
CARMS: Categorical-Antithetic-REINFORCE Multi-Sample Gradient Estimator

Anonymous Author(s)

Affiliation

Address

email

Abstract

1 Accurately backpropagating the gradient through categorical variables is a chal-
2 lenging task that arises in various domains, such as training discrete latent variable
3 models. To this end, we propose CARMS, an unbiased estimator for categorical
4 random variables based on multiple mutually negatively correlated (jointly anti-
5 thetic) samples. CARMS combines REINFORCE with copula based sampling
6 to avoid duplicate samples and reduce the variance, while keeping the estimator
7 unbiased using a simple multiplicative term. It generalizes both the ARMS anti-
8 thetic estimator for binary variables, which is CARMS for two categories, as well
9 as LOORF/VarGrad, the leave-one-out REINFORCE estimator, which is CARMS
10 with independent samples. We evaluate CARMS on several benchmark datasets on
11 a generative modeling task, as well as a structured output prediction task, and find
12 it to outperform competing methods including a strong self-control baseline. The
13 code is available in the supplementary material.

14 1 Introduction

15 When optimizing an expectation based objective of the form $\mathbb{E}_{z \sim q_\phi(z)}[f(z)]$, we sometimes require
16 the gradients with respect to the parameters ϕ of the distribution. This is challenging for discrete
17 variables, because the commonly used reparameterization gradient does not directly work, unless the
18 discrete distribution is approximated by a continuous and reparameterizable one [Jang et al., 2017,
19 Maddison et al., 2017]. A significant part of this field has thus been score function (REINFORCE)
20 based estimators [Glynn, 1990, Williams, 1992, Fu, 2006], which are general and do not require
21 differentiability of f . In this paper, we focus on the case when z is a high dimensional categorical
22 variable with logits ϕ . An common example of this form is the evidence lower bound (ELBO) [Jordan
23 et al., 1998], which arises in variational inference, which is used for training variational autoen-
24 coders [Kingma and Welling, 2014, Rezende et al., 2014]. Because their latent space consists of
25 a large number of categorical variables, they require Monte Carlo gradients with respect to the
26 parameters of the stochastic distribution, but have excellent performance, as a categorical VAE has in
27 practice achieved state of the art zero-shot image generation [Ramesh et al., 2021].

28 Our main contribution is a novel unbiased and low variance gradient estimator for categorical
29 variables. The Categorical-Antithetic-REINFORCE-Multi-Sample (CARMS) estimator uses a copula
30 to generate any number of antithetic (mutually negatively correlated) [Owen, 2013] categorical
31 samples, and constructs an unbiased estimator by combining them into a baseline for variance
32 reduction. This approach is inspired by the ARMS estimator [Dimitriev and Zhou, 2021], which also
33 uses multiple antithetic samples, but only works for binary variables. For two categories, CARMS
34 reduces to it, while for independent samples, CARMS reduces to the leave-one-out-REINFORCE
35 (LOORF) estimator. Our approach achieves higher ELBO with VAE models than the state of the art,
36 as well as higher log likelihood for conditional image completion.

Related work One widely used group of gradient estimators for categorical variables is based on trading off bias for lower variance. This includes the straight through (ST) estimator [Bengio et al., 2013], the direct argmax [Lorberbom et al., 2019], as well as the concurrently developed equivalent Gumbel-Softmax (GS) [Jang et al., 2017] or Concrete [Maddison et al., 2017] that uses a continuous relaxation. These were further improved by combining them with REINFORCE to obtain unbiased estimators, which includes REBAR [Tucker et al., 2017], as well as RELAX [Grathwohl et al., 2018], which uses a free form neural network.

In practice, REINFORCE is almost always augmented by baselines, e.g. in variational inference [Mnih and Gregor, 2014, Ranganath et al., 2014, Paisley et al., 2012, Ruiz et al., 2016, Kucukelbir et al., 2017]. MuProp [Gu et al., 2016] uses a first order mean field Taylor approximation, but requires f to be differentiable. Other estimators apply Rao-Blackwellization, *e.g.*, to one latent variable at a time [Titsias and Lázaro-Gredilla, 2015] or to the reparameterized Dirichlet vector in ARSM [Dong et al., 2021]. Besides REINFORCE and reparameterization type gradients, there is also the measure valued gradient [Rosca et al., 2019] and finite differences [Fu, 2006], but are less common in practice. A comprehensive review can be found in Mohamed et al. [2020].

The more recent approaches for categorical variables are unbiased and REINFORCE based, building on the idea of using multiple samples to construct a baseline for variance reduction. One such baseline is LOORF [Kool et al., 2019a], originally introduced in Salimans and Knowles [2014] and also known as VarGrad [Richter et al., 2020], where its theoretical properties are further analyzed. VIMCO [Mnih and Rezende, 2016] has a similar form, but is specific to the importance weighted multi sample bound [Burda et al., 2016]. A different approach is ARSM [Yin et al., 2019], which reparameterizes the gradient with a Dirichlet distribution and uses swaps to obtain multiple correlated samples. However, it adds some amount of variance due to the continuous reparameterization, and can require up to $C(C-1)/2$ function evaluations per step, which can be computationally expensive. More recently, the unordered set estimator (UNORD) [Kool et al., 2020] uses the Gumbel top-k trick to sample without replacement, and then constructs a baseline using a multiplicative term to preserve unbiasedness.

2 Background

Let $\mathbf{z} = (z_1, \dots, z_D)$ denote D independent categorical variables, where z_d is a one hot encoded categorical sample from $z_d \sim q_{\phi_d}(z_d) = \text{Cat}(\sigma(\phi_d))$, with $\sigma(\mathbf{x})_i = e^{x_i} / \sum_j e^{x_j}$ being the softmax function. Let also $\hat{\mathbf{i}}$ be the basis vector with its i^{th} coordinate set to one, and otherwise zero, such that $P(\mathbf{z} = \hat{\mathbf{i}}) = \sigma(\phi)_i$, or equivalently $\mathbb{E}[\mathbf{z}] = \sigma(\phi)$. Unless otherwise stated, superscripts denote the dimension, and subscripts denote different samples, with vectors and matrices being bold lowercase and bold uppercase symbols, respectively. We are interested in optimizing the following objective with respect to the logits $\phi = (\phi_1, \dots, \phi_D)$:

$$\mathcal{L}(\phi) = \mathbb{E}_{\mathbf{z} \sim q_{\phi}(\mathbf{z})}[f(\mathbf{z})], \quad q_{\phi}(\mathbf{z}) = \prod_{d=1}^D q_{\phi_d}(z_d).$$

Although the score function gradient contains two terms:

$$\nabla_{\phi} \mathcal{L}(\phi) = \mathbb{E}_{\mathbf{z} \sim q_{\phi}(\mathbf{z})} \left[\nabla_{\phi} f(\mathbf{z}) + f(\mathbf{z}) \nabla_{\phi} \ln q_{\phi}(\mathbf{z}) \right]$$

the first term is easily estimated, so we omit the subscript in f for notational clarity, and we focus on the latter term in this work. Since CARMS generalizes the multisample LOORF estimator beyond independent samples, we review it here. We also review the ARMS estimator, which uses a copula to generate antithetic samples, but is restricted to binary variables.

