

SIMPLE: A GRADIENT ESTIMATOR FOR k -SUBSET SAMPLING

Kareem Ahmed^{*1} Zhe Zeng^{*1} Mathias Niepert² Guy Van den Broeck¹

¹Dept. of Computer Science, UCLA ²Dept. of Computer Science, Stuttgart University

{ahmedk, zhezeng, guyvdb}@cs.ucla.edu

mathias.niepert@simtech.uni-stuttgart.de

ABSTRACT

k -subset sampling is ubiquitous in machine learning, enabling regularization and interpretability through sparsity. The challenge lies in rendering k -subset sampling amenable to end-to-end learning. This has typically involved *relaxing* the *reparameterized* samples to allow for backpropagation, with the risk of introducing high bias and high variance. In this work, we fall back to *discrete* k -subset sampling on the forward pass. This is coupled with using the gradient with respect to the *exact* marginals, computed *efficiently*, as a proxy for the true gradient. We show that our gradient estimator, SIMPLE, exhibits lower bias and variance compared to state-of-the-art estimators, including the straight-through Gumbel estimator when $k = 1$. Empirical results show improved performance on learning to explain and sparse linear regression. We provide an algorithm for computing the *exact* ELBO for the k -subset distribution, obtaining *significantly* lower loss compared to SOTA.

1 INTRODUCTION

k -subset sampling, sampling a subset of size k of n variables, is omnipresent in machine learning. It lies at the core of many fundamental problems that rely upon learning sparse features representations of input data, including stochastic high-dimensional data visualization (van der Maaten, 2009), parametric k -nearest neighbors (Grover et al., 2018), learning to explain (Chen et al., 2018), discrete variational auto-encoders (Rolfe, 2017), and sparse regression, to name a few. All such tasks involve optimizing an expectation of an objective function with respect to a latent *discrete* distribution parameterized by a neural network, which are often *assumed* intractable. Score-function estimators offer a cloyingly simple solution: rewrite the gradient of the expectation as an expectation of the gradient, which can subsequently be estimated using a finite number of samples offering an unbiased estimate of the gradient. Simple as it is, score-function estimators suffer from very high variance which can interfere with training. This provided the impetus for other, low-variance, gradient estimators, chief among them are those based on the reparameterization trick, which allows for biased, but low-variance gradient estimates. The reparameterization trick, however, does not allow for a direct application to discrete distributions thereby prompting continuous relaxations, e.g. Gumbel-softmax (Jang et al., 2017; Maddison et al., 2017), that allow for reparameterized gradients w.r.t the parameters of a *categorical* distribution. Reparameterizable subset sampling (Xie & Ermon, 2019) generalizes the Gumbel-softmax trick to k -subsets which while rendering k -subset sampling amenable to backpropagation at the cost of introducing bias in the learning by using relaxed samples.

In this paper, we set out with the goal of avoiding all such relaxations. Instead, we fall back to *discrete* sampling on the forward pass. On the backward pass, we reparameterize the gradient of the loss function with respect to the samples as a function of the *exact* marginals of the k -subset distribution. Computing the exact conditional marginals is, in general, intractable (Roth, 1996). We give an *efficient* algorithm for computing the k -subset probability, and show that the conditional marginals correspond to partial derivatives, and are therefore tractable for the k -subset distribution. We show that our proposed gradient estimator for the k -subset distribution, coined SIMPLE, is reminiscent of the straight-through (ST) Gumbel estimator when $k = 1$, with the gradients taken with respect to the unperturbed marginals. We empirically demonstrate that SIMPLE exhibits lower bias *and* variance compared to other known gradient estimators, including the ST Gumbel estimator in the case $k = 1$.

^{*}These authors contributed equally to this work.

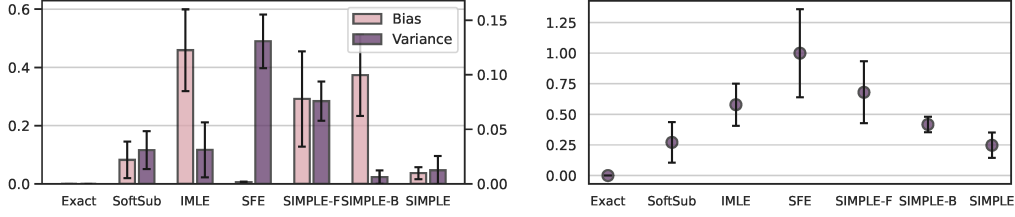


Figure 1: A comparison of the bias and variance of the gradient estimators (left) and the average and standard deviation of the cosine distance of a single-sample gradient estimate to the exact gradient. We used the cosine distance, defined as $(1 - \text{cosine similarity})$, in place of the euclidean distance as we only care about the direction of the gradient, not magnitude. The bias, variance and error were estimated using a sample of size 10,000. The details of this experiment are provided in Section 5.1.

We include an experiment on the task of learning to explain (L2X) using the BEERADVOCATE dataset (McAuley et al., 2012), where the goal is to select the subset of words that best explains the model’s classification of a user’s review. We also include an experiment on the task of stochastic sparse linear regression, where the goal is to learn the best sparse model, and show that we are able to recover the Kuramoto–Sivashinsky equation. Finally, we develop an efficient computation for the calculation of the exact variational evidence lower bound (ELBO) for the k -subset distribution, which when used in conjunction with SIMPLE leads to state-of-the-art discrete sparse VAE learning.

Contributions. In summary, we propose replacing *relaxed* sampling on the forward pass with *discrete* sampling. On the backward pass, we use the gradient with respect to the exact conditional marginals as a proxy for the true gradient, giving an algorithm for computing them efficiently. We empirically demonstrate that discrete samples on the forward pass, coupled with exact conditional marginals on the backward pass leads to a new gradient estimator, SIMPLE, with lower bias and variance compared to other known gradient estimators. We also provide an efficient computation of the exact ELBO for the k -subset distribution, leading to state-of-the-art discrete sparse VAE learning.

2 PROBLEM STATEMENT AND MOTIVATION

We consider models described by the equations

$$\theta = h_v(x), \quad z \sim p_\theta(z \mid \sum_i z_i = k), \quad \hat{y} = f_u(z, x), \quad (1)$$

where $x \in \mathcal{X}$ and $\hat{y} \in \mathcal{Y}$ denote feature inputs and target outputs, respectively, $h_v : \mathcal{X} \rightarrow \Theta$ and $f_u : \mathcal{Z} \times \mathcal{X} \rightarrow \mathcal{Y}$ are smooth, parameterized maps and θ are logits inducing a distribution over the latent binary vector z . The induced distribution $p_\theta(z)$ is defined as

$$p_\theta(z) = \prod_{i=1}^n p_{\theta_i}(z_i), \quad \text{with } p_{\theta_i}(z_i = 1) = \text{sigmoid}(\theta_i) \text{ and } p_{\theta_i}(z_i = 0) = 1 - \text{sigmoid}(\theta_i). \quad (2)$$

The goal of our stochastic latent layer is *not* to simply sample from $p_\theta(z)$, which would yield samples with a Hamming weight between 0 and n (i.e., with an arbitrary number of ones). Instead, we are interested in sampling from the distribution restricted to samples with a Hamming weight of k , for any given k . That is, we are interested in sampling from the conditional distribution $p_\theta(z \mid \sum_i z_i = k)$.

Conditioning the distribution $p_\theta(z)$ on this k -subset constraint introduces intricate dependencies between each of the z_i ’s. The probability of sampling any given k -subset vector z , therefore, becomes

$$p_\theta(z \mid \sum_i z_i = k) = p_\theta(z) / p_\theta(\sum_i z_i = k) \cdot \mathbb{I}[\sum_i z_i = k]$$

where $\mathbb{I}[\cdot]$ denotes the indicator function. In other words, the probability of sampling each k -subset is re-normalized by $p_\theta(\sum_i z_i = k)$ – the probability of sampling exactly k items from the *unconstrained* distribution induced by encoder h_v . The quantity $p_\theta(\sum_i z_i = k) = \sum_z p_\theta(z) \cdot \mathbb{I}[\sum_i z_i = k]$ appears to be intractable. We show that not to be the case, providing a tractable algorithm for computing it.

Given a set of samples $\mathcal{D}$, we are concerned with learning the parameters $\omega = (v, u)$ of the architecture in (1) through minimizing the training error L , which is the expected loss:

$$L(x, y; \omega) = \mathbb{E}_{z \sim p_\theta(z \mid \sum_i z_i = k)} [\ell(f_u(z, x), y)] \quad \text{with } \theta = h_v(x), \quad (3)$$

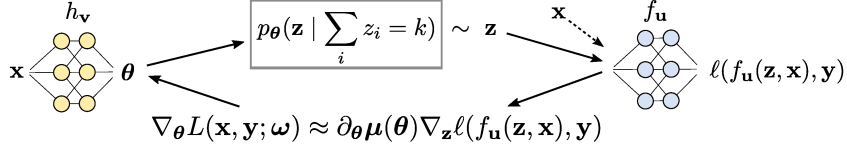


Figure 2: The problem setting considered in our paper. On the forward pass, a neural network h_v outputs θ parameterizing a *discrete* distribution over subsets of size k of n items, i.e., the k -subset distribution. We sample exactly, and efficiently, from this distribution, and feed the samples to a downstream neural network. On the backward pass, we approximate the true gradient by the product of the derivative of marginals and the gradient of the sample-wise loss.

where $\ell : \mathcal{Y} \times \mathcal{Y} \rightarrow \mathbb{R}^+$ is a point-wise loss function. This formulation, illustrated in Figure 2, is general and subsumes many settings. Different choices of mappings h_v and f_u , and sample-wise loss ℓ define various tasks. Table 1 presents some example settings used in our experimental evaluation.

