

Link Prediction Under Non-targeted Attacks: Do Soft Labels Always Help?

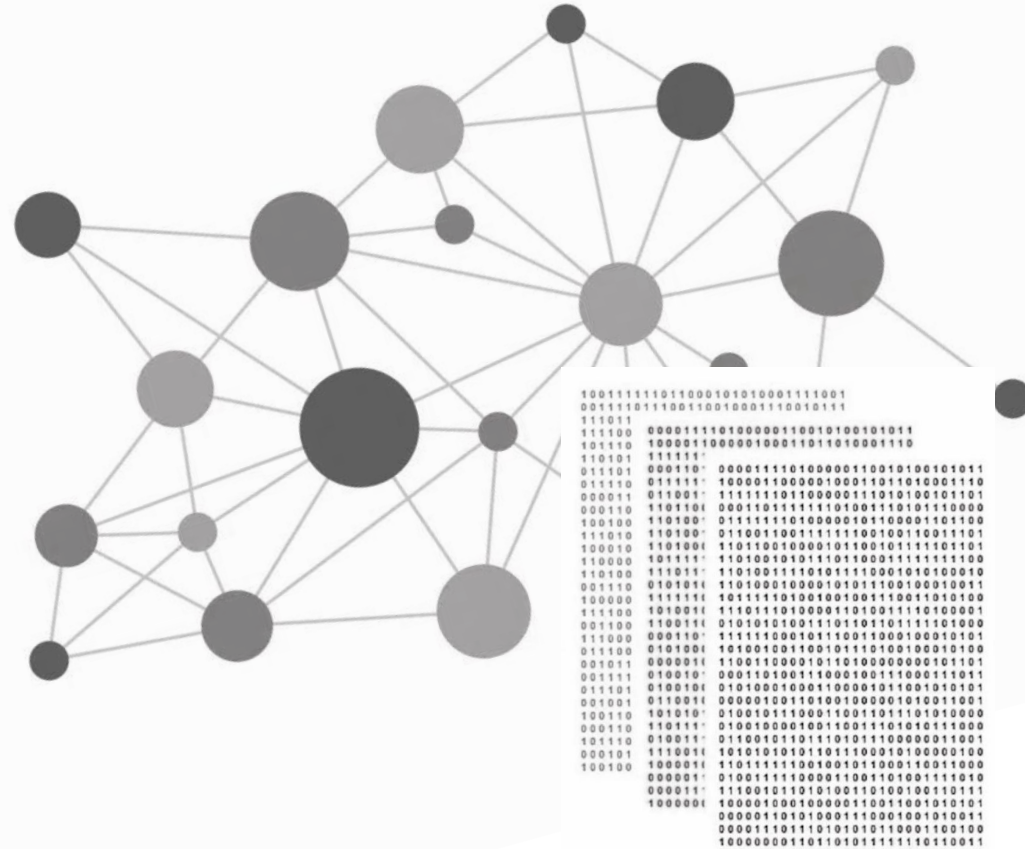
Adel Memariani, Michael Röder, Arnab Sharma, Caglar Demir, Axel-Cyrille Ngonga Ngomo

