
A Semantic Loss Function for Deep Learning with Symbolic Knowledge

Jingyi Xu¹ Zilu Zhang² Tal Friedman¹ Yitao Liang¹ Guy Van den Broeck¹

Abstract

This paper develops a novel methodology for using symbolic knowledge in deep learning. From first principles, we derive a semantic loss function that bridges between neural output vectors and logical constraints. This loss function captures how close the neural network is to satisfying the constraints on its output. An experimental evaluation shows that it effectively guides the learner to achieve (near-)state-of-the-art results on semi-supervised multi-class classification. Moreover, it significantly increases the ability of the neural network to predict structured objects, such as rankings and paths. These discrete concepts are tremendously difficult to learn, and benefit from a tight integration of deep learning and symbolic reasoning methods.

1. Introduction

The widespread success of representation learning raises the question of which AI tasks are amenable to deep learning, which tasks require classical model-based symbolic reasoning, and whether we can benefit from a tighter integration of both approaches. In recent years, significant effort has gone towards various ways of using representation learning to solve tasks that were previously tackled by symbolic methods. Such efforts include neural computers or differentiable programming (Weston et al., 2014; Reed & De Freitas, 2015; Graves et al., 2016; Riedel et al., 2016), relational embeddings or deep learning for graph data (Yang et al., 2014; Lin et al., 2015; Bordes et al., 2013; Neelakantan et al., 2015; Duvenaud et al., 2015; Niepert et al., 2016), neural theorem proving, and learning with constraints (Hu et al., 2016; Stewart & Ermon, 2017; Minervini et al., 2017; Wang et al., 2017).

¹Department of Computer Science, University of California Los Angeles, Los Angeles, CA, USA ²Peking University, Beijing, China. Correspondence to: Tal Friedman <tal@cs.ucla.edu>, Yitao Liang <yliang@cs.ucla.edu>, Guy Van den Broeck <guyvdb@cs.ucla.edu>.

This paper considers learning in domains where we have symbolic knowledge connecting the different outputs of a neural network. This knowledge takes the form of a constraint (or sentence) in Boolean logic. It can be as simple as an exactly-one constraint for one-hot output encodings, or as complex as a structured output prediction constraint for intricate combinatorial objects such as rankings, subgraphs, or paths. Our goal is to augment neural networks with the ability to learn how to make predictions subject to these constraints, and use the symbolic knowledge to improve the learning performance.

Most neuro-symbolic approaches aim to simulate or learn symbolic reasoning in an end-to-end deep neural network, or capture symbolic knowledge in a vector-space embedding. This choice is partly motivated by the need for smooth *differentiable* models; adding symbolic reasoning code (e.g., SAT solvers) to a deep learning pipeline destroys this property. Unfortunately, while making reasoning differentiable, the precise logical meaning of the knowledge is often lost. In this paper, we take a distinctly unique approach, and tackle the problem of differentiable but sound logical reasoning from first principles. Starting from a set of intuitive axioms, we derive the differentiable *semantic loss* which captures how well the outputs of a neural network match a given constraint. This function precisely captures the *meaning* of the constraint, and is independent of its *syntax*.

Next, we show how semantic loss gives significant *practical improvements* in semi-supervised classification. In this setting, semantic loss for the exactly-one constraint permits us to obtain a learning signal from vast amounts of unlabeled data. The key idea is that semantic loss helps us improve how confidently we are able to classify the unlabeled data. This simple addition to the loss function of standard deep learning architectures yields (near-)state-of-the-art performance in semi-supervised classification on MNIST, FASHION, and CIFAR-10 datasets.

Our final set of experiments study the benefits of semantic loss for learning tasks with highly structured output, such as preference learning and path prediction in a graph (Daumé et al., 2009; Chang et al., 2013; Choi et al., 2015; Graves et al., 2016). In these scenarios, the task is two-fold: learn both the structure of the output space, and

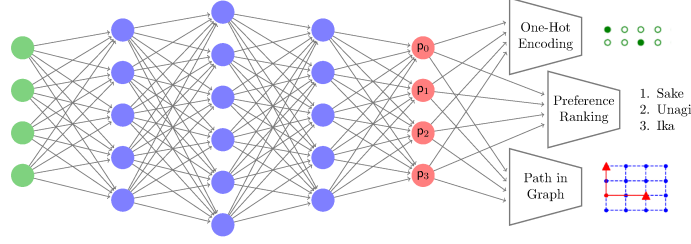


Figure 1: Outputs of a neural network feed into semantic loss functions for constraints representing a one-hot encoding, a total ranking of preferences, and paths in a grid graph.

the actual classification function within that space. By capturing the structure of the output space with logical constraints, and minimizing semantic loss for this constraint during learning, we are able to learn networks that are much more likely to correctly predict structured objects.

2. Background and Notation

To formally define semantic loss, we make use of concepts in propositional logic. We write uppercase letters (X, Y) for Boolean variables and lowercase letters (x, y) for their instantiation ($X = 0$ or $X = 1$). Sets of variables are written in bold uppercase ($\mathbf{X}, \mathbf{Y}$), and their joint instantiation in bold lowercase ($\mathbf{x}, \mathbf{y}$). A literal is a variable (X) or its negation ($\neg X$). A logical sentence (α or β) is constructed in the usual way, from variables and logical connectives ($\wedge, \vee$, etc.), and is also called a formula or constraint. A state or world $\mathbf{x}$ is an instantiation to all variables $\mathbf{X}$. A state $\mathbf{x}$ satisfies a sentence α , denoted $\mathbf{x} \models \alpha$, if the sentence evaluates to be true in that world, as defined in the usual way. A sentence α entails another sentence β , denoted $\alpha \models \beta$ if all worlds that satisfy α also satisfy β . A sentence α is logically equivalent to sentence β , denoted $\alpha \equiv \beta$, if both $\alpha \models \beta$ and $\beta \models \alpha$.

The output row vector of a neural net is denoted $\mathbf{p}$. Each value in $\mathbf{p}$ represents the probability of an output and falls in $[0, 1]$. We use both softmax and sigmoid units for our output activation functions. The notation for states $\mathbf{x}$ is used to refer the an assignment, the logical sentence enforcing that assignment, or the binary output vector capturing that same assignment, as these are all equivalent notions.

Figure 1 illustrates the three different concrete output constraints of varying difficulty that are studied in our experiments. First, we examine the exactly-one or *one-hot constraint* capturing the encoding used in multi-class classification. It states that for a set of indicators $\mathbf{X} = \{X_1, \dots, X_n\}$, one and exactly one of those indicators must be true, with the rest being false. This is enforced through a logical constraint α by conjoining sentences of the form $\neg X_1 \vee \neg X_2$ for all pairs of variables (at most one variable is true), and a single sentence $X_1 \vee \dots \vee X_n$ (at

least one variable is true). Our experiments further examine the *valid simple path constraint*. It states for a given source-destination pair and edge indicators that the edge indicators set to true must form a valid simple path from source to destination. Finally, we explore the *ordering constraint*, which requires that a set of n^2 indicator variables represent a total ordering over n variables, effectively encoding a permutation matrix. For a full description of the path and ordering constraints, we refer to Section 5.

3. Semantic Loss

In this section, we formally introduce semantic loss. We begin by giving the definition and our intuition behind it. This definition itself provides all of the necessary mechanics for enforcing constraints, and is sufficient for the understanding of our experiments in Sections 4 and 5. We also show that semantic loss is not just an arbitrary definition, but rather is defined uniquely by a set of intuitive assumptions. After stating the assumptions formally, we then provide an axiomatic proof of the uniqueness of semantic loss in satisfying these assumptions.