2.1 LOORF

Leave-one-out-REINFORCE (LOORF) [Salimans and Knowles, 2014, Kool et al., 2019a], also known as VarGrad [Richter et al., 2020], is a general score function based gradient estimator that uses N i.i.d. samples to construct a baseline for variance reduction, and is competitive with state of the art estimators. Its theoretical properties have recently been analyzed in Richter et al. [2020], where the same estimator results from minimizing the variance of the log ratio between the posterior and approximating distribution in variational inference. The only requirements for LOORF are the ability

84 to sample $\mathbf{z} \sim q_\phi(\mathbf{z})$ and evaluate $f(\mathbf{z})$. Given N samples $\mathbf{z}_1, \dots, \mathbf{z}_N \stackrel{iid}{\sim} \text{Cat}(\sigma(\phi))$, it has the form:
 85

$$g_{\text{LOORF}} = \frac{1}{N-1} \sum_{n=1}^N \left(f(\mathbf{z}_n) - \bar{f}(\mathbf{z}) \right) \nabla_\phi \ln q_\phi(\mathbf{z}_n) = \frac{1}{N-1} \sum_{n=1}^N \left(f(\mathbf{z}_n) - \bar{f}(\mathbf{z}) \right) (\mathbf{z}_n - \sigma(\phi_n)), \quad (1)$$

86 where $\bar{f}(\mathbf{z}) = \frac{1}{N} \sum_{n=1}^N f(\mathbf{z}_n)$. It is commonly used due to its simplicity and strong performance.

87 2.2 ARMS

88 The Antithetic-REINFORCE-MultiSample (ARMS) [Dimitriev and Zhou, 2021] estimator is a recent
 89 work that uses any number N of mutually negatively correlated samples. However, although unbiased,
 90 it is limited to binary variables, which motivated us to extend it to the categorical case. Since any
 91 antithetic copula in two dimensions reduces to the pair $(u, 1 - u)$, it generalizes DisARM [Dong
 92 et al., 2020], independently discovered as U2G [Yin et al., 2020], which use $N = 2$ samples. ARMS
 93 achieves this generalization by using a copula, which is any multivariate distribution whose marginals
 94 are uniform random variables:

$$\mathbf{u} = (u_1, \dots, u_N) \sim \mathcal{C}_N, \quad \forall i : u_i \sim \text{Unif}(0, 1).$$

95 More specifically, ARMS uses a Dirichlet or a Gaussian copula, which both have very strong negative
 96 dependence between each dimension of $\mathbf{u}$. However any copula can be used instead, with the only
 97 two requirements being the ability to generate samples easily, as well as being able to evaluate the
 98 bivariate CDF $\Phi(\mathbf{u}_i, \mathbf{u}_j)$ of the copula, so that the debiasing term can be calculated. Given a copula
 99 sample $\mathbf{u}$, it uses inverse CDF sampling to convert it into N antithetic $\text{Bern}(p)$ samples, which are
 100 simply $b_i = \mathbb{1}_{u_i < p}$, $\forall i$. For a D -dimensional vector of antithetic Bernoulli variables $\mathbf{b}_1, \dots, \mathbf{b}_N$ with
 101 probabilities $\sigma(\phi)$, the N -sample unbiased estimator has the following simple form:

$$g_{\text{ARMS}} = \frac{1}{N-1} \sum_{n=1}^N \left(f(\mathbf{b}_n) - \frac{1}{n} \sum_{m=1}^N f(\mathbf{b}_m) \right) \frac{\mathbf{b}_n - \sigma(\phi)}{1 - \rho}, \quad (2)$$

102 with $\rho = (\rho_1, \dots, \rho_D)$, and $\rho_d = \text{corr}(\mathbf{b}_i^d, \mathbf{b}_j^d)$ being the correlation of the d^{th} Bernoulli variable. It is
 103 easy to compute given the bivariate CDF of the copula.

104 3 CARMS

105 The CARMS estimator has a similar form to LOORF, but requires a multiplicative term to remain
 106 unbiased. It has two requirements: an easy way to sample antithetic categorical variables, and being
 107 able to compute the bivariate probability mass function (PMF), which should be identical for any
 108 pair. For clarity, we begin by assuming the ability to do this, and derive the univariate version for
 109 two samples, which we then extend to N samples. Next, we generalize CARMS to any number of
 110 categorical variables. Lastly, in Section 3.2, we show two different ways of sampling that satisfy
 111 both conditions: inverse CDF sampling with an easily computable analytical bivariate PMF, and the
 112 Gumbel max trick combined with an empirical estimate of the PMF.

113 The two sample version of CARMS can be obtained by replacing the two i.i.d. samples in 2-LOORF
 114 with an arbitrarily correlated pair of categorical variables and an added debiasing term. For $N = 2$
 115 samples $\mathbf{z}, \mathbf{z}' \sim \text{Cat}(\sigma(\phi))$, LOORF has the following simple form:

$$g_{\text{LOORF}}(\mathbf{z}, \mathbf{z}') = \frac{1}{2} \left(f(\mathbf{z}) - f(\mathbf{z}') \right) (\mathbf{z} - \mathbf{z}'), \quad (3)$$

116 which, importantly, is unbiased [Kool et al., 2019a, Dimitriev and Zhou, 2021]. This means that
 117 using a simple importance weight preserves its unbiasedness, for an arbitrary bivariate categorical
 118 distribution $(\mathbf{z}, \mathbf{z}') \sim \text{Cat}(\sigma(\phi), \sigma(\phi))$:

$$\begin{aligned} \mathbb{E} \left[g_{\text{LOORF}}(\mathbf{z}, \mathbf{z}') \frac{P(\mathbf{z} = \hat{\mathbf{i}})P(\mathbf{z}' = \hat{\mathbf{j}})}{P(\mathbf{z} = \hat{\mathbf{i}}, \mathbf{z}' = \hat{\mathbf{j}})} \right] &= \sum_{i,j} P(\mathbf{z} = \hat{\mathbf{i}}, \mathbf{z}' = \hat{\mathbf{j}}) \frac{P(\mathbf{z} = \hat{\mathbf{i}})P(\mathbf{z}' = \hat{\mathbf{j}})}{P(\mathbf{z} = \hat{\mathbf{i}}, \mathbf{z}' = \hat{\mathbf{j}})} g_{\text{LOORF}}(\hat{\mathbf{i}}, \hat{\mathbf{j}}) \\ &= \sum_{i,j} P(\mathbf{z} = \hat{\mathbf{i}})P(\mathbf{z}' = \hat{\mathbf{j}}) g_{\text{LOORF}}(\hat{\mathbf{i}}, \hat{\mathbf{j}}) = \mathbb{E}_{\mathbf{z}, \mathbf{z}' \stackrel{iid}{\sim} \text{Cat}(\sigma(\phi))} [g_{\text{LOORF}}(\mathbf{z}, \mathbf{z}')] = \nabla_\phi \mathcal{L}(\phi). \end{aligned}$$

119 We summarize the derivation of the two sample version of CARMS, which we denote as the
 120 Categorical-Antithetic-REINFORCE-Two-Sample (CARTS) estimator in the following theorem.