TASK	MAP h_v	MAP f_u	LOSS ℓ
Discrete VAE (Sec. 5.2)	Encoder	Decoder	ELBO
Learn To Explain (Sec. 5.3)	Embedding	Regression	RMSE
Sparse Regression (Sec. 5.4)	Identity	Linear Regression	RMSE

Table 1: Architectures of the three experiment settings.

Learning then requires computing the gradient of L w.r.t. $\omega = (v, u)$. The gradient of L w.r.t. u is

$$\nabla_u L(x, y; \omega) = \mathbb{E}_{z \sim p_\theta(z | \sum_i z_i = k)} [\partial_u f_u(z, x)^\top \nabla_y \ell(\hat{y}, y)], \quad (4)$$

where $\hat{y} = f_u(z, x)$ is the decoding of a latent sample z . Furthermore, the gradient of L w.r.t. v is

$$\nabla_v L(x, y; \omega) = \partial_v h_v(x)^\top \nabla_\theta L(x, y; \omega), \quad (5)$$

where $\nabla_\theta L(x, y; \omega) := \nabla_\theta \mathbb{E}_{z \sim p_\theta(z | \sum_i z_i = k)} [\ell(f_u(z, x), \hat{y})]$, the loss' gradient w.r.t. the encoder.

One challenge lies in computing the expectation in (3) and (4), which has no known closed-form solution. This necessitates a Monte-Carlo estimate through sampling from $p_\theta(z | \sum_i z_i = k)$.

A second, and perhaps more substantial hurdle lies in computing $\nabla_\theta L(x, y; \omega)$ in (5) due to the non-differentiable nature of discrete sampling. One could rewrite $\nabla_\theta L(x, y; \omega)$ as

$$\nabla_\theta L(x, y; \omega) = \mathbb{E}_{z \sim p_\theta(z | \sum_i z_i = k)} [\ell(f_u(z, x), y) \nabla_\theta \log p_\theta(z | \sum_i z_i = k)]$$

which is known as the REINFORCE estimator (Williams, 1992), or the score function estimator (SFE). It is typically avoided due to its notoriously high variance, despite its apparent simplicity. Instead, typical approaches (Xie & Ermon, 2019; Plötz & Roth, 2018) reparameterize the samples as a deterministic transformation of the parameters, and some independent standard Gumbel noise, and relaxing the deterministic transformation, the top- k function in this case, to allow for backpropagation.

3 SIMPLE: SUBSET IMPLICIT LIKELIHOOD ESTIMATION

Our goal is to build a gradient estimator for $\nabla_\theta L(x, y; \omega)$. We start by envisioning a hypothetical sampling-free architecture, where the downstream neural network f_u is a function of the marginals, $\mu := \mu(\theta) := \{p_\theta(z_j | \sum_i z_i = k)\}_{j=1}^n$, instead of a discrete sample z , resulting in a loss L_m s.t.

$$\nabla_\theta L_m(x, y; \omega) = \partial_\theta \mu(\theta)^\top \nabla_\mu \ell_m(f_u(\mu, x), y). \quad (6)$$

When the marginals $\mu(\theta)$ can be efficiently computed and differentiated, such a hypothetical pipeline can be trained end-to-end. Furthermore, Domke (2010) observed that, for an arbitrary loss function ℓ_m defined on the marginals, the Jacobian of the marginals w.r.t. the logits is symmetric, i.e.

$$\nabla_\theta L_m(x, y; \omega) = \partial_\theta \mu(\theta)^\top \nabla_\mu \ell_m(f_u(\mu, x), y) = \partial_\theta \mu(\theta) \nabla_\mu \ell_m(f_u(\mu, x), y). \quad (7)$$

Consequently, computing the gradient of the loss w.r.t. the logits, $\nabla_\theta L_m(x, y; \omega)$, reduces to computing the *directional derivative*, or the Jacobian-vector product, of the marginals w.r.t. the logits in the direction of the gradient of the loss. This offers an alluring opportunity: the conditional

Algorithm 1 PrExactlyk(θ, n, k)

Input: The logits θ of the distribution, the number of variables n , and the subset size k
Output: $p_\theta(\sum_i z_i = k)$
 // $a[i, j] = p_\theta(\sum_{m=1}^i z_m = j)$ for all i, j
 initialize a to be 0 everywhere
 $a[0, 0] = 1$ // $p_\theta(\sum_{m=1}^0 z_m = 0) = 1$
for $i = 1$ **to** n **do**
 for $j = 0$ **to** k **do**
 // cf. constructive proof of Prop. 1
 $a[i, j] = a[i - 1, j] \cdot p_{\theta_i}(z_i = 0)$
 $\quad + a[i - 1, j - 1] \cdot p_{\theta_i}(z_i = 1)$
return $a[n, k]$

Algorithm 2 Sample(θ, n, k)

Input: The logits θ of the distribution, the number of variables n , and the subset size k
Output: $\mathbf{z} = (z_1, \dots, z_n) \sim p_\theta(\mathbf{z} \mid \sum_i z_i = k)$
 $\text{sample} = [], j = k$
for $i = n$ **to** 1 **do**
 // cf. proof of Prop. 2
 $p = a[i - 1, j - 1]$
 $z_i \sim \text{Bernoulli}(p \cdot p_{\theta_i}(z_i = 1) / a[i, j])$
 // Pick next state based on value of sample
if $z_i = 1$ **then** $j = j - 1$
 $\text{sample.append}(z_i)$
return sample

marginals characterize the probability of each z_i in the sample, and could be thought of as a differentiable proxy for the samples. Specifically, by reparameterizing $\mathbf{z}$ as a function of the conditional marginal μ under approximation $\partial_\mu \mathbf{z} \approx \mathbf{I}$ as proposed by Niepert et al. (2021), and using the straight-through estimator for the gradient of the sample w.r.t. the marginals on the backward pass, we approximate our true $\nabla_\theta L(\mathbf{x}, \mathbf{y}; \omega)$ as

$$\nabla_\theta L(\mathbf{x}, \mathbf{y}; \omega) \approx \partial_\theta \mu(\theta) \nabla_{\mathbf{z}} L(\mathbf{x}, \mathbf{y}; \omega), \quad (8)$$

where the directional derivative of the marginals can be taken along *any downstream gradient*, rendering the whole pipeline end-to-end learnable, even in the presence of non-differentiable sampling.

Now, estimating the gradient of the loss w.r.t. the parameters can be thought of as decomposing into two sub-problems: **(P1)** Computing the derivatives of conditional marginals $\partial_\theta \mu(\theta)$, which requires the computation of the conditional marginals, and **(P2)** Computing the gradient of the loss w.r.t. the samples $\nabla_{\mathbf{z}} L(\mathbf{x}, \mathbf{y}; \omega)$ using sample-wise loss, which requires drawing exact samples. These two problems are complicated by conditioning on the k -subset constraint, which introduces intricate dependencies to the distribution, and is infeasible to solve naively, e.g. by enumeration. We will show simple, efficient, and exact solutions to each problem, at the heart of which is the insight that we need not care about the variables' order, only their sum, introducing symmetries that simplify the problem.

3.1 DERIVATIVES OF CONDITIONAL MARGINALS

In many probabilistic models, marginal inference is #P-hard (Roth, 1996; Zeng et al., 2020). However, we observe that it is not the case for the k -subset distribution. We notice that the conditional marginals correspond to the partial derivatives of the log-probability of the k -subset constraint. To see this, note that the derivative of a multi-linear function with respect to a single variable retains all the terms referencing that variable, and drops all other terms; this corresponds exactly to the unnormalized conditional marginals. By taking the derivative of the log-probability, this introduces the k -subset probability in the denominator, leading to the *conditional* marginals. Intuitively, the rate of change of the k -subset probability w.r.t. a variable only depends on that variable through its length- k subsets.

Theorem 1. Let $p_\theta(\sum_j z_j = k)$ be the probability of exactly- k of the unconstrained distribution parameterized by logits θ . For every variable Z_i , its conditional marginal is

$$p_\theta(z_i \mid \sum_j z_j = k) = \frac{\partial}{\partial \theta_i} \log p_\theta(\sum_j z_j = k). \quad (9)$$

We refer the reader to the appendix for a detailed proof of the above theorem. To establish the tractability of the above computation of the conditional marginals, we need to show that the probability of the exactly- k constraint $p_\theta(\sum_i z_i = k)$ can be obtained tractably, which we demonstrate next.

