



Universiteit
Leiden
The Netherlands

Estimating conditional mutual information for discrete-continuous mixtures using multi-dimensional adaptive histograms

Marx, A.; Yang, L.; Leeuwen, M. van; Demeniconi, C.; Davidson, I.

Citation

Marx, A., Yang, L., & Leeuwen, M. van. (2021). Estimating conditional mutual information for discrete-continuous mixtures using multi-dimensional adaptive histograms. *Proceedings Of The 2021 Siam International Conference On Data Mining (Sdm)*, 387-395.
doi:10.1137/1.9781611976700.44

Version: Publisher's Version

License: [Licensed under Article 25fa Copyright Act/Law \(Amendment Taverne\)](#)

Downloaded from: <https://hdl.handle.net/1887/3264126>

Note: To cite this publication please use the final published version (if applicable).

Estimating Conditional Mutual Information for Discrete-Continuous Mixtures using Multi-Dimensional Adaptive Histograms

Alexander Marx ^{*†}

Lincen Yang ^{*‡}

Matthijs van Leeuwen [‡]

Abstract

Estimating conditional mutual information (CMI) is an essential yet challenging step in many machine learning and data mining tasks. Estimating CMI from data that contains both discrete and continuous variables, or even discrete-continuous mixture variables, is a particularly hard problem. In this paper, we show that CMI for such mixture variables, defined based on the Radon-Nikodym derivative, can be written as a sum of entropies, just like CMI for purely discrete or continuous data. Further, we show that CMI can be consistently estimated for discrete-continuous mixture variables by learning an adaptive histogram model. In practice, we estimate such a model by iteratively discretizing the continuous data points in the mixture variables. To evaluate the performance of our estimator, we benchmark it against state-of-the-art CMI estimators as well as evaluate it in a causal discovery setting.

1 Introduction

In many research areas, such as classification [16], feature selection [34], and causal discovery [29], estimating the strength of a dependence plays a key role. A theoretically appealing way to measure dependencies is through mutual information (MI) since it has several important properties, such as the chain rule, the data processing inequality, and—last but not least—it is zero if (and only if) two random variables are independent of each other [4]. For structure identification, such as causal discovery, conditional mutual information (CMI) is even more interesting since it can help to distinguish between different graph structures. For instance, in a simple Markov chain $X \rightarrow Z \rightarrow Y$, X and Y may be dependent, but are rendered independent given Z . Vice versa, a collider structure such as $X \rightarrow Z \leftarrow Y$ may introduce a dependence between two marginally independent variables X and Y when conditioned on Z .

While estimating (conditional) mutual information for purely discrete or continuous data is a well-studied problem [4, 5, 8, 10, 22], many real-world settings concern a mix of discrete and continuous random variables,

such as age (in years) and height, or even random variables that can individually consist of a *mixture* of discrete and continuous components. Although several discretization-based approaches that can estimate MI for a mix of discrete and continuous random variables have recently emerged [2, 17, 31], so far only methods based on k -nearest neighbour (k NN) estimation were shown to work on *mixed variables*, which may consist of discrete-continuous mixture variables [7, 19, 24].

Regardless of the success of k NN-based estimators, discretization-based approaches have attractive properties, e.g., with regard to global interpretation. That is, a natural and understandable way to discretize a continuous random variable is via creating a histogram model, where we cut the sample space of the continuous variable in multiple non-overlapping parts called bins [27], or (hyper)rectangles for multi-dimensional variables. Within a bin, we consider the distribution to be constant, which allows us to estimate the density function via Riemann integration by making the bins smaller and smaller [4]. This definition, however, is less straightforward when mixed variables are involved.

In this paper, we approach the problem as follows: we first extend the definition of entropy for a univariate discrete-continuous mixture variable given by Politis [23] to multivariate variables. Using this definition, we show that CMI for mixed random variables can be written as a sum of entropies that are well-defined through the Radon-Nikodym derivative (see Sec. 2). Exploiting this property, we propose a consistent CMI estimator for such data that is based on adaptive histogram models in Sec. 3. To efficiently learn adaptive histograms from data, in Sec. 4 we define a model selection criterion based on the minimum description length (MDL) principle [9]. Subsequently, we propose an iterative greedy algorithm that aims to obtain the histogram model that minimizes the proposed MDL score in Sec. 5. We discuss related work in Sec. 6 and in Sec. 7, we empirically show that our method performs favourably to state-of-the-art estimators for mixed data and can be used in a causal discovery setting.¹

^{*}equal contribution (ordered alphabetically)

[†]CISPA Helmholtz Center for Information Security and Max Planck Institute for Informatics, Saarland University. amarx@mpi-inf.mpg.de

[‡]Leiden Institute of Advanced Computer Science, Leiden University. {l.yang,m.van.leeuwen}@liacs.leidenuniv.nl

¹Supplementary Material is published on arXiv:2101.05009.

2 Entropy for Mixed Random Variables

We consider multi-dimensional *mixed random variables*, of which any individual dimension can be discrete, continuous, or a discrete-continuous mixture. Further, we call a vector of such mixed random variables a *mixed random vector*. For a mixed random vector (X, Y) , where X and Y are possibly multivariate, we need to adopt the most general definition of mutual information (MI), i.e., the measure-theoretic definition:

$$I(X; Y) = \int_{\mathcal{X} \times \mathcal{Y}} \log \frac{dP_{XY}}{dP_X P_Y} dP_{XY},$$

where $dP_{XY}/(dP_X P_Y)$ is the Radon-Nikodym derivative, dP_{XY} the joint measure, and $P_X P_Y$ the product measure. It has been proven that $P_X P_Y$ is *absolutely continuous* with respect to P_{XY} [7], i.e., $P_{XY} = 0$ whenever $P_X P_Y = 0$; and therefore, such a Radon-Nikodym derivative always exists and $I(X, Y)$ is well-defined. This measure-theoretic definition can be extended to CMI using the chain rule: $I(X; Y|Z) = I(X; \{Y, Z\}) - I(X; Z)$.

