



What Can We Compute from a Representation?

Structure and Composition for Tractable Reasoning

YooJung Choi

ICML 2026 Workshop on Connecting Low-rank Representations in AI – *July 11, 2026*

*Data & Model
Compression*

Tensor Factorizations

Low-rank Adapters

*We exploit **structure** like low-rankness all the time*

Interpretability

*Factor Graphs,
Graphical Models*

Probabilistic Circuits

Structure for tractable computation

We can exploit structure (low-rank, sparsity, etc) in data & models for ...

Compact representations,

Efficient parameterization & training,

Efficient computation of models, ...

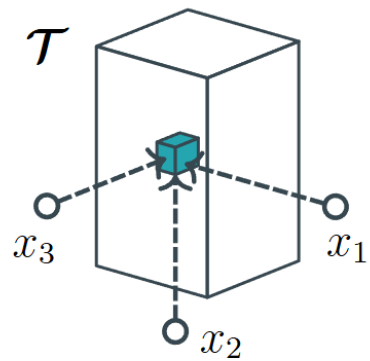
but also for...

*Tractable **reasoning** about models*

Outline

1. Probabilistic circuits – How structure enables tractable inference
2. A compositional view of inference – How tractable inference algorithms can be derived systematically
3. Tractable reasoning for distributional robustness

Imagine a *probability distribution*
 $P(x_1, x_2, x_3)$

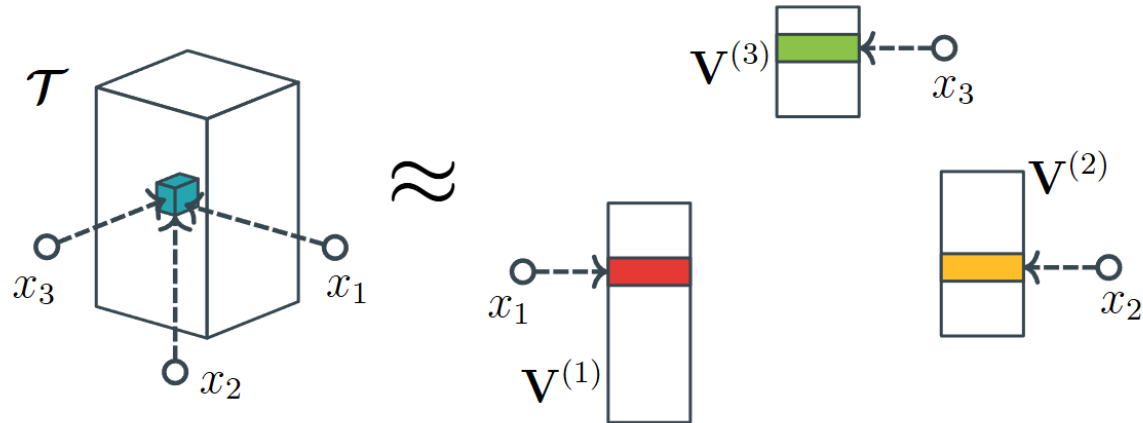


Tensor factorization

Imagine a *probability distribution*

$$P(x_1, x_2, x_3)$$

(e.g. CP)



$$t_{x_1 x_2 x_3} \approx \sum_{r=1}^R v_{x_1 r}^{(1)} v_{x_2 r}^{(2)} v_{x_3 r}^{(3)}$$

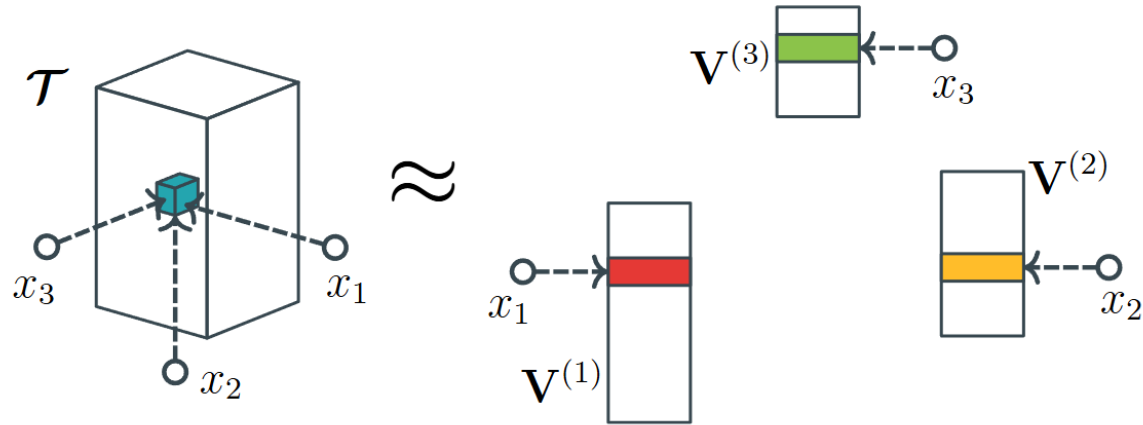
$$\mathbf{V}^{(1)} \in \mathbb{R}^{I_1 \times R}, \mathbf{V}^{(2)} \in \mathbb{R}^{I_2 \times R}, \mathbf{V}^{(3)} \in \mathbb{R}^{I_3 \times R}$$

Tensor factorization

(e.g. CP)

Imagine a *probability distribution*

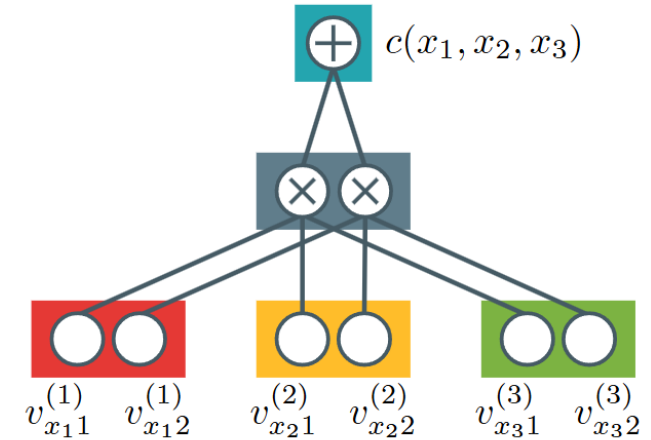
$$P(x_1, x_2, x_3)$$



$$t_{x_1 x_2 x_3} \approx \sum_{r=1}^R v_{x_1 r}^{(1)} v_{x_2 r}^{(2)} v_{x_3 r}^{(3)}$$