Proposition 1. The probability $p_\theta(\sum_i z_i = k)$ of sampling exactly k items from the unconstrained distribution $p_\theta(\mathbf{z})$ over n items as in Equation 2 can be computed exactly in time $\mathcal{O}(nk)$.

Proof. Our proof is constructive. As a base case, consider the probability of sampling $k = -1$ out of $n = 0$ items. We can see that the probability of such an event is 0. As a second base case, consider

the probability of sampling $k = 0$ out of $n = 0$ items. We can see that the probability of such an event is 1. Now assume that we are given the probability $p_{\theta}(\sum_i^{n-1} z_i = k')$, for $k' = 0, \dots, k$, and we are interested in computing $p_{\theta}(\sum_i^n z_i = k)$. By the partition theorem, we can see that

$$p_{\theta}(\sum_i^n z_i = k) = p_{\theta}(\sum_i^{n-1} z_i = k) \cdot p_{\theta_n}(z_n = 0) + p_{\theta}(\sum_i^{n-1} z_i = k-1) \cdot p_{\theta_n}(z_n = 1)$$

as events $\sum_i^{n-1} z_i = k$ and $\sum_i^{n-1} z_i = k-1$ are disjoint and, for any k , partition the sample space. Intuitively, for any k and n , we can sample k out of n items by choosing k of $n-1$ items, and not the n -th item, or choosing $k-1$ of $n-1$ items, and the n -th item. The above process gives rise to Algorithm 1, which returns $p_{\theta}(\sum_i z_i = k)$ in time $\mathcal{O}(nk)$. $\square$

By the construction described above, we obtain a closed-form $p_{\theta}(\sum_i^n z_i = k)$, which allows us to compute conditional marginals $p_{\theta}(z_i | \sum_j z_j = k)$ by Theorem 1 via auto-differentiation. This further allows the computation of the derivatives of conditional marginals $\partial_{\theta} \mu(\theta)_i = \partial_{\theta} p_{\theta}(z_i | \sum_j z_j = k)$ to be amenable to auto-differentiation, solving problem (P1) exactly and efficiently.

3.2 GRADIENTS OF LOSS W.R.T. SAMPLES

As alluded to in Section 3, we approximate $\nabla_{\theta} L(\mathbf{x}, \mathbf{y}; \omega)$ by the directional derivative of the marginals along the gradient of the loss w.r.t. discrete samples $\mathbf{z}$, $\nabla_{\mathbf{z}} L(\mathbf{x}, \mathbf{y}; \omega)$, where $\mathbf{z}$ is drawn from the k -subset distribution $p_{\theta}(\mathbf{z} | \sum_i z_i = k)$. What remains is to estimate the value of the loss, necessitating faithful sampling from the k -subset distribution, which might initially appear daunting.

Exact k -subset Sampling Next we show how to sample exactly from the k -subset distribution $p_{\theta}(\mathbf{z} | \sum_i z_i = k)$. We start by sampling the variables in reverse order, that is, we sample z_n through z_1 . The main intuition being that, having sampled $(z_n, z_{n-1}, \dots, z_{i+1})$ with a Hamming weight of $k-j$, we sample Z_i with a probability of choosing $k-j$ of $n-1$ variables and the n -th variable given that we choose $k-j+1$ of n variables. We formalize our intuition below.

Proposition 2. *Let Sample be defined as in Algorithm 2. Given n random variables $Z_1, \dots, Z_n$, a subset size k , and a k -subset distribution $p_{\theta}(\mathbf{z} | \sum_i z_i = k)$ parameterized by log probabilities θ , Algorithm 2 draws exact samples from $p_{\theta}(\mathbf{z} | \sum_i z_i = k)$ in time $\mathcal{O}(n)$.*

Proof. Assume that variables $Z_n, \dots, Z_{i+1}$ are sampled and have their values to be $z_n, \dots, z_{i+1}$ with $\sum_{m=i+1}^n z_m = k-j$. By Algorithm 2 we have that the probability with which to sample Z_i is

$$\begin{aligned} p_{\text{Sample}}(z_i = 1 | z_n, \dots, z_{i+1}) &= \frac{p_{\theta}(\sum_{m=i}^n z_m = k-j+1 | \sum_m z_m = k) p_{\theta_i}(z_i = 1)}{p_{\theta}(\sum_{m=i+1}^n z_m = k-j | \sum_m z_m = k)} \\ &= \frac{p_{\theta}(\sum_{m=i+1}^n z_m = k-j | z_i = 1, \sum_m z_m = k) p_{\theta_i}(z_i = 1)}{p_{\theta}(\sum_{m=i+1}^n z_m = k-j | \sum_m z_m = k)} \\ &= p_{\theta}(z_i = 1 | \sum_{m=i+1}^n z_m = k-j, \sum_m z_m = k) \text{ (by Bayes' theorem)} \end{aligned}$$

It follows that samples drawn from Algorithm 2 are distributed according to $p_{\theta}(\mathbf{z} | \sum_i z_i = k)$. $\square$

4 CONNECTION TO STRAIGHT-THROUGH GUMBEL-SOFTMAX

One might wonder if our gradient estimator reduces to the Straight-Through (ST) Gumbel-Softmax estimator, or relates to it in any way when $k = 1$. On the forward pass, the ST Gumbel Softmax estimator makes use of the Gumbel-Max trick (Maddison et al., 2014), which states that we can efficiently sample from a categorical distribution by perturbing each of the logits with standard Gumbel noise, and taking the MAP, or more formally $\mathbf{z} = \text{OneHot}(\arg \max_{i \in \{1, \dots, k\}} \theta_i + g_i) \sim p_{\theta}$ where the g_i 's are i.i.d Gumbel(0, 1) samples, and OneHot encodes the sample as a binary vector.

Since $\arg \max$ is non-differentiable, Gumbel-Softmax uses the *perturbed* relaxed samples, $\mathbf{y} = \text{Softmax}(\theta + g_i)$ as a proxy for discrete samples $\mathbf{z}$ on the backward pass, using differentiable Softmax in place of the non-differentiable $\arg \max$, with the entire function returning $(\mathbf{z} - \mathbf{y}).\text{detach}() + \mathbf{y}$ where detach ensures that the gradient flows only through the relaxed samples on the backward pass.

Algorithm 3 The proposed algorithm for the k -subset distribution

function FORWARDPASS(θ) <i>// $p_\theta(\sum_{m=1}^i z_m = j)$ for all i, j</i> $a = \text{PrExactlyk}(\theta, n, k)$ <i>// Sample from $p_\theta(z \mid \sum_i z_i = k)$</i> $z = \text{Sample}(\theta, n, k)$ save a for the backward pass return z	function BACKWARDPASS($\nabla_z \ell(f_u(z, x), y)$) load θ from the forward pass <i>// derivatives of $p_\theta(z \mid \sum_i z_i = k)$</i> $\mu = \nabla_\theta \log a[n, k]$ <i>// by auto-diff</i> <i>// Return the directional derivative of the</i> <i>// marginals along the downstream gradients</i> return JVP($\mu, \nabla_z \ell(f_u(z, x))$)
--	---

That is, just like SIMPLE, ST Gumbel-Softmax returns *exact, discrete* samples. However, whereas SIMPLE backpropagates through the exact marginals, ST Gumbel Softmax backpropagates through the *perturbed* marginals that result from applying the Gumbel-max trick. As can be seen in Figure 3, such a minor difference means that, empirically, SIMPLE exhibits lower bias and variance compared to ST Gumbel Softmax while being exactly as efficient.

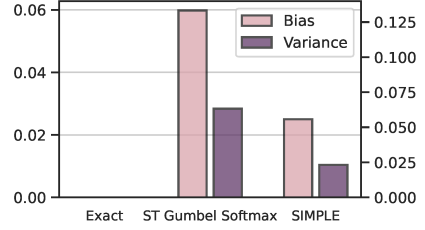


Figure 3: Bias and variance of SIMPLE and Gumbel Softmax over 10k samples

5 EXPERIMENTS

We conduct experiments on four different tasks: 1) A synthetic experiment designed to test the bias and variance, as well as the average deviation of SIMPLE compared to a variety of well-established estimators in the literature. 2) A discrete k -subset Variational Auto-Encoder (DVAE) setting, where the latent space models a probability distribution over k -subsets. We will show that we can compute the evidence lower bound (ELBO) exactly, and that, coupled with exact sampling and our SIMPLE gradient estimator, we attain a much lower loss compared to state of the art in sparse DVAEs. 3) The learning to explain (L2X) setting, where the aim is to select the k -subset of words that best describe the classifier’s prediction, where we show an improved mean-squared error, as well as precision, across the board. 4) A novel, yet simple task, sparse linear regression, where, in a vein similar to L2X, we wish to select a k -subset of features that give rise to a linear regression model, avoiding overfitting the spurious features present in the data. Table 1 details the architecture with the objective functions. Our code will be made publicly available at github.com/UCLA-StarAI/SIMPLE.

