in the surveyed papers is 10. As Dietterich notes in his original study, the probability of Type I error for random resampling procedures increases with the number of resampled folds [5]. Our simulation experiments confirm this finding, but we do not replicate the result here.

Half of the studies in our survey do not make use of an explicit significance test. However, of these, about half do report standard deviation, variance, or standard error (*StDev/Var/StErr* in Table 4). Finally, both accuracy and AUC are common measures of classifier performance, with precision/recall-based measures being much less common.

In addition to the related work discussed above, there have been other works on learning and generalization bounds from non-i.i.d. observations (e.g. [3, 4, 6, 14, 15, 22, 23, 30, 32, 33, 37, 39]). However, none of them address the problems of (1) dependence between related instances and (2) dependence between training and test set sizes. Usunier et al. [36] proposed a new framework to study the generalization properties of classifiers over data which can exhibit a suitable dependency structure. However, their focus was solely on dependent training sets.

4 Characteristics leading to bias in statistical tests

As discussed previously, there are two characteristics of within-network relational classification that can lead to increased levels of Type I error:

1. Inter-instance dependencies lead to correlated errors.
2. Small training sets lead to large test sets, increasing the dependence between samples.

Table 5 Error correlation and relational autocorrelation in real-world classification tasks

Data set	Task	Error corr.	Autocorr.
Enron email	Executive?	0.18	0.17
Citeseer	Neural nets?	0.23	0.59
Political books	Neutral?	0.25	0.22
Cora	Info. retrieval?	0.28	0.61
Reality mining	In study?	0.32	0.79
Reality mining	Student?	0.52	0.91

Here, we will discuss these two issues in more detail, providing both empirical and theoretical support for our claims.

First, we demonstrate that within-network classifiers produce correlated errors. To test the conjecture that within-network classifiers produce correlated errors, we experimented with several relational classifiers and real-world classification tasks, using the ϕ coefficient to measure the correlation of 0–1 errors over all pairs of related (i.e., linked) instances. We used a non-learning relational neighbor classifier [18] and a learning link-based classifier [17]. We ran each classifier both with and without collective inference on a number of prediction tasks. Table 5 shows the amount of measured error correlation for each task, averaged over all classifiers and proportions of labeled data (0.1, 0.3, 0.5, 0.7, 0.9). Although we report averages, we should note that all trials (i.e., tasks/classifiers) exhibited some degree of error correlation. We can observe that the level of error correlation is correlated with the level of relational autocorrelation (see e.g., [26]) in the class label. Since autocorrelation has been shown to be essentially ubiquitous in relational data, this suggests that error correlation is widespread as well.

Next, we show theoretically how error correlation increases the variance of average error—this increase in variance, if not accounted for in the statistical test will lead to increased Type I error.

Theorem 1 (Correlated errors increase variance) *Let X_i be a random variable that represents the classification error for an instance i and assume that each r.v. X_i is drawn from the same distribution with mean μ and variance σ^2 . Let $\bar{X} = \frac{1}{n} \sum_{i \in \mathbf{S}} X_i$ be a random variable that represents the average classification error in a sample $\mathbf{S} = \{X_i\}$ with n instances (i.e., $|\mathbf{S}| = n$). Assume that the sample $\mathbf{S}$ consists of n instances drawn equally from a set of groups $\mathbf{G}$, where $|\mathbf{G}| = m$ and each group $G_k \in \mathbf{G}$ has average size g (i.e., $m \cdot g = n$). Let ρ be the average correlation between the X_i that are members of the same group, otherwise assume the X_i are independent. Then the average classification error $\bar{X}$ has variance $\text{Var}(\bar{X}) = \frac{1}{n} \sigma^2 [1 + \rho(g - 1)]$.*

Proof

$$\begin{aligned} \text{Var}(\bar{X}) &= \text{Var}\left(\frac{1}{n} \sum_{X_i \in \mathbf{S}} X_i\right) \\ &= \frac{1}{n^2} \text{Var}\left(\sum_{X_i \in \mathbf{S}} X_i\right) \end{aligned}$$

$$\begin{aligned}
&= \frac{1}{n^2} \left(\sum_{X_i \in \mathbf{S}} \sum_{X_j \in \mathbf{S}} \text{Cov}(X_i, X_j) \right) \\
&= \frac{1}{n^2} \left(\sum_{X_i \in \mathbf{S}} \text{Var}(X_i) + \sum_{X_i \in \mathbf{S}} \sum_{X_j \in \mathbf{S}; j \neq i} \text{Cov}(X_i, X_j) \right) \\
&= \frac{1}{n^2} \left(\sum_{X_i \in \mathbf{S}} \sigma^2 + \sum_{G_k \in \mathbf{G}} \sum_{X_i \in G_k} \left[\sum_{X_j \in G_k; j \neq i} \text{Cov}(X_i, X_j) + \sum_{X_j \notin G_k} \text{Cov}(X_i, X_j) \right] \right) \\
&= \frac{1}{n^2} \left(\sum_{X_i \in \mathbf{S}} \sigma^2 + \sum_{G_k \in \mathbf{G}} \sum_{X_i \in G_k} \left[\sum_{X_j \in G_k; j \neq i} \rho \sqrt{\text{Var}(X_i)} \sqrt{\text{Var}(X_j)} + \sum_{X_j \notin G_k} 0 \right] \right) \\
&= \frac{1}{n^2} \left(\sum_{X_i \in \mathbf{S}} \sigma^2 + \sum_{G_k \in \mathbf{G}} \sum_{X_i \in G_k} \sum_{X_j \in G_k; j \neq i} \rho \sigma^2 \right) \\
&= \frac{1}{n^2} (n\sigma^2 + mg(g-1) [\rho\sigma^2]) \\
&= \frac{1}{n^2} (n\sigma^2 + n\sigma^2 \cdot \rho(g-1)) \\
&= \frac{1}{n} \sigma^2 [1 + \rho(g-1)]
\end{aligned}$$

□

We refer to this variance of the average error, when there is error correlation, as $\text{Var}_{\text{corr}}(\bar{X})$. Note that other than for very specific graph structures (e.g., bipartite graphs), if relational data are correlated, autocorrelation is positive and ρ will be greater than zero. Thus, as ρ or g (i.e., group size) increases, $\text{Var}_{\text{corr}}(\bar{X})$ also increases. If this correlation is ignored and error is estimated under the assumption of independence (i.e., with $\frac{1}{n}\sigma^2$), then the variance will be underestimated by $\rho(g-1) [\frac{1}{n}\sigma^2]$.

In Sect. 2.1, we described how existing resampling procedures create dependent (i.e., overlapping) test sets. Here, we show theoretically how overlapping samples decreases the variance of observed error. For the proof below, we consider the case of independent data instances and show the effect on variance when a set of instances appears repeatedly in each sample (i.e., test set). Then, in the next theorem, we will show how the effect of overlap combines with the effect due to correlated instances and together lead to bias in statistical tests.