$$\mathbf{V}^{(1)} \in \mathbb{R}^{I_1 \times R}, \mathbf{V}^{(2)} \in \mathbb{R}^{I_2 \times R}, \mathbf{V}^{(3)} \in \mathbb{R}^{I_3 \times R}$$

as a *circuit*

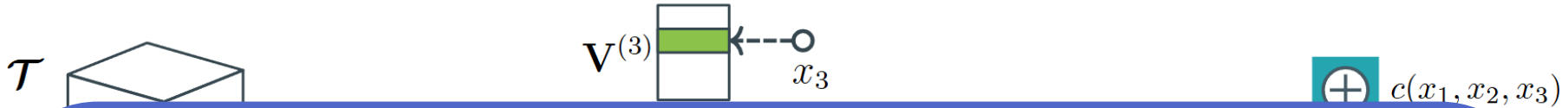


Tensor factorization

(e.g. CP)

as a **circuit**

Imagine a *probability distribution*
 $P(x_1, x_2, x_3)$



Check out the tutorial: <https://april-tools.github.io/aaai25-tf-pc-tutorial>

***from tensor factorizations
to circuits (and back)***



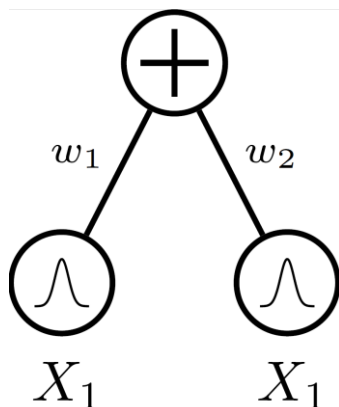
$v_{x_3 2}^{(3)}$

lorenzo loconte

antonio vergari

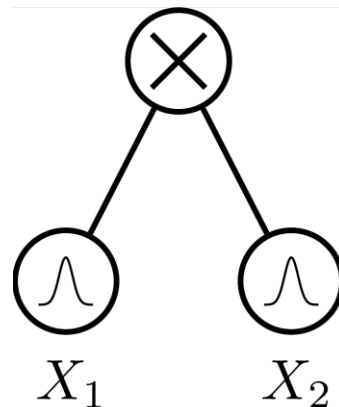
Probabilistic circuits (PCs)

computational graphs that recursively define probability mass/density functions



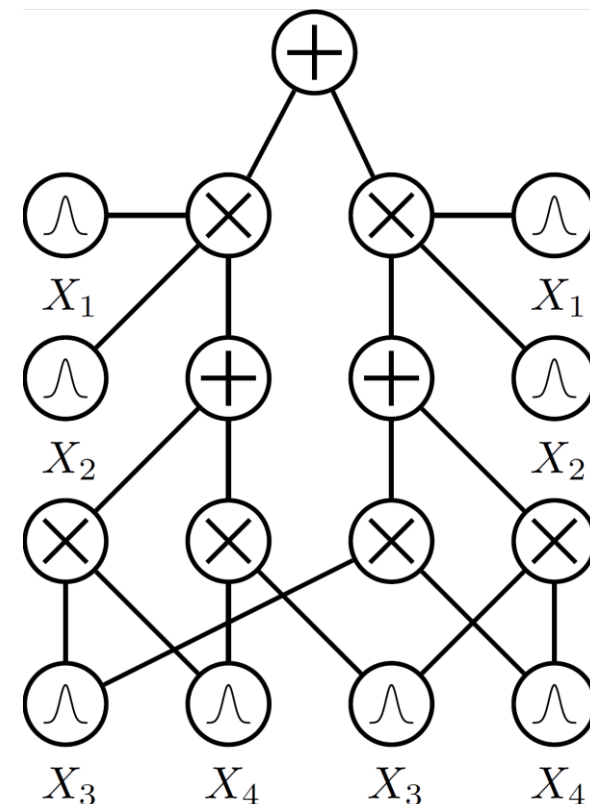
$$p(X_1) = w_1 p_1(X_1) + w_2 p_2(X_1)$$

$\Rightarrow$ mixtures



$$p(X_1, X_2) = p(X_1) \cdot p(X_2)$$

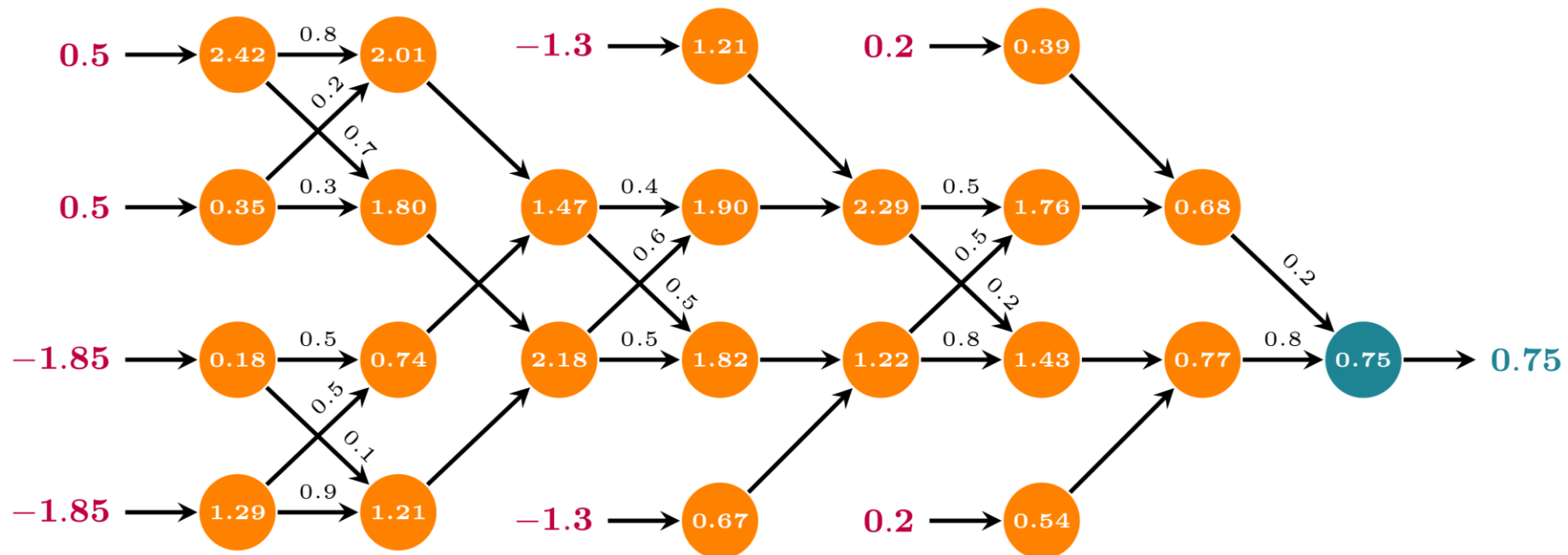
$\Rightarrow$ factorizations



Probabilistic circuits (PCs)