3.1. Definition

The semantic loss $L^s(\alpha, \mathbf{p})$ is a function of a sentence α in propositional logic, defined over variables $\mathbf{X} = \{X_1, \dots, X_n\}$, and a vector of probabilities $\mathbf{p}$ for the same variables $\mathbf{X}$. Element p_i denotes the predicted probability of variable X_i , and corresponds to a single output of the neural net. For example, the semantic loss between the one-hot constraint from the previous section, and a neural net output vector $\mathbf{p}$, is intended to capture how close the prediction $\mathbf{p}$ is to having exactly one output set to true (i.e. 1), and all others set to false (i.e. 0), regardless of which output is correct. The formal definition of this is as follows:

Definition 1 (Semantic Loss). Let $\mathbf{p}$ be a vector of probabilities, one for each variable in $\mathbf{X}$, and let α be a sentence over $\mathbf{X}$. The semantic loss between α and $\mathbf{p}$ is

$$L^s(\alpha, \mathbf{p}) \propto -\log \sum_{\mathbf{x} \models \alpha} \prod_{i: \mathbf{x} \models X_i} p_i \prod_{i: \mathbf{x} \models \neg X_i} (1 - p_i).$$

Intuitively, the semantic loss is proportional to a negative logarithm of the probability of generating a state that satisfies the constraint, when sampling values according to p . Hence, it is the self-information (or “surprise”) of obtaining an assignment that satisfies the constraint (Jones, 1979).

3.2. Derivation from First Principles

In this section, we begin with a theorem stating the uniqueness of semantic loss, as fixed by a series of axioms. The full set of axioms and the derivation of the precise semantic loss function is described in Appendix A¹.

Theorem 1 (Uniqueness). *The semantic loss function in Definition 1 satisfies all axioms in Appendix A and is the only function that does so, up to a multiplicative constant.*

In the remainder of this section, we provide a selection of the most intuitive axioms from Appendix A, as well as some key properties.

First, to retain logical meaning, we postulate that semantic loss is monotone in the order of implication.

Axiom 1 (Monotonicity). If $\alpha \models \beta$, then the semantic loss $L^s(\alpha, p) \geq L^s(\beta, p)$ for any vector p .

Intuitively, as we add stricter requirements to the logical constraint, going from β to α and making it harder to satisfy, the semantic loss cannot decrease. For example, when β enforces the output of an neural network to encode a subtree of a graph, and we tighten that requirement in α to be a path, the semantic loss cannot decrease. Every path is also a tree and any solution to α is a solution to β .

A direct consequence following the monotonicity axiom is that logically equivalent sentences must incur an identical semantic loss for the same probability vector p . Hence, the semantic loss is indeed a semantic property of the logical sentence, and *does not depend on its syntax*.

Proposition 2 (Semantic Equivalence). *If $\alpha \equiv \beta$, then the semantic loss $L^s(\alpha, p) = L^s(\beta, p)$ for any vector p .*

Another consequence is that semantic loss must be non-negative if we want the loss to be 0 for a true sentence.

Next, we state axioms establishing a correspondence between logical constraints and data. A state x can be equivalently represented as both a binary data vector, as well as a logical constraint that enforces a value for every variable in X . When both the constraint and the predicted vector represent the same state (for example, $X_1 \wedge \neg X_2 \wedge X_3$ vs. $[1\ 0\ 1]$), there should be no semantic loss.

Axiom 2 (Identity). For any state x , there is zero semantic loss between its representation as a sentence, and its representation as a deterministic vector: $\forall x, L^s(x, x) = 0$.

The axiom above together with the monotonicity axiom imply that any vector satisfying the constraint must incur zero loss. For example, when our constraint α requires that the output vector encodes an arbitrary total ranking, and the vector x correctly represents a single specific total ranking, there is no semantic loss.

Proposition 3 (Satisfaction). *If $x \models \alpha$, then the semantic loss $L^s(\alpha, x) = 0$.*

As a special case, logical literals (X or $\neg X$) constrain a single variable to take on a value, and thus play a role similar to the labels used in supervised learning. Such constraints require an even tighter correspondence: the semantic loss must act like a classical loss function (i.e., cross entropy).

Axiom 3 (Label-Literal Correspondence). The semantic loss of a single literal is proportionate to the cross-entropy loss for the equivalent data label: $L^s(X, p) \propto -\log(p)$ and $L^s(\neg X, p) \propto -\log(1 - p)$.

Appendix A states additional axioms that allow us to prove the following form of the semantic loss for a state x .

Lemma 4. *For state x and vector p , we have $L^s(x, p) \propto -\sum_{i: x_i = X_i} \log p_i - \sum_{i: x_i = \neg X_i} \log(1 - p_i)$.*

Lemma 4 falls short as a full definition of semantic loss for arbitrary sentences. One can define additional axioms to pin down L^s . For example, the following axiom is satisfied by Definition 1, and is highly desirable for learning.

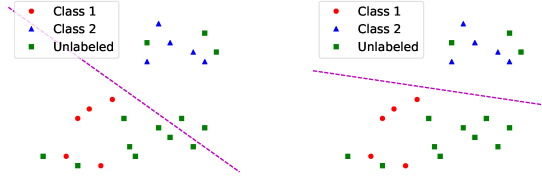
Axiom 4 (Differentiability). For any fixed α , the semantic loss $L^s(\alpha, p)$ is monotone in each probability in p , continuous and differentiable.

Appendix A makes the notion of semantic loss precise by stating one additional axiom. It is based on the observation that the state loss of Lemma 4 is proportionate to a log-probability. In particular, it corresponds to the probability of obtaining state x after independently sampling each X_i with probability p_i . We have now derived the semantic loss function from first principles, and arrived at Definition 1. Moreover, we can show that Theorem 1 holds - that it is the only choice of such a loss function.

4. Semi-Supervised Classification

The most straightforward constraint that is ubiquitous in classification is mutual exclusion over one-hot-encoded outputs. That is, for a given example, exactly one class and therefore exactly one binary indicator must be true. The machine learning community has made great strides on this task, due to the invention of assorted deep learning representations and their associated regularization terms (Krizhevsky et al., 2012; He et al., 2016). Many of these models take large amounts of labeled data for granted, and big data is indispensable for discovering accurate representations (Hastie et al., 2009). To sustain this progress, and

¹Appendices are included in the supplementary material.



(a) Trained w/o semantic loss (b) Trained with semantic loss

Figure 2: Binary classification toy example: a linear classifier without and with semantic loss.

alleviate the need for more labeled data, there is a growing interest into utilizing unlabeled data to augment the predictive power of classifiers (Stewart & Ermon, 2017; Bilenko et al., 2004). This section shows why semantic loss naturally qualifies for this task.

Illustrative Example To illustrate the benefit of semantic loss in the semi-supervised setting, we begin our discussion with a small toy example. Consider a binary classification task; see Figure 2. Ignoring the unlabeled examples, a simple linear classifier learns to distinguish the two classes by separating the labeled examples (Figure 2a). However, the unlabeled examples are also informative, as they must carry some properties that give them a particular label. This is the crux of semantic loss for semi-supervised learning: a model must confidently assign a consistent class even to unlabeled data. Encouraging the model to do so results in a more accurate decision boundary (Figure 2b).

4.1. Method

Our proposed method intends to be generally applicable and compatible with any feedforward neural net. Semantic loss is simply another regularization term that can directly be plugged into an existing loss function. More specifically, with some weight w , the new overall loss becomes

$$\text{existing loss} + w \cdot \text{semantic loss}.$$

When the constraint over the output space is simple (for example, there is a small number of solutions $\mathbf{x} \models \alpha$), semantic loss can be directly computed using Definition 1. Concretely, for the exactly-one constraint used in n -class classification, semantic loss reduces to

$$L^s(\text{exactly-one}, \mathbf{p}) \propto -\log \sum_{i=1}^n p_i \prod_{j=1, j \neq i}^n (1 - p_j),$$

where values p_i denote the probability of class i as predicted by the neural net. Semantic loss for the exactly-one constraint is efficient and causes no noticeable computational overhead in our experiments.