Theorem 2 (Overlap in samples reduces variance) *Let X_i , $\bar{X}$, and $\mathbf{S}$ be defined as in Theorem 1. Assume that every sample $\mathbf{S}$ has a subset $\mathbf{S}'$ of independent r.v.s and a subset $\mathbf{C}$ of common r.v.s, where the set $\mathbf{S}'$ is of size $n - c$ and the set $\mathbf{C}$ consists of c r.v.s with unknown values, but the values do not vary across samples (i.e., $\mathbf{S} = \mathbf{S}' + \mathbf{C} = \{X_j\}_{n-c} + \{x_i\}_c$). Then $\text{Var}(\bar{X}) = \frac{1}{n}\sigma^2 [1 - \frac{c}{n}]$.*

Proof

$$\begin{aligned}
 \text{Var}(\bar{X}) &= \text{Var}\left(\frac{1}{n} \sum_{X_i \in \mathbf{S}} X_i\right) \\
 &= \text{Var}\left(\frac{1}{n} \sum_{X_i \in \mathbf{C}} x_i + \frac{1}{n} \sum_{X_j \notin \mathbf{C}} X_j\right) \\
 &= \frac{1}{n^2} \text{Var}\left(\sum_{X_i \in \mathbf{C}} x_i\right) + \frac{1}{n^2} \text{Var}\left(\sum_{X_j \notin \mathbf{C}} X_j\right) \\
 &= 0 + \frac{1}{n^2} \sum_{X_j \notin \mathbf{C}} \text{Var}(X_j) \\
 &= \frac{1}{n^2} \sum_{j \notin \mathbf{C}} \sigma^2 \\
 &= \frac{1}{n^2} (n - c) \sigma^2 \\
 &= \frac{1}{n} \sigma^2 \left[1 - \frac{c}{n}\right]
 \end{aligned}$$

□

We refer to this variance of the average error, when there is overlap between samples, as $\text{Var}_{\text{ovlp}}(\bar{X})$. As c increases, $\text{Var}_{\text{ovlp}}(\bar{X})$ decreases. Note that for samples with n independent samples, the variance of the average performance is $\frac{1}{n} \sigma^2$, thus when overlap is ignored the variance will be underestimated by $\frac{c}{n^2} \sigma^2$.

Finally, we show how these two effects combine together to bias conventional statistical tests for network domains.

Theorem 3 (Sample overlap and correlation increase likelihood of Type I error) *Let algorithm A and algorithm B have equal classification error rates on data drawn from the same domain D. Let X_i be the classification error for an instance i in domain D and assume that X_i^A and X_i^B (the error made by algorithm A and B respectively) are drawn from the same distribution with mean μ and variance σ^2 . Furthermore, assume that a dataset of n instances from D consists of a set of groups $\mathbf{G}$, where $|\mathbf{G}| = m$ and each group $G_k \in \mathbf{G}$ has average size g (i.e., $m \cdot g = n$). Let ρ be the average correlation between the X_i that are members of the same group, otherwise assume the X_i are independent. Let $\bar{\mathbf{X}}^A = \{\bar{X}_1^A, \bar{X}_2^A, \dots, \bar{X}_k^A\}$ and $\bar{\mathbf{X}}^B = \{\bar{X}_1^B, \bar{X}_2^B, \dots, \bar{X}_k^B\}$ be the estimates of average error for $s = [1, k]$ samples drawn with replacement (i.e., repeated sampling) from a given dataset of size n in the domain D. Assume that there is a set instances $\mathbf{C}$, where $|\mathbf{C}| = c$, that is common to all k samples and assume otherwise the samples are independent. Then an unpaired t-test over $\bar{\mathbf{X}}^A$ and $\bar{\mathbf{X}}^B$ will underestimate the variance of the null distribution by: $\Delta = \frac{1}{n} \sigma^2 \left[\left(\frac{c}{n}\right) (1 + \rho(g - 1) \left[\frac{c-1}{c} + \frac{n-c-1}{n}\right]) \right]$*

Proof The observed variance $\text{Var}(\bar{X})$ will have effects due to both instance correlation and overlapping samples:

$$\text{Var}(\bar{X}) = \text{Var}\left(\frac{1}{n} \sum_{X_i \in \mathbf{S}} X_i\right)$$

$$\begin{aligned}
&= \frac{1}{n^2} \text{Var} \left(\sum_{X_i \in \mathbf{C}} x_i + \sum_{X_j \notin \mathbf{C}} X_j \right) \\
&= \frac{1}{n^2} \text{Var} \left(\sum_{X_j \notin \mathbf{C}} X_j \right) \\
&= \frac{1}{n^2} \left(\sum_{X_j \notin \mathbf{C}} \sum_{X_l \notin \mathbf{C}} \text{Cov}(X_j, X_l) \right) \\
&= \frac{1}{n^2} \left(\sum_{X_j \notin \mathbf{C}} \text{Var}(X_j) + \sum_{X_j \notin \mathbf{C}} \sum_{\substack{X_l \notin \mathbf{C} \\ l \neq j}} \text{Cov}(X_j, X_l) \right) \\
&= \frac{1}{n^2} \left(\sum_{X_j \notin \mathbf{C}} \sigma^2 + \sum_{X_j \notin \mathbf{C}} \sum_{\substack{X_l \notin \mathbf{C} \\ l \neq j}} p(X_l \in G_k | X_j \in G_k) \cdot \rho \sigma^2 + p(X_l \notin G_k | X_j \in G_k) \cdot 0 \right) \\
&= \frac{1}{n^2} \left(\sum_{X_j \notin \mathbf{C}} \sigma^2 + \sum_{X_j \notin \mathbf{C}} \sum_{\substack{X_l \notin \mathbf{C} \\ l \neq j}} \frac{g-1}{n-1} \cdot \rho \sigma^2 \right) \\
&= \frac{1}{n^2} \left((n-c) \sigma^2 + (n-c)(n-c-1) \frac{g-1}{n-1} \cdot \rho \sigma^2 \right) \\
&= \frac{1}{n} \sigma^2 \left(\left[1 - \frac{c}{n} \right] + \left[1 - \frac{c}{n} \right] \left[1 - \frac{c}{n-1} \right] \rho (g-1) \right) \\
&= \frac{1}{n} \sigma^2 \left[1 - \frac{c}{n} \right] \left(1 + \left[1 - \frac{c}{n-1} \right] \rho (g-1) \right)
\end{aligned}$$

We refer to this variance of the average error, when there is both overlap between samples and error correlation, as $\text{Var}_{\text{obs}}(\bar{X})$. The unpaired t -test uses the average (i.e., *pooled*) variance of $\bar{X}^A$ and $\bar{X}^B$ for the null distribution. Since the error distribution of A and B is equal, the average is equal to the variance of a single algorithm. When the samples have *independent* instances but c instances are common to each sample, we know from Theorem 2 that the variance of $\bar{X}$ will be the following: $\text{Var}_{\text{ovlp}}(\bar{X}) = \frac{1}{n} \sigma^2 \left[1 - \frac{c}{n} \right]$. However, when there is correlation ρ among the instances in the data but the samples themselves are independent, from Theorem 1, we know that the variance of $\bar{X}$ is the following: $\text{Var}_{\text{corr}}(\bar{X}) = \frac{1}{n} \sigma^2 [1 + \rho(g-1)]$.