121 **Theorem 1** Let (z, z') be a sample from an arbitrary bivariate Categorical distribution with
 122 marginal distributions $z, z' \sim \text{Cat}(\sigma(\phi))$, and a known bivariate PMF. An unbiased estimator
 123 of $\nabla_{\phi} \mathbb{E}[f(z)]$ is:

$$g_{\text{CARTS}}(z, z') = \frac{1}{2} \left(f(z) - f(z') \right) (z - z') z^T \mathcal{R} z', \quad \mathcal{R}_{ij} = \frac{\sigma(\phi)_i \sigma(\phi)_j}{P(z = \hat{i}, z' = \hat{j})}. \quad (4)$$

124 The intuition behind using negatively correlated variables is that we want to avoid the case when
 125 $z = z'$, because the sample is then “wasted.” We formalize this intuition below, and defer the proof
 126 to the appendix.

127 **Theorem 2** Let $(z, z') \sim \text{Cat}(\sigma(\phi), \sigma(\phi))$. If the bivariate PMF satisfies:

$$\forall i \neq j : P(z = \hat{i}, z' = \hat{j}) \geq P(z = \hat{i}) P(z' = \hat{j}),$$

128 with strict inequality for at least one pair, then $\text{Var}[g_{\text{CARTS}}(z, z')] < \text{Var}[g_{\text{LOORF}}(z, z')]$.

129 We now extend CARTS to N samples, using the following identity, the proof of which can be found
 130 in Dimitriev and Zhou [2021]. It states that N -sample LOORF is equivalent to averaging 2-sample
 131 LOORF over all $\binom{N}{2}$ pairs:

$$\begin{aligned} g_{\text{LOORF}}(z_1, \dots, z_N) &= \frac{1}{N} \sum_{n=1}^N \left(f(z_n) - \frac{1}{N} \sum_{m=1}^N f(z_m) \right) (z_n - \sigma(\phi)) \\ &= \frac{1}{N(N-1)} \sum_{n \neq m} \frac{1}{2} \left(f(z_n) - f(z_m) \right) (z_n - z_m) = \frac{1}{N(N-1)} \sum_{n \neq m} g_{\text{LOORF}}(z_n, z_m) \end{aligned}$$

132 With the above identity it easily follows that given N antithetic Categorical samples $\mathbf{Z} =$
 133 $[z_1, \dots, z_N]^T$, applying CARTS to all pairs results in an unbiased estimator, due to linearity of
 134 expectations:

$$\begin{aligned} \mathbb{E}[g_{\text{CARMS}}(\mathbf{Z}, \mathcal{R})] &= \mathbb{E} \left[\frac{1}{N(N-1)} \sum_{n \neq m} g_{\text{CARTS}}(z_n, z_m, \mathcal{R}) \right] = \mathbb{E} \left[\frac{1}{N(N-1)} \sum_{n \neq m} g_{\text{LOORF}}(z_n, z_m) \right] \\ &= \mathbb{E}[g_{\text{LOORF}}(\mathbf{Z})] = \nabla_{\phi} \mathcal{L}(\phi). \end{aligned}$$

135 We summarize CARMS in the next theorem. and we also rewrite it in a simpler matrix form used in
 136 our implementation. The matrix form can be obtained after some algebra, which can be found in the
 137 appendix.

138 **Theorem 3** Let $\mathbf{Z} = [z_1, \dots, z_N]^T$ be a sample from an arbitrary N -variate Categorical dis-
 139 tribution with identical marginal and bivariate distributions, such that $z_i \sim \text{Cat}(\sigma(\phi))$, and
 140 $\mathcal{R}_{ij} = \sigma(\phi)_i \sigma(\phi)_j / P(z_n = \hat{i}, z_m = \hat{j})$. An unbiased estimator of $\nabla_{\phi} \mathbb{E}[f(\mathbf{z})]$ is:

$$g_{\text{CARMS}}(\mathbf{Z}) = \frac{1}{N(N-1)} \sum_{n \neq m} \frac{1}{2} \left(f(z_n) - f(z_m) \right) (z_n - z_m) z_n^T \mathcal{R} z_m'$$

141 **Lemma 4** Let $f(\mathbf{Z}) = [f(z_1), \dots, f(z_N)]^T$, $\circ$ denote the Hadamard product, $\mathbf{1}_{N \times N}$ a matrix of
 142 ones, and $\mathbf{I}_N$ the identity matrix. Define $\mathcal{O} = \frac{1}{N-1} (\mathbf{1}_{N \times N} - \mathbf{I}_N) \circ (\mathbf{Z} \mathcal{R} \mathbf{Z}^T)$, and $\mathcal{D} = \text{diag}(\mathcal{O} \mathbf{1}_N)$,
 143 to be a diagonal matrix. The CARMS estimator can equivalently be written in the following form:

$$g_{\text{CARMS}}(\mathbf{Z}) = \frac{1}{N} f(\mathbf{Z})^T (\mathcal{D} - \mathcal{O}) (\mathbf{Z} - \mathbf{1}_N \sigma(\phi)^T). \quad (5)$$

144 Furthermore, for independent samples, $\mathbf{Z} \mathcal{R} \mathbf{Z}^T = \mathbf{1}_{N \times N}$ and LOORF has the form:

$$g_{\text{LOORF}}(\mathbf{Z}) = \frac{1}{N} f(\mathbf{Z})^T \left(\mathbf{I}_{N \times N} - \frac{1}{N-1} (\mathbf{1}_{N \times N} - \mathbf{I}_N) \right) (\mathbf{Z} - \mathbf{1}_N \sigma(\phi)^T) \quad (6)$$

145 3.1 Multivariate CARMS

146 In the univariate case, Monte Carlo estimators are not necessary, as the expectation has C terms
 147 and can simply be analytically summed, but for many categorical variables, stochastic gradients are
 148 required. For D dimensions, we can sample $\mathbf{Z}^d \sim \text{Cat}(\sigma(\phi^d))$ independently and combine them
 149 into a $D \times N \times C$ tensor $\mathbf{Z}$, with the corresponding $D \times C \times C$ importance ratio tensor $\mathbf{R}$. Below,
 150 we use superscripts and subscripts to index the first and second dimension of the tensors. Focusing
 151 on the d^{th} dimension, we have:

$$\begin{aligned} \nabla_{\phi^d} \mathbb{E}[f(\mathbf{Z})] &= \mathbb{E}_{\mathbf{Z}^{-d}} [\nabla_{\phi^d} \mathbb{E}_{\mathbf{Z}^d} [f(\mathbf{Z}^{-d}, \mathbf{Z}^d)]] \\ &= \mathbb{E}_{\mathbf{Z}^{-d}} \left[\mathbb{E}_{\mathbf{Z}^d} \left[\frac{1}{N(N-1)} \sum_{n \neq m} \frac{1}{2} (f(\mathbf{Z}_n^{-d}, \mathbf{Z}_n^d) - f(\mathbf{Z}_m^{-d}, \mathbf{Z}_m^d)) (\mathbf{Z}_n^d - \mathbf{Z}_m^d) \mathbf{Z}_n^{dT} \mathbf{R}^d \mathbf{Z}_m^d \right] \right] \\ &= \mathbb{E} \left[\frac{1}{N(N-1)} \sum_{n \neq m} \frac{1}{2} (f(\mathbf{Z}_n) - f(\mathbf{Z}_m)) (\mathbf{Z}_n^d - \mathbf{Z}_m^d) \mathbf{Z}_n^{dT} \mathbf{R}^d \mathbf{Z}_m^d \right]. \end{aligned}$$

152 Just like the univariate case, we only need N evaluations of f regardless of the dimensionality D .

153 3.2 Antithetic categorical variables

154 To be able to use CARMS in practice, we describe two ways to generate antithetic categorical
 155 variables in this section. Both methods are based on transformations of uniform random variables,
 156 which makes them amenable to antithetic copulas, such as the Gaussian or Dirichlet copula [Dimitriev
 157 and Zhou, 2021].