5.1 SYNTHETIC EXPERIMENTS

We carried out a series of experiments with a 5-subset distribution, and a latent space of dimension 10. We set the loss to $L(\theta) = \mathbb{E}_{z \sim p_\theta(z \mid \sum_i z_i = k)} [\|z - \mathbf{b}\|^2]$, where $\mathbf{b}$ is the groundtruth logits sampled from $\mathcal{N}(0, \mathbf{I})$. Such a distribution is tractable: we only have $\binom{10}{5} = 252$ k -subsets, which are easily enumerable and therefore, the exact gradient, the golden standard, can be computed in closed form.

In this experiment, we are interested in three metrics: bias, variance, and the average error of each gradient estimator, where the latter is measured by averaging the deviation of each single-sample gradient estimate from the exact gradient. We used the cosine distance, defined as $1 - \text{cosine similarity}$ as the measure of deviation in our calculation of the metrics above, as we only care about direction.

We compare against four different baselines: *exact*, which denotes the exact gradient; *SoftSub* (Xie & Ermon, 2019), which uses an extension of the Gumbel-Softmax trick to sample *relaxed* k -subsets on the forward pass; I-MLE, which denotes the IMLE gradient estimator (Niepert et al., 2021), where *approximate* samples are obtained using perturb-and-map (PAM) on the forward pass, approximating the marginals using PAM samples on the backward pass; and score function estimator, denoted *SFE*.

We tease apart SIMPLE’s improvements by comparing three different flavors: SIMPLE-F, which only uses *exact* sampling, falling back to estimating the marginals using *exact* samples; SIMPLE-B, which uses *exact* marginals on the backward pass with *approximate* PAM samples on the forward pass; and SIMPLE, coupling *exact* samples on the forward pass with *exact* marginals on the backward pass.

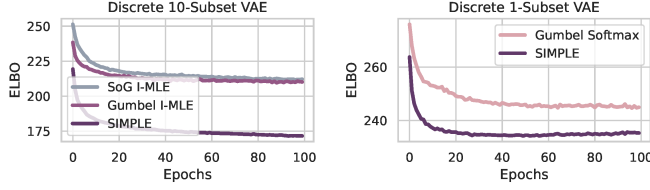


Figure 4: ELBO against # of epochs. (Left) Comparison of SIMPLE against different flavors of IMLE on the 10-subset DVAE, and (Right) against ST Gumbel Softmax on the 1-subset DVAE.

Our results are shown in Figure 1. As expected, we observe that *SFE* exhibits no bias, but high variance whereas *SoftSub* suffers from both bias and variance, due to the Gumbel noise injection into the samples to make them differentiable. We observe that I-MLE exhibits very high bias, as well as very low variance. This can be attributed to the PAM sampling, which in the case of k -subset distribution does not sample faithfully from the distribution, but is instead biased to sampling only the mode of the distribution. This also means that, by approximating the marginals using PAM samples, there is a lot less variance to our gradients. On to our SIMPLE gradient estimator, we see that it exhibits *less bias as well as less variance* compared to all the other gradient estimators. We also see that each estimated gradient is, on average, much more aligned with the exact gradient. To understand why that is, we compare SIMPLE, SIMPLE-F, and SIMPLE-B. As hypothesized, we observe that *exact sampling*, SIMPLE-F, reduces the bias, but increases the variance compared to I-MLE, this is since, unlike the PAM samples, our exact sample span the entire sample space. We also observe that, even compared to I-MLE, SIMPLE-B, reduces the variance by marginalizing over all possible samples.

5.2 DISCRETE VARIATIONAL AUTO-ENCODER

Next, we test our SIMPLE gradient estimator in the k -subset discrete variational auto-encoder (DVAE) setting, where the latent variables model a probability distribution over k -subsets, and has a dimensionality of 20. Similar to prior work (Jang et al., 2017; Niepert et al., 2021), the encoding and decoding functions of the VAE consist of three dense layers (encoding: 512-256-20x20; decoding: 256-512-784). The DVAE is trained to minimize the sum of reconstruction loss and KL-divergence of the k -subset distribution and the constrained uniform distribution, known as the ELBO, on MNIST.

In prior work, the KL-divergence was approximated using the unconditional marginals, obtained simply through a Softmax layer. Instead we show that the KL-divergence between the k -subset distribution and the uniform distribution can be computed exactly. First note that, through simple algebraic manipulations, the KL-divergence between the k -subset distribution and the constrained uniform distribution can be rewritten as the sum of negative entropy, $-H(\mathbf{z})$, where $\mathbf{z} \sim p_{\theta}(\mathbf{z} \mid \sum_i z_i = k)$ and $\log$ the number of k -subsets, $\log \binom{n}{k}$ (see appendix for details), reducing the hardness of computing the KL-divergence, to computing the entropy of a k -subset distribution, for which Algorithm 4 gives a tractable algorithm. Intuitively, the uncertainty in the distribution over a sequence of length n , k of which are true, decomposes as the uncertainty over Z_n , and the average of the uncertainties over the remainder of the sequence. We refer the reader to the appendix for the proof of the below proposition.

Proposition 3. Let Entropy be defined as in Algorithm 4. Given variables, $Z_1, \dots, Z_n$, and a k -subset distribution $p_{\theta}(\mathbf{z} \mid \sum_i z_i = k)$, Algorithm 4 computes entropy of $p_{\theta}(\mathbf{z} \mid \sum_i z_i = k)$.

We plot the loss ELBO against the number of epochs, as seen in Figure 4. We compared against I-MLE using sum-of-gamma noise as well as Gumbel noise for PAM sampling, on the 10-subset DVAE, and against ST Gumbel Softmax on the 1-subset DVAE. We observe a *significantly* lower loss on the test set on the 10-subset DVAE, partly attributable to the exact ELBO computation, but also on the 1-subset DVAE compared to ST Gumbel Softmax, where the sole difference is the backward pass.

5.3 LEARNING TO EXPLAIN

The BEERADVOCATE dataset (McAuley et al., 2012) consists of free-text reviews and ratings for 4 different aspects of beer: appearance, aroma, palate, and taste. The training set has 80k reviews for the aspect APPEARANCE and 70k reviews for all other aspects. In addition to the ratings for all reviews, each sentence in the test set contains annotations of the words that best describe the review

Algorithm 4 Entropy(θ, n, k)

Input: The logits θ of the distribution, the number of variables n , and the subset size k

Output: $H(\mathbf{z}) = -\mathbb{E}_{\mathbf{z} \sim p_{\theta}(\mathbf{z} \mid \sum_i z_i = k)} [\log p(\mathbf{z})]$

```

 $h = \text{zeros}(n, k)$ 
for  $i = k$  to  $n$  do
  for  $j = 0$  to  $k$  do
    //  $p(z_i \mid \sum_{m=1}^i z_m = j)$ 
     $p = a[i-1, j-1] * p_{\theta_i}(z_i = 1) / a[i, j]$ 
    // cf. proof of Prop. 3 in Appendix
     $h[i, j] = H_b(p) + p * h[i-1, j] +$ 
       $(1-p) * h[i-1, j+1]$ 
  return  $h$ 
```

Method	Appearance		Palate		Taste	
	Test MSE	Precision	Test MSE	Precision	Test MSE	Precision
SIMPLE (Ours)	2.35 $\pm$ 0.28	66.81 $\pm$ 7.56	2.68 $\pm$ 0.06	44.78 $\pm$ 2.75	2.11 $\pm$ 0.02	42.31 $\pm$ 0.61
L2X ($t = 0.1$)	10.70 $\pm$ 4.82	30.02 $\pm$ 15.82	6.70 $\pm$ 0.63	50.39 $\pm$ 13.58	6.92 $\pm$ 1.61	32.23 $\pm$ 4.92
SoftSub ($t = 0.5$)	2.48 $\pm$ 0.10	52.86 $\pm$ 7.08	2.94 $\pm$ 0.08	39.17 $\pm$ 3.17	2.18 $\pm$ 0.10	41.98 $\pm$ 1.42
I-MLE ($\tau = 30$)	2.51 $\pm$ 0.05	65.47 $\pm$ 4.95	2.96 $\pm$ 0.04	40.73 $\pm$ 3.15	2.38 $\pm$ 0.04	41.38 $\pm$ 1.55

Table 2: Results for three aspects with $k = 10$: test MSE and subset precision, both $\times 100$

Method	$k = 5$		$k = 10$		$k = 15$	
	Test MSE	Precision	Test MSE	Precision	Test MSE	Precision
SIMPLE (Ours)	2.27 $\pm$ 0.05	57.30 $\pm$ 3.04	2.23 $\pm$ 0.03	47.17 $\pm$ 2.11	3.20 $\pm$ 0.04	53.18 $\pm$ 1.09
L2X ($t = 0.1$)	5.75 $\pm$ 0.30	33.63 $\pm$ 6.91	6.68 $\pm$ 1.08	26.65 $\pm$ 9.39	7.71 $\pm$ 0.64	23.49 $\pm$ 10.93
SoftSub ($t = 0.5$)	2.57 $\pm$ 0.12	54.06 $\pm$ 6.29	2.67 $\pm$ 0.14	44.44 $\pm$ 2.27	2.52 $\pm$ 0.07	37.78 $\pm$ 1.71
I-MLE ($\tau = 30$)	2.62 $\pm$ 0.05	54.76 $\pm$ 2.50	2.71 $\pm$ 0.10	47.98 $\pm$ 2.26	2.91 $\pm$ 0.18	39.56 $\pm$ 2.07

Table 3: Results for aspect Aroma: test MSE and subset precision, both $\times 100$, for $k \in \{5, 10, 15\}$.

score with respect to the various aspects. We address the problem introduced by the L2X paper (Chen et al., 2018) of learning a k -subset distribution over words that best explain a given rating. We follow the architecture suggested in the L2X paper, consisting of four convolutional and one dense layer.