Since the two algorithms are equal, the variance of the null distribution should correspond to the variance with independent samples $\text{Var}_{\text{corr}}(\bar{X})$. However, the observed variance $\text{Var}_{\text{obs}}(\bar{X})$ will be used in the t -test, resulting in an underestimate of Δ :

$$\begin{aligned}
\Delta &= \text{Var}_{\text{corr}}(\bar{X}) - \text{Var}_{\text{obs}}(\bar{X}) \\
&= \frac{1}{n} \sigma^2 [1 + \rho(g-1)] - \frac{1}{n} \sigma^2 \left[1 - \frac{c}{n} \right] \left(1 + \left[1 - \frac{c}{n-1} \right] \rho (g-1) \right) \\
&= \frac{1}{n} \sigma^2 \left[1 + \rho(g-1) - \left[1 - \frac{c}{n} \right] - \left(\left[1 - \frac{c}{n} \right] \left[1 - \frac{c}{n-1} \right] \rho (g-1) \right) \right]
\end{aligned}$$

$$\begin{aligned}
 &= \frac{1}{n} \sigma^2 \left[1 - \left[1 - \frac{c}{n} \right] + \rho(g-1) \left(1 - \left[1 - \frac{c}{n} \right] \left[1 - \frac{c}{n-1} \right] \right) \right] \\
 &= \frac{1}{n} \sigma^2 \left[\left(\frac{c}{n} \right) + \rho(g-1) \left(\frac{c}{n} + \frac{c}{n-1} - \frac{c}{n} \frac{c}{n-1} \right) \right] \\
 &= \frac{1}{n} \sigma^2 \left[\left(\frac{c}{n} \right) + \rho(g-1) \left(\frac{c}{n} \right) \left(1 + \frac{n}{n-1} - \frac{c}{n-1} \right) \right] \\
 &= \frac{1}{n} \sigma^2 \left[\left(\frac{c}{n} \right) + \rho(g-1) \left(\frac{c}{n} \right) \left(1 + \frac{n-c}{n-1} \right) \right] \\
 &= \frac{1}{n} \sigma^2 \left[\left(\frac{c}{n} \right) \left(1 + \rho(g-1) \left[1 + \frac{n-c}{n-1} \right] \right) \right] \\
 &\geq 0
 \end{aligned}$$

□

As ρ (the amount of error correlation) or c (the amount of overlap between samples) increases, the amount of underestimation increases (i.e., Δ). This increases the probability of a Type I error in the following way. For unpaired tests, the t -statistic is:

$$\hat{t} = \frac{\bar{X}^A - \bar{X}^B}{\sqrt{\text{Var}(\bar{X})} \cdot \sqrt{\frac{2}{n}}}$$

where $\text{Var}(\bar{X})$ is the pooled variance. Since $\text{Var}_{\text{obs}}(\bar{X}) < \text{Var}_{\text{corr}}(\bar{X})$, the result will be that $\hat{t}_{\text{obs}} > \hat{t}_{\text{corr}}$ and thus $P(\hat{t}_{\text{obs}}|T) < P(\hat{t}_{\text{corr}}|T)$, where T is the t distribution with $2k - 2$ degrees of freedom. Thus, using $\text{Var}_{\text{obs}}(\bar{X})$ instead of $\text{Var}_{\text{corr}}(\bar{X})$, it is more likely that the null hypothesis will be rejected even when it holds, and as such Type I error will increase.

This effect will impact paired t -tests in a similar way, as the decrease in observed variance of $\bar{X}^A$ and $\bar{X}^B$ will also result in an underestimate of the *difference variance* $\text{Var}(\bar{X}^A - \bar{X}^B)$, which is used instead of the pooled variance.

5 Empirical investigation with simulated classifiers

Type I errors occur when a statistical test incorrectly rejects the null hypothesis (i.e, the test concludes that there is a significant difference between two classifiers when there is none). We run simulation experiments in order to assess the contributions of two key factors in Type I errors for within-network classification: (1) correlation of errors among related data instances and (2) dependence between samples. We also assess the potential of various statistical tests to produce Type I errors in a within-network classification setting.

Our method preserves the basic structure of Dietterich's [5], but introduces a group-based model to more easily represent sets of related instances and vary the degree of error correlation among them. We also ran Dietterich's original procedure with qualitatively similar results.

5.1 Methodology

As we have seen, real classifiers exhibit correlated errors on sets of related instances. To simulate this behavior, we divide all data instances into disjoint groups such that classification errors are more likely to be correlated on instances within a group than on instances from different groups.

Table 6 Simulation algorithm

```

for simulation 1 to 10
  for trial 1 to 1000
    (NCV*) Create sample and split into 10 disjoint folds
    for propLabeled  $\in$  (0.1, 0.3, 0.5, 0.7, 0.9)
      (RS*) Create sample and resample 30 folds
      for each fold
        run groupBasedClassification(fold)
      end for each
      Apply significance test to all folds in trial/propLabeled
      Accept or reject null hypothesis
      Measure error correlation for this trial/propLabeled
    end for
  end for
  Calculate null hypotheses rejection rate for simulation
  Calculate mean error correlation for this simulation
end for
Calculate final mean rejection rate and error correlation
NCV*: Performed for NCV (see Table 3) only.
RS*: Performed for RRS and ERS (see Tables 1–2) only.

```

See Table 7 for *groupBasedClassification* procedure

Table 6 outlines our simulation algorithm for both the resampling procedures and the network cross-validation procedure. The two procedures differ only in that: (1) they use different resampling algorithms to create their test sets and (2) NCV uses the same samples and folds across all proportions of labeled data, whereas the resampling procedures choose a different sample and random split for each trial and proportion labeled.