158 3.2.1 Inverse CDF sampling

159 The inverse transform sampling is a commonly used identity to transform uniform random variables,
 160 which are easy to generate, to another distribution:

$$x \sim F_x(x) \iff u \sim \text{Unif}(0, 1), x = F_x^{-1}(u),$$

161 where $F_x(x)$ denotes the CDF of the desired distribution, and is simple for categorical variables. Let
 162 $\mathbf{p}^{(o)} = (p_1^{(o)}, \dots, p_C^{(o)})$ be a vector of probabilities, which are reordered according to some ordering o .
 163 Define the left and right boundaries:

$$\mathbf{l}_i^{(o)} = \sum_{j=1}^{i-1} p_j^{(o)}, \quad \mathbf{r}_j^{(o)} = \sum_{j=1}^i p_j^{(o)}. \quad (7)$$

164 Then, we can transform $u \sim \text{Unif}(0, 1)$ by setting $z = \hat{\mathbf{j}}$, where j is such that $u \in [\mathbf{l}_j^{(o)}, \mathbf{r}_j^{(o)}]$. To
 165 obtain N antithetic categorical variables, we can sample $\mathbf{u} \sim \mathcal{C}_N$ for a given copula and set $\mathbf{z}_n$
 166 in a vectorized manner. Importantly, this sampling approach allows us to analytically evaluate the
 167 bivariate PMF:

$$\begin{aligned} P(\mathbf{z}^{(o)} = \hat{\mathbf{i}}, \mathbf{z}'^{(o)} = \hat{\mathbf{j}}) &= P(u \in [\mathbf{l}_i^{(o)}, \mathbf{r}_i^{(o)}], u' \in [\mathbf{l}_j^{(o)}, \mathbf{r}_j^{(o)}]) = \\ &= \Phi_{\mathcal{C}}(\mathbf{l}_i^{(o)}, \mathbf{l}_j^{(o)}) + \Phi_{\mathcal{C}}(\mathbf{r}_i^{(o)}, \mathbf{r}_j^{(o)}) - \Phi_{\mathcal{C}}(\mathbf{l}_i^{(o)}, \mathbf{r}_j^{(o)}) - \Phi_{\mathcal{C}}(\mathbf{r}_i^{(o)}, \mathbf{l}_j^{(o)}), \end{aligned} \quad (8)$$

168 where $\Phi_{\mathcal{C}}$ denotes the bivariate CDF of the copula. Unfortunately, keeping to one specific ordering
 169 does not allow for all possible pairs to have non-zero probabilities, so we randomly reorder $\mathbf{p}$ at every
 170 sampling step in the following manner. An ordering o consists of two indexes: i, j uniformly sampled
 171 from $\{1, \dots, C\}$. First we translate $\mathbf{p}$ i elements to the right, then swap the last element with the j^{th} :

$$\forall k : \mathbf{p}_k = \mathbf{p}_{(k+i) \bmod C}, \quad \text{and} \quad \mathbf{p}_j \leftrightarrow \mathbf{p}_C. \quad (9)$$

172 This guarantees at least one of the $C(C-1)/2$ orderings has i and j as the first and last elements,
 173 respectively. And since i, j are uniformly chosen, this allows us to easily compute the needed bivariate

Algorithm 1 Antithetic inverse CDF categorical sampling

Input: Number of samples N , probabilities $\mathbf{p} = \sigma(\phi)$, copula $\mathcal{C}$.

Sample $\mathbf{u} = (u_1, \dots, u_N) \sim \mathcal{C}_N$.

Shuffle $\mathbf{p}$ according to Eq. 9, and define the left and right boundaries $\mathbf{l}, \mathbf{r}$ according to Eq. 7.

Set $\forall n: z_n = \hat{j}$, where j is such that $u_n \in [l_j, r_j]$.

For $i, j \in \{1, \dots, C\}$: compute $\mathcal{R}_{ij} = \mathbf{p}_i \mathbf{p}_j / P(z = \hat{i}, z' = \hat{j})$ according to Eq. 8.

return: $(z_1, \dots, z_N), \mathcal{R}$.

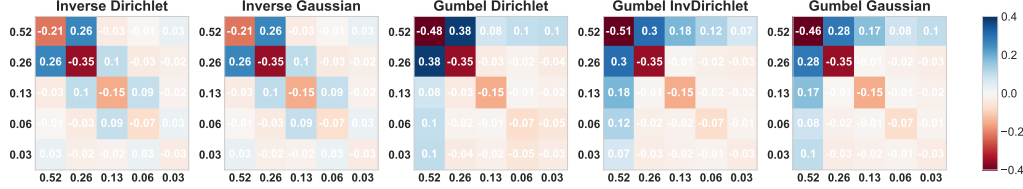


Figure 1: Correlation matrix for each pair of variables $\rho_{ij} = \text{corr}(z_i, z_j)$ using both types of categorical sampling and different copulas.

174 PMF for $\mathcal{R}$, given all of the orderings:

$$P(z = \hat{i}, z' = \hat{j}) = \frac{2}{C(C-1)} \sum_{o \in O} P(z^{(o)} = \hat{i}, z'^{(o)} = \hat{j}) \quad (10)$$

175 Although it would be simpler to permute the elements randomly, a sum over $C!$ orderings would
 176 quickly become prohibitive as C increases. We summarize the inverse CDF sampling method in
 177 Algorithm 1.

178 3.2.2 Gumbel max sampling

179 It is not strictly necessary to analytically calculate the bivariate PMF. If the variance is not too large,
 180 a Monte Carlo estimation suffices. In such a case, we can use a simpler sampling approach using
 181 the well known Gumbel max trick [Gumbel, 1954, Jang et al., 2017, Maddison et al., 2017]. Since
 182 a Gumbel distribution $g \sim \text{Gumbel}(0, 1) \iff g = -\ln(-\ln(u))$, $u \sim \text{Unif}(0, 1)$ we can again
 183 use copulas to encode negative correlations between samples. Given $\mathbf{u} \sim \mathcal{C}_N$, an antithetic Gumbel
 184 categorical sample is $\mathbf{z} = \hat{j}$ such that $j = \text{argmax}_i \phi_i - \ln(-\ln(u_i))$, where ϕ are the logits of the
 185 desired categorical distribution. This is also equivalent to the following exponential racing [Yin et al.,
 186 2019, Zhang and Zhou, 2018] sampling: $\mathbf{z} = \hat{j}$, such that $j = \text{argmin}_i \epsilon_i$, where $\epsilon_i \sim \text{Exp}(e^{\phi_i})$.

187 If we let $\mathbf{Z} = [z_1, \dots, z_N]^T$ as before, an simple empirical estimate of the bivariate PMF computed
 188 using all pairs and only matrix operations is:

$$P(z = \hat{i}, z' = \hat{j}) \approx \frac{1}{N(N-1)} \mathbf{Z}^T (\mathbf{1}_{N \times N} - I_N) \mathbf{Z}.$$

189 In Fig. 1 we show what bivariate PMF both approaches produce. The results are similar for both
 190 a(n inverted) Dirichlet or Gaussian copula, with larger differences between the categorical sampling
 191 method.