It is common knowledge that for purely discrete or continuous random variables, CMI can be written as a sum of entropies, i.e., $I(X; Y|Z) = H(X, Z) + H(Y, Z) - H(X, Y, Z) - H(Z)$. What is not clear, however, is if this formula also holds when (X, Y, Z) contains discrete-continuous mixture random variables. We investigate this problem in two steps. We first define the measure-theoretic entropy for a (possibly multi-dimensional) discrete-continuous mixture random variable and prove it to be well-defined, though previous work claimed the opposite [7]. Second, using this definition, we prove that (conditional) MI for a mixed random vector can be written as the sum of measure-theoretic entropies, just like purely continuous or discrete random vectors.

2.1 A Generalized Definition of Entropy The measure-theoretic entropy is defined only for one-dimensional random variables [23]. Building upon this definition, we give an explicit proof that such a one-dimensional measure-theoretic entropy is well-defined, and then extend this definition to the multi-dimensional case, which we prove is also well-defined.

2.1.1 Generalized One-Dimensional Entropy

We start off by reviewing the existing definition for the one-dimensional case [23]. Given a one-dimensional random variable X , entropy H is defined as

$$H(X) = \int_{\mathbb{R}} \frac{dP_X(x)}{dv(x)} \log \frac{dP_X(x)}{dv(x)} dv(x),$$

where $v(\cdot)$ is a measure defined on all one-dimensional Borel sets [23]. If $v(\cdot)$ is the Lebesgue measure, which we

denote as $u(\cdot)$, $H(X)$ becomes the differential entropy. Alternatively, if $v(\cdot)$ is a counting measure, $H(X)$ becomes the common (discrete) entropy.

If, however, X is a discrete-continuous mixture variable, v is defined as follows. We split $\mathbb{R}$ into three disjoint subsets s.t. $\mathbb{R} = S_d \cup S_c \cup S_o$. First, S_o is the subset of $\mathbb{R}$ on which X has zero probability measure, i.e., $P_X(S_o) = 0$. Second, the set S_d contains all discrete points, i.e., S_d is countable and $\forall x \in S_d, P_X(x) > 0$. Third, S_c covers the continuous points, hence $P_X(S_c) + P_X(S_d) = 1$ and for any Borel set $A \subseteq S_c$ satisfying $u(A) = 0$, we have $P_X(A) = 0$. Based on these three subsets S_d, S_c , and S_o , we can define v as

$$(2.1) \quad v(A) = u(A \cap S_c) + |A \cap S_d|,$$

where $|A \cap S_d|$ is the cardinality of this intersection.

To show that the generalized one-dimensional entropy is well-defined, we need to prove that the Radon-Nikodym derivative dP_X/dv always exists. This we show in the following lemma.

LEMMA 2.1. *Given a one-dimensional discrete-continuous random variable X with probability measure P_X , P_X is absolutely continuous w.r.t. v , i.e., $P_X = 0$ whenever $v = 0$, and hence dP_X/dv always exists.*

We provide the proof of Lemma 2.1, as well as for Lemmas 2.2 and 2.3 in Supplementary Material S.1.

2.1.2 Generalized Multi-Dimensional Entropy

In the following, we extend the measure-theoretic entropy definition to a mixed k -dimensional random vector $W = (W_1, \dots, W_k)$. For each W_i , we define S_d^i, S_c^i, S_o^i and measure v^i as above, and also define the *product measure* v for the k -dimensional random vector as $v = v^1 \times \dots \times v^k$. Then, define the entropy for W as

$$(2.2) \quad H(W) = \int_{\mathbb{R}^k} \frac{dP_W(w)}{dv(w)} \log \frac{dP_W(w)}{dv(w)} dv(w).$$

To prove that such entropy is well-defined, we show that dP_W/dv always exists.

LEMMA 2.2. *Given a mixed k -dimensional random vector $W = (W_1, \dots, W_k)$ with probability measure P_W , dP_W/dv always exists.*

Last, based on Lemma 2.1 and 2.2, we can prove that just like for a purely continuous or discrete random vector, conditional mutual information for a mixed random vector can be written as a sum of entropies.

LEMMA 2.3. *Given a mixed random vector (X, Y, Z) with joint probability measure P_{XYZ} , we can write $I(X; Y|Z) = H(X, Z) + H(Y, Z) - H(Z) - H(X, Y, Z)$, where each entropy can be defined as in Eq. (2.2).*

As a direct implication of the above proof, it follows that mutual information can also be written as the sum of entropies, since it is a special case of CMI with $Z = \emptyset$. With this generalized definition, we can now show how to estimate CMI using adaptive histogram models.

3 Adaptive Histogram Models

Adaptive histogram models have been thoroughly studied for continuous random variables [27]; however, to the best of our knowledge, there exists no rigorous definition of histograms for mixed random variables. Thus, to use histogram models as a foundation to estimate the measure-theoretic (conditional) MI, we need to rigorously define histograms for mixed random variables. We start with the one-dimensional case.

3.1 One-Dimensional Histogram Models A histogram model is typically defined based on a set of consecutive intervals called *bins* [27]. However, to deal with discrete-continuous mixture random variables, we define the set of bins, denoted as B , such that each bin is either an interval or a set containing only a single point. That is, $B = B' \cup B''$, where B' and B'' are sets of subsets of $\mathbb{R}$, with B' consisting of countable consecutive intervals and B'' consisting of countable single point sets. Last, we define the “width” of a bin using the measure v as defined in Eq. 2.1, i.e., for a bin $B_j \in B$ we have

$$v(B_j) = u(B_j \cap B') + |B_j \cap B''|.$$

As any $B_j \in B''$ contains only a single discrete point, $v(B_j) = 1$ for all $B_j \in B''$.