We simulate drawing a network sample s from an underlying population by creating 300 instances and assigning one of 10 groups to each instance uniformly at random (a skewed group-size distribution produced qualitatively similar results). We then resample s and run a simulated classification experiment on each resampled train/test split (see Table 7). For each trial, we apply a significance test to either accept or reject the null hypothesis. We calculate the proportion of trials for which the null hypothesis was rejected. Since the simulation is designed so each classifier has the same error rate in the underlying population, any rejections of the null hypothesis represent Type I errors. In addition, we measure the degree of error correlation for each trial. We use the ϕ coefficient to measure the pairwise correlation of 0–1 errors among all instances in the same group, averaged over all groups.

5.2 Results

For all experiments, we present average Type I error rates for various statistical tests over 10 simulations of 1000 trials each, on data samples of size 300 instances. Unless otherwise noted, our default experimental parameters are classifier error rate $P(err) = 0.1$, error correlation parameter $errCorr = 0.9$, and proportion of labeled instances $propLabeled = 0.9$. The $errCorr$ parameter determines the likelihood of misclassifying instances in the chosen misclassification group versus instances in other groups.

Table 7 Group-based classifier simulation algorithm

```

groupBasedClassification(fold)
  NG = total number of groups
  M = round(NG * P(err))
  P(MG) = P(err) + errCorr * (1 - P(err))
  P(MG') = P(err) +  $\frac{1 - P(MG)}{1 - P(err)}$ 
  // a real classifier would normally train here, but there is no training phase since we are
  // simulating classification
  // choose groups to misclassify
  MGA = random set of M groups from 1 to  $\frac{NG}{2}$ 
  MGB = random set of M groups from  $\frac{NG}{2} + 1$  to NG
  for each instance  $i \in fold$ 
    // simulate application of classifier A
    if group(i)  $\in$  MGA
      i misclassified by classifier A with P(MG)
    else
      i misclassified by classifier A with P(MG')
    end if
    // simulate application of classifier B
    if group(i)  $\in$  MGB
      i misclassified by classifier B with P(MG)
    else
      i misclassified by classifier B with P(MG')
    end if
  end for each

```

This method ensures that classifiers A and B have the same error rate, while still making different kinds of errors (i.e., A misclassifies different groups from B). M is the number of groups chosen for misclassification by each classifier. $P(MG)$ is the misclassification probability for instances of the chosen groups (MG_A/MG_B) and $P(MG')$ is the misclassification probability for instances of all other groups. $P(err)$ is the overall error rate of both classifiers A and B and $errCorr$ controls the degree of error correlation among instances within MG_A/MG_B .

Figure 2a shows the effects of varying the proportion of labeled data available (0.1, 0.3, 0.5, 0.7, 0.9). For both resampling procedures, the Type I error rate increases as *propLabeled* decreases. This result is expected since the degree of overlap between test sets increases as the test sets become larger due to the larger number of unlabeled instances. Since NCV disallows overlapping test sets by design, it is not susceptible to this problem, achieving low Type I error rates across the range of *propLabeled* values.

Figure 2b shows the effects of measured error correlation on Type I error rate as we vary $P(err) = [0.1, 0.2, 0.3, 0.4]$ and $errCorr = [0, 0.2, 0.4, \dots, 1.0]$. We can observe that the Type I error rate of the resampling procedures increase as the error correlation increases. This is not surprising since increased error correlation is expected to lead to increased imbalance between samples. In addition, we note that our experiments showed, for a fixed value of $P(err)$ ($errCorr$), both the measured type I error rate and the measured error correlation increased monotonically with $errCorr$ ($P(err)$). Overall, NCV is less affected by imbalanced samples since test sets do not overlap, so it exhibits much lower levels of Type I error. The Type I error rates of equal-instance resampling (ERS) are lower than simple random

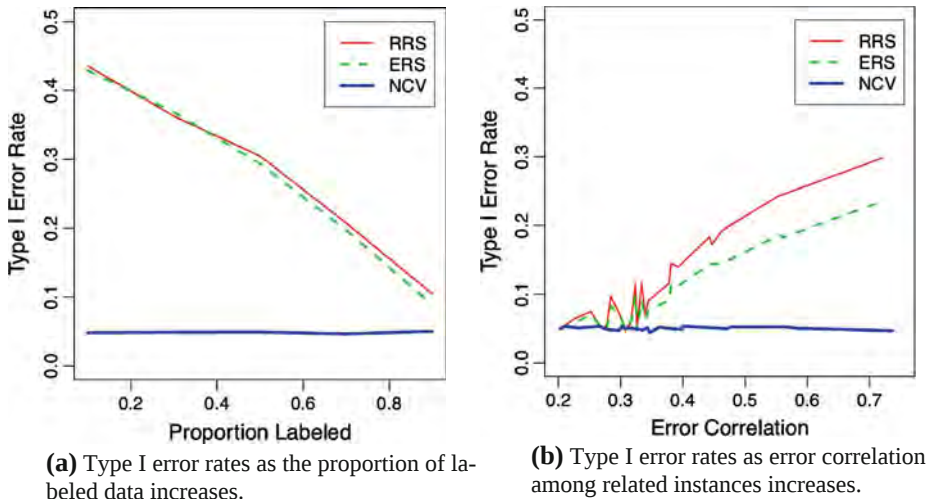


Fig. 2 Results for simulated classifiers on synthetic data

resampling; however, since the improvement is not sufficient to make it competitive with NCV, we do not consider ERS further.

6 Empirical investigation with real classifiers

This section describes our investigation of the characteristics of statistical tests when comparing real relational learning algorithms. We consider two collective inference models and compare their performance on synthetic data. The synthetic data generation enables the simulation of multiple draws of networks from the same distribution. We evaluate the Type I error and power of statistical tests, as the performance of the two models is varied.

6.1 Models

In order to run experiments at the scale needed to assess Type I error and power rates, we choose to investigate two simple and efficient collective models described in Macskassy and Provost [19].

The first model is the weighted-vote relational neighbor (wvRN), which estimates class label probabilities by assuming the existence of homophily. Given the unlabeled nodes in a network $v_i \in V_U$, wvRN estimates $P(y_i | N_i)$ as the average of the class probabilities of the instances in N_i (v_i 's neighbors): $P(y_i = + | N_i) = \frac{1}{Z} \sum_{v_j \in N_i} P(y_j = + | N_j)$, where Z is a normalizing constant.

The second model is the network-only Bayes classifier (nBC), which estimates class label probabilities for v_i with a multinomial naive Bayesian model, based on the classes of v_i 's neighbors N_i : $P(y_i = + | N_i) = \frac{P(N_i | +)P(+)}{P(N_i)} \propto [\frac{1}{Z} \prod_{v_j \in N_i} P(y_j = \tilde{y}_j | y_i = +)]P(+)$, where Z is a normalizing constant and $\tilde{y}_j$ is the class observed at node v_j .