Compute $p(\mathbf{X})$ by simple feedforward evaluation

$$p(X_1 = -1.85, X_2 = 0.5, X_3 = -1.3, X_4 = 0.2)$$



Tractable inference

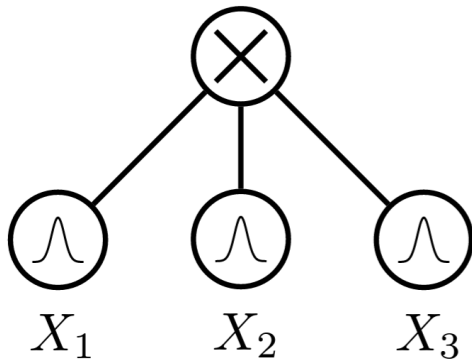
- **Structural properties** enable **tractable** inference of different queries

polytime in the size of circuit (# edges)
- E.g., *marginal & conditional inference* takes *linear time* in the circuit size for *smooth* and *decomposable* PCs

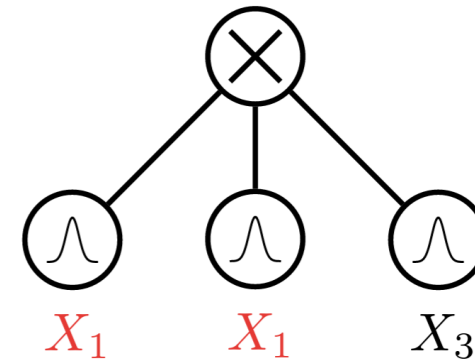
“What’s the diagnosis if some test results are missing?”

Decomposability

- A product node is *decomposable* if its children are defined on disjoint sets of variables (i.e. their scopes are mutually exclusive)



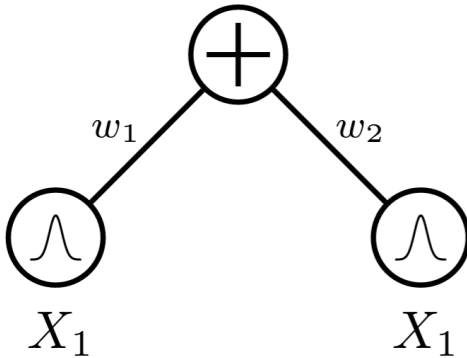
decomposable circuit



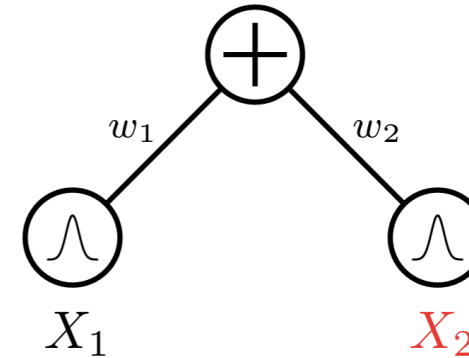
non-decomposable circuit

Smoothness

- A sum node is *smooth* if its children are defined on the same variables



smooth circuit

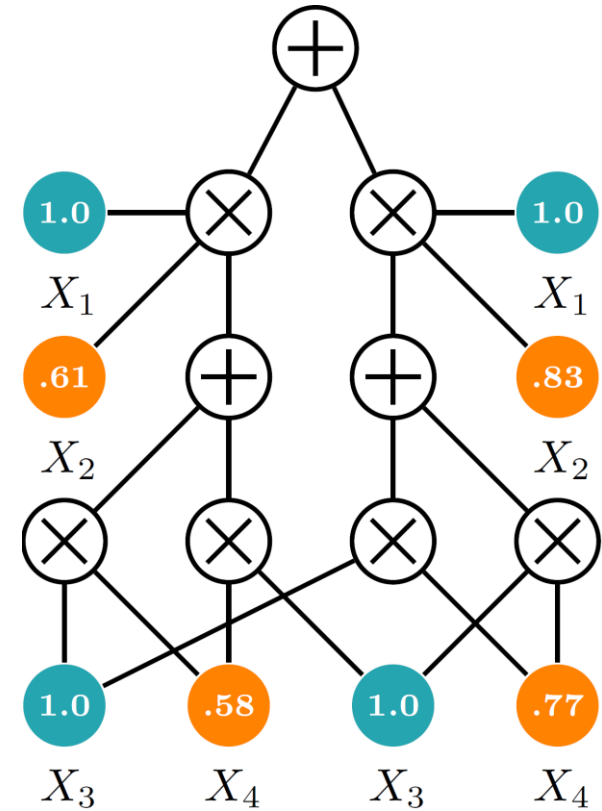


*non-smooth
circuit*

Tractable marginals

smooth and *decomposable* PC =>
compute any marginal by a single evaluation

$$\text{e.g. } p(x_2, x_4) = \int p(x_1, x_2, x_3, x_4) dX_1 X_3$$



Outline

1. Probabilistic circuits – How structure enables tractable computation
2. A compositional view of inference – How tractable inference algorithms can be derived systematically
3. Tractable reasoning for distributional robustness

Given a new inference task, e.g.,

$$-\sum_{\mathbf{x}} p(\mathbf{x}) \log p(\mathbf{x})$$

Entropy

$$\sum_{\mathbf{x}} p(\mathbf{x}) \log \frac{p(\mathbf{x})}{q(\mathbf{x})}$$

KL divergence

$$\max_{\mathbf{x}} \sum_{\mathbf{y} \sim e} p(\mathbf{x}, \mathbf{y})$$

Marginal MAP inference

$$\sum_{\mathbf{x}} p(\mathbf{z}) p(\mathbf{y} | \mathbf{x}, \mathbf{z})$$

Causal inference
(backdoor adjustment)