Data Science Group (DICE), Heinz Nixdorf Institute, Paderborn University

03.11.2025

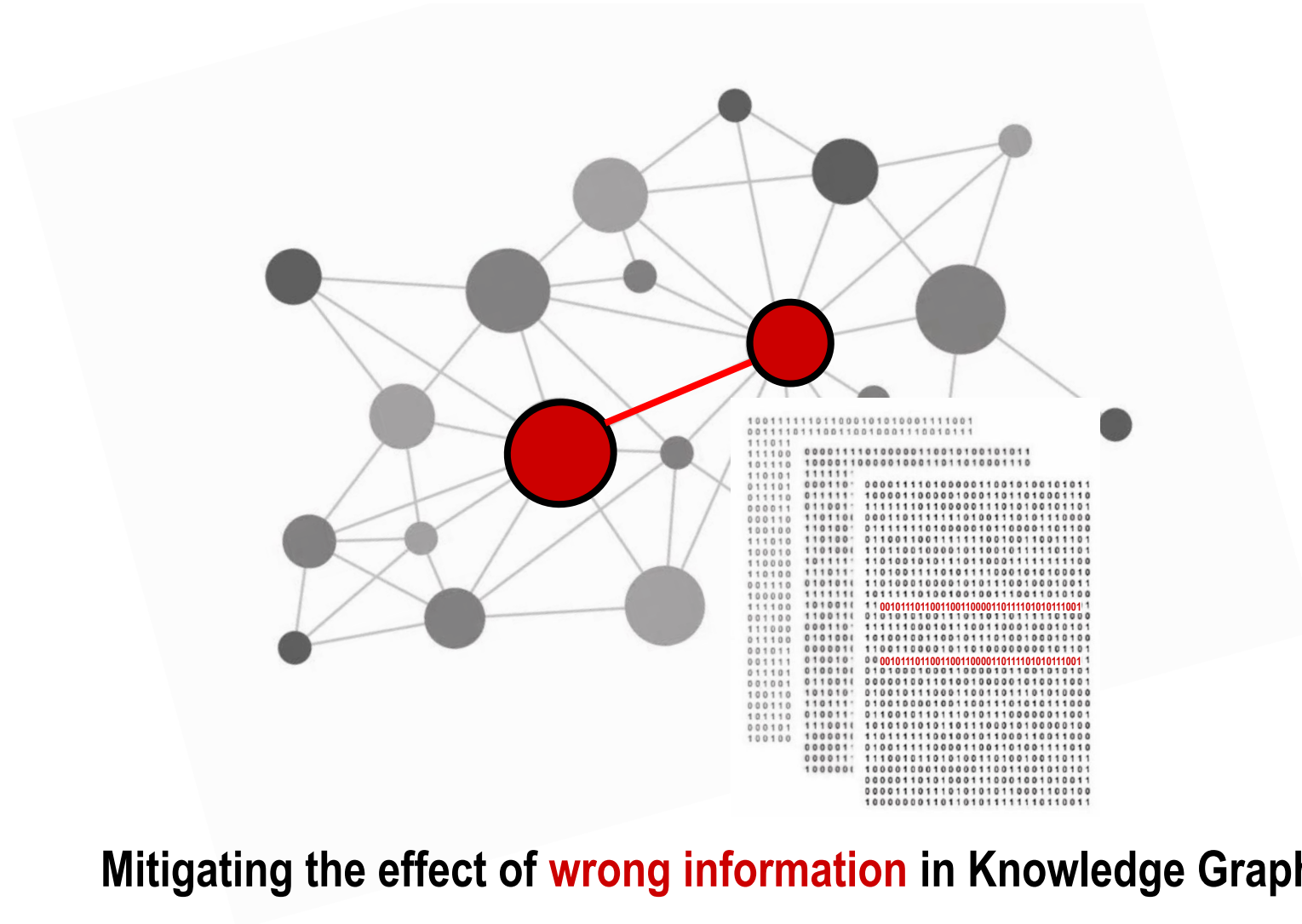


Motivation



Knowledge graphs and their embeddings

Motivation



Mitigating the effect of **wrong information** in Knowledge Graphs

Motivation

- **Real-world Knowledge Graphs (KGs) may contain wrong information due to:**
 - Human curation mistakes
 - Automated construction
 - Data poisoning attacks

Motivation

- **Real-world Knowledge Graphs (KGs) may contain wrong information due to:**
 - Human curation mistakes
 - Automated construction
 - Data poisoning attacks
- **Fact checking is costly because:**
 - Rely on trusted reference knowledge
 - Evidence retrieval in large KGs is computationally expensive

Motivation

- **Real-world Knowledge Graphs (KGs) may contain wrong information due to:**
 - Human curation mistakes
 - Automated construction
 - Data poisoning attacks
- **Fact checking is costly because:**
 - Rely on trusted reference knowledge
 - Evidence retrieval in large KGs is computationally expensive
- **Unreliable facts can distort knowledge graph embeddings**
 - Compromise the trustworthiness of downstream!

How do data-poisoning attacks work?

- **Targeted:** Add or remove triples to change specific inferences
 - Facts and the AI model are assumed to be known in advance!
 - Less practical

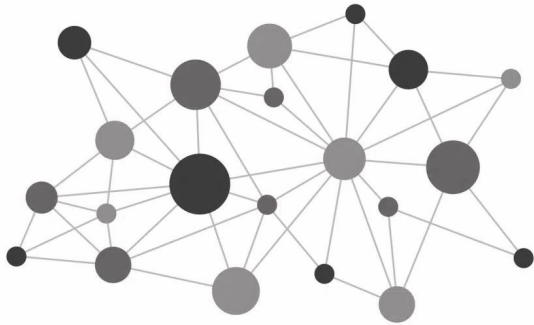


How do data-poisoning attacks work?

- **Targeted:** Add or remove triples to change specific inferences
 - Facts and the AI model are assumed to be known in advance!
 - Less practical
- **Non-targeted:** Add noise/remove facts to lower overall performance
 - Our focus in this research

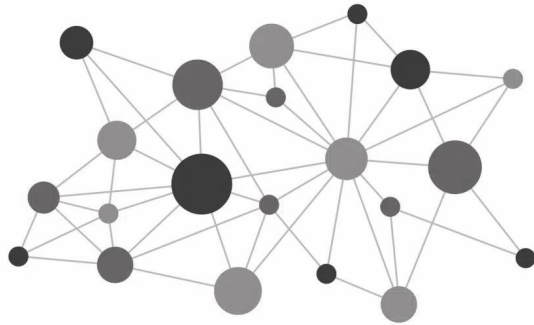


How do data-poisoning attacks work?



Knowledge Graph

How do data-poisoning attacks work?



Knowledge Graph

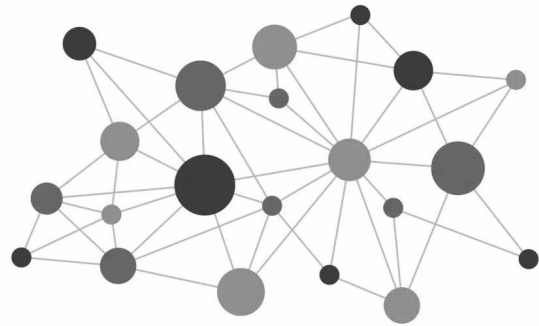


```

1001111101100010101000111001
001111011100110010001110010111
111011
111100 000011110100000110010100101011
101110 100001100000100011011010001110
110101 111111
011101 000110 00011110100000110010100101011
011110 011111 100001100000100011011010001110
000011 011001 11111101100000111010100101101
000110 110110 001101111110100111010110000
100100 110100 01111110100000101100001101100
111010 110100 01100110011111100100110011101
100010 110100 11011001000010110010111101101
110000 101111 11010010101110110001111111100
110100 111011 110100111101011110001010100010
001110 010100 11010001000101011100100010001
100000 111111 101111101001001001110011010100
111100 101001 111011101000011010011110100001
001100 110011 010101010011101101101111101000
111000 000110 111111000101110011000100010101
011100 010100 101001001100101110100100010100
001011 000011 110011000010110100000000101101
001111 010010 000110100111000100111000111011
011101 010010 01010001000110001011001010101
001001 011001 0000010011010011000001010011001
100110 101010 010010111000110011011101010000
000110 110111 0100100001001100111010111000
101110 010011 01100101101110101110000011001
111001 101010101011011100010100000100
100001 1101111110010010110001100110011000
000001 010011111000111001101001111010
000011 111001011010100110100100110111
100001 100010001000001100110001010101
00001101010001110001001010001
000011101110101010110001100100
100000011011010111111110110011
    
```

Embedding Vectors

How do data-poisoning attacks work?