192 4 Experimental results

193 In this section, we first illustrate the variance reduction that CARMS offers on a toy example. We
 194 also optimize a categorical variational autoencoder (VAE) [Kingma and Welling, 2014, Rezende
 195 et al., 2014], and a stochastic network for structured output prediction, which are standard tasks [Jang
 196 et al., 2017] for categorical variables, done on three different benchmark datasets. Each experiment
 197 uses both antithetic categorical approaches for CARMS: inverse CDF sampling and the Gumbel max
 198 trick, denoted as CARMS-I and CARMS-G, respectively. For CARMS-G we clip the empirical ratio
 199 values to avoid numerical instabilities, and we use the Dirichlet copula for both. We compare our
 200 approach to three state-of-the-art unbiased estimators: LOORF/VarGrad [Kool et al., 2019a, Richter

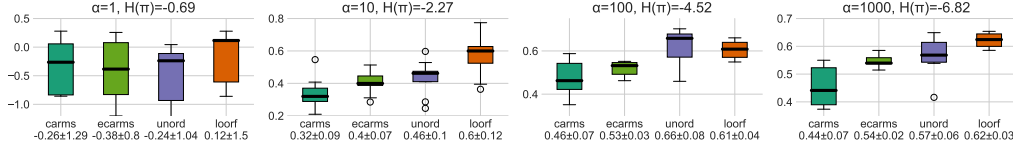


Figure 2: Log variance of the gradient of different estimators with respect to the logits on a toy problem. Columns correspond to different entropy levels of the logits.

et al., 2020], the unordered set estimator (UNORD) [Kool et al., 2020] and ARSM [Yin et al., 2019]. The code for all experiments is available in the supplementary material.

4.1 Toy example

We first showcase the variance reduction over other methods in a simple toy example, where we take the gradient with respect to the logits ϕ of:

$$\mathcal{L}(\phi) = \mathbb{E}[f(z_1, \dots, z_D)], \quad z_d \sim \text{Cat}(\sigma(\phi_d)), \quad f(z_1, \dots, z_D) = \sum_{d=1}^D \sum_{c=1}^C d \cdot c \cdot z_{dc}.$$

For simplicity, let the number of categories, dimensions, and samples be $C = D = N = 3$. The probabilities are randomly sampled from a Dirichlet distribution: $\sigma(\phi) \sim \text{Dir}(\mathbf{1}_C \cdot \alpha)$, but are identical for all methods for a given α . We vary the entropy of the probabilities from high ($\alpha = 1$) to low ($\alpha = 1000$). In a high entropy setting, there is little difference between the estimators, as the variance itself is very high, but differences emerge as we increase α and lower the entropy. The combined log variance of the gradient of each logit, for different methods and different α , is shown in Fig 2.

4.2 Categorical variational autoencoder

For this task, we follow the experimental setting from ARMS [Dimitriev and Zhou, 2021], except we use categorical instead of binary latent variables, and maximize the ELBO:

$$\text{ELBO}(\phi) = \mathbb{E} \left[\ln \frac{p(\mathbf{x}|\mathbf{z})p(\mathbf{z})}{q_\phi(\mathbf{z}|\mathbf{x})} \right] \approx \sum_{n=1}^N \ln \frac{p(\mathbf{x}|\mathbf{z}_n)p(\mathbf{z}_n)}{q_\phi(\mathbf{z}_n|\mathbf{x})}, \quad \mathbf{z}_1, \dots, \mathbf{z}_N \sim \text{Cat}_N(\sigma(\phi)),$$

where $\text{Cat}_N(\sigma(\phi))$ denotes an N -variate categorical distribution with identical marginals. The number of categories is $C \in \{3, 5, 10\}$ with $D = \lfloor 200/C \rfloor$ latent variables, respectively, to make the total computational effort similar. In the binary case $C = 2$, CARMS reduces to ARMS, for which a thorough comparison has already been produced. The task is training a categorical VAE using either a linear or nonlinear encoder/decoder pair on three different datasets: Dynamic(ally binarized) MNIST [LeCun et al., 2010], Fashion MNIST [Xiao et al., 2017], and Omniglot [Lake et al., 2015]. All datasets are freely available under the MIT license, and do not contain any personally identifiable information or offensive content. For a fair comparison, all methods use the same learning rate, optimizer, model architecture, and number of samples. Since ARSM uses a variable number of function evaluations per step, we use one sample per step, for which ARSM uses around twice as many evaluations as the other methods. The results are combined from five independent runs for each experimental configuration.

The VAE consists of a stochastic layer with $\lfloor 200/C \rfloor$ units, each of which is a C -way categorical variable. For the nonlinear case, there are additionally two layers of 200 units with LeakyReLU [Maas et al., 2013] activations. The prior logits are optimized using SGD with a learning rate of 10^{-2} , whereas the encoder and decoder are optimized using Adam [Kingma and Ba, 2015] with a learning rate of 10^{-4} , following Yin et al. [2019]. The optimization is run for 10^6 steps, with a batch size of 50, from which the global dataset mean is subtracted. The models are trained on an Intel Xeon Platinum 8280 2.7GHz CPU, and an individual run took 5-8 hours on one core of the machine, with a total carbon emissions estimated to be 28.19 kg of CO_2 [Lacoste et al., 2019]. An exception is UNORD for 10 samples, which was significantly more computation heavy. Although the paper states that a step can be performed in $O(2^C)$, the provided code requires $O(C!)$ evaluations per step.

Table 1: Final training 100 sample log likelihood of VAEs using different estimators, where the stochastic layer contains $C=3, 5$, or 10 categories, with $\lfloor 200/C \rfloor$ latent variables and C samples per gradient step, respectively. Results are reported on three datasets: Dynamic MNIST, Fashion MNIST, and Omniglot over 5 runs, with the best performing methods in bold.