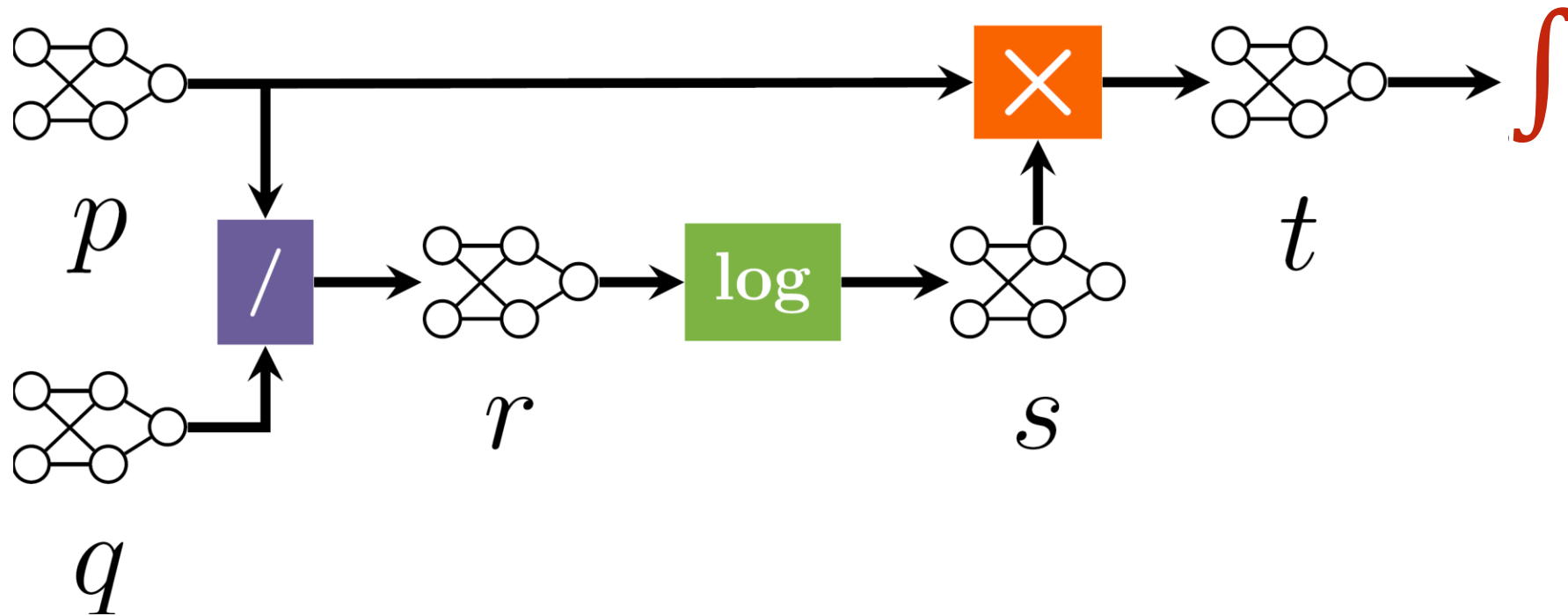
...and more...

Can we systematically derive efficient algorithms and circuit conditions for each query?

*Inference as **composition of operations!***

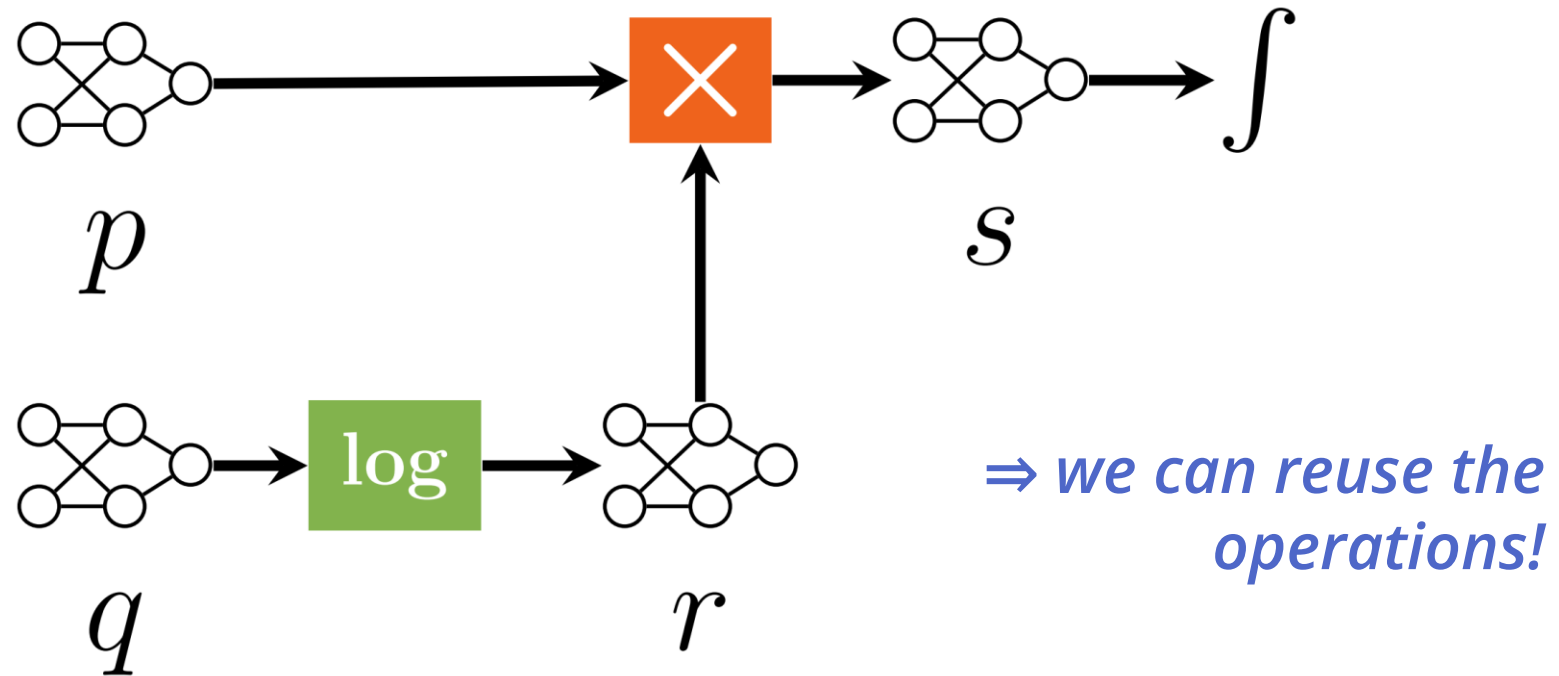
Queries as pipelines

$$\text{KL}\mathbb{D}(p \parallel q) = \int p(\mathbf{x}) \times \log((p(\mathbf{x})/q(\mathbf{x})))d\mathbf{X}$$