In general, for arbitrary constraints α , semantic loss is not efficient to compute using Definition 1, and more advanced automated reasoning is required. Section 5 discusses this issue in more detail. For example, using automated reasoning can reduce the time complexity to compute semantic loss for the exactly-one constraint from $O(n^2)$ (as shown above), to $O(n)$.

4.2. Experimental Evaluation

In this section, we evaluate semantic loss in the semi-supervised setting by comparing it with several competitive models.² As most semi-supervised learners build on a supervised learner, changing the underlying model significantly affects the semi-supervised learner’s performance. For comparison, we add semantic loss to the same base models used in ladder nets (Rasmus et al., 2015), which currently achieves state-of-the-art results on semi-supervised MNIST and CIFAR-10 (Krizhevsky, 2009). Specifically, the MNIST base model is a fully-connected multilayer perceptron (MLP), with layers of size 784-1000-500-250-250-10. On CIFAR-10, it is a 10-layer convolutional neural network (CNN) with 3-by-3 padded filters. After every 3 layers, features are subject to a 2-by-2 max-pool layer with strides of 2. Furthermore, we use ReLu (Nair & Hinton, 2010), batch normalization (Ioffe & Szegedy, 2015), and Adam optimization (Kingma & Ba, 2015) with a learning rate of 0.002. We refer to Appendix B and C for a specification of the CNN model and additional details about hyper-parameter tuning.

For all semi-supervised experiments, we use the standard 10,000 held-out test examples provided in the original datasets and randomly pick 10,000 from the standard 60,000 training examples (50,000 for CIFAR-10) as validation set. For values of N that depend on the experiment, we retain N randomly chosen labeled examples from the training set, and remove labels from the rest. We balance classes in the labeled samples to ensure no particular class is over-represented. Images are preprocessed for standardization and Gaussian noise is added to every pixel ($\sigma = 0.3$).

MNIST The permutation invariant MNIST classification task is commonly used as a test-bed for general semi-supervised learning algorithms. This setting does not use any prior information about the spatial arrangement of the input pixels. Therefore, it excludes many data augmentation techniques that involve geometric distortion of images, as well as convolutional neural networks.

When evaluating on MNIST, we run experiments for 20 epochs, with a batch size of 10. Experiments are repeated

²The code to reproduce all the experiments in this paper can be found at <https://github.com/UCLA-StarAI/Semantic-Loss/>.

Table 1: MNIST. Previously reported test accuracies followed by baselines and semantic loss results ($\pm$ stddev)

Accuracy % with # of used labels	100	1000	ALL
AtlasRBF (Pitelis et al., 2014)	91.9 (± 0.95)	96.32 (± 0.12)	98.69
Deep Generative (Kingma et al., 2014)	96.67 (± 0.14)	97.60 (± 0.02)	99.04
Virtual Adversarial (Miyato et al., 2016)	97.67	98.64	99.36
Ladder Net (Rasmus et al., 2015)	98.94 (± 0.37)	99.16 (± 0.08)	99.43 (± 0.02)
Baseline: MLP, Gaussian Noise	78.46 (± 1.94)	94.26 (± 0.31)	99.34 (± 0.08)
Baseline: Self-Training	72.55 (± 4.21)	87.43 (± 3.07)	
Baseline: MLP with Entropy Regularizer	96.27 (± 0.64)	98.32 (± 0.34)	99.37 (± 0.12)
MLP with Semantic Loss	98.38 (± 0.51)	98.78 (± 0.17)	99.36 (± 0.02)

Table 2: FASHION. Test accuracy comparison between MLP with semantic loss and ladder nets.

Accuracy % with # of used labels	100	500	1000	ALL
Ladder Net (Rasmus et al., 2015)	81.46 (± 0.64)	85.18 (± 0.27)	86.48 (± 0.15)	90.46
Baseline: MLP, Gaussian Noise	69.45 (± 2.03)	78.12 (± 1.41)	80.94 (± 0.84)	89.87
MLP with Semantic Loss	86.74 (± 0.71)	89.49 (± 0.24)	89.67 (± 0.09)	89.81

10 times with different random seeds. Table 1 compares semantic loss to three baselines and state-of-the-art results from the literature. The first baseline is a purely supervised MLP, which makes no use of unlabeled data. The second is the classic self-training method for semi-supervised learning, which operates as follows. After every 1000 iterations, the unlabeled examples that are predicted by the MLP to have more than 95% probability of belonging to a single class, are assigned a pseudo-label and become labeled data.

Additionally, we constructed a third baseline by replacing the semantic loss term with the entropy regularizer described in Grandvalet & Bengio (2005) as a direct comparison for semantic loss. With the same amount of parameter tuning, we found that using entropy achieves an accuracy of 96.27% with 100 labeled examples, and 98.32% with 1000 labelled examples, both are slightly worse than the accuracies reached by semantic loss. Furthermore, to our best knowledge, there is no straightforward method to generalize entropy loss to the settings of complex constraints, where semantic loss is clearly defined and can be easily deployed. We will discuss this more in Section 5.

Lastly, We attempted to create a fourth baseline by constructing a constraint-sensitive loss term in the style of Hu et al. (2016), using a simple extension of Probabilistic Soft Logic (PSL) (Kimmig et al., 2012). PSL translates logic into continuous domains by using soft truth values, and defines functions in the real domain corresponding to each Boolean function. This is normally done for Horn clauses, but since they are not sufficiently expressive for our constraints, we apply fuzzy operators to arbitrary sentences instead. We are forced to deal with a key difference between semantic loss and PSL: encodings in fuzzy logic are highly sensitive to the syntax used for the constraint (and therefore violate Proposition 2). We selected two reasonable encod-

ings detailed in Appendix E. The first encoding results in a constant value of 1, and thus could not be used for semi-supervised learning. The second encoding empirically deviates from 1 by < 0.01 , and since we add Gaussian noise to the pixels, no amount of tuning was able to extract meaningful supervision. Thus, we do not report these results.

When given 100 labeled examples ($N = 100$), MLP with semantic loss gains around 20% improvement over the purely supervised baseline. The improvement is even larger (25%) compared to self-training. Considering *the only change is an additional loss term*, this result is very encouraging. Comparing to the state of the art, ladder nets slightly outperform semantic loss by 0.5% accuracy. This difference may be an artifact of the excessive tuning of architectures, hyper-parameters and learning rates that the MNIST dataset has been subject to. In the coming experiments, we extend our work to more challenging datasets, in order to provide a clearer comparison with ladder nets. Before that, we want to share a few more thoughts on how semantic loss works. A classical softmax layer interprets its output as representing a categorical distribution. Hence, by normalizing its outputs, softmax enforces the same mutual exclusion constraint enforced in our semantic loss function. However, there does not exist a natural way to extend softmax loss to unlabeled samples. In contrast, semantic loss does provide a learning signal on unlabeled samples, by forcing the underlying classifier to make an decision and construct a confident hypothesis for all data. However, for the fully supervised case ($N = \text{all}$), semantic loss does not significantly affect accuracy. Because the MLP has enough capacity to almost perfectly fit the training data, where the constraint is always satisfied, semantic loss is almost always zero. This is a direct consequence of Proposition 3.

Table 3: CIFAR. Test accuracy comparison between CNN with Semantic Loss and ladder nets.