We compare to relaxation-based baselines L2X (Chen et al., 2018) and SoftSub (Xie & Ermon, 2019) as well as to I-MLE which uses perturb-and-MAP to both compute an approximate sample in the forward pass and to estimate the marginals. Prior work has shown that the straight-through estimator (STE) did not work well and we omit it here. We used the standard hyperparameter settings of Chen et al. (2018) and choose the temperature parameter $t \in \{0.1, 0.5, 1.0, 2.0\}$ for all methods. We used the standard Adam settings and trained separate models for each aspect using MSE as point-wise loss ℓ . Table 3 lists results for $k \in \{5, 10, 15\}$ for the AROMA aspect. The mean-squared error (MSE) of SIMPLE is almost always lower and its subset precision never significantly exceeded by those of the baselines. Table 2 shows results on the remaining aspects Appearance, Palate, and Taste for $k = 10$.

5.4 SPARSE LINEAR REGRESSION

Given a library of feature functions, the task of sparse linear regression aims to learn from data which feature subset best describes the nonlinear partial differential equation (PDE) that the data are sampled from. We propose to tackle this task by learning a k -subset distribution over the feature functions. During learning, we first sample from the k -subset distribution to decide which feature function subset to choose. With k chosen features, we perform linear regression to learn the coefficients of the features from data, and then update the k -subset distribution logit parameters by minimizing RMSE.

To test our proposed approach, we follow the experimental setting in PySINDy (de Silva et al., 2020; Kaptanoglu et al., 2022) and use the dataset collected by PySINDy where the samples are collected from the Kuramoto–Sivashinsky (KS) equation, a fourth-order nonlinear PDE known for its chaotic behavior. This PDE takes the form $v_t = -v_{xx} - v_{xxxx} - vv_x$, which can be seen as a linear combination of feature functions $\mathcal{V} = \{v_{xx}, v_{xxxx}, vv_x\}$ with the coefficients all set to a value of -1 . At test time, we use the MAP estimation of the learned k -subset distribution to choose the k feature functions. For $k = 3$, our proposed method achieves the same performance as the state-of-the-art solver on this task, PySINDy. It identifies the KS PDE from data by choosing exactly the ground truth feature function subset $\mathcal{V}$, obtaining an RMSE of 0.00622 after applying linear regression on $\mathcal{V}$.

6 COMPLEXITY ANALYSIS

In Proposition 1, we prove that computing the marginal probability of the exactly- k constraint can be done tractably in time $\mathcal{O}(nk)$. In the context of deep learning, we often care about vectorized com-

plexity. We demonstrate an optimized algorithm achieving a vectorized complexity $\mathcal{O}(\log k \log n)$, assuming perfect parallelization. The optimization is possible by computing the marginal probability in a divide-and-conquer way: it partitions the variables into two subsets and compute their marginals respectively such that the complexity $\mathcal{O}(n)$ is reduced to $\mathcal{O}(\log n)$; the summation over the k terms also has its complexity reduced to $\mathcal{O}(\log k)$ in a similar manner. We refer the readers to Algorithm 5 in Appendix for the optimized algorithm. We further modify Algorithm 2 to perform divide-and-conquer such that sampling k -subsets achieves a vectorized complexity being $\mathcal{O}(\log n)$, shown as Algorithm 6 in the Appendix. As a comparison, SoftSub (Xie & Ermon, 2019) has its complexity to be $\mathcal{O}(nk)$ due to the relaxed top- k operation and its vectorized complexity to be $\mathcal{O}(k \log n)$ stemming from the fact that softmax layers need $\mathcal{O}(\log n)$ rounds of communication for normalization.

7 RELATED WORK

There is a large body of work on gradient estimation for categorical random variables. Maddison et al. (2017); Jang et al. (2017) propose the Gumbel-softmax distribution (named the concrete distribution by the former) to relax categorical random variables. For more complex distributions, such as the k -subset distribution which we are concerned with in this paper, existing approaches either use the straight-through and score function estimators or propose tailor-made relaxations (see for instance Kim et al. (2016); Chen et al. (2018); Grover et al. (2018)). We directly compare to the score function and straight-through estimator as well as the tailored relaxations of Chen et al. (2018); Grover et al. (2018) and show that we are competitive and obtain a lower bias and/or variance than these other estimators. Tucker et al. (2017); Grathwohl et al. (2018) develop parameterized control variates based on continuous relaxations for the score-function estimator. Lastly, Paulus et al. (2020) offers a comprehensible work on relaxed gradient estimators, deriving several extensions of the softmax trick. All of the above works, ours included, assume the independence of the selected items, beyond there being k of them. That is with the exception of Paulus et al. (2020) which make use of a relaxation using pairwise embeddings, but do not make their code available. We leave that to future work.

A related line of work has developed and analyzed sparse variants of the softmax function, motivated by their potential computational and statistical advantages. Representative examples are Blondel et al. (2020a); Peters et al. (2019); Correia et al. (2019); Martins & Astudillo (2016). SparseMAP (Nicolae et al., 2018) has been proposed in the context of structured prediction and latent variable models, also replacing the softmax with a sparser distribution. LP-SparseMAP (Nicolae & Martins, 2020) is an extension that uses a relaxation of the optimization problem rather than a MAP solver. Sparsity can also be exploited for efficient marginal inference in latent variable models (Correia et al., 2020). Contrary to our work, they cannot control the sparsity level exactly through a k -subset constraint or guarantee a sparse output. Also, we aim at cases where samples in the forward pass are required.

Integrating specialized discrete algorithms into neural networks is a line of research with increasing popularity. Examples are sorting algorithms (Cuturi et al., 2019; Blondel et al., 2020b; Grover et al., 2018), ranking (Rolinek et al., 2020; Kool et al., 2019), dynamic programming (Mensch & Blondel, 2018; Corro & Titov, 2019), and solvers for combinatorial optimization problems Berthet et al. (2020); Rolínek et al. (2020); Shirobokov et al. (2020); Niepert et al. (2021); Minervini et al. (2023); Zeng et al. or even probabilistic circuits over structured output spaces (Ahmed et al., 2022; Blondel, 2019). There has also been work on making common programming language expression such as conditional statements, loops, and indexing differentiable through relaxations (Petersen et al., 2021). Xie et al. (2020) proposes optimal transport as a way to obtain differentiable sorting methods for top- k classification. In contrast, we focus on the k -subset sampling problem and provide exact discrete sampling and marginal inference algorithms, obtaining a gradient estimator for the k -subset distribution with a favorable bias-variance trade-off.

8 CONCLUSION

We introduced a gradient estimator for the k -subset distribution which replaces *relaxed* and *approximate* sampling on the forward pass with *exact* sampling. It sidesteps the non-differentiable nature of discrete sampling by estimating the gradients as a function of our distribution’s marginals, for which we prove a simple characterization, showing that we can compute them exactly and efficiently. We demonstrated improved empirical results on a number of tasks: L2X, DVAEs, and sparse regression.

ACKNOWLEDGMENTS

We would like to thank all the reviewers whose insightful comments helped improve this work. We would especially like to thank reviewers for pointing out that we can forgo the method of finite differences in favor of computing the directional derivative, alleviating the need for λ , and enabling us to compute exact derivative along the downstream gradient. This work was funded in part by the DARPA Perceptually-enabled Task Guidance (PTG) Program under contract number HR00112220005, NSF grants #IIS-1943641, #IIS-1956441, #CCF-1837129, Samsung, CISCO, a Sloan Fellowship, and by Deutsche Forschungsgemeinschaft under Germany’s Excellence Strategy - EXC 2075. ZZ is supported by an Amazon Doctoral Student Fellowship.