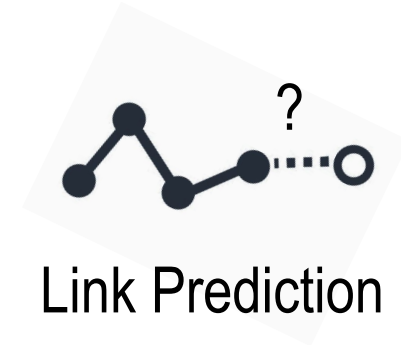
Knowledge Graph



```

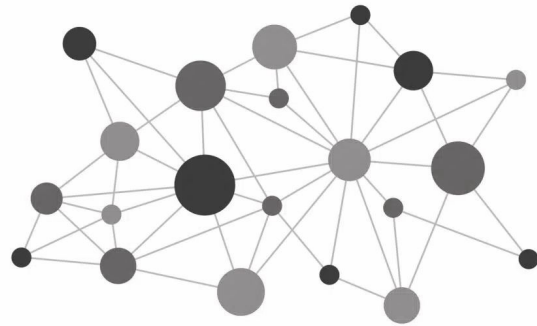
1001111101100010101000111001
001111011100110010001110010111
111011
111100 0000111010000011001010010101
101110 100001100000100011011010001110
110101 111111
01101 000110 0001111010000011001010010101
011110 011111 100001100000100011011010001110
000011 011001 11111101100000111010100101101
000110 110101 00010111111010011101110000
100100 110100 01111110100000101100001101100
111010 110100 01100110011111100100110011101
100010 110100 11011001000010110010111101101
110000 101111 11010010101110110001111111100
110100 111011 110100111101011110001010100010
001110 010100 11010001000101011100100010011
100000 111111 10111101001001001110011010100
111100 101001 11101101000011010011110100001
001100 110011 010101010011101101101111101000
111000 000110 1111110001011100110001000101
011100 010100 101001001100101110100010100
001011 000011 110011000010110100000000101101
001111 010010 000110100111000100111000111011
011101 010010 0101000100011000010110010101
001001 011001 0000010011010001000001010011001
100110 101010 01001011000110011011101010000
000110 110111 0100100001001100111010111000
101110 010011 01100101101110101110000011001
111001 101010101011011100010100000100
100001 110111110010010110001100110011000
000001 010011111000011001101001111010
000011 111001011010100110100100110111
100000 100010000001100110001010101
00001101010001110001001001001
000011101110101010110001100100
10000001101101011111110110011
    
```

Embedding Vectors



Link Prediction

How do data-poisoning attacks work?



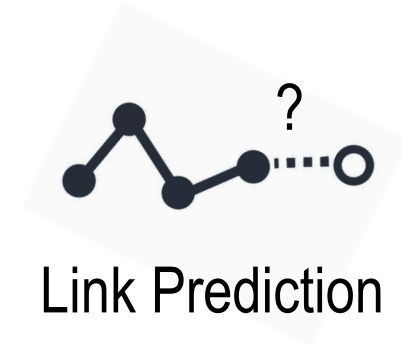
Knowledge Graph



```

10011111011000101010001111001
001111011100110010001110010111
111011
111100 0000111010000011001010010101
101110 1000011000001000110110001110
110101 111111
01101 000110 00011110100000110010010101
011110 011111 100001100000100011011010001110
000011 011001 11111101100000111010010010101
000110 110110 00110111111010011101110000
100100 110100 0111111010000010110000110100
111010 110100 01100110011111100100110011101
100010 110100 1101100100001011001011110101
110000 101111 1101001010111011000111111100
110100 111011 110100111101011110001010100010
001110 010100 11010001000101011100100010011
100000 111111 101111101001001001110011010100
111100 101001 111011101000011010011110100001
001100 110011 0101010100111011011011111010000
111000 000110 111111000101110011000100010101
011100 010100 101001001100101110100100010100
001011 000001 110011000010110100000000101101
001111 010010 000110100111000100111000111011
011101 010010 010100010001100001011001010101
001001 011001 0000010011010011000001010011001
100110 101010 010010111000110011011101010000
000110 110111 010010000100110011101010111000
101110 010011 01100101101110101110000011001
111001 101010101011011100010100000100
100001 1101111110010010110001100110011000
000001 010011111000110011010101111010
000011 11100101010100110100100110111
1000010001000001100110001010101
00001101010001110001001010011
000011101110101010110001100100
10000001101101011111110110011
    
```

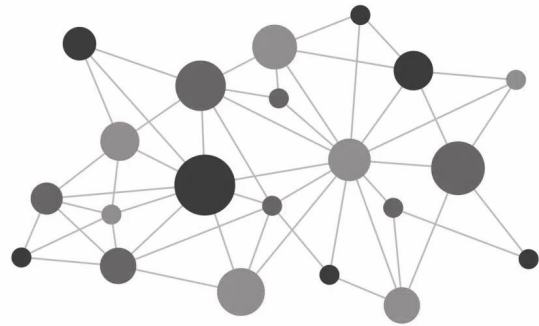
Embedding Vectors



Link Prediction

(Jupiter → surfaceType → ?)

How do data-poisoning attacks work?