Accuracy % with # of used labels	4000	ALL
CNN Baseline in Ladder Net	76.67 (± 0.61)	90.73
Ladder Net (Rasmus et al., 2015)	79.60 (± 0.47)	
Baseline: CNN, Whitening, Cropping	77.13	90.96
CNN with Semantic Loss	81.79	90.92

FASHION The FASHION (Xiao et al., 2017) dataset consists of Zalando’s article images, aiming to serve as a more challenging drop-in replacement for MNIST. Arguably, it has not been overused and requires more advanced techniques to achieve good performance. As in the previous experiment, we run our method for 20 epochs, whereas ladder nets need 100 epochs to converge. Again, experiments are repeated 10 times and Table 2 reports the classification accuracy and its standard deviation (except for $N = \text{all}$ where it is close to 0 and omitted for space).

Experiments show that utilizing semantic loss results in a very large 17% improvement over the baseline when only 100 labels are provided. Moreover, our method compares favorably to ladder nets, except when the setting degrades to be fully supervised. Note that our method already nearly reaches its maximum accuracy with 500 labeled examples, which is only 1% of the training dataset.

CIFAR-10 To show the general applicability of semantic loss, we evaluate it on CIFAR-10. This dataset consisting of 32-by-32 RGB images in 10 classes. A simple MLP would not have enough representation power to capture the huge variance across objects within the same class. To cope with this spike in difficulty, we switch our underlying model to a 10-layer CNN as described earlier. We use a batch size of 100 samples of which half are unlabeled. Experiments are run for 100 epochs. However, due to our limited computational resources, we report on a single trial. Note that we make slight modifications to the underlying model used in ladder nets to reproduce similar baseline performance. Please refer to Appendix B for further details.

As shown in Table 3, our method compares favorably to ladder nets. However, due to the slight difference in performance between the supervised base models, a direct comparison would be methodologically flawed. Instead, we compare the net improvements over baselines. In terms of this measure, our method scores a gain of 4.66% whereas ladder nets gain 2.93%.

4.3. Discussion

The experiments so far have demonstrated the competitiveness and general applicability of our proposed method on semi-supervised learning tasks. It surpassed the previous

state of the art (ladder nets) on FASHION and CIFAR-10, while being close on MNIST. Considering the simplicity of our method, such results are encouraging. Indeed, a key advantage of semantic loss is that it only requires a simple additional loss term, and thus incurs almost no computational overhead. Conversely, this property makes our method sensitive to the underlying model’s performance.

Without the underlying predictive power of a strong supervised learning model, we do not expect to see the same benefits we observe here. Recently, we became aware that Miyato et al. (2016) extended their work to CIFAR-10 and achieved state-of-the-art results (Miyato et al., 2017), surpassing our performance by 5%. In future work, we plan to investigate whether applying semantic loss on their architecture would yield an even stronger performance.

Figure 5 in the appendix illustrates the effect of semantic loss on FASHION pictures whose correct label was hidden from the learner. Pictures 5a and 5b are correctly classified by the supervised base model, and on the first set it is confident about this prediction ($p_i > 0.8$). Semantic loss rarely diverts the model from these initially correct labels. However, it bootstraps these unlabeled examples to achieve higher confidence in the learned concepts. With this additional learning signal, the model changes its beliefs about Pictures 5c, which it was previously uncertain about. Finally, even on confidently misclassified Pictures 5d, semantic loss is able to remedy the mistakes of the base model.

5. Learning with Complex Constraints

While much of current machine learning research is focused on problems such as multi-class classification, there remain a multitude of difficult problems involving highly constrained output domains. As mentioned in the previous section, semantic loss has little effect on the fully-supervised exactly-one classification problem. This leads us to seek out more difficult problems to illustrate that semantic loss can also be highly informative in the supervised case, provided the output domain is a sufficiently complex space. Because semantic loss is defined by a Boolean formula, it can be used on any output domain that can be fully described in this manner. Here, we develop a framework for making semantic loss tractable on highly complex constraints, and evaluate it on some difficult examples.

5.1. Tractability of Semantic Loss

Our goal here is to develop a general method for computing both semantic loss and its gradient in a tractable manner. Examining Definition 1 of semantic loss, we see that the right-hand side is a well-known automated reasoning task called weighted model counting (WMC) (Chavira & Darwiche, 2008; Sang et al., 2005).

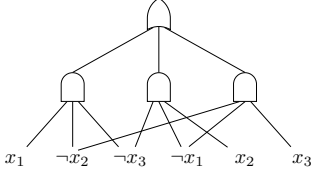


Figure 3: A compiled decomposable and deterministic circuit for the exactly-one constraint with 3 variables.

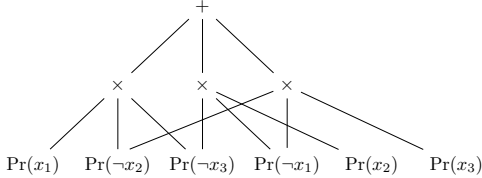


Figure 4: The corresponding arithmetic circuit for the exactly-one constraint with 3 variables.

Furthermore, we know of circuit languages that compute WMCs, and that are amenable to backpropagation (Darwiche, 2003). We use the circuit compilation techniques in Darwiche (2011) to build a Boolean circuit representing semantic loss. We refer to the literature for details of this compilation approach. Due to certain properties of this circuit form, we can use it to compute both the values and the gradients of semantic loss in time linear in the size of the circuit (Darwiche & Marquis, 2002). Once constructed, we can add it to our standard loss function as described in Section 4.1.

Figure 3 shows an example Boolean circuit for the exactly-one constraint with 3 variables. We begin with the standard logical encoding for the exactly-one constraint $(x_1 \vee x_2 \vee x_3) \wedge (\neg x_1 \vee \neg x_2) \wedge (\neg x_1 \wedge \neg x_3) \wedge (\neg x_2 \wedge \neg x_3)$, and then compile it into a circuit that can perform WMC efficiently (Chavira & Darwiche, 2008). The cost of this step depends on the type of the constraint: for bounded-treewidth constraints it can be done efficiently, and for some constraints exact compilation is theoretically hard. In that case, we have to rely on advanced knowledge compilation algorithms to still perform this step efficiently in practice. Our semantic loss framework can be applied regardless of how the circuit gets compiled. On our example, following the circuit bottom up, the logical function can be read as $(x_1 \wedge \neg x_2 \wedge \neg x_3) \vee (\neg x_1 \wedge x_2 \wedge \neg x_3) \vee (\neg x_1 \wedge \neg x_2 \wedge x_3)$. Once this Boolean circuit is built, we can convert it to an arithmetic circuit, by simply changing AND gates into $*$, and OR gates into $+$, as shown in Figure 4. Now, by pushing the probabilities up through the arithmetic circuit, evaluating the root gives the probability of the logical formula described by the Boolean circuit – this is precisely the exponentiated semantic loss. Notice that this computation was not possible with the Boolean formula we began with:

it is a direct result of our circuit having two key properties called determinism and decomposability. Finally, we can similarly do another pass down on the circuit to compute partial derivatives (Darwiche & Marquis, 2002).

5.2. Experimental Evaluation

Our ambition when evaluating semantic loss’ performance on complex constraints is not to achieve state-of-the-art performance on any particular problem, but rather to highlight its effect. To this end, we evaluate our method on problems with a difficult output space, where the model could no longer be fit directly from data, and purposefully use simple MLPs for evaluation. We want to emphasize that the constraints used in this evaluation are intentionally designed to be very difficult; much more so than the simple implications that are usually studied (e.g., Hu et al. (2016)). Hyper-parameter tuning details are again in Appendix C.