Further, we define a histogram model M as a set of bins equipped with a parameter vector of length K , where $K = |B|$ is the number of bins. That is, a histogram model M is a family of probability distributions $P_{X,\theta}$, parametrized by the vector $\theta = (\theta_1, \dots, \theta_K)$. Each element of θ represents the Radon-Nikodym derivative (or density) of each bin. Note that this definition generalizes to purely continuous random variables when $B'' = \emptyset$ and also to discrete random variables if $B' = \emptyset$. For the latter case, the histogram model degenerates to a multinomial model.

3.2 Multi-Dimensional Histograms First, we define the set of multi-dimensional bins. For a mixed k -dimensional random vector $W = (W_1, \dots, W_k)$, we define the set of *bins* for each W_i as in Sec. 3.1, denoted as B^i . Consequently, we can define a set of k -dimensional *bins*, denoted B , by the Cartesian product $B = B^1 \times \dots \times B^k$.

Since each B^i is countable, B is also countable, and we can hence assume B is indexed by j . Then, we split B in a similar way as in the one-dimensional

case, i.e., $B = B' \cup B''$, where B'' contains *only discrete values*. That is, for any k -dimensional bin $B_j \in B''$, each dimension of B_j is a set that contains a single one-dimensional point. Note that, however, for any $B_j \in B'$, each dimension of B_j can either be a one-dimensional interval or a one-dimensional single-point set. Further, we define the volume of a multi-dimensional bin $B_j \in B$ using the product measure $v(B_j)$ (see Sec. 2.1.2).

Similar to one-dimensional histograms, a multi-dimensional histogram model M can be described by a probability distribution $P_{W,\theta}$ parametrized by the vector $\theta = (\theta_1, \dots, \theta_K)$, where K is the number of bins and θ_i is the Radon-Nikodym derivative for each bin.

3.3 Maximum Likelihood Estimator Given a possibly multi-dimensional histogram with K bins, we denote the Radon-Nikodym derivative $dP_{W,\theta}/dv$ as f_θ^h and its maximum likelihood estimator as $\hat{f}_\theta^h$. Observe that for any parameter $\theta_j \in \theta$, the product $\theta_j v(B_j)$ follows a multinomial distribution. Thus, given a dataset $D = \{D_i\}_{i=1, \dots, n}$, with D_i representing a row, the maximum log-likelihood is denoted as and equal to

$$(3.3) \quad l_M(D) = \log f_{\hat{\theta}(D)}^h(D) = \log \prod_{j=1}^K \left(\frac{c_j}{n \cdot v(B_j)} \right)^{c_j},$$

where c_j and $v(B_j)$ are respectively the number of data points and the bin volumes of bin $j \in \{1 \dots K\}$. Notice that this maximum likelihood generalizes to the purely discrete case (i.e., multinomial distribution) where all $v(B_j) = 1$, and to the purely continuous case [27] where v becomes the Lebesgue measure.

3.4 Conditional Mutual Information Estimator Combining all previous theoretical discussions, we can now estimate conditional mutual information for three (possibly multivariate) random variables X, Y and Z by

$$I^h(X; Y|Z) = H^h(X, Z) + H^h(Y, Z) - H^h(X, Y, Z) - H^h(Z).$$

The corresponding measure-theoretic entropies are estimated from k -dimensional data over (X, Y, Z) , where k_X, k_Y and k_Z are the corresponding number of dimensions of X, Y and Z . We estimate the entropies as

$$\begin{aligned} H^h(X, Y, Z) &= - \int_{\mathbb{R}^k} f_\theta^h(x, y, z) \log(f_\theta^h(x, y, z)) dv \\ H^h(X, Z) &= - \int_{\mathbb{R}^{k_X + k_Z}} f_\theta^h(x, z) \log(f_\theta^h(x, z)) dv \\ H^h(Y, Z) &= - \int_{\mathbb{R}^{k_Y + k_Z}} f_\theta^h(y, z) \log(f_\theta^h(y, z)) dv \\ H^h(Z) &= - \int_{\mathbb{R}^{k_Z}} f_\theta^h(z) \log(f_\theta^h(z)) dv \end{aligned}$$

in which $f_{\hat{\theta}}^h(x, y, z)$ is the maximum likelihood estimator given the data, while we obtain $f_{\hat{\theta}}^h(x, z)$, $f_{\hat{\theta}}^h(y, z)$, and $f_{\hat{\theta}}^h(z)$ via marginalization from $f_{\hat{\theta}}^h(x, y, z)$. Next, we will prove that I^h is a strongly consistent estimator for conditional mutual information on mixed data.

THEOREM 3.1. *Given a mixed random vector (X, Y, Z) with probability measure P_{XYZ} ,*

$$\lim_{v' \rightarrow 0} \lim_{n \rightarrow \infty} I^h(X; Y|Z) = I(X; Y|Z)$$

almost surely, where n refers to the sample size and v' refers to the maximum of the histogram volumes for bins in B' (defined in Sec. 3.2).

The proof is provided in Supplementary Material S.1. Informally, our proof is based on the following key aspects: 1) All volume-related terms in I^h cancel out, 2) discrete empirical entropy converges to the true entropy almost surely [1], and 3) in the limit, differential entropy can be obtained by discretizing a continuous random variable into “infinitely” small bins [4, Theorem 8.3.1]. Notably, the order of the double limit in Theorem 3.1 inherently indicates that n should grow faster than the number of bins [26], which is also required for histograms on purely continuous data to converge [27].

4 Learning Adaptive Histograms from Data

To efficiently estimate a histogram model that inherits the consistency guarantees from Theorem 3.1 we need to consider the following requirements. First of all, we need to ensure that we learn a joint histogram model over (X, Y, Z) . This is due to the fact that we obtain the lower-dimensional entropies such as $H^h(X, Z)$ by marginalization over the likelihood estimator $f_{\hat{\theta}}^h(x, y, z)$. If we would not learn a joint model, the volume-related terms in $H^h(X, Y, Z)$, $H^h(X, Z)$, $H^h(Y, Z)$, and $H^h(Z)$ would not cancel out. In addition, we need to make sure that the number of bins is in $o(n)$ and increases if we were to increase the number of samples n , while at the same time the size of the bins decreases.