Queries as pipelines

$$H(p, q) = \int p(\mathbf{x}) \times \log(q(\mathbf{x})) d\mathbf{X}$$



Tractable circuit operations

Operation		Tractability		Hardness
		Input properties	Output properties	
SUM	$\theta_1 p + \theta_2 q$	(+Cmp)	(+SD)	NP-hard for Det output
PRODUCT	$p \cdot q$	Cmp (+Det, +SD)	Dec (+Det, +SD)	#P-hard w/o Cmp
POWER	$p^n, n \in \mathbb{N}$	SD (+Det)	SD (+Det)	#P-hard w/o SD
	$p^\alpha, \alpha \in \mathbb{R}$	Sm, Dec, Det (+SD)	Sm, Dec, Det (+SD)	#P-hard w/o Det
QUOTIENT	p/q	Cmp; q Det (+ p Det, +SD)	Dec (+Det, +SD)	#P-hard w/o Det
→ LOG	$\log(p)$	Sm, Dec, Det	Sm, Dec	#P-hard w/o Det
EXP	$\exp(p)$	linear	SD	#P-hard

*smooth,
decomposable,
deterministic*



*smooth,
decomposable*

**Which conditions
are needed...**

**...to efficiently perform
the operation...**

**...while maintaining some
circuit properties?**

Inference by tractable operations

systematically derive tractable inference algorithm of complex queries

	Query	Tract. Conditions	Hardness
CROSS ENTROPY	$-\int p(\mathbf{x}) \log q(\mathbf{x}) d\mathbf{X}$	Cmp, q Det	#P-hard w/o Det
SHANNON ENTROPY	$-\sum p(\mathbf{x}) \log p(\mathbf{x})$	Sm, Dec, Det	coNP-hard w/o Det
RÉNYI ENTROPY	$(1 - \alpha)^{-1} \log \int p^\alpha(\mathbf{x}) d\mathbf{X}, \alpha \in \mathbb{N}$	SD	#P-hard w/o SD
	$(1 - \alpha)^{-1} \log \int p^\alpha(\mathbf{x}) d\mathbf{X}, \alpha \in \mathbb{R}_+$	Sm, Dec, Det	#P-hard w/o Det
MUTUAL INFORMATION	$\int p(\mathbf{x}, \mathbf{y}) \log(p(\mathbf{x}, \mathbf{y}) / (p(\mathbf{x})p(\mathbf{y})))$	Sm, SD, Det*	coNP-hard w/o SD
KULLBACK-LEIBLER DIV.	$\int p(\mathbf{x}) \log(p(\mathbf{x}) / q(\mathbf{x})) d\mathbf{X}$	Cmp, Det	#P-hard w/o Det
RÉNYI'S ALPHA DIV.	$(1 - \alpha)^{-1} \log \int p^\alpha(\mathbf{x}) q^{1-\alpha}(\mathbf{x}) d\mathbf{X}, \alpha \in \mathbb{N}$	Cmp, q Det	#P-hard w/o Det
	$(1 - \alpha)^{-1} \log \int p^\alpha(\mathbf{x}) q^{1-\alpha}(\mathbf{x}) d\mathbf{X}, \alpha \in \mathbb{R}$	Cmp, Det	#P-hard w/o Det
ITAKURA-SAITO DIV.	$\int [p(\mathbf{x}) / q(\mathbf{x}) - \log(p(\mathbf{x}) / q(\mathbf{x})) - 1] d\mathbf{X}$	Cmp, Det	#P-hard w/o Det
CAUCHY-SCHWARZ DIV.	$-\log \frac{\int p(\mathbf{x}) q(\mathbf{x}) d\mathbf{X}}{\sqrt{\int p^2(\mathbf{x}) d\mathbf{X} \int q^2(\mathbf{x}) d\mathbf{X}}}$	Cmp	#P-hard w/o Cmp
SQUARED LOSS	$\int (p(\mathbf{x}) - q(\mathbf{x}))^2 d\mathbf{X}$	Cmp	#P-hard w/o Cmp

Compositional inference algorithms

```
function kld(p, q)
    r = quotient(p, q)
    s = log(r)
    t = product(p, s)
    return integrate(t)
end
```

```
function xent(p, q)
    r = log(q)
    s = product(p, r)
    return -integrate(s)
end
```

```
function ent(p)
    q = log(p)
    r = product(p, q)
    return -integrate(s)
end
```

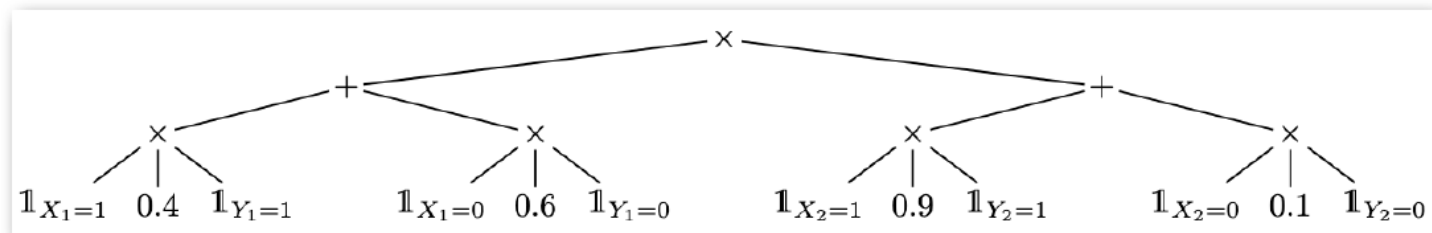
```
function alphadiv(p, q, alpha=1.5)
    r = real_pow(p, alpha)
    s = real_pow(q, 1.0-alpha)
    t = product(r, s)
    return log(integrate(t)) / (1.0-alpha)
end
```