The wvRN model does not require any learning. To estimate the parameters of the nBC model, we use maximum likelihood estimation over the labeled part of the network. For collective inference, we use relaxation labeling with both models.

6.2 Data

The synthetic datasets are generated with a latent group model [25]. Each network is generated with 300 nodes (instances). The nodes are generated as members of (hidden) groups, and group membership determines the binary class label values and link existence for each node. The average group size is 10, and there are two types of groups: A and B . The network is skewed towards A groups, $P(A) = 0.75$. Members of A groups are more likely to have a positive class label, $P(+|A) = 0.9$. Members of A groups also have higher intra-group linkage with $P_A(e_{ij} = 1 | i \in g_k^A \wedge j \in g_k^A) = 0.6$ and lower inter-group linkage with $P_A(e_{ij} = 1 | i \in g_k^A \wedge j \notin g_k^A) = 0.003$. Members of B groups are more likely to have a negative class label, $P(+|B) = 0.1$. Members of B groups have relatively lower intra-group linkage with $P_B(e_{ij} = 1 | i \in g_k^B \wedge j \in g_k^B) = 0.4$ and higher inter-group linkage with $P_B(e_{ij} = 1 | i \in g_k^B \wedge j \notin g_k^B) = 0.013$. The resulting networks have an average autocorrelation of 0.40, an average error correlation of 0.20, and a class prior of $P(+) = 0.70$.

Our choice of data generation parameters was designed to create networks where the wvRN and nBC would make *different* classification errors. Many of the nodes in type B groups have more links to nodes in type A groups so the networks do not fully meet the assumption of homophily which underlies the wvRN model. The nBC should thus more accurately *learn* how to classify the type B nodes, while the wvRN will likely be more accurate on the type A nodes. However, since wvRN does not *learn* the concept of homophily, it will not experience variance due to small labeled sets.

6.3 Methodology

To estimate Type I error, we need two models with equal performance on partially labeled networks of the same size, drawn from the same domain. To achieve this, we measured the average accuracy of the wvRN and the nBC models on the synthetic data and handicapped the better model (wvRN) until the performance difference of the models was ≤ 0.005 . The better performing model was handicapped by randomly selecting $p\%$ of its predictions and perturbing those probabilities toward the opposite class. To set p , we generated 50 networks for use as a *calibration* set. Each of the 50 networks was sampled into 10-fold network cross-validation sets, resulting in 500 training/test set splits on which we measured average accuracy of each model. Using this calibration set, we searched for a value of p that resulted in a performance difference of ≤ 0.005 between the two models.

To estimate power, we need to vary the performance difference between the two models. To achieve this, we perturbed the predictions of the *worse* performing model (nBC) to increase the mean difference in performance between the two models. For the power experiments, we used perturbation rates of $p = [0.025, 0.075, 0.15, 0.3]$.

6.4 Results

To measure Type I error rates and power of the statistical tests, we used four synthetic networks (in addition to the calibration set). On each network, we considered four levels of labeling: [0.1, 0.2, 0.3, 0.4]. At each level of labeling, we sampled the network 10 times, either by repeated sampling or by cross-validation. On each of the ten samples, we learned the nBC model on the labeled portion of the network, and then, we applied both models to the unlabeled portion of the network with the perturbation rate p . We measured the accuracy of each model on the ten test sets and then assessed the difference in performance using a t -test.

Fig. 3 Type I error rates for synthetic data and real classifiers—as the proportion of labeled data increases

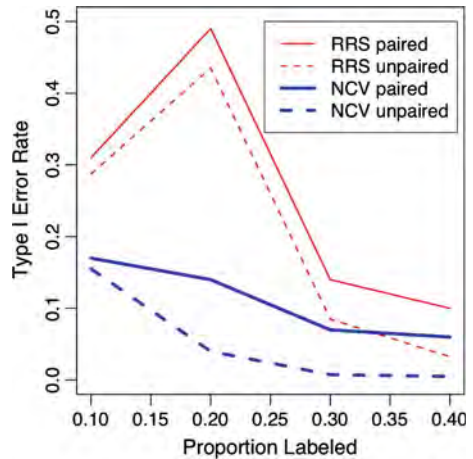


Figure 3 plots the Type I error rates for four combinations of sampling method (RS, NCV) and statistical test (paired and unpaired t -test). Type I error rates for each dataset are measured over 100 trials and averaged. Recall that we use the calibration set to choose a value for p that makes the average performance of the wvRN and nBC equal at the given level of labeling. Thus, any trial in which the t -test assesses the observed performance difference as *significant* corresponds to a Type I error.

All the tests have high Type I error rates with 10% of instances labeled. This error generally decreases as the amount of labeled data increases (and thus the size of the test set decreases). Since we are using relatively simple classifiers, as the number of labeled data increase model performance does indeed converge, and the two models make similar classification errors at labeling rates greater than 40%. In reality, when comparing relational models with different representations and different complexity, we expect Type I error to occur at all levels of labeling.

Notably, the repeated sampling approach experiences as high as 50% Type I errors (at 20% labeling). This means that *half* the time the method is concluding a significant performance difference between the two models, when in fact there is none. For data mining and machine learning researchers that are investigating the trade-offs between learning algorithms, this is an unacceptable level of error. On the same data, the network cross-validation procedure error rate is only 15% with the paired t -test—a 70% reduction in Type I error. Clearly, the network cross-validation approach results in a more accurate comparison of model performance.

Note that in the simulated classifier experiments (see Fig. 2a,b), NCV across the proportions labeled is equivalent to 10-fold CV at 90% labeled, since performance in the simulated classifier experiments does not depend on: (1) the number of labeled neighbors available during inference and (2) the number of instances available during training. These additional dependencies contribute to the higher Type I error observed for NCV with real classifiers.

Although the Type I error of NCV combined with a paired t -test is much lower than resampling, it is still higher than the expected level of 5%. To investigate this behavior, we examined the estimates used in the t -test calculation: (1) the estimate of the mean performance difference between the two models: $\mu_{diff} = \mu_{A_{acc}} - \mu_{B_{acc}}$ and (2) the estimate of the variance of the differences: $Var_{diff} = Var(\{A_{acc} - B_{acc}\}_i)$. The error correlation increases the variance of the estimated μ_{diff} , but the estimates are not biased. On the other hand, the error correlation decreases the variance of the estimated differences, resulting in