Categories		CARMS-I	CARMS-G	LOORF	UNORD	ARSM	
Dynamic MNIST	Linear	3	-105.34 ± 0.25	-105.36 ± 0.24	-105.64 ± 0.23	-105.23 ± 0.23	-107.35 ± 0.56
		5	-103.53 ± 0.13	-103.35 ± 0.18	-103.54 ± 0.15	-103.50 ± 0.13	-106.13 ± 0.53
		10	-103.22 ± 0.05	-103.12 ± 0.06	-103.48 ± 0.06	-103.56 ± 0.06	-106.71 ± 0.58
	Nonlinr	3	-94.85 ± 0.28	-94.60 ± 0.28	-95.12 ± 0.21	-95.21 ± 0.22	-99.62 ± 0.50
		5	-93.05 ± 0.14	-92.60 ± 0.12	-92.91 ± 0.16	-92.98 ± 0.12	-98.89 ± 0.43
		10	-92.13 ± 0.05	-92.42 ± 0.10	-92.44 ± 0.04	-93.05 ± 0.14	-97.76 ± 0.41
Fashion MNIST	Linear	3	-245.44 ± 0.22	-245.69 ± 0.19	-245.8 ± 0.19	-245.90 ± 0.21	-247.51 ± 0.45
		5	-242.06 ± 0.13	-241.90 ± 0.10	-242.17 ± 0.09	-242.43 ± 0.12	-244.63 ± 0.47
		10	-240.44 ± 0.03	-240.52 ± 0.04	-240.91 ± 0.04	-241.08 ± 0.06	-243.29 ± 0.36
	Nonlinr	3	-233.13 ± 0.16	-233.20 ± 0.16	-233.75 ± 0.10	-233.34 ± 0.12	-237.93 ± 0.20
		5	-231.72 ± 0.09	-231.67 ± 0.13	-232.01 ± 0.05	-232.19 ± 0.09	-237.29 ± 0.34
		10	-230.77 ± 0.05	-231.16 ± 0.05	-231.35 ± 0.03	-231.74 ± 0.02	-235.88 ± 0.17
Omniglot	Linear	3	-114.53 ± 0.12	-114.73 ± 0.15	-114.90 ± 0.13	-114.77 ± 0.15	-116.34 ± 0.42
		5	-114.01 ± 0.10	-113.93 ± 0.11	-114.01 ± 0.10	-114.00 ± 0.09	-115.72 ± 0.35
		10	-114.97 ± 0.06	-114.99 ± 0.06	-115.17 ± 0.05	-115.55 ± 0.08	-117.48 ± 0.35
	Nonlinr	3	-110.19 ± 0.23	-110.16 ± 0.21	-110.33 ± 0.27	-110.31 ± 0.21	-115.13 ± 0.51
		5	-109.14 ± 0.09	-109.35 ± 0.13	-109.39 ± 0.16	-109.55 ± 0.14	-115.07 ± 0.37
		10	-108.66 ± 0.08	-108.62 ± 0.07	-108.98 ± 0.06	-109.54 ± 0.15	-115.11 ± 0.34

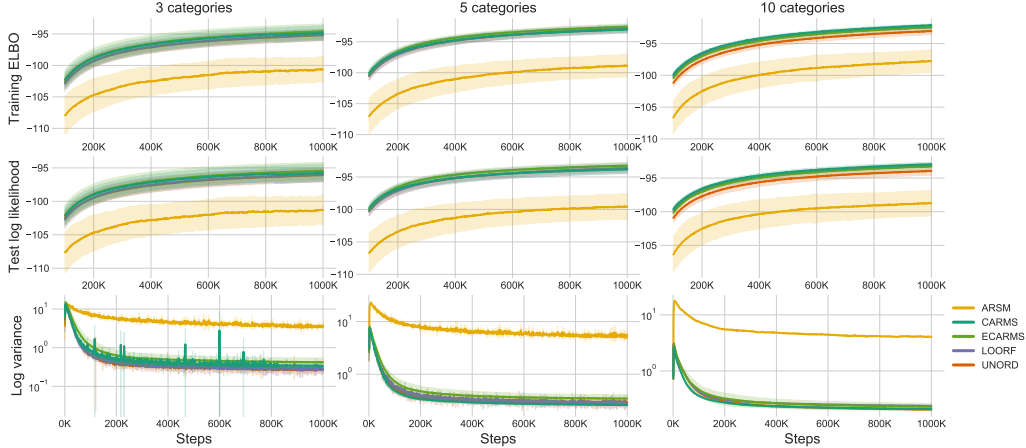


Figure 3: Training a nonlinear categorical VAE with different estimators on Dynamic MNIST using ELBO. Columns correspond to $C \in \{3, 5, 10\}$ categories with C samples per gradient step, respectively. Rows correspond to the 100 sample training and test log likelihood, and the variance of the gradient with respect to the logits of the encoder network. Results for different datasets and other networks can be found in the Appendix.

In Fig 3, we plot the training and test log likelihood using 100 samples, and gradient variance w.r.t the logits over time, for a nonlinear network on dynamic MNIST. Similar plots for other datasets and network types can be found in the appendix. Also shown in Table 1 is the final training log likelihood using 100 samples for all three datasets, network types, categories, and gradient estimators. The corresponding table with the final test log likelihood can be found in the appendix. In general, both versions of CARMS perform comparably, and result in slightly higher log likelihood than other methods. Because Gumbel CARMS uses an empirical estimate of the debiasing ratio, it has higher variance and slightly worse performance. When limiting the number of function evaluations, there is a gap between ARSM (which used on average $2C$ evaluations per step, out of a maximum of $C(C-1)/2$ per step) compared to the other methods, which used C evaluations. This is possibly due to the continuous reparameterization that ARSM uses, which adds variance.

Table 2: Final training log likelihood of a categorical network for conditional estimation using different gradient estimators, where the stochastic layer contains $C = 3, 5$, or 10 categories, with $\lfloor 200/C \rfloor$ latent variables and C samples per gradient step, respectively. Results are reported on three datasets: Dynamic MNIST, Fashion MNIST, and Omniglot over 5 runs, with the best performing methods in bold.

	Categories	CARMS-I	CARMS-G	LOORF	UNORD	ARSM
Dynamic MNIST	3	57.98 $\pm$ 0.10	58.35 $\pm$ 0.14	58.19 $\pm$ 0.14	58.06 $\pm$ 0.04	60.22 $\pm$ 0.19
	5	57.57 $\pm$ 0.05	57.85 $\pm$ 0.12	57.78 $\pm$ 0.08	57.6 $\pm$ 0.05	59.38 $\pm$ 0.13
	10	58.17 $\pm$ 0.26	58.33 $\pm$ 0.13	58.20 $\pm$ 0.14	58.18 $\pm$ 0.17	59.32 $\pm$ 0.11
Fashion MNIST	3	132.83 $\pm$ 0.08	132.90 $\pm$ 0.08	133.10 $\pm$ 0.05	133.06 $\pm$ 0.06	134.56 $\pm$ 0.35
	5	132.68 $\pm$ 0.05	132.81 $\pm$ 0.12	132.91 $\pm$ 0.07	132.94 $\pm$ 0.14	134.09 $\pm$ 0.12
	10	133.32 $\pm$ 0.15	133.43 $\pm$ 0.23	133.54 $\pm$ 0.17	133.38 $\pm$ 0.10	134.02 $\pm$ 0.21
Omniglot	3	65.57 $\pm$ 0.10	66.05 $\pm$ 0.14	65.92 $\pm$ 0.26	65.81 $\pm$ 0.09	68.00 $\pm$ 0.09
	5	65.65 $\pm$ 0.18	66.16 $\pm$ 0.32	65.92 $\pm$ 0.23	65.78 $\pm$ 0.06	67.99 $\pm$ 0.32
	10	66.76 $\pm$ 0.31	66.94 $\pm$ 0.24	66.87 $\pm$ 0.07	66.66 $\pm$ 0.24	68.35 $\pm$ 0.16

4.3 Structured prediction with stochastic categorical networks

We also compare all methods on the standard benchmark task of predicting the lower half of an image from the upper half, i.e. the conditional distribution $p(\mathbf{x}_l|\mathbf{x}_u)$, where $\mathbf{x}_u$ and $\mathbf{x}_l$ denote the upper and lower half of an image, respectively. We use a stochastic categorical network to estimate this distribution, with the objective: $\mathbb{E}_{\mathbf{z} \sim p(\mathbf{z}_m|\mathbf{x}_u)} \left[\frac{1}{M} \sum_{m=1}^M \ln p(\mathbf{x}_l|\mathbf{z}_m) \right]$, where $\mathbf{z}$ denotes a stochastic categorical layer. The encoder/decoder pair each contain one hidden layer with $\lfloor 200/C \rfloor$ latent variables and a LeakyReLU activation, with the optimization performed for $C \in \{3, 5, 10\}$ categories, on all three datasets, with identical settings for each gradient estimator for a fair comparison. We use $M = 1$ for training, and $M = 1000$ for evaluation on both the train and test set. In Table 2, we show the final training set log likelihood, with the corresponding test log likelihood table in the appendix, though the results are qualitatively similar. The results are similar to the VAE experiment, with the inverse CDF CARMS having a slightly higher log likelihood. However, the differences between estimators are less pronounced, with the unordered set estimator is being mostly on par with CARMS, and ARSM only slightly trailing the other methods, with a much smaller gap.