One way to achieve those properties would be to fix the bin width or the number of bins depending on the number of samples. However, such an approach is not very flexible and does not allow for variable bin widths. To allow for a more flexible model, we formally consider the problem of constructing an adaptive multi-dimensional histogram as a model selection problem and employ a selection criterion based on the minimum description length (MDL) principle [25]. MDL-based model selection has been successfully used for learning one-dimensional [13] and two-dimensional histograms [11, 36], demonstrating adaptivity to both local density changes and sample size.

We now briefly introduce the MDL principle and define the MDL-optimal histogram model. Specifically, while previous work [11, 13, 36] only considers purely continuous data (or more precisely, data with arbitrarily small precision), we apply the MDL principle to mixed-type data, based on our rigorous definition of histogram models for mixed random variables. On top of that, we empirically show that our score fulfils the desired properties—i.e. the number of bins grows as $o(n)$.

4.1 MDL and Stochastic Complexity The minimum description length principle is arguably one of the best off-the-shelf model selection criteria [9], which has been successfully applied to many machine learning and data mining tasks. The general idea is to assign a code length to data D compressed by a model M , e.g., a histogram model. Given a collection of candidate models, denoted as $\mathcal{M}$, MDL selects the model M^* that minimizes the joint code length of the model and the data. Formally, our goal is to find

$$M^* = \operatorname{argmin}_{M \in \mathcal{M}} L(D|M) + L(M),$$

where $L(D|M)$ denotes the code length² of the data given the model, while $L(M)$ denotes the code length needed to encode the model.

The optimal way of encoding data D given M , in the sense that it will result in minimax *regret*, is to use the *normalized maximum likelihood (NML)* code [9]. Accordingly, the code length of the data is called *stochastic complexity (SC)*, which is defined as the sum of the negative log-likelihood $-l_M(D)$, defined in Eq. 3.3, and the *parametric complexity* (also called *regret*) $\log R(n, K)$ [9]. The parametric complexity of a histogram model with K bins is given by [13, 36]

$$R(n, K) = \sum_{c_1 + \dots + c_K = n} \frac{n!}{c_1! \dots c_K!} \prod_{i=1}^K \left(\frac{c_i}{n} \right)^{c_i},$$

and can be computed in sub-linear time [20].

4.2 Code Length of the Model Given a dataset D with n rows and k individual columns D^j , we now define the model class $\mathcal{M}$. First, we create fixed bins according to B'' (as defined in Sec. 3.2) per discrete value that occurs in D_j . Next, we enumerate all possible bins for B' with fixed precision ϵ . To this end, denote the remaining non-discrete data points in D_j as D_j^c . If D_j^c is empty D_j corresponds to a discrete variable and we can

²The code length L denotes the number of bits needed to describe the given object. Hence, all logarithms are to base 2 and $0 \log 0 = 0$.

stop here. Otherwise, we create all possible cut points for D_j^c as $C_j^0 = \{\min(D_j^c), \min(D_j^c) + \epsilon, \dots, \max(D_j^c)\}$. By selecting a subset of cut points $C_j \subseteq C_j^0$, we get a valid solution for B' . We can enumerate all possible segmentations by enumerating each $C_j \subseteq C_j^0$.

By repeating this process for each dimension, we obtain our model class $\mathcal{M}$. Further, we get the code length for a model $M \in \mathcal{M}$ by encoding all combinations of cut points for each dimension [13], i.e.,

$$L(M) = \sum_{j \in \{1, \dots, k\}} L(C_j) = \sum_{j \in \{1, \dots, k\}} \log \left(\frac{|C_j^0|}{|C_j|} \right).$$

This completes the definition of our final optimization score $L(D|M) + L(M)$.

To proof consistency for this score, we need to show that the number of selected bins grows at rate $o(n)$. Since the theoretical analysis is rather difficult, we instead empirically demonstrate this property for Gaussian distributed data in Supplementary Material S.3. In the next section, we present an iterative greedy algorithm that optimizes our MDL score.

5 Implementation

In this section, we describe our algorithm to estimate the joint entropy $H(X_1, \dots, X_k)$ for a k -dimensional discrete-continuous mixture random vector.

5.1 Algorithm To discretize a one-dimensional random variable X , we first create bins for the discrete values of X and then discretize the continuous values. We detect discrete data points by checking if a single value x in the domain $\mathcal{X}$ of X occurs multiple times. If a user-defined threshold t , e.g., 5 is reached, we create a bin for this point. To discretize the remaining continuous values, we start by splitting $\mathcal{X}$ into K_{init} equi-width bins, which we can safely choose from the complexity class $o(\sqrt{n})$ (see Supplementary Material S.3). Using dynamic programming, we compute the variable-width histogram model M that minimizes $L(D, M)$ in quadratic time w.r.t. K_{init} [13].

Since the runtime complexity to compute the optimal variable-width histogram over a multi-dimensional random variable would grow exponentially w.r.t. k , we opt for an iterative greedy algorithm (we provide the pseudocode in Supplementary Material S.2). We start by initializing the optimization: for every dimension, we fix bins for the discrete values and put the remaining continuous values into a single bin. Then, in each iteration, we compute a candidate discretization for each dimension and keep the discretization of that dimension that provides the highest gain in compression. To compute a candidate discretization for a dimension X_j , we

extend the one-dimensional algorithm described above. That is, we determine those cut points for X_j that provide the highest gain in $L(D, M)$, while keeping the bins for the remaining dimensions fixed. We repeat this until the maximum number of iterations i_{max} is reached or we cannot further decrease $L(D, M)$.