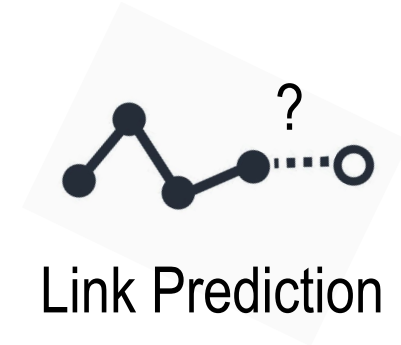
Knowledge Graph



```

1001111101100010101000111001
001111011100110010001110010111
111011
111100 0000111010000011001010010101
101110 1000011000001000110110001110
110101 111111
011101 000110 0001111010000011001010010101
011110 011111 100001100000100011011010001110
000011 011001 11111101100000111010100101101
000110 110101 00011011111010011101110000
100100 110100 01111110100000101100001101100
111010 110100 01100110011111100100110011101
100010 110100 11011001000010110010111101101
110000 101111 11010010101110110001111111100
110100 111011 110100111101011110001010100010
001110 010101 110100010000101011100100010011
100000 111111 101111101001001001110011010100
111100 101001 111011101000011010011110100001
001100 110011 0101010100111011011011111010000
111000 000110 111111000101110011000100010101
011100 010100 101001001100101110100100010100
001011 000011 110011000010110100000000101101
001111 010010 000110100111000100111000111011
011101 010010 010100010001100001011001010101
001001 011001 000010011010001000001010011001
100110 101010 01001011000110011011101010000
000110 110111 010010000100110011101010111000
101110 010011 01100101101110101110000011001
111001 101010101011011100010100000100
100001 110111110010010110001100110011000
000001 010011111000110011010101111010
000011 11100101010100110100100110111
10000100100000110011001010101
00001101010100011100010010001
000011101110101010110001100100
10000001101101011111110110011
    
```