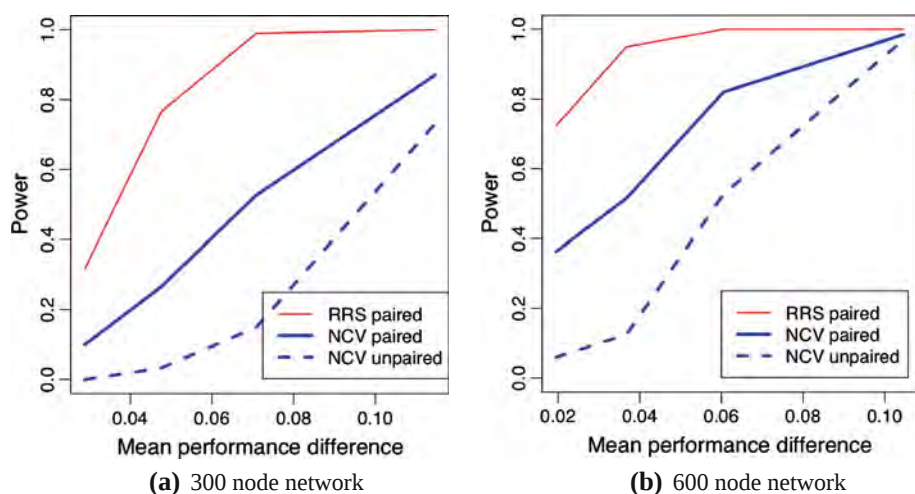


Fig. 4 Statistical power on synthetic data as the performance difference between classifiers increases. Results are shown with 30% of instances labeled

a biased underestimate of Var_{diff} . We considered the unpaired t -test to adjust for this bias. The unpaired t -test uses a *pooled* estimate of variance over the performance measurements $\{A_{acc}\}$ and $\{B_{acc}\}$, instead of an estimate of variance of the differences. The pooled estimate of variance is higher than the variance of the differences, so it can offset the bias in variance estimation due to error correlation. This is indeed the case—combining NCV with unpaired t -tests results in Type I error rates of less than 0.05 (when proportion labeled is greater than 10%).

Figure 4a plots the statistical power of each test as we varied the average performance difference between wvRN and nBC. More specifically, we used $p = [0.025, 0.075, 0.15, 0.3]$ and measured the average performance difference between the two models on the calibration set. This gives us the mean performance difference that is plotted on the x -axis. Then, for each of the four evaluation networks, we sampled the network 10 times and learned/applied/evaluated the models as described above. Again, the results are measured over 100 trials. Since the two models *do* perform differently, any trial in which the t -test *does not* conclude that the observed performance difference is significant corresponds to a Type II error. Statistical power is defined as the proportion of trials in which the t -test correctly concludes that the two models are different (i.e., 1–Type II error). We only plot the results for 30% labeled. The results for other levels of labeling are qualitatively similar.

The power results illustrate an additional challenge in evaluating the performance difference between the models. Even when there is 5% difference in the mean performance of the two algorithms, it is sobering to note that repeated sampling can detect this difference less than 80% of the time. Network cross-validation is significantly worse—the paired t -test can detect the difference less than 30% of the time and the unpaired t -test less than 5% of the time. This may be due to the difference in test set size used by the two approaches. Recall that repeated sampling uses *all* the unlabeled data for evaluation, so at 30% labeling, this corresponds to 210 nodes. On the other hand, network cross-validation uses only 10% of the nodes for evaluation (i.e., 30 nodes) regardless of the level of labeling in the network.

To explore this issue, we increased the dataset size to 600 and measured the power of each approach again. Figure 4b graphs the resulting power rates. In general, power of any test

will be increased as the sample size is increased. However, here, we can see that the gains for network cross-validation are relatively larger than for repeated sampling. It is difficult to compare across the two sets of results due to different mean performance of the models. However, if we interpolate between the results in Fig. 4b to assess the power at 5% mean performance difference and compare to Fig. 4a at 5%, we can see that doubling the dataset size reduced the Type II error of repeated sampling by 10%, but the Type II error for network cross-validation was reduced by 45% (paired *t*-test) and 70% (unpaired), respectively.

7 Empirical investigation with real-world data

This section describes our investigation of the characteristics of statistical tests when comparing relational learning algorithms on real-world relational data. To confirm the behavior we observed in the synthetic datasets, we compared the performance of wvRN and nBC on data from the National Longitudinal Study of Adolescent Health [13] as well as data from the Internet Movie Database (imdb.com).

The Adolescent Health (Add Health) data consist of survey information from 144 middle and high schools, collected in 1994–1995. The survey questions queried for the students' social networks along with myriad behavioral and academic attributes. In this paper, we consider the social networks of six schools with similar autocorrelation and link patterns. The classification task is to predict whether the student *smokes* based on the behavior of their friends in the social network. The six schools we selected have sizes ranging from 300–700 nodes, average degree of 7–8, autocorrelation in the range [0.25, 0.35], and average error correlation of 0.14.

Our second dataset is drawn from the Internet Movie Database (IMDb). We collected a sample of 1,543 movies released in the United States between 2003 and 2007, with their associated producers and studios. To create six disjoint network samples, we partitioned the movies by their associated studios, using stratified sampling by studio size. Within each partition, we created links among movies with a common producer, ignoring any links across the partitions. The resulting networks have an average size of 257 nodes, and the movies have average degree of 16. The classification task is to predict whether the movie will make more than \$60mil (inflation adjusted) in total box office receipts. The average autocorrelation in these networks is 0.35, and the average error correlation is 0.22.

To assess the Type I error characteristics of the models, we used a procedure similar to the one described in Sect. 6. Each trial considers one network as the evaluation set, then we calibrate the models on the remaining 5 networks under the assumption that these networks were drawn from the same distribution. We sampled each of these five networks 10 times into 10-fold NCV sets, producing a calibration set of 500 training/test splits. As described previously, we searched for a value of p that resulted in a performance difference of ≤ 0.005 between the two models. We considered five levels of labeling: [0.1, 0.2, 0.3, 0.4, 0.5], calibrated the models at each level of labeling, and measured the Type I error on the held out network. The models converge in performance at 50–60% labeling (i.e., $p = 0$) so we do not consider larger labeling rates.

Figure 5 shows Type I error for each combination of sampling method and statistical test, measured over 50 trials. Figure 5a and b show the results for the AddHealth and IMDb datasets, respectively.