5.2 Complexity The complexity of discretizing a univariate random variable is in $\mathcal{O}(K_{\text{max}} \cdot (K_{\text{init}})^2)$ and depends on the number of initial bins K_{init} and the maximum number of bins K_{max} , which we typically chose as a fraction of K_{init} (both in $o(\sqrt{n})$). In a multi-dimensional setting we have to multiply this complexity by the current domain size of the remaining variables, since we have to update each bin conditioned on those. In the worst case, this number is equal to $(K_{\text{max}})^{k-1}$. Overall, we apply this procedure—if all variables are continuous— $i_{\text{max}} \cdot k$ times.

6 Related Work

We discuss related methods for adaptive histograms and (conditional) mutual information estimation.

Both theoretical properties and practical issues of density estimation using histograms have been studied for decades [27]. Various algorithms have been proposed for the challenging task of constructing an adaptive one-dimensional histogram, among which the MDL-based histogram [13] is considered to be the state-of-the-art, as it is self-adaptive to both local density structure and sample size and does not have any hyperparameters. Learning adaptive multivariate histograms is even harder due to the combinatorial explosion of the search space. One approach is to resort to the dyadic CART algorithm [12]; various methods designed for specific tasks also exist [11, 35]. Our algorithm is similar to that of Kameya [11], but they only consider the two-dimensional case.

For discrete data, (conditional) mutual information estimation is a well-studied problem [4, 10, 18, 21, 33] and it has been shown that mutual information can be estimated using the $3H$ principle [10]. An important observation is that for discrete data, the empirical estimator for entropy is sub-optimal [21], which encouraged the design of more efficient entropy estimators with sub-linear sample complexity [10, 33].

For estimating (conditional) mutual information on continuous data or a mix of discrete and continuous data, three classes of approaches exist. The first class concerns kernel density estimation (KDE) methods [8, 22], which perform well on continuous data; however, no KDE-based MI and CMI estimation methods exist that are designed for discrete-continuous mixture random variables. Moreover, bandwidth tuning

for KDE can be computationally expensive, which becomes even worse for mixed data, as different bandwidths may be needed for discrete random variables. The second class of methods relies on k -nearest neighbour (k NN) estimates [6, 14, 15], which have been established as the state of the art [7, 24]. k NN approaches can be applied not only to a mix of discrete and continuous variables, but can also be used as consistent MI [7] and CMI [19, 24] estimators for discrete-continuous mixtures. The third class of methods first discretizes the continuous random variables and then calculates mutual information from the discretized variables [4, 5, 31]. Two recent approaches based on adaptive partitioning for mixed random variables have been proposed [2, 17]. While Mandros et al. [17] focus on mutual information and its application to functional dependency discovery, Cabeli et al. [2], similar to us, build upon an MDL-based score to estimate MI and CMI, to which we compare in Sec. 7. The key difference is that Cabeli et al. [2] compute $I(X; Y|Z)$ as $(I(X; \{Y, Z\}) - I(X; Z) + I(Y; \{X, Z\}) - I(Y; Z))/2$ and maximize each of the four terms (with penalty terms) directly, while we first learn a joint histogram.

To the best of knowledge, we are the first to propose a CMI estimator for discrete-continuous mixture variables based on discretization or histogram density estimation. Our method can consistently estimate CMI on mixed random variables containing discrete-continuous mixtures. We focus on histogram-based models instead of k NN estimation, since histograms are more interpretable [27] and do not require tuning of the parameter k , which can have a large impact on the outcome.

7 Experiments

In this section, we empirically evaluate the performance of our approach. First, we will benchmark our estimator against state-of-the-art CMI estimators on different data types. After that, we evaluate how well our estimator is suited to test for conditional independence in a causal discovery setup. For reproducibility, we make our code available online.³

7.1 Mutual Information Estimation On the mutual information estimation task, we compare our approach to the state-of-the-art MI estimators. In particular, we compare against FP [6], RAVK [24] and MS [19], which all rely on k NN estimates, and MIIC [2], which is a discretization-based method. All of those can be applied to our setup, but only the authors of RAVK and MS specifically consider discrete-continuous mixture variables. We apply MIIC using the default param-

eters and use $k = 10$ for all k NN-based approaches.⁴ For our algorithm, we set the maximum number of iterations and the threshold to detect discrete points in a mixture variable to 5, set $K_{\text{init}} = 20 \log n$ and $K_{\text{max}} = 5 \log n$. To comply with the literature, we compute all entropies in this section using the *natural logarithm*.

7.1.1 Experiment I-IV As a sanity check, we start with an experiment on purely continuous data. That is, for **Experiment I**, let X and Y be Gaussian distributed random variables with mean 0, variance 1, and covariance 0.6. Consequently, the correlation ρ between X and Y is 0.6 and true MI can be calculated as $I(X; Y) = -\frac{1}{2} \log(1 - \rho^2)$. In **Experiment II**, X is discrete and drawn from $\text{Unif}(0, m - 1)$, with $m = 5$ and Y is continuous with $Y \sim \text{Unif}(x, x + 2)$ for $X = x$. Therefore, $I(X; Y) = \log(m) - \frac{(m-1) \log 2}{m}$ [7]. Next, for **Experiment III**, X is exponentially distributed with rate 1 and Y is a zero-inflated Poissonization of X —i.e., $Y = 0$ with probability $p = 0.15$ and $Y \sim \text{Pois}(x)$ for $X = x$ with probability $1 - p$. The ground truth is $I(X; Y) = (1 - p)(2 \log 2 - \gamma - \sum_{k=1}^{\infty} \log k \cdot 2^{-k}) \approx (1 - p)0.3012$, where γ is the Euler-Mascheroni constant [7]. Last, in **Experiment IV**, we generate the data according to the Markov chain $X \rightarrow Z \rightarrow Y$ (see Mesner and Shalizi [19]). In particular, X is exponentially distributed with rate $\frac{1}{2}$, $Z \sim \text{Pois}(x)$ for $X = x$ and Y is binomial distributed with size $n = z$ for $Z = z$ and probability $p = \frac{1}{2}$. Due to the Markov chain structure, the ground truth $I(X; Y | Z) = 0$.