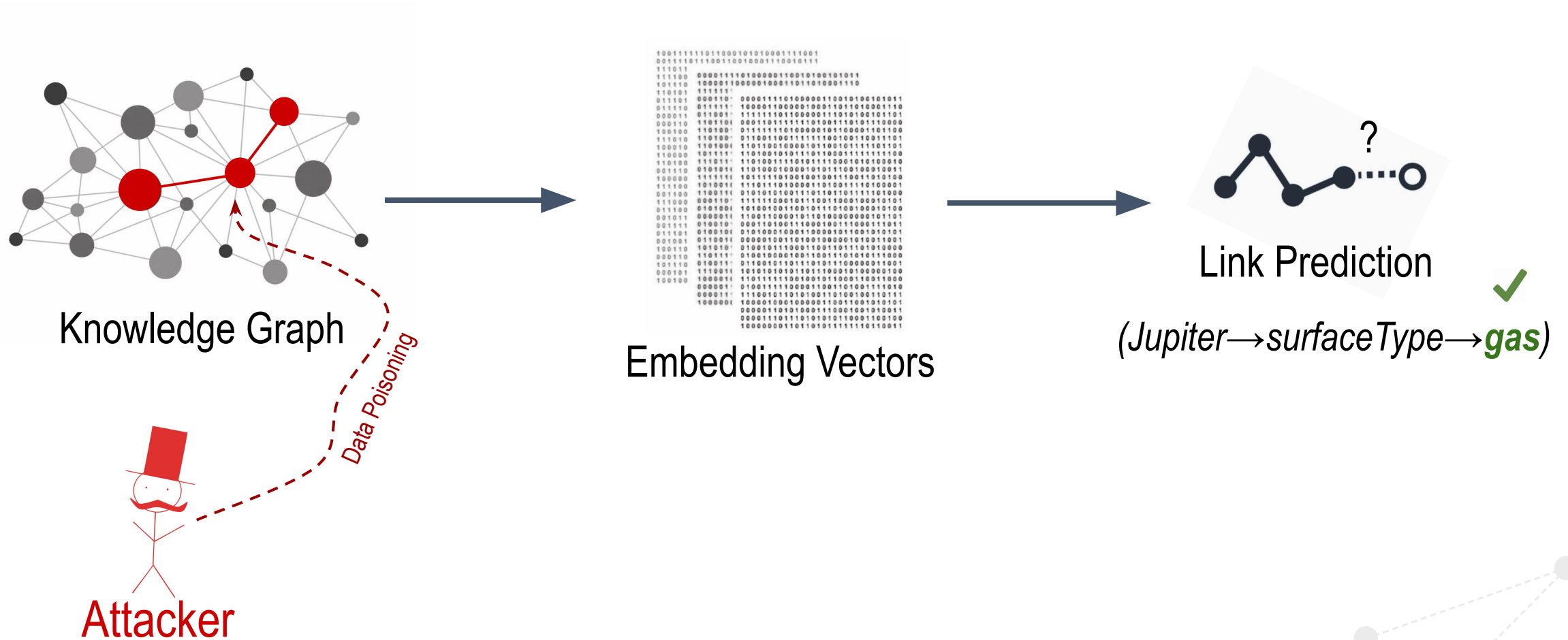
Embedding Vectors



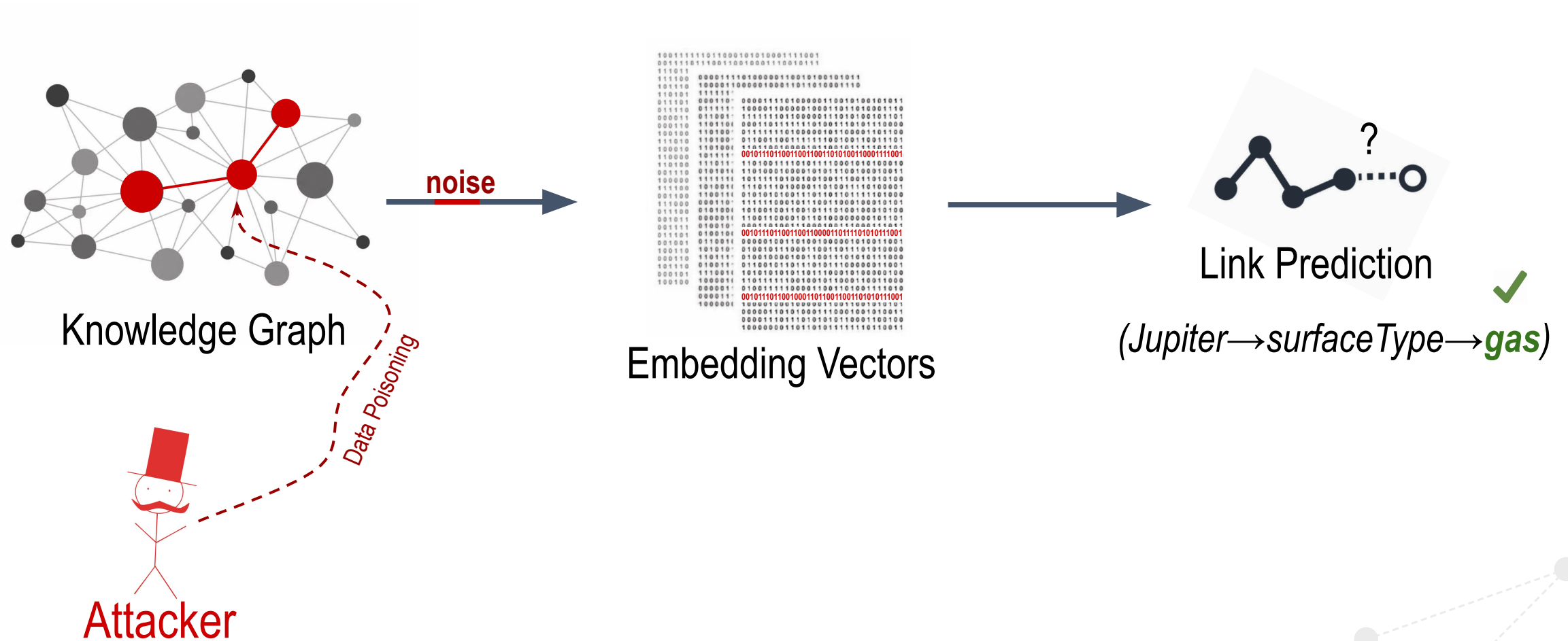
Link Prediction

(Jupiter → surfaceType → **gas**) ✓

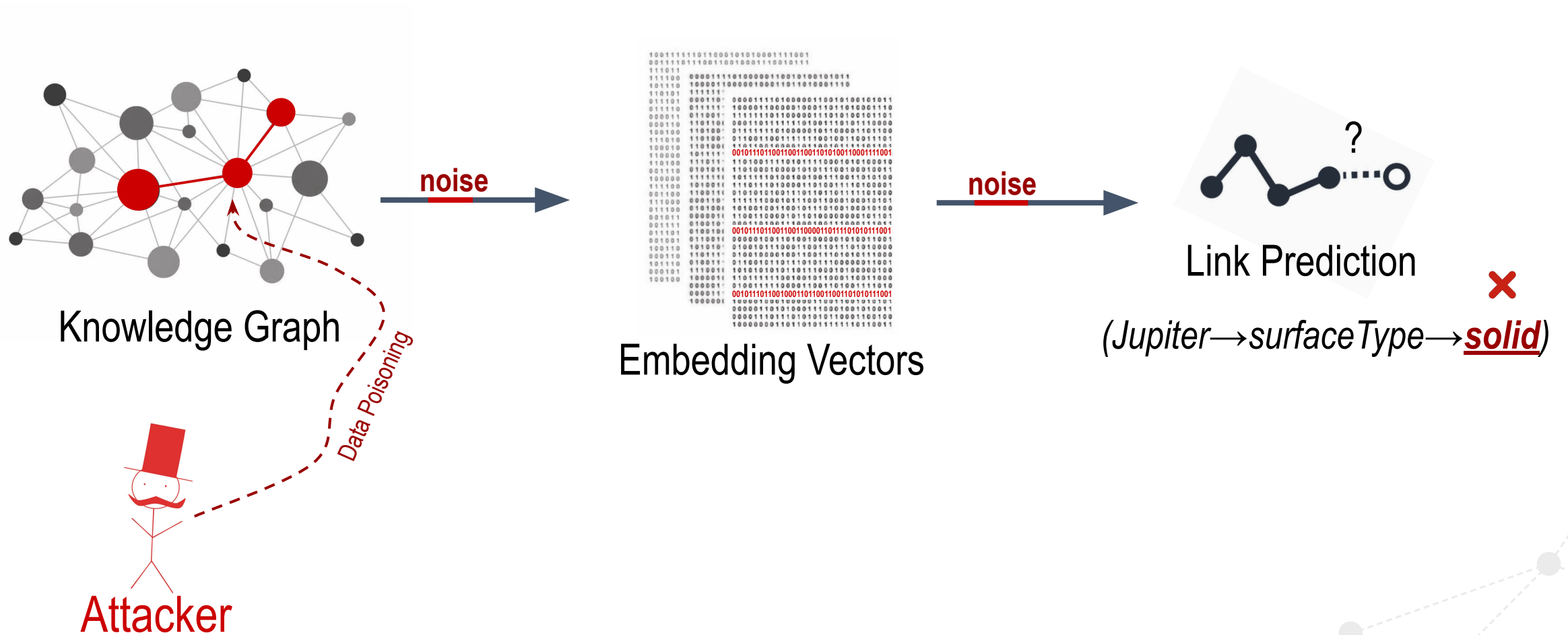
How do data-poisoning attacks work?