REFERENCES

- Kareem Ahmed, Stefano Teso, Kai-Wei Chang, Guy Van den Broeck, and Antonio Vergari. Semantic probabilistic layers for neuro-symbolic learning. In *Advances in Neural Information Processing Systems 35 (NeurIPS)*, dec 2022.
- Quentin Berthet, Mathieu Blondel, Olivier Teboul, Marco Cuturi, Jean-Philippe Vert, and Francis R. Bach. Learning with differentiable perturbed optimizers. In *NeurIPS*, 2020.
- Mathieu Blondel. Structured prediction with projection oracles. *Advances in neural information processing systems*, 32, 2019.
- Mathieu Blondel, André FT Martins, and Vlad Niculae. Learning with fenchel-young losses. *J. Mach. Learn. Res.*, 21(35):1–69, 2020a.
- Mathieu Blondel, Olivier Teboul, Quentin Berthet, and Josip Djolonga. Fast differentiable sorting and ranking. In *International Conference on Machine Learning*, pp. 950–959. PMLR, 2020b.
- Jianbo Chen, Le Song, Martin Wainwright, and Michael Jordan. Learning to explain: An information-theoretic perspective on model interpretation. In Jennifer Dy and Andreas Krause (eds.), *Proceedings of the 35th International Conference on Machine Learning*, volume 80 of *Proceedings of Machine Learning Research*, pp. 883–892. PMLR, 10–15 Jul 2018.
- Gonalo M Correia, Vlad Niculae, and André FT Martins. Adaptively sparse transformers. In *Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP)*, pp. 2174–2184, 2019.
- Gonalo M Correia, Vlad Niculae, Wilker Aziz, and André FT Martins. Efficient marginalization of discrete and structured latent variables via sparsity. *Advances in Neural Information Processing Systems*, 2020.
- Caio Corro and Ivan Titov. Differentiable perturb-and-parse: Semi-supervised parsing with a structured variational auto-encoder. In *International Conference on Learning Representations*, 2019.
- Marco Cuturi, Olivier Teboul, and Jean-Philippe Vert. Differentiable ranking and sorting using optimal transport. *Advances in neural information processing systems*, 32, 2019.
- Brian de Silva, Kathleen Champion, Markus Quade, Jean-Christophe Loiseau, J. Kutz, and Steven Brunton. Pysindy: A python package for the sparse identification of nonlinear dynamical systems from data. *Journal of Open Source Software*, 5(49):2104, 2020. doi: 10.21105/joss.02104. URL <https://doi.org/10.21105/joss.02104>.
- Justin Domke. Implicit differentiation by perturbation. In *Advances in Neural Information Processing Systems 23*, pp. 523–531. 2010.
- Will Grathwohl, Dami Choi, Yuhuai Wu, Geoffrey Roeder, and David Duvenaud. Backpropagation through the void: Optimizing control variates for black-box gradient estimation. *ICLR*, 2018.

- Aditya Grover, Eric Wang, Aaron Zweig, and Stefano Ermon. Stochastic optimization of sorting networks via continuous relaxations. In *International Conference on Learning Representations*, 2018.
- Eric Jang, Shixiang Gu, and Ben Poole. Categorical reparameterization with gumbel-softmax. In *5th International Conference on Learning Representations, ICLR 2017, Toulon, France, April 24-26, 2017, Conference Track Proceedings*. OpenReview.net, 2017.
- Alan A. Kaptanoglu, Brian M. de Silva, Urban Fasel, Kadierdan Kaheman, Andy J. Goldschmidt, Jared Callahan, Charles B. Delahunt, Zachary G. Nicolaou, Kathleen Champion, Jean-Christophe Loiseau, J. Nathan Kutz, and Steven L. Brunton. Pysindy: A comprehensive python package for robust sparse system identification. *Journal of Open Source Software*, 7(69):3994, 2022. doi: 10.21105/joss.03994. URL <https://doi.org/10.21105/joss.03994>.
- Carolyn Kim, Ashish Sabharwal, and Stefano Ermon. Exact sampling with integer linear programs and random perturbations. In *Thirtieth AAAI Conference on Artificial Intelligence*, 2016.
- Wouter Kool, Herke Van Hoof, and Max Welling. Stochastic beams and where to find them: The gumbel-top-k trick for sampling sequences without replacement. In *International Conference on Machine Learning*, pp. 3499–3508. PMLR, 2019.
- Tao Lei, R. Barzilay, and T. Jaakkola. Rationalizing neural predictions. In *EMNLP*, 2016.
- Chris J Maddison, Daniel Tarlow, and Tom Minka. A* sampling. In Z. Ghahramani, M. Welling, C. Cortes, N. Lawrence, and K.Q. Weinberger (eds.), *Advances in Neural Information Processing Systems*, volume 27. Curran Associates, Inc., 2014.
- Chris J. Maddison, Andriy Mnih, and Yee Whye Teh. The concrete distribution: A continuous relaxation of discrete random variables. In *5th International Conference on Learning Representations, ICLR 2017, Toulon, France, April 24-26, 2017, Conference Track Proceedings*. OpenReview.net, 2017.
- Andre Martins and Ramon Astudillo. From softmax to sparsemax: A sparse model of attention and multi-label classification. In *International conference on machine learning*, pp. 1614–1623. PMLR, 2016.
- Julian McAuley, J. Leskovec, and Dan Jurafsky. Learning attitudes and attributes from multi-aspect reviews. *2012 IEEE 12th International Conference on Data Mining*, pp. 1020–1025, 2012.
- Arthur Mensch and Mathieu Blondel. Differentiable dynamic programming for structured prediction and attention. In *International Conference on Machine Learning*, pp. 3462–3471. PMLR, 2018.
- Pasquale Minervini, Luca Franceschi, and Mathias Niepert. Adaptive perturbation-based gradient estimation for discrete latent variable models. In *AAAI*, 2023.
- Vlad Niculae and André F. T. Martins. Lp-sparsemap: Differentiable relaxed optimization for sparse structured prediction. In *ICML*, 2020.
- Vlad Niculae, André F. T. Martins, Mathieu Blondel, and Claire Cardie. Sparsemap: Differentiable sparse structured inference. In *ICML*, 2018.
- Mathias Niepert, Pasquale Minervini, and Luca Franceschi. Implicit mle: Backpropagating through discrete exponential family distributions. *Advances in Neural Information Processing Systems*, 34, 2021.
- Max B Paulus, Dami Choi, Daniel Tarlow, Andreas Krause, and Chris J Maddison. Gradient estimation with stochastic softmax tricks. *arXiv preprint arXiv:2006.08063*, 2020.
- Ben Peters, Vlad Niculae, and André FT Martins. Sparse sequence-to-sequence models. In *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics*, pp. 1504–1519, 2019.
- Felix Petersen, Christian Borgelt, Hilde Kuehne, and Oliver Deussen. Learning with algorithmic supervision via continuous relaxations. *Advances in Neural Information Processing Systems*, 34, 2021.

- Tobias Plötz and Stefan Roth. Neural nearest neighbors networks. In Samy Bengio, Hanna M. Wallach, Hugo Larochelle, Kristen Grauman, Nicolò Cesa-Bianchi, and Roman Garnett (eds.), *Advances in Neural Information Processing Systems 31: Annual Conference on Neural Information Processing Systems 2018, NeurIPS 2018, December 3-8, 2018, Montréal, Canada*, pp. 1095–1106, 2018.
- Jason Tyler Rolfe. Discrete variational autoencoders. In *5th International Conference on Learning Representations, ICLR 2017, Toulon, France, April 24-26, 2017, Conference Track Proceedings*. OpenReview.net, 2017.
- Michal Rolínek, Vít Musil, Anselm Paulus, Marin Vlastelica, Claudio Michaelis, and Georg Martius. Optimizing rank-based metrics with blackbox differentiation. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pp. 7620–7630, 2020.
- Michal Rolínek, Paul Swoboda, Dominik Zietlow, Anselm Paulus, Vít Musil, and Georg Martius. Deep graph matching via blackbox differentiation of combinatorial solvers. In *ECCV*, 2020.
- Dan Roth. On the hardness of approximate reasoning. *Artif. Intell.*, 82(1-2):273–302, 1996. doi: 10.1016/0004-3702(94)00092-1. URL [https://doi.org/10.1016/0004-3702\(94\)00092-1](https://doi.org/10.1016/0004-3702(94)00092-1).
- Sergey Shirobokov, Vladislav Belavin, Michael Kagan, Andrei Ustyuzhanin, and Atilim Gunes Baydin. Black-box optimization with local generative surrogates. *Advances in Neural Information Processing Systems*, 33:14650–14662, 2020.
- George Tucker, Andriy Mnih, Chris J Maddison, Dieterich Lawson, and Jascha Sohl-Dickstein. Rebar: Low-variance, unbiased gradient estimates for discrete latent variable models. *Advances in Neural Information Processing Systems*, 2017.
- Laurens van der Maaten. Learning a parametric embedding by preserving local structure. In David van Dyk and Max Welling (eds.), *Proceedings of the Twelfth International Conference on Artificial Intelligence and Statistics*, volume 5 of *Proceedings of Machine Learning Research*, pp. 384–391, Hilton Clearwater Beach Resort, Clearwater Beach, Florida USA, 16–18 Apr 2009. PMLR.
- Ronald J. Williams. Simple statistical gradient-following algorithms for connectionist reinforcement learning. *Mach. Learn.*, 8:229–256, 1992. doi: 10.1007/BF00992696.
- Sang Michael Xie and Stefano Ermon. Reparameterizable subset sampling via continuous relaxations. In *Proceedings of the Twenty-Eighth International Joint Conference on Artificial Intelligence, IJCAI-19*, pp. 3919–3925. International Joint Conferences on Artificial Intelligence Organization, 7 2019. doi: 10.24963/ijcai.2019/544.
- Yujia Xie, Hanjun Dai, Minshuo Chen, Bo Dai, Tuo Zhao, Hongyuan Zha, Wei Wei, and Tomas Pfister. Differentiable top-k with optimal transport. *Advances in Neural Information Processing Systems*, 33:20520–20531, 2020.
- Zhe Zeng, Paolo Morettin, Fanqi Yan, Andrea Passerini, Guy Van den Broeck, et al. Is parameter learning via weighted model integration tractable? In *The 4th Workshop on Tractable Probabilistic Modeling*.
- Zhe Zeng, Paolo Morettin, Fanqi Yan, Antonio Vergari, and Guy Van den Broeck. Probabilistic inference with algebraic constraints: Theoretical limits and practical approximations. *Advances in Neural Information Processing Systems*, 33:11564–11575, 2020.