For each of the above experiments, we sample data with sample size $n \in \{100, 200, \dots, 1000\}$ and generate 100 data sets per sample size. We run each of the estimators on the generated data and show the mean squared error (MSE) of each estimator in Fig. 1. Overall, our estimator performs best or very close to the best throughout the experiments and reaches an MSE lower than 0.001 with at most 1000 samples. The best competitors are MS and MIIC; however, both are biased when we consider discrete-continuous mixture variables, as we show in Experiment V.

7.1.2 Experiment V Next, we generate data according to a discrete-continuous mixture [7]. Half of the data points are continuous, with X and Y being standard Gaussian with correlation $\rho = 0.8$, while the other half follows a discrete distribution with $P(1, 1) = P(-1, -1) = 0.4$ and $P(1, -1) = P(-1, 1) = 0.1$. In addition, we generate Z independently with $Z \sim$

³<https://github.com/ylicen/CMI-adaptive-hist.git>

⁴We evaluated all approaches with $k = 5, 10, 20$. Since $k = 10$ had the best trade-off and is close to $k = 7$ as used by Mesner and Shalizi [19], we report those results.

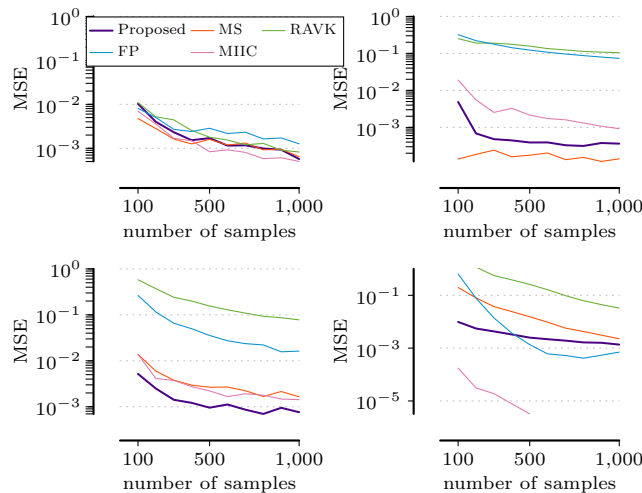


Figure 1: Synthetic data with known ground truth. Ordered from top-left to bottom-right, we show the MSE for Experiments I-IV, for our estimator and competing algorithms MS, RAVK, FP and MIIC.

Binomial(3, 0.2). Hence the ground truth is equal to $I(X; Y) = I(X; Y | Z) = 0.4 \cdot \log \frac{0.4}{0.5^2} + 0.1 \cdot \log \frac{0.1}{0.5^2} - \frac{1}{4} \log(1 - 0.8^2) \approx 0.352$.

In Fig. 2 (top) we show the mean and MSE for this experiment. We see that our estimator starts by overestimating the true value, but its average quickly converges to the true value, while the competing estimators seem to have a slightly positive or negative bias. Especially FP and MIIC, which were not designed for this setup, have a clear bias even for 1000 data points. The same trend can be observed for MSE.

7.1.3 Experiment VI Last, we test how sensitive our method is to dimensionality. We generate X and Y as in Experiment II, but fix n to 2000 and add k independent random variables, $Z_k \sim \text{Binomial}(3, 0.5)$.

Fig. 2 (bottom) shows the mean and MSE. Our estimator recovers the true CMI up to a negligible error up to $k = 2$. After that, it starts to slowly underestimate the true CMI. This can be explained by the fact that the model costs increase linearly with the domain size and hence, we will fit fewer bins to the continuous variable for large k . We validated this conjecture by repeating the experiment for $n = 10000$. On this larger sample size, the MSE for our estimator remained below 0.001 even for $k = 4$. While MIIC is slightly more stable for $k \geq 3$, the competing k NN-based estimators deviate quite a bit from the true estimate for higher dimensions.

Overall, we are on par with or outperform the best competitor throughout Experiments I-VI. Especially on

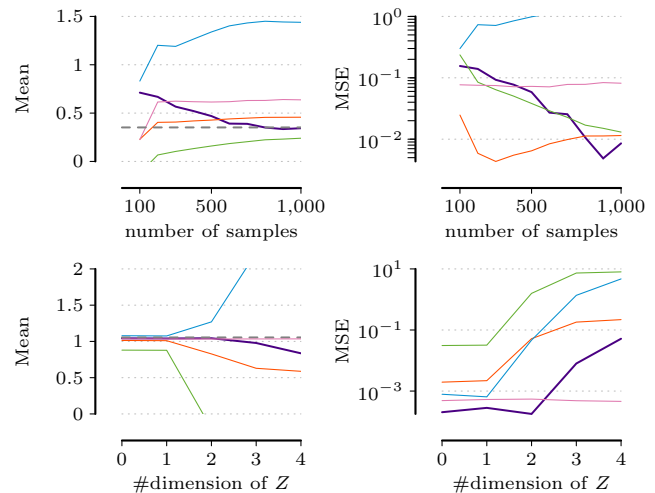


Figure 2: Top row: Experiment V, where we show the mean of the estimators (left) with the true CMI as a dashed gray line and the MSE (right). Bottom row: Experiment VI, where the sample size is constant at 2000 and the x-axis refers to the number of dimensions of Z . We show the mean (left) and MSE (right). The color coding is chosen as in Fig. 1.

mixture data, which is our main focus, our method is the only one that converges to the true estimate.