As expected, the statistical tests exhibit similar behavior on the real relational data and the synthetic data. Again resampling produces unacceptable levels of Type I error (up to 40%) and network cross-validation has more reasonable error rates. Overall Type I error decreases

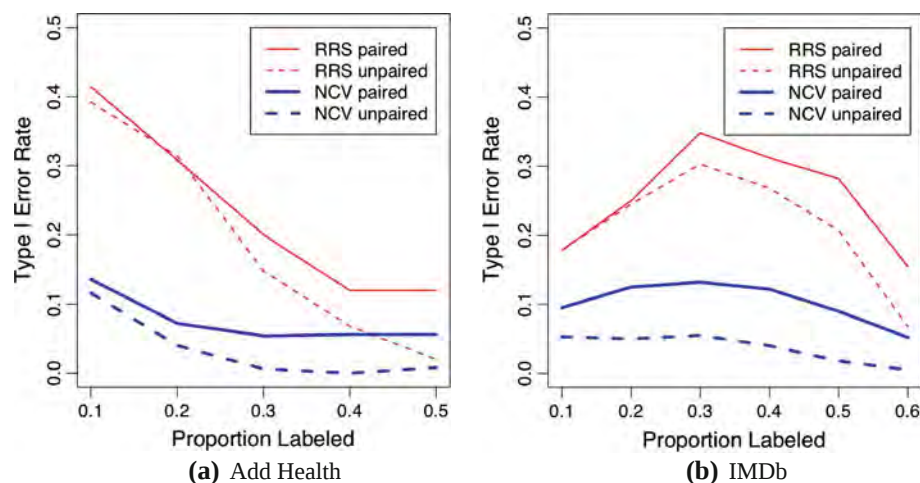


Fig. 5 Type I error rates on real relational datasets, as the proportion of labeled data increases

as the proportion of labeled data increases to 50–60% (i.e., test set size decreases). Recall, however, that we are investigating simple models that are nearly equivalent on a restricted task involving only the class label and no other attribute/link features. In practice, as the complexity of models and concepts increase, Type I errors are likely to occur at all levels of labeling.

8 Conclusions and discussion

In this paper, we examined the characteristics of statistical tests for comparing *within-network* classification algorithms. We presented three resampling procedures and three significance tests, performed experiments on both real and synthetic data using real and simulated classifiers.

Our analysis shows that a commonly used form of evaluation in relational learning (paired *t*-tests on overlapping network samples) can result in unacceptably high levels of Type I error (as high as 50%). High Type I error indicates that many algorithm differences will be judged incorrectly as significant when in fact performance is equivalent. Although for efficiency reasons we considered relatively simple relational models for this work, our findings apply to evaluations of more complex relational models as well—since any relational model that attempts to exploit relational autocorrelation is likely to produce correlated errors.

Furthermore, we demonstrated, both theoretically and empirically, that Type I error increases as (1) the correlation among instances increases and (2) the overlap among the evaluation sets increase (i.e., the proportion of labeled nodes in the network decreases).

Although we investigated the properties of significance tests for within-network classification, the findings are also applicable to across-network tasks and other forms of hypothesis testing (e.g., standard error bars will be underestimated). The extent of the effect will depend on the level of observed autocorrelation (which will cause error correlation), as well as the amount of overlap between samples.

We proposed a method for *network cross-validation* that reduces the overlap between samples. We note that although the method creates disjoint test sets, the predictions for those test set instances will be influenced by other predictions in the unlabeled *inference* set (due

to the collective inference process). This means that there is still some dependency between test sets (since the inference sets overlap) which could increase the Type I error of NCV.

Our empirical evaluation shows that NCV combined with *unpaired t*-tests results in low levels of Type I error. However, this low error is achieved at the expense of statistical power (i.e., Type II error). NCV combined with *paired t*-tests produces more acceptable levels of Type I error while still providing reasonable levels of statistical power.

Promising future directions for this work include: (1) using patterns (e.g., communities) in relational data to split train/test data (e.g., stratified by community, or biased by community); (2) investigating non-random labeling patterns and their impact on error correlation for different collective inference methods; and (3) investigating how characteristics of relational data affect the power of statistical tests (i.e., Type II error).

Acknowledgments We thank Rongjing Xiang for her assistance in experimental implementation. This work was performed under the auspices of the US Department of Energy by Lawrence Livermore National Laboratory under contract DE-AC52-07NA27344 (LLNL-JRNL-455699). This work was also supported by DARPA and NSF under contract numbers NBCH1080005 and SES-0823313.

Open Access This article is distributed under the terms of the Creative Commons Attribution Noncommercial License which permits any noncommercial use, distribution, and reproduction in any medium, provided the original author(s) and source are credited.

References

1. Chakrabarti S, Dom B, Indyk P (1998) Enhanced hypertext categorization using hyperlinks. In: Proceedings of the ACM SIGMOD int'l conference on management of data. pp 307–318
2. Cohen P (1995) Empirical methods for artificial intelligence. MIT Press, Cambridge
3. Dhurandhar A, Dobra A (2008) Probabilistic characterization of random decision trees. *J Mach Learn Res* 9:2321–2348
4. Dhurandhar A, Dobra A (2008) Study of classification models and model selection measures based on moment analysis. In: NIPS'08 workshop on new challenges in theoretical machine learning: learning with data-dependent concept spaces
5. Dietterich T (1998) Approximate statistical tests for comparing supervised classification learning algorithms. *Neural Comput* 10:1895–1923
6. Dunder M, Krishnapuram B, Bi J, Rao RB (2007) Learning classifiers when the training data is not iid. In: Proceedings of the 20th int'l joint conference on artificial intelligence
7. Gallagher B, Eliassi-Rad T (2007) An examination of experimental methodology for classifiers of relational data. In: Workshop proceedings of the seventh IEEE int'l conference on data mining. pp 411–416
8. Gallagher B, Eliassi-Rad T (2008) Leveraging label-independent features for classification in sparsely labeled networks: an empirical study. In: Proceedings of the second ACM SIGKDD workshop on social network mining and analysis (SNA-KDD'08)
9. Gallagher B, Tong H, Eliassi-Rad T, Faloutsos C (2008) Using ghost edges for classification in sparsely labeled networks. In: Proceedings of the 14th ACM SIGKDD Int'l conference on knowledge discovery and data mining. ACM, New York, NY, USA, pp 256–264
10. Getoor L, Friedman N, Koller D, Taskar B (2001) Learning probabilistic models of relational structure. In: Proceedings of the 18th Int'l conference on machine learning. Morgan Kaufmann, San Francisco, CA, pp 170–177
11. Getoor L, Friedman N, Koller D, Taskar B (2002) Learning probabilistic models with link uncertainty. *J Mach Learn Res* 3:679–707
12. Getoor L, Segal E, Taskar B, Koller D (2001) Probabilistic models of text and link structure for hypertext classification. In: In IJCAI'01 workshop on text learning: beyond supervision. Morgan Kaufmann, San Francisco, CA, pp 170–177
13. Harris K (2008) The national longitudinal study of adolescent health (add health), waves i & ii, 1994–1996; wave iii, 2001–2002 [machine-readable data file and documentation]. Carolina Population Center, University of North Carolina at Chapel Hill, Chapel Hill, NC