Grids We begin with a classic algorithmic problem: finding the shortest path in a graph. Specifically, we use a 4-by-4 grid $G = (V, E)$ with uniform edge weights. We randomly remove edges for each example to increase difficulty. Formally, our input is a binary vector of length $|V| + |E|$, with the first $|V|$ variables indicating sources and destinations, and the next $|E|$ which edges are removed. Similarly, each label is a binary vector of length $|E|$ indicating which edges are in the shortest path. Finally, we require through our constraint α that the output form a valid simple path between the desired source and destination. To compile this constraint, we use the method of Nishino et al. (2017) to encode pairwise simple paths, and enforce the correct source and destination. For more details on the constraint and data generation process, see Appendix D.

To evaluate, we use a dataset of 1600 examples, with a 60/20/20 train/validation/test split. Table 4 compares test accuracy between a 5-layer MLP baseline, and the same model augmented with semantic loss. We report three different accuracies that illustrate the effect of semantic loss: “Coherent” indicates the percentage of examples for which the classifier gets the entire configuration right, while “Incoherent” measures the percentage of individually correct binary labels, which as a whole may not constitute a valid path at all. Finally, “Constraint” describes the percentage of predictions given by the model that satisfy the constraint associated with the problem. In the case of incoherent accuracy, semantic loss has little effect, and in fact slightly reduces the accuracy as it combats the standard sigmoid cross entropy. In regard to coherent accuracy however, semantic loss has a very large effect in guiding the network to jointly learn true paths, rather than optimizing each binary output individually. We further see this by observing the large increase in the percentage of predictions that really are paths between the desired nodes in the graph.

Table 4: Grid shortest path test results: coherent, incoherent and constraint accuracy.

Test accuracy %	Coherent	Incoherent	Constraint
5-layer MLP	5.62	85.91	6.99
+ Semantic loss	28.51	83.14	69.89

Table 5: Preference prediction test results: coherent, incoherent and constraint accuracy.

Test accuracy %	Coherent	Incoherent	Constraint
3-layer MLP	1.01	75.78	2.72
+ Semantic loss	13.59	72.43	55.28

Preference Learning The next problem is that of predicting a complete order of preferences. That is, for a given set of user features, we want to predict how the user ranks their preference over a fixed set of items. We encode a preference ordering over n items as a flattened binary matrix $\{X_{ij}\}$, where for each $i, j \in \{1, \dots, n\}$, X_{ij} denotes that item i is at position j (Choi et al., 2015). Clearly, not all configurations of outputs correspond to a valid ordering, so our constraint allows only for those that are.

We use preference ranking data over 10 types of sushi for 5000 individuals, taken from PREFLIB (Mattei & Walsh, 2013). We take the ordering over 6 types of sushi as input features to predict the ordering over the remaining 4 types, with splits identical to those in Shen et al. (2017). We again split the data 60/20/20 into train/test/split, and employ a 3-layer MLP as our baseline. Table 5 compares the baseline to the same MLP augmented with semantic loss for valid total orderings. Again, we see that semantic loss has a marginal effect on incoherent accuracy, but significantly improves the network’s ability to predict valid, correct orderings. Remarkably, without semantic loss, the network is only able to output a valid ordering on 1% of examples.

6. Related Work

Incorporating symbolic background knowledge into machine learning is a long-standing challenge (Srinivasan et al., 1995). It has received considerable attention for structured prediction in natural language processing, in both supervised and semi-supervised settings. For example, *constrained conditional models* extend linear models with constraints that are enforced through integer linear programming (Chang et al., 2008; 2013). Constraints have also been studied in the context of probabilistic graphical models (Mateescu & Dechter, 2008; Ganchev et al., 2010). Kisa et al. (2014) utilize a circuit language called the *probabilistic sentential decision diagram* to induce distributions over arbitrary logical formulas. They learn generative models that satisfy preference and path constraints (Choi et al.,

2015; 2016), which we study in a discriminative setting.

Various deep learning techniques have been proposed to enforce either arithmetic constraints (Pathak et al., 2015; Márquez-Neila et al., 2017) or logical constraints (Rocktäschel et al., 2015; Hu et al., 2016; De-meester et al., 2016; Stewart & Ermon, 2017; Minervini et al., 2017; Diligenti et al., 2017; Donadello et al., 2017) on the output of a neural network. The common approach is to reduce logical constraints into differentiable arithmetic objectives by replacing logical operators with their fuzzy t-norms and logical implications with simple inequalities. A downside of this fuzzy relaxation is that the logical sentences lose their precise meaning. The learning objective becomes a function of the syntax rather than the semantics (see Section 4). Moreover, these relaxations are often only applied to Horn clauses. One alternative is to encode the logic into a factor graph and perform loopy belief propagation to compute a loss function (Naradowsky & Riedel, 2017), which is known to have issues in the presence of complex logical constraints (Smith & Gogate, 2014).

Several specialized techniques have been proposed to exploit the rich structure of real-world labels. Deng et al. (2014) propose hierarchy and exclusion graphs that jointly model hierarchical categories. It is a method invented to address examples whose labels are not provided at the most specific level. Finally, the objective of semantic loss to increase the confidence of predictions on unlabeled data is related to information-theoretic approaches to semi-supervised learning (Grandvalet & Bengio, 2005; Erkan & Altun, 2010), and approaches that increase robustness to output perturbation (Miyato et al., 2016). A key difference between semantic loss and these information-theoretic losses is that semantic loss generalizes to arbitrary logical output constraints that are much more complex.

7. Conclusions & Future Work

Both reasoning and semi-supervised learning are often identified as key challenges for deep learning going forward. In this paper, we developed a principled way of combining automated reasoning for propositional logic with existing deep learning architectures. Moreover, we showed that semantic loss provides significant benefits during semi-supervised classification, as well as deep structured prediction for highly complex output spaces.

An interesting direction for future work is to come up with effective approximations of semantic loss, for settings where even the methods we have described are not sufficient. There are several potential ways to proceed with this, including hierarchical abstractions, relaxations of the constraints, or projections on random subsets of variables.

Acknowledgements

This research was conducted while Zilu Zhang was a visiting student at StarAI Lab, UCLA. The authors thank Arthur Choi and Yujia Shen for helpful discussions. This work is partially supported by NSF grants #IIS-1657613, #IIS-1633857 and DARPA XAI grant #N66001-17-2-4032.