A PROOFS

Proposition 2. Let $p_{\theta}(\sum_i z_i = k)$ be the probability of exactly- k of the unconstrained distribution parameterized by log probabilities θ . For every Z_i , the conditional marginal is

$$p_{\theta}(z_i | \sum_j z_j = k) = \frac{\partial}{\partial \theta_i} \log p_{\theta}(\sum_j z_j = k). \quad (10)$$

Proof. We first rewrite the marginal $p_{\theta}(\sum_i z_i = k)$ into a summation as the probability for all possible events by definition as follows.

$$p_{\theta}(\sum_j z_j = k) = \sum_{\mathbf{z}: \sum_j z_j = k} \prod_{j: z_j = 1} \exp(\theta_j) \prod_{j: z_j = 0} (1 - \exp(\bar{\theta}_j)) \quad (11)$$

Here we assume that the probability of $z_j = 0$ is a constant term w.r.t. parameter θ_j , i.e., $\frac{\partial}{\partial \theta_j} (1 - \exp(\bar{\theta}_j)) = 0$.^{*} Further, the derivative of $p_{\theta}(\sum_j z_j = k)$ w.r.t. θ_i is as follows,

$$\begin{aligned} \frac{\partial}{\partial \theta_i} p_{\theta}(\sum_j z_j = k) &= \frac{\partial}{\partial \theta_i} \sum_{\mathbf{z}: \sum_j z_j = k \wedge z_i = 1} \prod_{j: z_j = 1} \exp(\theta_j) \prod_{j: z_j = 0} (1 - \exp(\bar{\theta}_j)) \\ &= \frac{\partial}{\partial \theta_i} \exp(\theta_i) \sum_{\mathbf{z}: \sum_j z_j = k \wedge z_i = 1} \prod_{j: z_j = 1, j \neq i} \exp(\theta_j) \prod_{j: z_j = 0} (1 - \exp(\bar{\theta}_j)) \\ &= \exp(\theta_i) \sum_{\mathbf{z}: \sum_j z_j = k \wedge z_i = 1} \prod_{j: z_j = 1, j \neq i} \exp(\theta_j) \prod_{j: z_j = 0} (1 - \exp(\bar{\theta}_j)) \\ &= p_{\theta}(\sum_j z_j = k \wedge z_i = 1), \end{aligned}$$

where the first equality holds since terms corresponding to $z_i \neq 1$ has their derivative to be zero w.r.t. θ_i . It further holds that

$$\begin{aligned} \frac{\partial}{\partial \theta_i} \log p_{\theta}(\sum_j z_j = k) &= \frac{\frac{\partial}{\partial \theta_i} p_{\theta}(\sum_j z_j = k)}{p_{\theta}(\sum_j z_j = k)} \\ &= \frac{p_{\theta}(\sum_j z_j = k \wedge z_i = 1)}{p_{\theta}(\sum_j z_j = k)} \\ &= p_{\theta}(z_i | \sum_j z_j = k) \end{aligned}$$

which finishes our proof. $\square$

Proposition 3. Let Entropy be defined as in Algorithm 4. Given variables $Z_1, \dots, Z_n$ and a k -subset distribution $p_{\theta}(\mathbf{z} | \sum_i z_i = k)$ parameterized by θ , Algorithm 4 computes entropy of $p_{\theta}(\mathbf{z} | \sum_i z_i = k)$.

Proof. In a slight abuse of notation, let z_n denote $z_n = 1$, and let $\bar{z}_n$ denote $z_n = 0$. Furthermore, we denote by $\sigma_n^k, \sigma_{n-1}^k$ and σ_{n-1}^{k-1} the events $\sum_{i=0}^n = k, \sum_{i=0}^{n-1} = k$ and $\sum_{i=0}^{n-1} = k-1$, respectively.

The entropy of the k -subset distribution is given by

$$H(\mathbf{Z}) = -\mathbb{E}_{\mathbf{z} \sim p_{\theta}(\mathbf{z} | \sigma_n^k)} [\log p(\mathbf{z})] = - \sum_{\mathbf{z}: \sigma_n^k} p_{\theta}(\mathbf{z} | \sigma_n^k) \log p_{\theta}(\mathbf{z} | \sigma_n^k)$$

We start by simplifying the expression for $p_{\theta}(\mathbf{z} | \sigma_n^k)$, where, by the chain rule, the above is

$$\sum_{\mathbf{z}: \sigma_n^k} p_{\theta}(\bar{z}_n | \sigma_n^k) \cdot p_{\theta}(\sigma_{n-1}^k | \sigma_n^k, \bar{z}_n) + p_{\theta}(z_n | \sigma_n^k) \cdot p_{\theta}(\sigma_{n-1}^{k-1} | \sigma_n^k, z_n)$$

^{*}In practice, this can be easily implemented. For example, in framework `Tensorflow`, it can be done by setting `tf.stop_gradients`.

Plugging the above in the expression for the entropy, distributing the sum over the product, we get

$$\begin{aligned}
&= - \sum_{\mathbf{z}:\sigma_n^k} p_{\theta}(\bar{z}_n | \sigma_n^k) \cdot p_{\theta}(\sigma_{n-1}^k | \sigma_n^k, \bar{z}_n) \\
&\quad \cdot \log [p_{\theta}(\bar{z}_n | \sigma_n^k) \cdot p_{\theta}(\sigma_{n-1}^k | \sigma_n^k, \bar{z}_n) + p_{\theta}(z_n | \sigma_n^k) \cdot p_{\theta}(\sigma_{n-1}^{k-1} | \sigma_n^k, z_n)] \\
&+ p_{\theta}(z_n | \sigma_n^k) \cdot p_{\theta}(\sigma_{n-1}^{k-1} | \sigma_n^k, z_n) \\
&\quad \cdot \log [p_{\theta}(\bar{z}_n | \sigma_n^k) \cdot p_{\theta}(\sigma_{n-1}^k | \sigma_n^k, \bar{z}_n) + p_{\theta}(z_n | \sigma_n^k) \cdot p_{\theta}(\sigma_{n-1}^{k-1} | \sigma_n^k, z_n)],
\end{aligned}$$

where, since the two events $\bar{z}_n$ and z_n are mutually exclusive, we can simplify the above to

$$\begin{aligned}
&- \sum_{\mathbf{z}:\sigma_n^k} p_{\theta}(\bar{z}_n | \sigma_n^k) \cdot p_{\theta}(\sigma_{n-1}^k | \sigma_n^k, \bar{z}_n) \cdot \log [p_{\theta}(\bar{z}_n | \sigma_n^k) \cdot p_{\theta}(\sigma_{n-1}^k | \sigma_n^k, \bar{z}_n)] \\
&\quad + p_{\theta}(z_n | \sigma_n^k) \cdot p_{\theta}(\sigma_{n-1}^{k-1} | \sigma_n^k, z_n) \cdot \log [p_{\theta}(z_n | \sigma_n^k) \cdot p_{\theta}(\sigma_{n-1}^{k-1} | \sigma_n^k, z_n)].
\end{aligned}$$

Expanding the logarithms, rearranging terms, and using that conditional probabilities sum to 1 we get

$$\begin{aligned}
&p_{\theta}(\bar{z}_n | \sigma_n^k) \log p_{\theta}(\bar{z}_n | \sigma_n^k) + p_{\theta}(z_n | \sigma_n^k) \log p_{\theta}(z_n | \sigma_n^k) \\
&+ p_{\theta}(\bar{z}_n | \sigma_n^k) \cdot p_{\theta}(\sigma_{n-1}^k | \sigma_n^k, \bar{z}_n) \log p_{\theta}(\sigma_{n-1}^k | \sigma_n^k, \bar{z}_n) \\
&+ p_{\theta}(z_n | \sigma_n^k) \cdot p_{\theta}(\sigma_{n-1}^{k-1} | \sigma_n^k, z_n) \log p_{\theta}(\sigma_{n-1}^{k-1} | \sigma_n^k, z_n) \\
&= -\mathbb{E}_{z_n \sim p_{\theta}(z_n | \sigma_n^k)} [-\log p_{\theta}(z_n | \sigma_n^k)] + \mathbb{E}_{z_n \sim p_{\theta}(z_n | \sigma_n^k)} [H(\mathbf{Z}_{:n-1} | \sigma_n^k, z_n)] \\
&= H_b(Z_n | \sigma_n^k) + \mathbb{E}_{z_n \sim p_{\theta}(z_n | \sigma_n^k)} [H(\mathbf{Z}_{:n-1} | \sigma_n^k, z_n)].
\end{aligned}$$

That is, simply stated, the entropy of the k -subset distribution decomposes as the entropy of the *constrained* distribution over Z_n , and average entropy of the distribution on the remaining variables.