7.2 Independence Testing In theory, two random variables X and Y are conditionally independent given a set of random variables Z , denoted as $X \perp\!\!\!\perp Y | Z$, if $I(X; Y | Z) = 0$. Vice versa, X and Y are dependent given Z , if $I(X; Y | Z) > 0$. In practice, we cannot simply rely on our estimator to conclude independence: due to the monotonicity of mutual information, i.e., $I(X; Y) \leq I(X; Y \cup Z)$, estimates will rarely be *exactly zero* in the limited sample regime, but only *close to zero* [18, 34]. To address this problem, we use our algorithm to discretize X, Y and Z , and compute $I_C(X; Y | Z) := \max\{0, I_n(X_d; Y_d | Z_d) + C_n(X_d; Y_d | Z_d)\}$, where C_n is a correction term calculated from the discretized variables, which is negative. In the following, we evaluate our estimator with two different correction criteria. The first one is a correction for mutual information based on the Chi-squared distribution, with C_n equal to $-\chi_{\alpha, l}/2n$ [34], where $\chi_{\alpha, l}$ refers to the critical value of the Chi-squared distribution with significance level α and degrees of freedom l . We can compute the degrees of freedom l from the domain sizes of the discretized variables for the conditional case as $l = (|\mathcal{X}| - 1)(|\mathcal{Y}| - 1)|\mathcal{Z}|$ [32], and for the unconditional case as $l = (|\mathcal{X}| - 1)(|\mathcal{Y}| - 1)$. For the second correction, we replace each empirical entropy in $I_n(X_d; Y_d | Z_d)$

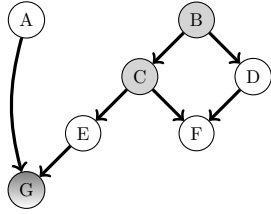


Figure 3: Synthetic network with continuous (white), discrete (gray) and mixed (shaded) random variables consisting of different causal structures, such as colliders, a chain ($C \rightarrow E \rightarrow G$), and a fork ($C \leftarrow B \rightarrow D$).

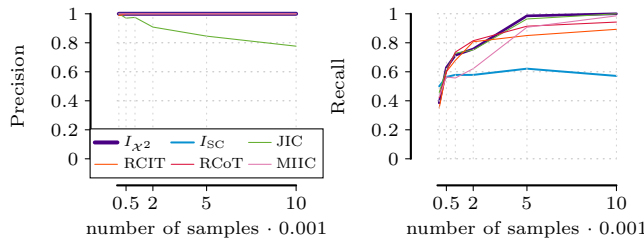


Figure 4: Precision (left) and recall (right) on undirected graphs inferred using the PC-stable algorithm equipped with the corresponding independence test. The data is generated from the graph shown in Fig. 3.

with its corresponding stochastic complexity term as defined in Sec. 4.1. If we subtract the regret terms for $H_n(X_d, Y_d, Z_d)$ and $H_n(Z_d)$ from those for $H_n(X_d, Z_d)$ and $H_n(Y_d, Z_d)$, we are guaranteed to get a negative value, thus a valid regret term [18]. In the following, we refer to the test using the Chi-squared correction as I_{χ^2} and to the one based on stochastic complexity as I_{SC} .

To test how well I_{χ^2} and I_{SC} perform on mixed-type and continuous data, we benchmark both against state-of-the-art kernel-based tests RCIT and RCoT [30], as well as JIC [31], and MIIC [2], which are both discretization-based methods using a correction based on stochastic complexity.⁵ To apply RCIT and RCoT on mixed data, we treat the discrete data points as integers. In the following, we evaluate the performance of each test in a causal discovery setup. In addition, we provide a more detailed description of the data generation and experiments on individual collider and non-collider structures in Supplementary Material S.3.

7.2.1 Causal Discovery To evaluate our test in a causal discovery setting, we generate data according to a small synthetic network—shown in Fig. 3—that consists

⁵Note that MIIC calculates stochastic complexity based on factorized NML and JIC uses an asymptotic approximation of stochastic complexity, while we use quotient NML for I_{SC} [18, 28].

of a mixture of generating mechanisms that we used in experiments I-IV and includes continuous and discrete (ordinal) random variables and one mixture variable, which is partially Gaussian and partially Poisson distributed (for details see Supplementary Material S.3). To evaluate how well the ground truth graph can be recovered, we apply the PC-stable algorithm [3, 29] equipped with the different independence tests, where we use $\alpha = 0.01$ for I_{χ^2} , RCIT and RCoT.

Fig 4 shows recovery precision and recall for the undirected graph, averaged over 20 draws per sample size $n \in \{100, 500, 1\,000, 2\,000, 5\,000, 10\,000\}$.

We see that overall I_{χ^2} performs best and is the only method that reaches both a perfect accuracy and recall. While JIC also reaches a perfect recall, it finds too many edges leading to a precision of only 80%. Although also MIIC, RCIT and RCoT have a perfect precision, their recall is worse than for I_{χ^2} . Neither of the kernel-based tests manages to recall all the edges even for 10 000 samples. After a closer inspection, this is due to the edge $E \rightarrow G$ that involves the discrete-continuous variable G . If we compare I_{χ^2} to I_{SC} , we clearly see that the latter is too conservative, which leads to a bad recall.

8 Conclusion

We proposed a novel approach for the estimation of conditional mutual information from data that may contain discrete, continuous, and mixture variables. To be able to deal with discrete-continuous mixture variables, we defined a class of generalized adaptive histogram models. Based on our observation that CMI for mixture-variables can be written as a sum of entropies, we presented a CMI estimator based on such histograms, for which we proved that it is consistent.

Further, we used the minimum description length principle to formally define optimal histograms, and proposed a greedy algorithm to practically learn good histograms from data. Finally, we demonstrated that our algorithm outperforms state-of-the-art (conditional) mutual information estimation methods, and that it can be successfully used as a conditional independence test in causal graph structure learning. Notably, for both setups, we observe that our approach performs especially well when mixture variables are present.

Acknowledgements

This work is part of the research programme ‘Human-Guided Data Science by Interactive Model Selection’ with project number 612.001.804, which is (partly) financed by the Dutch Research Council (NWO).