References

- Bilenko, M., Basu, S., and Mooney, R. J. Integrating constraints and metric learning in semi-supervised clustering. In *ICML*, pp. 11. ACM, 2004.
- Bordes, A., Usunier, N., Garcia-Duran, A., Weston, J., and Yakhnenko, O. Translating embeddings for modeling multi-relational data. In *NIPS*, pp. 2787–2795, 2013.
- Chang, K.-W., Samdani, R., and Roth, D. A constrained latent variable model for coreference resolution. In *EMNLP*, 2013.
- Chang, M.-W., Ratnoff, L.-A., Rizzolo, N., and Roth, D. Learning and inference with constraints. In *AAAI*, pp. 1513–1518, 2008.
- Chavira, M. and Darwiche, A. On probabilistic inference by weighted model counting. *JAIR*, 2008.
- Choi, A., Van den Broeck, G., and Darwiche, A. Tractable learning for structured probability spaces: A case study in learning preference distributions. In *IJCAI*, 2015.
- Choi, A., Tavabi, N., and Darwiche, A. Structured features in naive Bayes classification. In *AAAI*, pp. 3233–3240, 2016.
- Darwiche, A. A differential approach to inference in bayesian networks. *J. ACM*, 50(3):280–305, May 2003. ISSN 0004-5411. doi: 10.1145/765568.765570.
- Darwiche, A. SDD: A new canonical representation of propositional knowledge bases. In *IJCAI*, 2011.
- Darwiche, A. and Marquis, P. A knowledge compilation map. *JAIR*, 17:229–264, 2002.
- Daumé, H., Langford, J., and Marcu, D. Search-based structured prediction. *Machine learning*, 75(3):297–325, 2009.
- Demeester, T., Rocktäschel, T., and Riedel, S. Lifted rule injection for relation embeddings. In *EMNLP*, pp. 1389–1399, 2016.
- Deng, Ding, N., Jia, Y., Frome, A., Murphy, K., Bengio, S., Li, Y., Neven, H., and Adam, H. Large-scale object classification using label relation graphs. In *ECCV*, volume 8689, 2014.
- Diligenti, M., Gori, M., and Sacca, C. Semantic-based regularization for learning and inference. *JAIR*, 244:143–165, 2017.
- Donadello, I., Serafini, L., and Garcez, A. d. Logic tensor networks for semantic image interpretation. In *IJCAI*, pp. 1596–1602, 2017.
- Duvenaud, D. K., Maclaurin, D., Iparraguirre, J., Bombarell, R., Hirzel, T., Aspuru-Guzik, A., and Adams, R. P. Convolutional networks on graphs for learning molecular fingerprints. In *NIPS*, pp. 2224–2232, 2015.
- Erkan, A. and Altun, Y. Semi-supervised learning via generalized maximum entropy. In *AISTATS*, volume PMLR, pp. 209–216, 2010.
- Ganchev, K., Gillenwater, J., Taskar, B., et al. Posterior regularization for structured latent variable models. *JMLR*, 11(Jul):2001–2049, 2010.
- Grandvalet, Y. and Bengio, Y. Semi-supervised learning by entropy minimization. In *NIPS*, 2005.
- Graves, A., Wayne, G., Reynolds, M., Harley, T., Danihelka, I., Grabska-Barwińska, A., Colmenarejo, S. G., Grefenstette, E., Ramalho, T., Agapiou, J., et al. Hybrid computing using a neural network with dynamic external memory. *Nature*, 538(7626):471–476, 2016.
- Hastie, T., Tibshirani, R., and Friedman, J. Overview of supervised learning. In *The elements of statistical learning*, pp. 9–41. Springer, 2009.
- He, K., Zhang, X., Ren, S., and Sun, J. Deep residual learning for image recognition. In *CVPR*, June 2016.
- Hu, Z., Ma, X., Liu, Z., Hovy, E., and Xing, E. Harnessing deep neural networks with logic rules. In *ACL*, 2016.
- Ioffe, S. and Szegedy, C. Batch normalization: Accelerating deep network training by reducing internal covariate shift. In *ICML*, pp. 448–456, 2015.
- Jones, D. S. *Elementary information theory*. Clarendon Press, 1979.
- Kimmig, A., Bach, S. H., Broecheler, M., Huang, B., and Getoor, L. A short introduction to probabilistic soft logic. In *NIPS (Workshop Track)*, 2012.
- Kingma, D. P. and Ba, J. L. Adam: A method for stochastic optimization. In *ICLR*, 2015.
- Kingma, D. P., Mohamed, S., Jimenez Rezende, D., and Welling, M. Semi-supervised learning with deep generative models. In Ghahramani, Z., Welling, M., Cortes, C., Lawrence, N. D., and Weinberger, K. Q. (eds.), *NIPS*, pp. 3581–3589. Curran Associates, Inc., 2014.

- Kisa, D., Van den Broeck, G., Choi, A., and Darwiche, A. Probabilistic sentential decision diagrams. In *KR*, 2014.
- Krizhevsky, A. Learning multiple layers of features from tiny images. 2009.
- Krizhevsky, A., Sutskever, I., and Hinton, G. E. Imagenet classification with deep convolutional neural networks. In Pereira, F., Burges, C. J. C., Bottou, L., and Weinberger, K. Q. (eds.), *NIPS*, pp. 1097–1105, 2012.
- Lin, Y., Liu, Z., Sun, M., Liu, Y., and Zhu, X. Learning entity and relation embeddings for knowledge graph completion. In *AAAI*, 2015.
- Márquez-Neila, P., Salzmann, M., and Fua, P. Imposing hard constraints on deep networks: Promises and limitations. *arXiv preprint arXiv:1706.02025*, 2017.
- Mateescu, R. and Dechter, R. Mixed deterministic and probabilistic networks. *Annals of mathematics and artificial intelligence*, 54(1-3):3, 2008.
- Mattei, N. and Walsh, T. Preflib: A library of preference data [HTTP://PREFLIB.ORG](http://preflib.org). In *ADT*, 2013.
- Minervini, P., Demeester, T., Rocktäschel, T., and Riedel, S. Adversarial sets for regularising neural link predictors. *arXiv preprint arXiv:1707.07596*, 2017.
- Miyato, T., Maeda, S.-i., Koyama, M., Nakae, K., and Ishii, S. Distributional smoothing with virtual adversarial training. In *ICLR*, 2016.
- Miyato, T., Maeda, S.-i., Koyama, M., Nakae, K., and Ishii, S. Virtual adversarial training: a regularization method for supervised and semi-supervised learning. *ArXiv e-prints*, 2017.
- Nair, V. and Hinton, G. E. Rectified linear units improve restricted boltzmann machines. In *ICML*, pp. 807–814. Omnipress, 2010.
- Naradowsky, J. and Riedel, S. Modeling exclusion with a differentiable factor graph constraint. In *ICML (Workshop Track)*, 2017.
- Neelakantan, A., Roth, B., and McCallum, A. Compositional vector space models for knowledge base inference. In *ACL-IJCNLP*, pp. 156–166, 2015.
- Niepert, M., Ahmed, M., and Kutzkov, K. Learning convolutional neural networks for graphs. In *ICML*, pp. 2014–2023, 2016.
- Nishino, M., Yasuda, N., Minato, S., and Nagata, M. Compiling graph substructures into sentential decision diagrams. In *AAAI*, 2017.
- Pathak, D., Krahenbuhl, P., and Darrell, T. Constrained convolutional neural networks for weakly supervised segmentation. In *ICCV*, pp. 1796–1804, 2015.
- Pitelis, N., Russell, C., and Agapito, L. Semi-supervised learning using an unsupervised atlas. In *ECML-PKDD*, pp. 565–580. Springer, 2014.
- Rasmus, A., Berglund, M., Honkala, M., Valpola, H., and Raiko, T. Semi-supervised learning with ladder networks. In *NIPS*, 2015.
- Reed, S. and De Freitas, N. Neural programmer-interpreters. *arXiv preprint arXiv:1511.06279*, 2015.
- Riedel, S., Bosnjak, M., and Rocktäschel, T. Programming with a differentiable forth interpreter. *CoRR*, abs/1605.06640, 2016.
- Rocktäschel, T., Singh, S., and Riedel, S. Injecting logical background knowledge into embeddings for relation extraction. In *HLT-NAACL*, 2015.
- Sang, T., Beame, P., and Kautz, H. A. Performing bayesian inference by weighted model counting. In *AAAI*, volume 5, pp. 475–481, 2005.
- Shen, Y., Choi, A., and Darwiche, A. A tractable probabilistic model for subset selection. In *UAI*, 2017.
- Smith, D. and Gogate, V. Loopy belief propagation in the presence of determinism. In *AISTATS*, 2014.
- Srinivasan, A., Muggleton, S., and King, R. Comparing the use of background knowledge by inductive logic programming systems. In *ILP*, pp. 199–230, 1995.
- Stewart, R. and Ermon, S. Label-free supervision of neural networks with physics and domain knowledge. In *AAAI*, pp. 2576–2582, 2017.
- Wang, M., Tang, Y., Wang, J., and Deng, J. Premise selection for theorem proving by deep graph embedding. *arXiv preprint arXiv:1709.09994*, 2017.
- Weston, J., Chopra, S., and Bordes, A. Memory networks. *arXiv preprint arXiv:1410.3916*, 2014.
- Xiao, H., Rasul, K., and Vollgraf, R. Fashion-mnist: a novel image dataset for benchmarking machine learning algorithms. *CoRR*, abs/1708.07747, 2017.
- Yang, M.-C., Duan, N., Zhou, M., and Rim, H.-C. Joint relational embeddings for knowledge-based question answering. In *EMNLP*, 2014.