As the base case, the entropy of the k -subset distribution when $k = n$ is 0; there is only one way in which to pick to choose n of n variables, and the k -subset distribution is therefore deterministic. $\square$

B OPTIMIZED ALGORITHMS

Algorithm 5 is the optimized version of Algorithm 1, both of which compute the marginal probability of the exactly- k constraint. Algorithm 6 is the optimized version of Algorithm 2, both of which sample faithfully from the k -subset distribution.

Algorithm 5 PrExactlyk(θ, l, u, k)

Input: The logits θ of the distribution, range of variable indices $[l, u]$, and the subset size k

Output: The exact marginal probability of variables summing up to k , $P(\sum_{i=l}^u X_i = k)$

if $l > u$ **then return** 0

if $l = u$ **then return** $p_{\theta}(X_l = k)$

for $m = 0$ to k **do**

$p_m = \text{PrExactlyk}(\theta, l, \lfloor u/2 \rfloor, m)$

$\text{PrExactlyk}(\theta, \lfloor u/2 \rfloor + 1, u, k - m)$

return $\sum_{m=0}^k p_m$

Algorithm 6 Sample(θ, l, u, k)

Input: The logits θ of the distribution, range of variable indices $[l, u]$, and the subset size k

Output: A sample $\mathbf{z} = (z_1, \dots, z_n)$ from $p_{\theta}(\mathbf{z} | \sum_i z_i = k)$

define $p(x = m) = p_m, m = 0, \dots, k$

// with p_m as defined in Algorithm 5

sample m^* from p

$\mathbf{z}_{l:\lfloor u/2 \rfloor} = \text{Sample}(\theta, l, \lfloor u/2 \rfloor, m^*)$

$\mathbf{z}_{\lfloor u/2 \rfloor + 1:u} = \text{Sample}(\theta, \lfloor u/2 \rfloor + 1, u, k - m^*)$

return Concat($\mathbf{z}_{l:\lfloor u/2 \rfloor}, \mathbf{z}_{\lfloor u/2 \rfloor + 1:u}$)

C EXPERIMENTAL DETAILS

C.1 SYNTHETIC EXPERIMENTS

In this experiment we analyzed the behavior of various discrete gradient estimators for the k -subset distribution. We were interested in three different metrics: the bias of the the gradients estimators, the variance of the gradient estimators, as well as the average deviation of each estimated gradient from the exact gradient. We used cosine distance, defined as $1 - \text{cosine similarity}$ as our measure of

distance, as we typically care about the direction, not the magnitude of the gradient; the latter can be recovered using an appropriate learning rate. Following Niepert et al. (2021), we chose a tractable 5-subset distribution, where $n = 10$, and were therefore limited to $\binom{10}{5} = 252$ possible subsets. We set the loss to $L(\theta) = \mathbb{E}_{z \sim p_\theta(z | \sum_i z_i = k)} [\|z - \mathbf{b}\|^2]$, where $\mathbf{b}$ is the groundtruth logits sampled from $\mathcal{N}(0, \mathbf{I})$. We used a sample size of 10000 to estimate each of our metrics.

C.2 DISCRETE VARIATIONAL AUTO-ENCODER

We tested our SIMPLE gradient estimator in the discrete k -subset Variational Auto-Encoder (VAE) setting, where the latent variables model a probability distribution over k -subsets, and has a dimensionality of 20. The experimental setup is similar to those used in prior work on the Gumbel softmax tricks (Jang et al., 2017) and IMLE (Niepert et al., 2021). The encoding and decoding functions of the VAE consist of three dense layers (encoding: 512-256-20x20; decoding: 256-512-784). As is commonplace in discrete VAEs, the loss is the sum of the reconstruction loss (binary cross-entropy loss on output pixels) and KL divergence of the k -subset distribution and the uniform distribution, known as the evidence lower bound, or the ELBO. The task being to learn a *sparse* generative model of MNIST. As in prior work, we use a batch size of 100 and train for 100 epochs, plotting the test loss after each epoch. We use the standard Adam settings in Tensorflow 2.x, and do not employ any learning rate scheduling. The encoder network consists of an input layer with dimension 784 (we flatten the images), a dense layer with dimension 512 and ReLU activation, a dense layer with dimension 256 and ReLU activation, and a dense layer with dimension 400 (20×20) which outputs θ and no non-linearity SIMPLE takes θ as input and outputs a discrete latent code of size 20×20 . The decoder network, which takes this discrete latent code as input, consists of a dense layer with dimension 256 and ReLU activation, a dense layer with dimension 512 and ReLU activation, and finally a dense layer with dimension 784 returning the logits for the output pixels. Sigmoids are applied to these logits and the binary cross-entropy loss is computed. To obtain the best performing model of each of the compared methods, we performed a grid search over the learning rate in the range $[1 \times 10^{-3}, 5 \times 10^{-4}]$, λ in the range $[1 \times 10^{-3}, 1 \times 10^{-2}, 1 \times 10^{-1}, 1 \times 10^0, 1 \times 10^1, 1 \times 10^2, 1 \times 10^3]$, and for SoG I-MLE, the temperature τ in the range $[1 \times 10^{-1}, 1 \times 10^0, 1 \times 10^1, 1 \times 10^2]$.

We will now present a formal proof on how to compute the KL-divergence between the k -subset distribution and a uniform distribution tractably and exactly.

Proposition 4. *Let $p_\theta(z | \sum_i z_i = k)$ be a k -subset distribution parameterized by θ and $\mathcal{U}(z)$ be a uniform distribution on the constrained space $\mathcal{C} = \{z | \sum_i z_i = k\}$. Then the KL-divergence between distribution $p_\theta(z | \sum_i z_i = k)$ and $\mathcal{U}(z)$ can be computed by*

$$D(p_\theta(z | \sum_i z_i = k) || \mathcal{U}(z)) = -H(z) + \log \binom{n}{k},$$

where H denote the entropy of distribution $p_\theta(z | \sum_i z_i = k)$.

Proof. By the definition of KL divergence, it holds that

$$\begin{aligned} D(p_\theta(z | \sum_i z_i = k) || \mathcal{U}(z)) &= \sum_{z \in \mathcal{C}} p_\theta(z | \sum_i z_i = k) \cdot \log \frac{p_\theta(z | \sum_i z_i = k)}{U(z)} \\ &= \left(\sum_{z \in \mathcal{C}} p_\theta(z | \sum_i z_i = k) \log p_\theta(z | \sum_i z_i = k) \right) - \sum_{z \in \mathcal{C}} p_\theta(z | \sum_i z_i = k) \log U(z) \\ &= -H(z) + \log \binom{n}{k}. \end{aligned}$$

The last equality holds since $U(z) \equiv 1/\binom{n}{k}$. $\square$

C.3 LEARNING TO EXPLAIN

The BEERADVOCATE dataset (McAuley et al., 2012) consists of free-text reviews and ratings for 4 different aspects of beer: appearance, aroma, palate, and taste. The training set has 80k reviews for

the aspect APPEARANCE and 70k reviews for all other aspects. The maximum review length is 350 tokens. We follow Niepert et al. (2021) in computing 10 different evenly sized validation/test splits of the 10k held out set and compute mean and standard deviation over 10 models, each trained on one split. In addition to the ratings for all reviews, each sentence in the test set contains annotations of the words that best describe the review score with respect to the various aspects. Following the experimental setup of recent work (Paulus et al., 2020; Niepert et al., 2021), we address the problem introduced by the L2X paper (Chen et al., 2018) of learning a k -subset distribution over words that best explain a given aspect rating. Subset precision was computed using a set of 993 annotated reviews. We use pre-trained word embeddings from Lei et al. (2016)[†] We use the standard neural network architecture from prior work Chen et al. (2018); Paulus et al. (2020) with 4 convolutional and one dense layer. This neural network outputs the parameters θ of the k -subset distribution over k -hot binary latent masks with $k \in \{5, 10, 15\}$. We train for 20 epochs using the standard Adam settings in Tensorflow 2.x, and no learning rate schedule. We always evaluate the model with the best validation MSE among the 20 epochs.

[†]<http://people.csail.mit.edu/taolei/beer/>.