14. Herbster M, Lever G, Pontil M (2008) Exploiting cluster-structure to predict the labeling of a graph. In: NIPS'08 workshop on new challenges in theoretical machine learning: learning with data-dependent concept spaces
15. Herbster M, Lever G, Pontil M (2008) Online prediction on large diameter graphs. In: NIPS'08 workshop on new challenges in theoretical machine learning: learning with data-dependent concept spaces
16. Jensen D, Neville J, Gallagher B (2004) Why collective inference improves relational classification. In: Proceedings of the 10th ACM SIGKDD int'l conference on knowledge discovery and data mining. pp 593–598
17. Lu Q, Getoor L (2003) Link-based classification. In: Proceedings of the 20th int'l conference on machine learning. pp 496–503
18. Macskassy S, Provost F (2003) A simple relational classifier. In: Proceedings of the 2nd workshop on multi-relational data mining, KDD2003. pp 64–76
19. Macskassy S, Provost F (2007) Classification in networked data: a toolkit and a univariate case study. *J Mach Learn Res* 8(May):935–983
20. Macskassy SA (2007) Classification in networked data: a toolkit and a univariate case study. In: Proceedings of the twenty-second conference on artificial intelligence. pp 590–595
21. McDowell L, Gupta K, Aha D (2007) Cautious inference in collective classification. In: Proceedings of the 22nd AAAI conference on artificial intelligence
22. Mohri M, Rostamizadeh A (2007) Stability bounds for non-i.i.d. processes. In: Proceedings of the neural information processing systems conference, 20
23. Mohri M, Rostamizadeh A (2010) Stability bounds for stationary ϕ -mixing and β -mixing processes. *J Mach Learn Res* 11:789–814
24. Neville J, Gallagher B, Eliassi-Rad T (2009) Evaluating statistical tests for within-network classifiers of relational data. In: Proceedings of the 9th IEEE int'l conference on data mining. pp 397–406
25. Neville J, Jensen D (2005) Leveraging relational autocorrelation with latent group models. In: Proceedings of the 5th IEEE int'l conference on data mining. pp 322–329
26. Neville J, Jensen D (2007) Relational dependency networks. *J Mach Learn Res* 8:653–692
27. Neville J, Jensen D, Friedland L, Hay M (2003) Learning relational probability trees. In: Proceedings of the 9th ACM SIGKDD int'l conference on knowledge discovery and data mining. pp 625–630
28. Neville J, Jensen D, Gallagher B (2003) Simple estimators for relational Bayesian classifiers. In: Proceedings of the 3rd IEEE int'l conference on data mining. pp 609–612
29. Perlch C, Provost F (2006) Acora: distribution-based aggregation for relational learning from identifier attributes. *Mach Learn* 62(1/2):65–105
30. Ralaivola L, Szafranski M, Stempfel G (2008) Chromatic pac-bayes bounds for non-iid data. In: NIPS'08 workshop on new challenges in theoretical machine learning: learning with data-dependent concept spaces
31. Sen P, Namata G, Bilgic M, Getoor L, Gallagher B, Eliassi-Rad T (2008) Collective classification in network data. *AI Mag* 29(3):93–106
32. Steinwart I, Christmann A (2009) Fast learning from non-i.i.d. observations. In: Proceedings of the neural information processing systems conference, 22
33. Taskar B (2009) Structured prediction cascades. In: NIPS'09 workshop on approximate learning of large scale graphical models: theory and applications
34. Taskar B, Abbeel P, Koller D (2002) Discriminative probabilistic models for relational data. In: Proceedings of the 18th conference on uncertainty in artificial intelligence. pp 485–492
35. Taskar B, Segal E, Koller D (2001) Probabilistic classification and clustering in relational data. In: Proceedings of the 17th int'l joint conference on artificial intelligence. pp 870–878
36. Usunier N, Amini MR, Gallinari P (2005) Generalization error bounds for classifiers trained with interdependent data. In: Proceedings of the neural information processing systems conference, 18
37. Vitale F, Cesa-Bianchi N, Gentile C (2008) Online graph prediction with random trees. In: NIPS'08 workshop on new challenges in theoretical machine learning: learning with data-dependent concept spaces
38. Xu Z, Kersting K, Tresp V (2009) Multi-relational learning with gaussian processes. In: Proceedings of the 21st int'l joint conference on artificial intelligence
39. Zhang X, Song L, Gretton A, Smola A (2008) Kernel measures of independence for non-iid data. In: Proceedings of the neural information processing systems conference, 21
40. Zhu X, Ghahramani Z, Lafferty J (2003) Semi-supervised learning using gaussian fields and harmonic functions. In: Proceedings of the 20th int'l conference on machine learning

Author Biographies



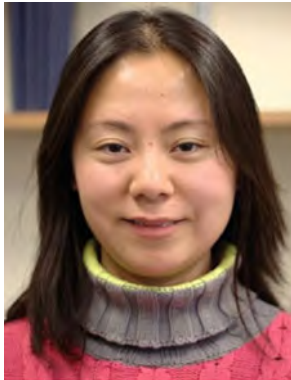
Jennifer Neville is an assistant professor at Purdue University with a joint appointment in the Departments of Computer Science and Statistics. She received her Ph.D. in computer science from the University of Massachusetts Amherst in 2006. She received a DARPA IPTO Young Investigator Award in 2003, was selected as a member of the DARPA Computer Science Study Group in 2007, and was chosen by IEEE as one of “AI’s 10 to watch” for 2008. Her research focuses on data mining techniques for relational and network domains.



Brian Gallagher is a computer scientist at the Center for Applied Scientific Computing at Lawrence Livermore National Laboratory. He received a B.A. in computer science from Carleton College in 1998 and an M.S. in computer science from the University of Massachusetts Amherst in 2004. His research interests include artificial intelligence, machine learning, and knowledge discovery and data mining. His current research focuses on classification and analysis in network-structured data.



Tina Eliassi-Rad is an assistant professor at the Department of Computer Science at Rutgers University. Until September 2010, she was a Member of Technical Staff at Lawrence Livermore National Laboratory. Tina earned her Ph.D. in Computer Sciences (with a minor in Mathematical Statistics) at the University of Wisconsin-Madison in 2001. Broadly speaking, Tina’s research interests include machine learning, data mining, and artificial intelligence. Her work has been applied to the World-Wide Web, text corpora, large-scale scientific simulation data, and complex networks. Tina is an action editor for the Data Mining and Knowledge Discovery Journal. She received a US DOE Office of Science Outstanding Mentor Award in 2010.



Tao Wang is a postdoctoral research associate in the Department of Computer Science at Purdue University. She received her Ph.D. in Computing Science from the University of Alberta in 2007. From 2007–2008, she was a research fellow at Australian National University. In 2009 she was a visiting assistant professor in the Department of Statistics at Purdue University. Her research interests are in the broad area of statistical machine learning and data mining. In particular, she is interested in the development of statistical tools and methodology for solving problems involving automated learning and decision-making in dynamical worlds.