5 Discussion

We have presented a novel approach for training categorical variables, which extends the ARMS estimator for the binary case. It goes beyond i.i.d. samples by using a copula to generate antithetic categorical samples, but preserves unbiasedness by including a multiplicative term. For i.i.d. samples, the form of the estimator reduces to LOORF. We showcase its usefulness on several datasets, a different number of categories and types of deep neural networks. In variational inference tasks, and conditional estimation, CARMS outperforms other state of the art estimators. The main limitation of this work is its specificity to categorical (including binary) variables. We hope to extend this, e.g. to Plackett-Luce models for top-k sampling [Kool et al., 2019b, Grover et al., 2019], which has important applications in ranking [Dadaneh et al., 2020]. There is also a general limitation that CARMS shares with other state-of-the-art estimators, which is higher complexity than LOORF, a very simple but strong baseline. Future work includes investigating theoretical properties and large scale applications, as well as possible general antithetic gradient estimators, and we plan to investigate adaptive correlations that take into account the properties of f to further reduce the variance.

Potential societal impact This work is focused on better optimization for categorical variables, which includes any network containing a stochastic categorical layer. In particular, generative models such as VAEs are widely used, and are sometimes trained on more human centered datasets. These models can sometimes be used to impersonate a person’s face or voice. We only used non-human datasets such as MNIST, but with enough knowledge and using the available code, anyone, including bad actors, can train the same model on human content for malicious purposes. However, we strongly believe that wide knowledge dissemination and open source code is crucial for reproducibility in all of science.

References

- Y. Bengio, N. Léonard, and A. Courville. Estimating or propagating gradients through stochastic neurons for conditional computation. *arXiv preprint arXiv:1308.3432*, 2013.
- Y. Burda, R. B. Grosse, and R. Salakhutdinov. Importance weighted autoencoders. In *4th International Conference on Learning Representations, ICLR*, 2016.
- S. Z. Dadaneh, S. Boluki, M. Zhou, and X. Qian. Arsm gradient estimator for supervised learning to rank. In *ICASSP 2020-2020 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pages 3157–3161. IEEE, 2020.
- A. Dimitriev and M. Zhou. Arms: Antithetic-reinforce-multi-sample gradient for binary variables. *International Conference on Machine Learning*, 2021.
- Z. Dong, A. Mnih, and G. Tucker. Disarm: An antithetic gradient estimator for binary latent variables. In *Advances in Neural Information Processing Systems 33*, 2020.
- Z. Dong, A. Mnih, and G. Tucker. Coupled gradient estimators for discrete latent variables. *Third Symposium on Advances in Approximate Bayesian Inference*, 2021.
- M. C. Fu. Gradient estimation. *Handbooks in operations research and management science*, 13: 575–616, 2006.
- P. W. Glynn. Likelihood ratio gradient estimation for stochastic systems. *Communications of the ACM*, 33(10):75–84, 1990.
- W. Grathwohl, D. Choi, Y. Wu, G. Roeder, and D. Duvenaud. Backpropagation through the void: Optimizing control variates for black-box gradient estimation. In *6th International Conference on Learning Representations, ICLR*, 2018.
- A. Grover, E. Wang, A. Zweig, and S. Ermon. Stochastic optimization of sorting networks via continuous relaxations. *arXiv preprint arXiv:1903.08850*, 2019.
- S. Gu, S. Levine, I. Sutskever, and A. Mnih. Muprop: Unbiased backpropagation for stochastic neural networks. In *4th International Conference on Learning Representations, ICLR*, 2016.
- E. J. Gumbel. *Statistical theory of extreme values and some practical applications: a series of lectures*, volume 33. US Government Printing Office, 1954.
- E. Jang, S. Gu, and B. Poole. Categorical reparameterization with gumbel-softmax. In *5th International Conference on Learning Representations, ICLR*. OpenReview.net, 2017.
- M. I. Jordan, Z. Ghahramani, T. S. Jaakkola, and L. K. Saul. An introduction to variational methods for graphical models. In *Learning in graphical models*, pages 105–161. Springer, 1998.
- D. P. Kingma and J. Ba. Adam: A method for stochastic optimization. In *3rd International Conference on Learning Representations, ICLR*, 2015.
- D. P. Kingma and M. Welling. Auto-encoding variational bayes. In *2nd International Conference on Learning Representations, ICLR*, 2014.
- W. Kool, H. van Hoof, and M. Welling. Buy 4 REINFORCE samples, get a baseline for free! In *Workshop, Deep Reinforcement Learning Meets Structured Prediction, ICLR*, 2019a.
- W. Kool, H. Van Hoof, and M. Welling. Stochastic beams and where to find them: The gumbel-top-k trick for sampling sequences without replacement. In *International Conference on Machine Learning*, pages 3499–3508. PMLR, 2019b.
- W. Kool, H. van Hoof, and M. Welling. Estimating gradients for discrete random variables by sampling without replacement. In *International Conference on Learning Representations*, 2020. URL <https://openreview.net/forum?id=rk1Ej2EFvB>.
- A. Kucukelbir, D. Tran, R. Ranganath, A. Gelman, and D. M. Blei. Automatic differentiation variational inference. *The Journal of Machine Learning Research*, 18(1):430–474, 2017.

330 A. Lacoste, A. Luccioni, V. Schmidt, and T. Dandres. Quantifying the carbon emissions of machine
331 learning. *arXiv preprint arXiv:1910.09700*, 2019.

332 B. M. Lake, R. Salakhutdinov, and J. B. Tenenbaum. Human-level concept learning through proba-
333 bilistic program induction. *Science*, 350(6266):1332–1338, 2015.

334 Y. LeCun, C. Cortes, and C. Burges. Mnist handwritten digit database. *ATT Labs [Online]*. Available:
335 <http://yann.lecun.com/exdb/mnist>, 2, 2010.

336 G. Lorberbom, T. S. Jaakkola, A. Gane, and T. Hazan. Direct optimization through arg max for
337 discrete variational auto-encoder. In *Advances in Neural Information Processing Systems 32*, pages
338 6200–6211, 2019.

339 A. L. Maas, A. Y. Hannun, and A. Y. Ng. Rectifier nonlinearities improve neural network acoustic
340 models. In *ICML Workshop on Deep Learning for Audio, Speech and Language Processing*.
341 Citeseer, 2013.

342 C. J. Maddison, A. Mnih, and Y. W. Teh. The concrete distribution: A continuous relaxation of
343 discrete random variables. In *5th International Conference on Learning Representations, ICLR*.
344 OpenReview.net, 2017.