References

- [1] A. Antos and I. Kontoyiannis, "Convergence properties of functional estimates for discrete distributions," *Random Structures & Algorithms*, vol. 19, no. 3-4, pp. 163–193, 2001.
- [2] V. Cabeli, L. Verny, N. Sella, G. Uguzzoni, M. Verny, and H. Isambert, "Learning clinical networks from medical records based on information estimates in mixed-type data," *PLOS Comput. Biol.*, vol. 16, no. 5, p. e1007866, 2020.
- [3] D. Colombo and M. H. Maathuis, "A modification of the pc algorithm yielding order-independent skeletons," *CoRR*, abs/1211.3295, 2012.
- [4] T. M. Cover and J. A. Thomas, *Elements of information theory*. John Wiley & Sons, 2012.
- [5] G. A. Darbellay and I. Vajda, "Estimation of the information by an adaptive partitioning of the observation space," *IEEE Trans. Inf. Theorie*, vol. 45, no. 4, pp. 1315–1321, 1999.
- [6] S. Frenzel and B. Pompe, "Partial mutual information for coupling analysis of multivariate time series," *Physical Review Letters*, vol. 99, no. 20, p. 204101, 2007.
- [7] W. Gao, S. Kannan, S. Oh, and P. Viswanath, "Estimating mutual information for discrete-continuous mixtures," in *NIPS*, 2017, pp. 5986–5997.
- [8] W. Gao, S. Oh, and P. Viswanath, "Breaking the bandwidth barrier: Geometrical adaptive entropy estimation," in *NIPS*. Curran Associates, Inc., 2016, pp. 2460–2468.
- [9] P. Grünwald, *The Minimum Description Length Principle*. MIT Press, 2007.
- [10] Y. Han, J. Jiao, and T. Weissman, "Adaptive estimation of shannon entropy," in *ISIT*. IEEE, 2015, pp. 1372–1376.
- [11] Y. Kameya, "Time series discretization via mdl-based histogram density estimation," in *ICTAI*. IEEE, 2011, pp. 732–739.
- [12] J. Klemelä, "Multivariate histograms with data-dependent partitions," *Statistica sinica*, pp. 159–176, 2009.
- [13] P. Kontkanen and P. Myllymäki, "MDL histogram density estimation," in *AISTATS*, 2007.
- [14] L. F. Kozachenko and N. N. Leonenko, "Sample estimate of the entropy of a random vector," *Probl. Peredachi Inf.*, vol. 23, pp. 9–16, 1987.
- [15] A. Kraskov, H. Stögbauer, and P. Grassberger, "Estimating mutual information," *Phys. Rev. E*, vol. 69, no. 6, p. 066138, 2004.
- [16] J. Lee and D.-W. Kim, "Feature selection for multi-label classification using multivariate mutual information," *Pattern Recognition Letters*, vol. 34, no. 3, pp. 349–357, 2013.
- [17] P. Mandros, D. Kaltenpoth, M. Boley, and J. Vreeken, "Discovering functional dependencies from mixed-type data," in *KDD*, 2020, pp. 1404–1414.
- [18] A. Marx and J. Vreeken, "Testing conditional independence on discrete data using stochastic complexity," in *AISTATS*, 2019, pp. 496–505.
- [19] O. C. Mesner and C. R. Shalizi, "Conditional Mutual Information Estimation for Mixed, Discrete and Continuous, Data," *IEEE Transactions on Information Theory*, 2020.
- [20] T. Mononen and P. Myllymäki, "Computing the multinomial stochastic complexity in sub-linear time," in *PGM*, 2008, pp. 209–216.
- [21] L. Paninski, "Estimation of entropy and mutual information," *Neural Computation*, vol. 15, no. 6, pp. 1191–1253, 2003.
- [22] L. Paninski and M. Yajima, "Undersmoothed kernel entropy estimators," *IEEE Trans. Inf. Theorie*, vol. 54, no. 9, pp. 4384–4388, 2008.
- [23] D. N. Politis, *On the entropy of a mixture distribution*. Purdue University. Department of Statistics, 1991.
- [24] A. Rahimzamani, H. Asnani, P. Viswanath, and S. Kannan, "Estimators for multivariate information measures in general probability spaces," in *NIPS*, 2018, pp. 8664–8675.
- [25] J. Rissanen, "Modeling by shortest data description," *Automatica*, vol. 14, no. 1, pp. 465–471, 1978.
- [26] W. Rudin et al., *Principles of mathematical analysis*. McGraw-hill New York, 1964, vol. 3.
- [27] D. W. Scott, *Multivariate density estimation: theory, practice, and visualization*. John Wiley & Sons, 2015.
- [28] T. Silander, J. Leppä-aho, E. Jääsaari, and T. Roos, "Quotient normalized maximum likelihood criterion for learning bayesian network structures," in *AISTATS*, vol. 84. PMLR, 2018, pp. 948–957.
- [29] P. Spirtes, C. N. Glymour, R. Scheines, and D. Heckerman, *Causation, prediction, and search*. MIT press, 2000.
- [30] E. V. Strobl, K. Zhang, and S. Visweswaran, "Approximate kernel-based conditional independence tests for fast non-parametric causal discovery," *JCI*, vol. 7, no. 1, 2019.
- [31] J. Suzuki, "An estimator of mutual information and its application to independence testing," *Entropy*, vol. 18, no. 4, p. 109, 2016.
- [32] —, "Mutual information estimation: Independence detection and consistency," in *ISIT*. IEEE, 2019, pp. 2514–2518.
- [33] G. Valiant and P. Valiant, "Estimating the unseen: an $n/\log(n)$ -sample estimator for entropy and support size, shown optimal via new clts," in *STOC*, 2011, pp. 685–694.
- [34] N. X. Vinh, J. Chan, and J. Bailey, "Reconsidering mutual information based feature selection: A statistical significance view," in *AAAI*, 2014.
- [35] D. Weiler and J. Eggert, "Multi-dimensional histogram-based image segmentation," in *ICONIP*. Springer, 2007, pp. 963–972.
- [36] L. Yang, M. Baratchi, and M. van Leeuwen, "Unsupervised discretization by two-dimensional mdl-based histogram," *arXiv preprint arXiv:2006.01893*, 2020.