A. Axiomatization of Semantic Loss: Details

This appendix provides further details on our axiomatization of semantic loss. We detail here a complete axiomatization of semantic loss, which will involve restating some axioms and propositions from the main paper.

The first axiom says that there is no loss when the logical constraint α is always true (it is a logical tautology), independent of the predicted probabilities p .

Axiom 5 (Truth). The semantic loss of a true sentence is zero: $\forall p, L^s(\text{true}, p) = 0$.

Next, when enforcing two constraints on disjoint sets of variables, we want the ability to compute semantic loss for the two constraints separately, and sum the results for their joint semantic loss.

Axiom 6 (Additive Independence). Let α be a sentence over $\mathbf{X}$ with probabilities p . Let β be a sentence over $\mathbf{Y}$ disjoint from $\mathbf{X}$ with probabilities q . The semantic loss between sentence $\alpha \wedge \beta$ and the joint probability vector $[p \ q]$ decomposes additively: $L^s(\alpha \wedge \beta, [p \ q]) = L^s(\alpha, p) + L^s(\beta, q)$.

It directly follows from Axioms 5 and 6 that the probabilities of variables that are not used on the constraint do not affect the semantic loss.

Proposition 5 formalizes this intuition.

Proposition 5 (Locality). Let α be a sentence over $\mathbf{X}$ with probabilities p . For any $\mathbf{Y}$ disjoint from $\mathbf{X}$ with probabilities q , the semantic loss $L^s(\alpha, [p \ q]) = L^s(\alpha, p)$.

Proof. Follows from the additive independence and truth axioms. Set $\beta = \text{true}$ in the additive independence axiom, and observe that this sets $L^s(\beta, q) = 0$ because of the truth axiom. $\square$

To maintain logical meaning, we postulate that semantic loss is monotone in the order of implication.

Axiom 7 (Monotonicity). If $\alpha \models \beta$, then the semantic loss $L^s(\alpha, p) \geq L^s(\beta, p)$ for any vector p .

Intuitively, as we add stricter requirements to the logical constraint, going from β to α and making it harder to satisfy, semantic loss cannot decrease. For example, when β enforces the output of an neural network to encode a subtree of a graph, and we tighten that requirement in α to be a path, semantic loss cannot decrease. Every path is also a tree and any solution to α is a solution to β .

A first consequence following the monotonicity axiom is that logically equivalent sentences must incur an identical semantic loss for the same probability vector p . Hence, the semantic loss is indeed a semantic property of the logical

sentence, and *does not depend on the syntax* of the sentence.

Proposition 6. If $\alpha \equiv \beta$, then the semantic loss $L^s(\alpha, p) = L^s(\beta, p)$ for any vector p .

A second consequence is that semantic loss must be non-negative.

Proposition 7 (Non-Negativity). Semantic loss is non-negative.

Proof. Because $\alpha \models \text{true}$ for all α , the monotonicity axiom implies that $\forall p, L^s(\alpha, p) \geq L^s(\text{true}, p)$. By the truth axiom, $L^s(\text{true}, p) = 0$, and therefore $L^s(\alpha, p) \geq 0$ for all choices of α and p . $\square$

A state $\mathbf{x}$ is equivalently represented as a data vector, as well as a logical constraint that enforces a value for every variable in $\mathbf{X}$. When both the constraint and the predicted vector represent the same state (for example, $X_1 \wedge \neg X_2 \wedge X_3$ vs. $[1 \ 0 \ 1]$), there should be no semantic loss.

Axiom 8 (Identity). For any state $\mathbf{x}$, there is zero semantic loss between its representation as a sentence, and its representation as a deterministic vector: $\forall \mathbf{x}, L^s(\mathbf{x}, \mathbf{x}) = 0$.

The axioms above together imply that any vector satisfying the constraint must incur zero loss. For example, when our constraint α requires that the output vector encodes an arbitrary total ranking, and the vector $\mathbf{x}$ correctly represents a single specific total ranking, there is no semantic loss.

Proposition 8 (Satisfaction). If $\mathbf{x} \models \alpha$, then the semantic loss $L^s(\alpha, \mathbf{x}) = 0$.

Proof of Proposition 8. The monotonicity axiom specializes to say that if $\mathbf{x} \models \alpha$, we have that $\forall p, L^s(\mathbf{x}, p) \geq L^s(\alpha, p)$. By choosing p to be $\mathbf{x}$, this implies $L^s(\mathbf{x}, \mathbf{x}) \geq L^s(\alpha, \mathbf{x})$. From the identity axiom, $L^s(\mathbf{x}, \mathbf{x}) = 0$, and therefore $0 \geq L^s(\alpha, \mathbf{x})$. Proposition 7 bounds the loss from below as $L^s(\alpha, \mathbf{x}) \geq 0$. $\square$

As a special case, logical literals (x or $\neg x$) constrain a single variable to take on a single value, and thus play a role similar to the labels used in supervised learning. Such constraints require an even tighter correspondence: semantic loss must act like a classical loss function (i.e., cross entropy).

Axiom 9 (Label-Literal Correspondence). The semantic loss of a single literal is proportionate to the cross-entropy loss for the equivalent data label: $L^s(x, p) \propto -\log(p)$ and $L^s(\neg x, p) \propto -\log(1 - p)$.

Next, we have the symmetry axioms.

Axiom 10 (Value Symmetry). For all p and α , we have that $L^s(\alpha, p) = L^s(\bar{\alpha}, 1 - p)$ where $\bar{\alpha}$ replaces every variable in α by its negation.

Axiom 11 (Variable Symmetry). Let α be a sentence over $\mathbf{X}$ with probabilities $\mathbf{p}$. Let π be a permutation of the variables $\mathbf{X}$, let $\pi(\alpha)$ be the sentence obtained by replacing variables x by $\pi(x)$, and let $\pi(\mathbf{p})$ be the corresponding permuted vector of probabilities. Then, $L^s(\alpha, \mathbf{p}) = L^s(\pi(\alpha), \pi(\mathbf{p}))$.

The value and variable symmetry axioms together imply the equality of the multiplicative constants in the label-literal duality axiom for all literals.

Lemma 9. *There exists a single constant K such that $L^s(X, \mathbf{p}) = -K \log(p)$ and $L^s(\neg X, \mathbf{p}) = -K \log(1 - p)$ for any literal x .*

Proof. Value symmetry implies that $L^s(X_i, \mathbf{p}) = L^s(\neg X_i, 1 - \mathbf{p})$. Using label-literal correspondence, this implies $K_1 \log(p_i) = K_2 \log(1 - (1 - p_i))$ for the multiplicative constants K_1 and K_2 that are left unspecified by that axiom. This implies that the constants are identical. A similar argument based on variable symmetry proves equality between the multiplicative constants for different i . $\square$

Finally, this allows us to prove the following form of semantic loss for a state $\mathbf{x}$.

Lemma 10. *For state $\mathbf{x}$ and vector $\mathbf{p}$, we have $L^s(\mathbf{x}, \mathbf{p}) \propto -\sum_{i:\mathbf{x} \models X_i} \log p_i - \sum_{i:\mathbf{x} \models \neg X_i} \log(1 - p_i)$.*

Proof of Lemma 10. A state $\mathbf{x}$ is a conjunction of independent literals, and therefore subject to the additive independence axiom. Each literal’s loss in this sum is defined by Lemma 9. $\square$

The following and final axiom requires that semantic loss is proportionate to the logarithm of a function that is additive for mutually exclusive sentences.