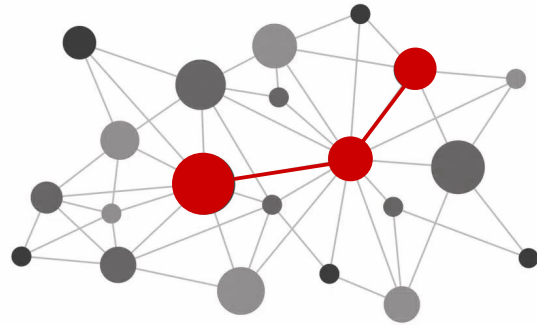


How do data-poisoning attacks work?



How do data-poisoning attacks work?





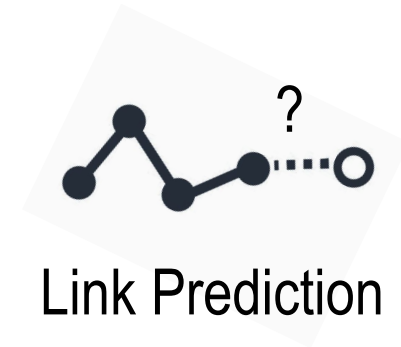
Knowledge Graph

noise

```
10011111011000101010001111001
001111011100110010001110010111
111011
111100 0000111101000001100101001011
101110 1000011000001000110110001110
110101 111111
011101 000110 00011110100001100101001011
011110 011111 100001100000100011011010001110
000011 011001 1111110110000011100100101101
000110 110101 00010111111010011101110000
100100 110100 0111111010000010110001101100
111010 110100 01100110011111100100110011101
100010 110100 11011001000010110010111101101
110000 101111 101111
110100 111011 1101001111010111100010100010
001110 010101 110100010000101011100100010011
100000 111111 101111010010011001110011010100
111100 101001 111011101000011010011110100001
001100 110011 010101010011101100101111010000
111000 000110 1111110001011100110001000101
011100 010100 101001001100101110100100010100
001011 000011 110011000010110100000000101101
001111 010010 0001010001110000001110001110014
011101 010010 0010111001100001101110101011001
001001 011001 0000100110100100000100011001
100110 101010 01001011000110011011101010000
000110 110111 0100100001001100111010111000
101110 010011 01100101101110101110000011001
111001 101010101011011100010100000100
100011 110111110100010110001100110011000
000011 010011110000110011100111101
000011 000011 000011
001011 0010110000011001100110101011001
00001101010100011100010010001
00011101101010101100011001100100
```

Embedding Vectors

noise



Link Prediction

(Jupiter → surfaceType → solid)