*Efficient inference algorithms in a couple lines of Julia code!*¹

¹<https://github.com/UCLA-StarAI/circuit-ops-atlas>

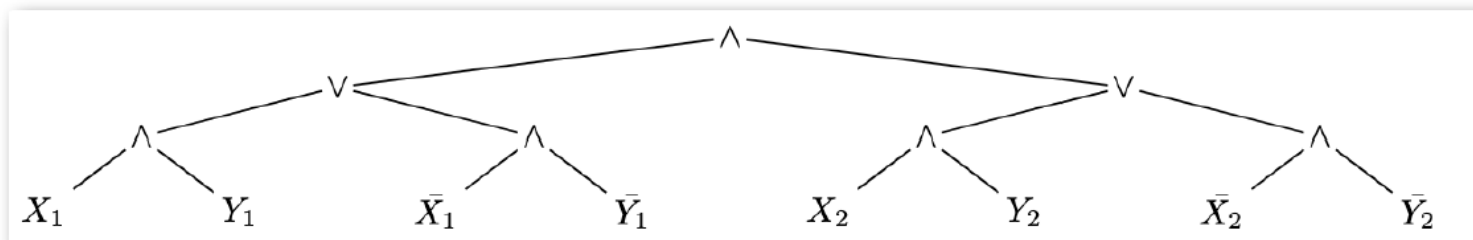
Algebraic circuits

Let $\mathcal{S} = (S, \oplus, \otimes, 0_{\mathcal{S}}, 1_{\mathcal{S}})$ be a commutative semiring. An algebraic circuit $\mathcal{C}$ over $\mathcal{S}$ recursively defines $f: \text{Assign}(\mathbf{X}) \rightarrow S$ as a computation graph of $\oplus$ and $\otimes$



Probabilistic Semiring
 $(\mathbb{R}_{\geq 0}, +, \times, 0, 1)$

probabilistic circuits: PC, SPN, PSDD, Cutset Networks, etc.



Boolean Semiring
 $(\{\perp, \top\}, \vee, \wedge, \perp, \top)$

logical circuits: BDD, DNNF, DFA, etc.

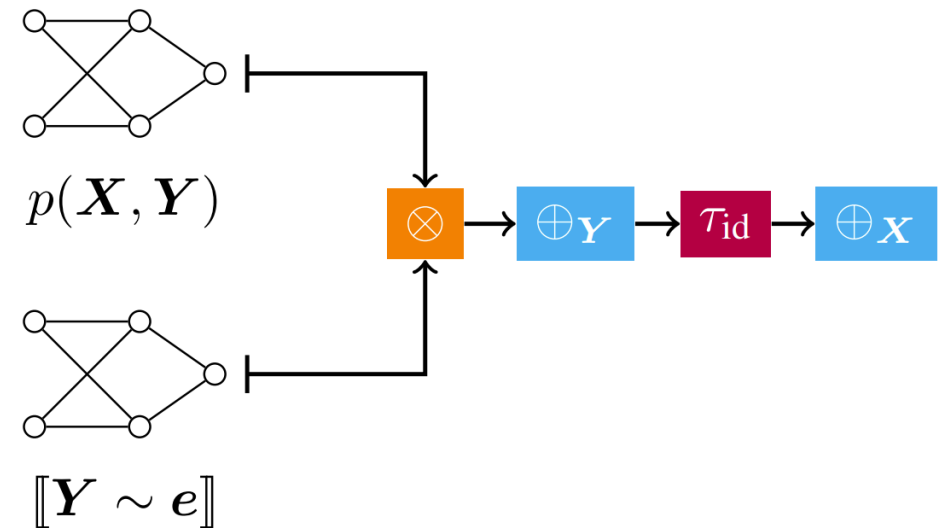
Algebraic queries...

$$\max_x \sum_{y \sim e} p(\mathbf{x}, \mathbf{y})$$

Marginal MAP
for a PC $p(\mathbf{X}, \mathbf{Y})$

- Anatomy of MMAP:
- $y \sim e$: **product** with logical circuit $\llbracket Y \sim e \rrbracket$
 - $\sum_y$: **summation** over semiring $(\mathbb{R}_{\geq 0}, +, \times, 0, 1)$
 - $\max_x$: **summation** over semiring $(\mathbb{R}_{\geq 0}, \max, \times, 0, 1)$

...as composition of operations...

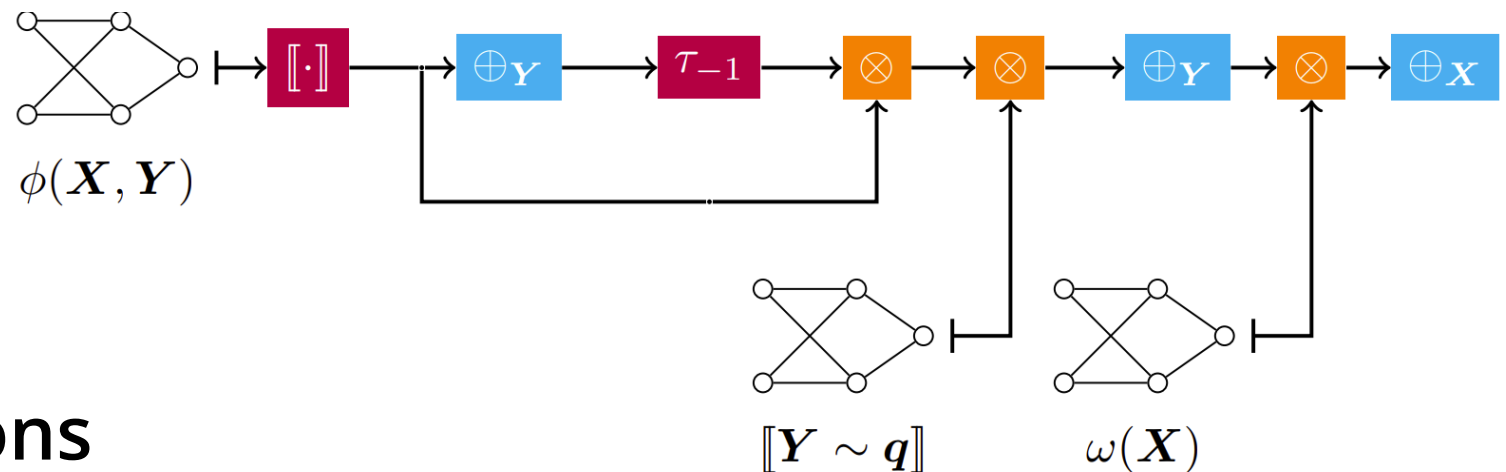


Algebraic queries...

$$\sum_x w(x) \sum_{y \sim q} \frac{\phi(x, y)}{\sum_{y'} \phi(x, y')}$$

Probabilistic answer set programming inference
(maximum entropy semantics)

...as composition of operations...