Axiom 12 (Exponential Additivity). Let α and β be mutually exclusive sentences (i.e., $\alpha \wedge \beta$ is unsatisfiable), and let $f^s(K, \alpha, \mathbf{p}) = K^{-L^s(\alpha, \mathbf{p})}$. Then, there exists a positive constant K such that $f^s(K, \alpha \vee \beta, \mathbf{p}) = f^s(K, \alpha, \mathbf{p}) + f^s(K, \beta, \mathbf{p})$.

We are now able to state and prove the main uniqueness theorem.

Theorem 11 (Uniqueness). *The semantic loss function in Definition 1 satisfies all axioms in Appendix A and is the only function that does so, up to a multiplicative constant.*

Proof of Theorem 11. The truth axiom states that $\forall \mathbf{p}, f^s(K, \text{true}, \mathbf{p}) = 1$ for all positive constants K . This is the first Kolmogorov axiom of probability. The second Kolmogorov axiom for $f^s(K, \cdot, \mathbf{p})$ follows from the additive independence axiom of semantic loss. The

third Kolmogorov axiom (for the finite discrete case) is given by the exponential additivity axiom of semantic loss. Hence, $f^s(K, \cdot, \mathbf{p})$ is a probability distribution for some choice of K , which implies the definition up to a multiplicative constant. $\square$

B. Specification of the Convolutional Neural Network Model

Table 6 shows the slight architectural difference between the CNN used in ladder nets and ours. The major difference lies in the choice of ReLU. Note we add standard padded cropping to preprocess images and an additional fully connected layer at the end of the model, neither is used in ladder nets. We only make those slight modification so that the baseline performance reported by Rasmus et al. (2015) can be reproduced.

Table 6: Specifications of CNNs in Ladder Net and our proposed method.

CNN in Ladder Net	CNN in this paper
Input 32×32 RGB image	
	Resizing to 36 × 36 with padding; Cropping Back
Whitening Contrast Normalization Gaussian Noise with std. of 0.3	
3×3 conv. 96 BN LeakyReLU	3×3 conv. 96 BN ReLU
3×3 conv. 96 BN LeakyReLU	3×3 conv. 96 BN ReLU
3×3 conv. 96 BN LeakyReLU	3×3 conv. 96 BN ReLU
2×2 max-pooling stride 2 BN	
3×3 conv. 192 BN LeakyReLU	3×3 conv. 192 BN ReLU
3×3 conv. 192 BN LeakyReLU	3×3 conv. 192 BN ReLU
3×3 conv. 192 BN LeakyReLU	3×3 conv. 192 BN ReLU
2×2 max-pooling stride 2 BN	
3×3 conv. 192 BN LeakyReLU	3×3 conv. 192 BN ReLU
1×1 conv. 192 BN LeakyReLU	3×3 conv. 192 BN ReLU
1×1 conv. 10 BN LeakyReLU	1×1 conv. 10 BN ReLU
Global meanpool BN	
	Fully connected BN
10-way softmax	

C. Hyper-parameter Tuning Details

Validation sets are used for tuning the weight associated with semantic loss, the only hyper-parameter that causes noticeable difference in performance for our method. For our semi-supervised classification experiments, we perform a grid search over $\{0.001, 0.005, 0.01, 0.05, 0.1\}$ to find the optimal value. Empirically, 0.005 always gives the

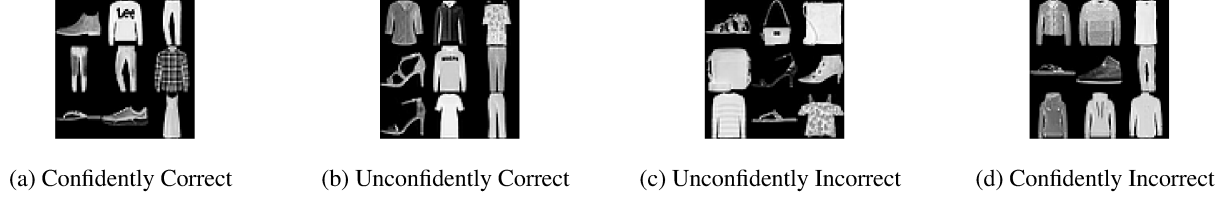


Figure 5: FASHION pictures grouped by how confidently the supervised base model classifies them correctly. With semantic loss, the final semi-supervised model predicts all correctly and confidently.

best or nearly the best results and we report its results on all experiments.

For the FASHION dataset specifically, because MNIST and FASHION share the same image size and structure, methods developed in MNIST should be able to directly perform on FASHION without heavy modifications. Because of this, we use the same hyper-parameters when evaluating our method. However, for the sake of fairness, we subject ladder nets to a small-scale parameter tuning in case its performance is more volatile.

For the grids experiment, the only hyper parameter that needed to be tuned was again the weight given to semantic loss, which after trying $\{0.001, 0.005, 0.01, 0.05, 0.1, 0.5, 1\}$ was selected to be 0.5 based on validation results. For the preference learning experiment, we initially chose the semantic loss weight from $\{0.001, 0.005, 0.01, 0.05, 0.1, 0.5, 1\}$ to be 0.1 based on validation, and then further tuned the weight to 0.25.

D. Complex Constraint Specification

Grids To compile our grid constraint, we first use Nishino et al. (2017) to generate a constraint for each source destination pair. Then, we conjoin each of these with indicators specifying which source and destination pair must be used, and finally we disjoin all of these together to form our constraint.

To generate the data, we begin by randomly removing one third of edges. We then filter out connected components with fewer than 5 nodes to reduce degenerate cases, and proceed with randomly selecting pairs of nodes to create data points.

The predictive model we employ as our baseline is a 5 layer MLP with 50 hidden sigmoid units per layer. It is trained using Adam Optimizer, with full data batches (Kingma & Ba, 2015). Early stopping with respect to validation loss is used as a regularizer.

Preference Learning We split each user’s ordering into their ordering over sushis 1,2,3,5,7,8, which we use as the features, and their ordering over 4,6,9,10 which are the la-

bels we predict. The constraint is compiled directly from logic, as this can be done in a straightforward manner for an n-item ordering.

The predictive model we use here is a 3 layer MLP with 25 hidden sigmoid units per layer. It is trained using Adam Optimizer with full data batches (Kingma & Ba, 2015). Early stopping with respect to validation loss is used as a regularizer.

E. Probabilistic Soft Logic Encodings

We here give both encodings on the exactly-one constraint over three x_1, x_2, x_3 . The first encoding is:

$$(\neg x_1 \wedge x_2 \wedge x_3) \vee (x_1 \wedge \neg x_2 \wedge x_3) \vee (x_1 \wedge x_2 \wedge \neg x_3)$$

The second encoding is:

$$(x_1 \vee x_2 \vee x_3) \wedge (\neg x_1 \vee \neg x_2) \wedge (\neg x_1 \vee \neg x_3) \wedge (\neg x_2 \vee \neg x_3)$$

Both encodings extend to cases whether the number of variables is arbitrary.

The norm functions used for these experiments are as described in Kimmig et al. (2012), with the loss for an interpretation I being defined as follows:

$$\begin{aligned} x_1 \wedge x_2 &= \max\{0, I(x_1) + I(x_2) - 1\} \\ x_1 \vee x_2 &= \min\{I(x_1) + I(x_2), 1\} \\ \neg x_1 &= 1 - I(x_1) \end{aligned}$$