×

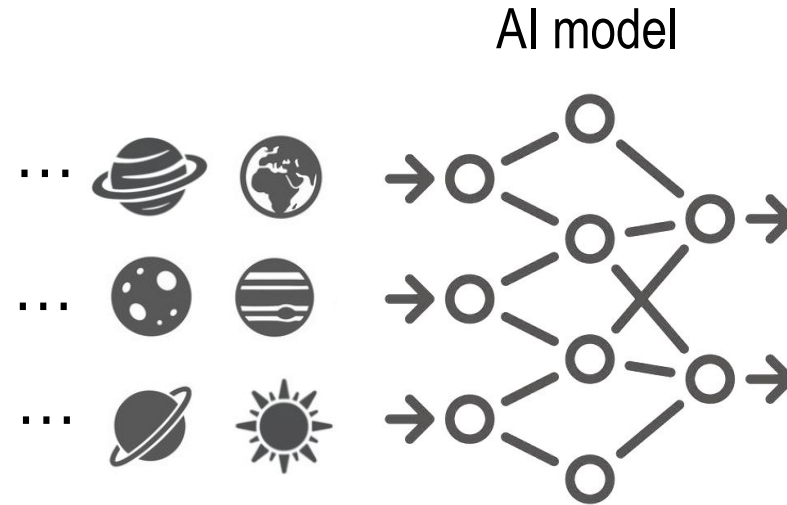
Goal



Robust Knowledge Graph Embedding
Reliable even with *noisy* data

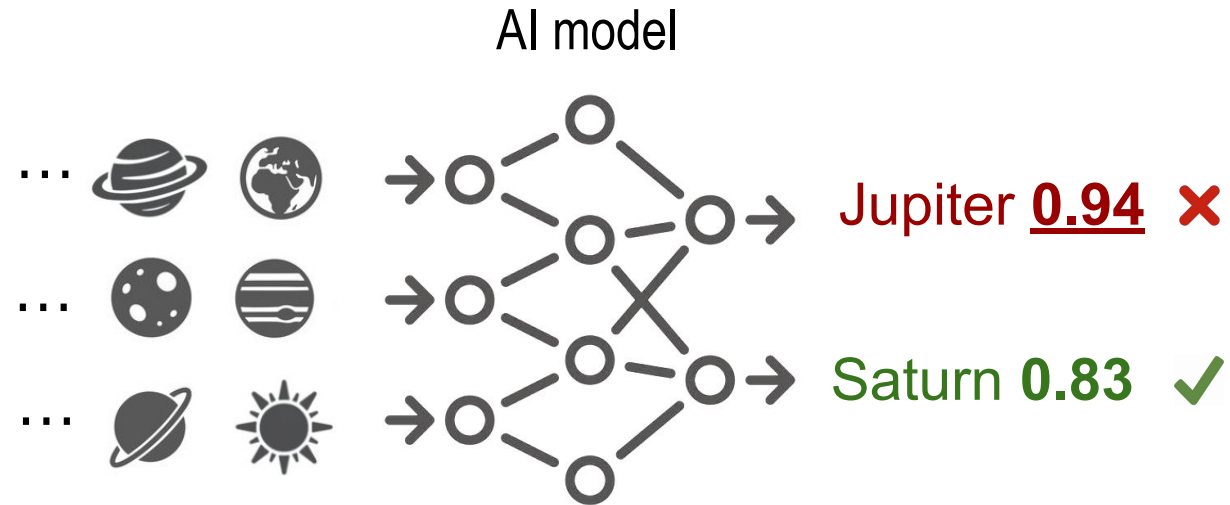
Model confidence

Model confidence



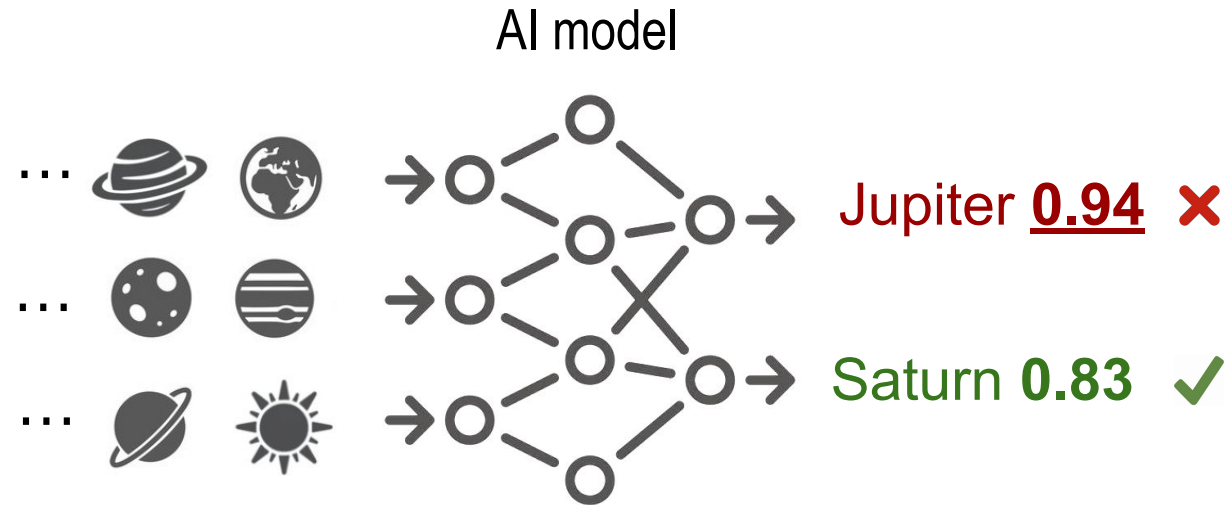
(Titan, orbits, ?)

Model confidence



(Titan, orbits, **Jupiter**)

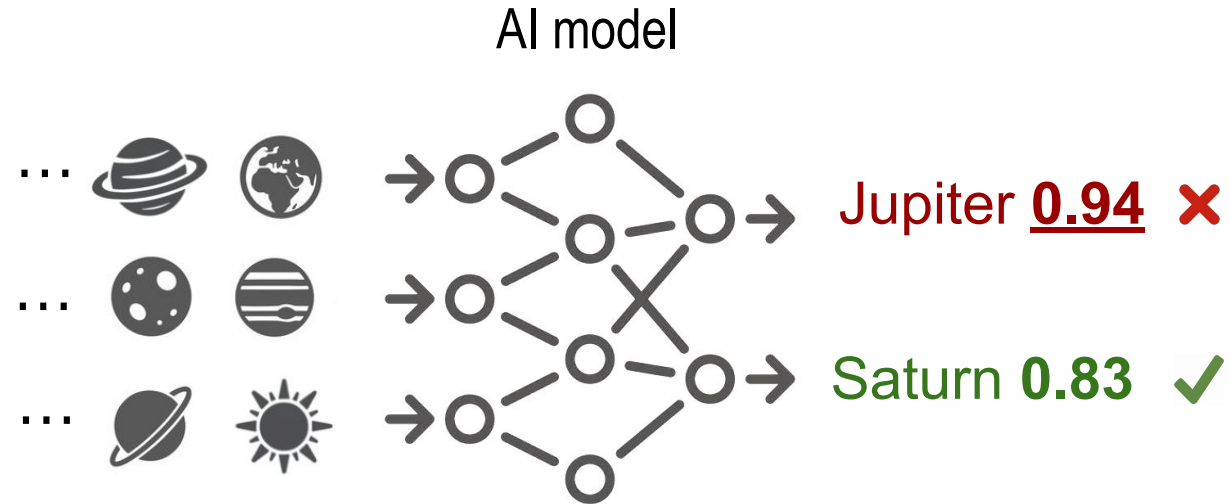
Model confidence



(Titan, orbits, **Jupiter**)

A highly confident wrong prediction 🤔

Model confidence



(Titan, orbits, **Jupiter**)

A highly confident wrong prediction 🤔

➡ Add a bit of uncertainty during training

The Challenge: Noisy training triples

The Challenge: Noisy training triples

Our goal: Avoid over-confidence

The Challenge: Noisy training triples

Our goal: Avoid over-confidence

Our approach: Use soft labels during training

Soft labels

Binary Class Entropy (BCE) Loss

$$\mathcal{L}(\hat{y}, y) = -\frac{1}{|\mathcal{E}|} \sum_{i=1}^{|\mathcal{E}|} [y_i \log(\hat{y}_i) + (1 - y_i) \log(1 - \hat{y}_i)]$$

Binary Class Entropy (BCE) Loss

$$\mathcal{L}(\hat{y}, y) = -\frac{1}{|\mathcal{E}|} \sum_{i=1}^{|\mathcal{E}|} [y_i \log(\hat{y}_i) + (1 - y_i) \log(1 - \hat{y}_i)]$$