...and reuse operations

2AMC [Kiesel et al., 2022]

Class of inference problems
inc. MMAP, PLP inference, SDP, etc.

$$\bigoplus_x \left(\tau_{\mathcal{S}_Y \rightarrow \mathcal{S}_X} \left(\bigoplus_y [\phi(x, y)]_{\mathcal{B} \rightarrow \mathcal{S}_Y} \otimes \omega(y) \right) \otimes \omega'(x) \right)$$

Corollary 2AMC can be computed in $O(|\phi|)$ for deterministic, $\mathbf{X}$ -deterministic, and $\mathbf{X}$ -first circuits ϕ

Further, we can relax these conditions for specific 2AMC problems.

Causal inference

Adjustment formulae

$$p(\mathbf{y} | do(\mathbf{x})) = \sum_z p(z) p(\mathbf{y} | \mathbf{x}, z),$$

$$p(\mathbf{y} | do(\mathbf{x})) = \sum_z p(z | \mathbf{x}) \sum_{\mathbf{x}'} p(\mathbf{x}') p(\mathbf{y} | \mathbf{x}', z).$$

Corollary Backdoor adjustment can be computed in $O(|p|)$ for a $\mathbf{Z}$ -det and $(\mathbf{X} \cup \mathbf{Z})$ -det circuit p

Corollary Frontdoor adjustment can be computed in $O(|p|^2)$ for a $\mathbf{X}$ -det and $(\mathbf{X} \cup \mathbf{Z})$ -det circuit

Improves on previously best-known
custom algorithms!

Outline

1. Probabilistic circuits – How structure enables tractable computation
2. A compositional view of inference – How tractable inference algorithms can be derived systematically
3. Tractable reasoning for distributional robustness

Robust Maximum-Likelihood Learning

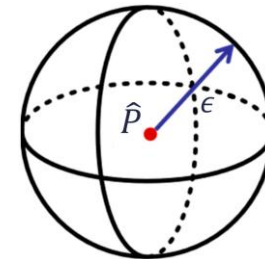
Maximum-likelihood estimation: given a set of samples $\{\mathbf{x}^i\}_{i=1,\dots,N}$ forming an empirical distribution $\hat{P}(\mathbf{X})$, learn the model P that maximizes (log-)likelihood

$$\max_P \mathbb{E}_{\hat{P}} [\log P(\mathbf{X})]$$

However, often *not reliable* under *data noise/perturbations* or *distribution shifts*

Robust Maximum-likelihood learning: maximize likelihood wrt the *worst-case distribution* Q within ϵ -Wasserstein ball of $\hat{P}$

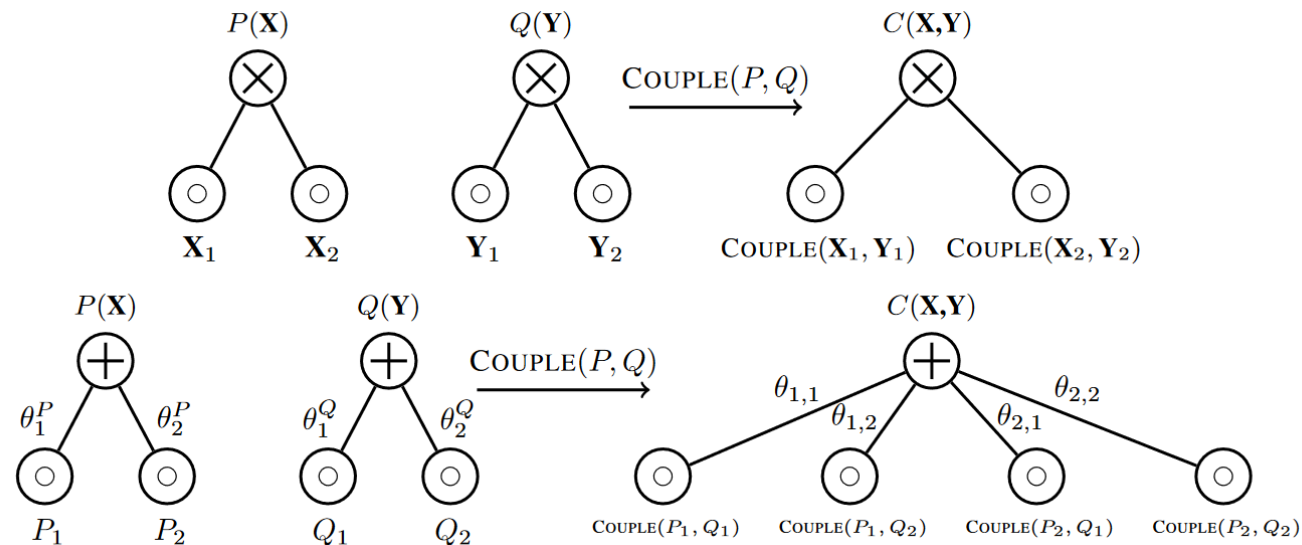
$$\begin{aligned} \max_P \min_Q \mathbb{E}_Q [\log P(\mathbf{X})] \\ \text{s.t. } \mathbf{W}(\hat{P}, Q) \leq \epsilon \end{aligned}$$



Circuit-Wasserstein Distance

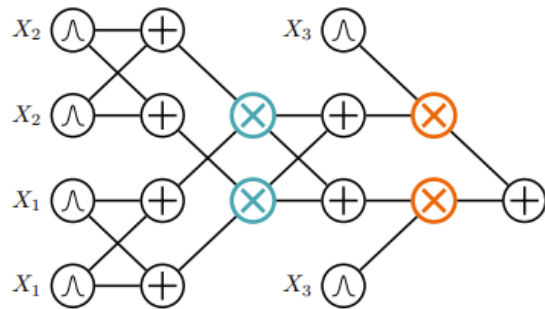
$$W_p^p(P, Q) = \inf_{\gamma \in \Gamma(P, Q)} \mathbb{E}_{\gamma(\mathbf{x}, \mathbf{y})} [\|\mathbf{x} - \mathbf{y}\|_p^p]$$

If P and Q are PCs with *compatible structure*, we can construct *coupling circuits*, and minimize the Wasserstein objective in *quadratic time* (solving a series of small LPs)