345 A. Mnih and K. Gregor. Neural variational inference and learning in belief networks. In *International*
346 *Conference on Machine Learning*, pages 1791–1799. PMLR, 2014.

347 A. Mnih and D. J. Rezende. Variational inference for monte carlo objectives. In *Proceedings of*
348 *the 33rd International Conference on Machine Learning, ICML*, volume 48, pages 2188–2196.
349 JMLR.org, 2016.

350 S. Mohamed, M. Rosca, M. Figurnov, and A. Mnih. Monte carlo gradient estimation in machine
351 learning. *J. Mach. Learn. Res.*, 21:132:1–132:62, 2020.

352 A. B. Owen. Monte carlo theory, methods and examples. , 2013.

353 J. W. Paisley, D. M. Blei, and M. I. Jordan. Variational bayesian inference with stochastic search. In
354 *Proceedings of the 29th International Conference on Machine Learning, ICML*, 2012.

355 A. Ramesh, M. Pavlov, G. Goh, S. Gray, C. Voss, A. Radford, M. Chen, and I. Sutskever. Zero-shot
356 text-to-image generation. *arXiv preprint arXiv:2102.12092*, 2021.

357 R. Ranganath, S. Gerrish, and D. Blei. Black box variational inference. In *Artificial intelligence and*
358 *statistics*, pages 814–822. PMLR, 2014.

359 D. J. Rezende, S. Mohamed, and D. Wierstra. Stochastic backpropagation and approximate inference
360 in deep generative models. In *International conference on machine learning*, pages 1278–1286.
361 PMLR, 2014.

362 L. Richter, A. Boustati, N. Nüsken, F. J. Ruiz, and Ö. D. Akyildiz. Vargrad: A low-variance gradient
363 estimator for variational inference. In *Advances in Neural Information Processing Systems*,
364 volume 33, pages 13481–13492, 2020.

365 M. Rosca, M. Figurnov, S. Mohamed, and A. Mnih. Measure-valued derivatives for approximate
366 bayesian inference. *4th workshop on Bayesian Deep Learning, NeurIPS*, 2019.

367 F. J. R. Ruiz, M. K. Titsias, and D. M. Blei. The generalized reparameterization gradient. In *Advances*
368 *in Neural Information Processing Systems 29*, pages 460–468, 2016.

369 T. Salimans and D. A. Knowles. On using control variates with stochastic approximation for
370 variational bayes and its connection to stochastic linear regression. *arXiv preprint arXiv:1401.1022*,
371 2014.

372 M. Titsias and M. Lázaro-Gredilla. Local expectation gradients for black box variational inference.
373 In *Advances in neural information processing systems*, pages 2620–2628. Citeseer, 2015.

- 374 G. Tucker, A. Mnih, C. J. Maddison, D. Lawson, and J. Sohl-Dickstein. REBAR: low-variance,
 375 unbiased gradient estimates for discrete latent variable models. In *Advances in Neural Information*
 376 *Processing Systems 30*, pages 2627–2636, 2017.
- 377 R. J. Williams. Simple statistical gradient-following algorithms for connectionist reinforcement
 378 learning. *Machine learning*, 8(3-4):229–256, 1992.
- 379 H. Xiao, K. Rasul, and R. Vollgraf. Fashion-mnist: a novel image dataset for benchmarking machine
 380 learning algorithms. *arxiv.org/abs/1708.07747*, 2017.
- 381 M. Yin, Y. Yue, and M. Zhou. ARSM: augment-reinforce-swap-merge estimator for gradient
 382 backpropagation through categorical variables. In *Proceedings of the 36th International Conference*
 383 *on Machine Learning*, volume 97, pages 7095–7104. PMLR, 2019.
- 384 M. Yin, N. Ho, B. Yan, X. Qian, and M. Zhou. Probabilistic Best Subset Selection by Gradient-Based
 385 Optimization. *arXiv e-prints*, 2020.
- 386 Q. Zhang and M. Zhou. Nonparametric bayesian lomax delegate racing for survival analysis with
 387 competing risks. *arXiv preprint arXiv:1810.08564*, 2018.

388 Checklist

- 389 1. For all authors...
- 390 (a) Do the main claims made in the abstract and introduction accurately reflect the paper’s
 391 contributions and scope?
 392 [Yes] Our work focuses specifically on improving training that includes stochastic
 393 categorical variables, where we offer improvement over the state of the art.
- 394 (b) Did you describe the limitations of your work?
 395 [Yes] We describe the limitations in the discussion section.
- 396 (c) Did you discuss any potential negative societal impacts of your work?
 397 [Yes] We discuss societal impacts in the discussion section.
- 398 (d) Have you read the ethics review guidelines and ensured that your paper conforms to
 399 them?
 400 [Yes] We read the guidelines and ensured the paper conforms to them.
- 401 2. If you are including theoretical results...
- 402 (a) Did you state the full set of assumptions of all theoretical results?
 403 [Yes] The theorems/lemmas contain all assumptions.
- 404 (b) Did you include complete proofs of all theoretical results?
 405 [Yes] Proofs are given either in the main text or the supplemental material (appendix).
- 406 3. If you ran experiments...
- 407 (a) Did you include the code, data, and instructions needed to reproduce the main experi-
 408 mental results (either in the supplemental material or as a URL)?
 409 [Yes] The code, data, and instructions are available in the supplementary material.
- 410 (b) Did you specify all the training details (e.g., data splits, hyperparameters, how they
 411 were chosen)?
 412 [Yes] All of information is available in the supplementary material.
- 413 (c) Did you report error bars (e.g., with respect to the random seed after running experi-
 414 ments multiple times)?
 415 [Yes] We provide the mean and variance over 5 runs for all experiments.
- 416 (d) Did you include the total amount of compute and the type of resources used (e.g., type
 417 of GPUs, internal cluster, or cloud provider)?
 418 [Yes] As stated in the results section, in total it took 5-8 hours on one CPU machine for
 419 a full run of all methods on all datasets, categories and network types, and an estimated
 420 28kg of CO₂.
- 421 4. If you are using existing assets (e.g., code, data, models) or curating/releasing new assets...

422 (a) If your work uses existing assets, did you cite the creators?
 423 [Yes] We included the proper dataset citations. The code was written by the authors.

424 (b) Did you mention the license of the assets?
 425 [Yes] All datasets are available under the MIT license, mentioned in the results section.

426 (c) Did you include any new assets either in the supplemental material or as a URL?
 427 [N/A] There are no new assets used. The existing datasets and code are either part of
 428 the supplemental material, or are downloaded when training starts.

429 (d) Did you discuss whether and how consent was obtained from people whose data you're
 430 using/curating?
 431 [N/A] Under the MIT license, permission is granted to any person to use, copy, modify,
 432 merge, publish etc.

433 (e) Did you discuss whether the data you are using/curating contains personally identifiable
 434 information or offensive content?
 435 [Yes] Yes, it is stated in the results section that no personally identifiable information
 436 or offensive content is present in any of the datasets.

437 5. If you used crowdsourcing or conducted research with human subjects...

438 (a) Did you include the full text of instructions given to participants and screenshots, if
 439 applicable? [N/A]

440 (b) Did you describe any potential participant risks, with links to Institutional Review
 441 Board (IRB) approvals, if applicable? [N/A]

442 (c) Did you include the estimated hourly wage paid to participants and the total amount
 443 spent on participant compensation? [N/A]