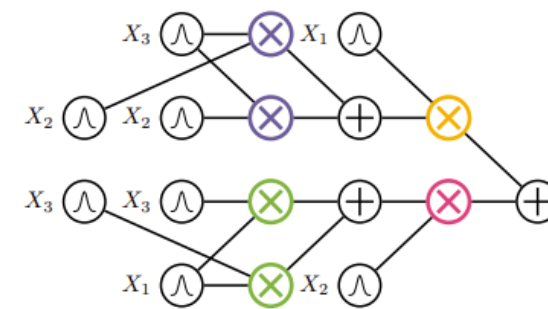
Compatibility

Two circuits are *compatible* if they have the same *hierarchical scope partitioning*



compatible circuits

support polytime products



non-compatible circuits

PeTeR: Post-training Robustification of PCs

Given a *pre-trained* PC $\hat{P}$, learn the PC parameters P_θ that guarantees high likelihood on the worst-case distribution Q_ϕ within ϵ -Circuit-Wasserstein ball of $\hat{P}$.

$$\begin{aligned} \max_{P_\theta} \min_{Q_\phi} \mathbb{E}_{Q_\phi} [\log P_\theta(\mathbf{X})] \\ \text{s.t. } \text{CW}(\hat{P}, Q_\phi) \leq \epsilon \end{aligned}$$

Note: updates a pre-trained PC. No training data needed!

Solve the equivalent Lagrangian problem by gradient descent-ascent

Tractable expectation, circuit-Wasserstein, & gradient computation thanks to *compatibility*!

		<i>adversarially-perturbed</i>						<i>randomly-perturbed</i>		
Dataset	h	$\mathcal{T}$			$\mathcal{T}_a$			$\mathcal{T}_r$		
		MLE-PC	RL-TPM	PeTeR	MLE-PC	RL-TPM	PeTeR	MLE-PC	RL-TPM	PeTeR
NLTCs	1		-9.36	<u>-6.79</u>	-11.26	<u>-10.91</u>	-9.31	<u>-8.41</u> ± 0.03	<u>-10.13</u> ± 0.03	-8.04 ± 0.02
	3	-6.09	-10.34	<u>-7.79</u>	-18.43	-11.14	<u>-11.56</u>	-11.46 ± 0.05	<u>-11.22</u> ± 0.02	-9.79 ± 0.02
	5		-10.86	<u>-9.90</u>	-23.02	-10.61	<u>-11.75</u>	-13.24 ± 0.09	<u>-11.44</u> ± 0.02	-10.76 ± 0.01
MSNBC	1		-9.14	<u>-7.05</u>	-12.01	<u>-10.33</u>	-8.45	<u>-8.31</u> ± 0.01	-9.88 ± 0.01	-8.13 ± 0.00
	3	-6.43	-8.47	-8.69	-19.88	<u>-11.52</u>	-10.96	-11.29 ± 0.01	<u>-10.17</u> ± 0.01	-9.99 ± 0.00
	5		<u>-9.92</u>	-10.76	-25.39	-11.85	<u>-11.95</u>	-13.52 ± 0.02	-10.97 ± 0.00	<u>-11.29</u> ± 0.00
Plants	1		-21.01	<u>-14.98</u>	-23.39	<u>-22.59</u>	-19.62	<u>-18.78</u> ± 0.03	-23.65 ± 0.04	-18.03 ± 0.02
	3	-14.46	-22.30	<u>-16.16</u>	-38.81	-23.22	<u>-26.69</u>	<u>-26.26</u> ± 0.09	-28.76 ± 0.05	-23.31 ± 0.05
	5		-23.64	<u>-17.23</u>	-52.40	-27.54	<u>-33.30</u>	<u>-32.36</u> ± 0.11	-32.63 ± 0.04	-27.49 ± 0.07
Netflix	1		-58.82	<u>-57.66</u>	-61.13	<u>-60.67</u>	-60.61	-58.27 ± 0.02	-59.37 ± 0.01	<u>-58.32</u> ± 0.02
	3	-57.59	-59.27	<u>-58.09</u>	-67.12	-63.08	<u>-65.06</u>	-59.56 ± 0.03	-60.55 ± 0.02	<u>-59.69</u> ± 0.03
	5		-60.10	<u>-59.21</u>	-72.17	-65.32	<u>-66.48</u>	-60.78 ± 0.06	-61.88 ± 0.04	<u>-61.28</u> ± 0.04
DNA	1		-82.69	<u>-81.05</u>	-87.68	<u>-86.67</u>	-85.12	<u>-83.44</u> ± 0.09	-84.65 ± 0.05	-83.42 ± 0.05
	3	-79.97	-84.31	<u>-82.62</u>	-102.88	<u>-93.38</u>	-90.70	-90.33 ± 0.15	<u>-88.80</u> ± 0.05	-88.07 ± 0.07
	5		-85.85	<u>-84.23</u>	-117.82	<u>-98.24</u>	-96.08	-97.00 ± 0.23	<u>-92.00</u> ± 0.08	-91.18 ± 0.10
Movie	1		-55.02	<u>-54.07</u>	-61.93	<u>-61.21</u>	-58.59	-58.53 ± 0.09	-59.77 ± 0.07	<u>-58.99</u> ± 0.08
	3	-52.81	-57.06	<u>-53.38</u>	-79.51	<u>-69.61</u>	-67.62	-69.77 ± 0.20	-68.82 ± 0.12	<u>-68.91</u> ± 0.17
	5		-59.34	<u>-56.09</u>	-96.75	-78.24	<u>-79.51</u>	-80.66 ± 0.16	<u>-77.02</u> ± 0.11	-75.07 ± 0.16
BBC	1		-250.72	<u>-250.30</u>	-258.72	-254.99	<u>-256.75</u>	-253.04 ± 0.12	-253.43 ± 0.10	<u>-253.10</u> ± 0.11
	3	-250.20	-251.45	<u>-250.55</u>	-275.07	-262.65	<u>-266.11</u>	-258.74 ± 0.15	-259.25 ± 0.12	<u>-258.82</u> ± 0.13
	5		-252.52	<u>-250.58</u>	-290.24	-269.24	<u>-274.58</u>	<u>-264.30</u> ± 0.30	-265.01 ± 0.26	-264.16 ± 0.29

PeTeR improves the performance of pre-trained PCs on adversarially- and randomly-perturbed datasets, and frequently outperforms baseline robust learning method

Takeaways

1. Probabilistic circuits – How structure enables tractable computation
2. A compositional view of inference – How tractable inference algorithms can be derived systematically
3. Tractable reasoning for distributional robustness

What was not in this talk:

- How to build/learn the circuits
- Expressive efficiency vs tractability trade-off (for both exact & approximate)



Thank you!

Thanks to my students and collaborators!