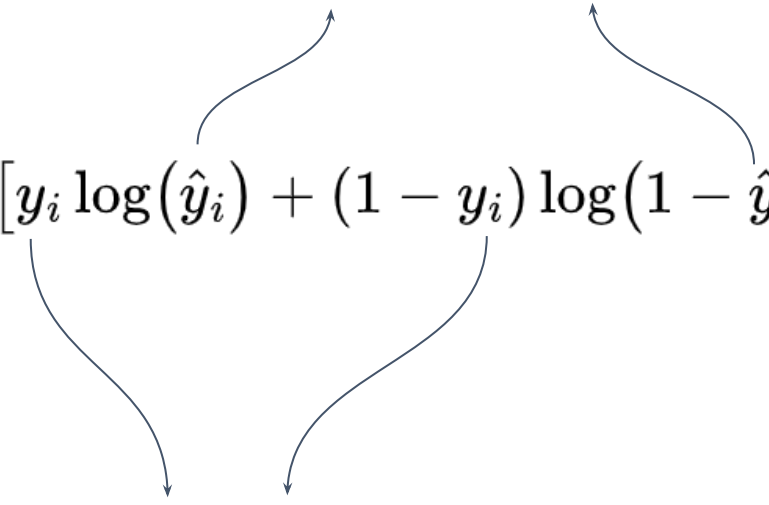
Predicted

Binary Class Entropy (BCE) Loss

$$\mathcal{L}(\hat{y}, y) = -\frac{1}{|\mathcal{E}|} \sum_{i=1}^{|\mathcal{E}|} [y_i \log(\hat{y}_i) + (1 - y_i) \log(1 - \hat{y}_i)]$$

Predicted

Label



Binary: 1.0, 0.0, 1.0, 0.0 ...

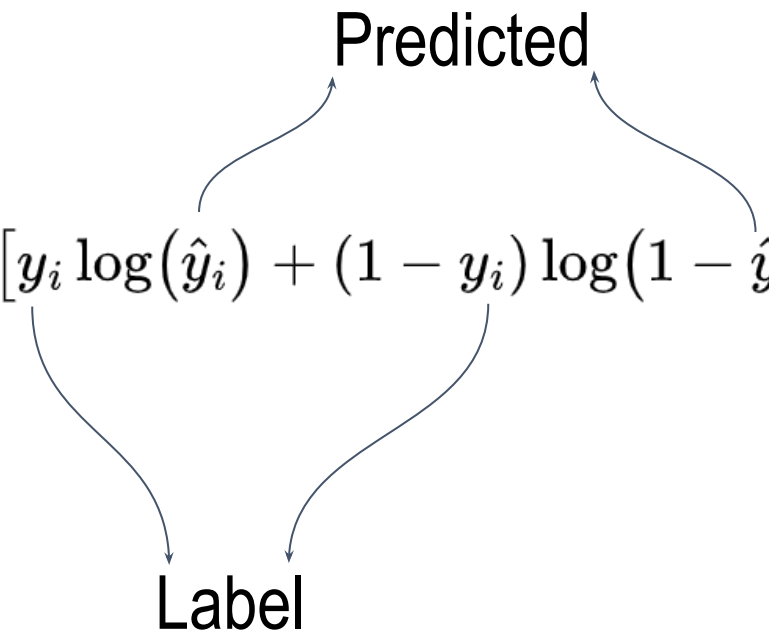
Label Smoothing (LS)

→ Is a training-time regularizer:

$$\mathcal{L}(\hat{y}, y) = -\frac{1}{|\mathcal{E}|} \sum_{i=1}^{|\mathcal{E}|} [y_i \log(\hat{y}_i) + (1 - y_i) \log(1 - \hat{y}_i)]$$

Predicted

Label



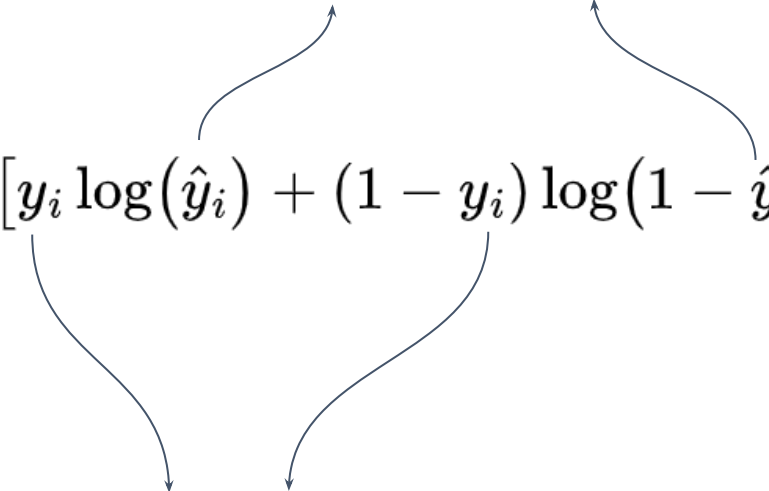
The diagram shows two arrows originating from the term $\hat{y}_i$ in the loss function, both pointing to the word "Predicted". Another two arrows originate from the term y_i in the loss function, both pointing to the word "Label".

Smoothed: 0.95, 0.05, 0.95, 0.05 ...

Label Smoothing (LS)

→ Is a training-time regularizer:

$$\mathcal{L}(\hat{y}, y) = -\frac{1}{|\mathcal{E}|} \sum_{i=1}^{|\mathcal{E}|} [y_i \log(\hat{y}_i) + (1 - y_i) \log(1 - \hat{y}_i)]$$


 Predicted
 Label

Smoothed: 0.95, 0.05, 0.95, 0.05 ...

Tail prediction: $p_s(t \mid h, r) = \begin{cases} 1 - \alpha & \text{if } t' = t \\ \frac{\alpha}{(|\mathcal{E}| - 1)} & \text{otherwise} \end{cases}$

Label Smoothing (LS)

→ Is a training-time regularizer:

$$\mathcal{L}(\hat{y}, y) = -\frac{1}{|\mathcal{E}|} \sum_{i=1}^{|\mathcal{E}|} [y_i \log(\hat{y}_i) + (1 - y_i) \log(1 - \hat{y}_i)]$$

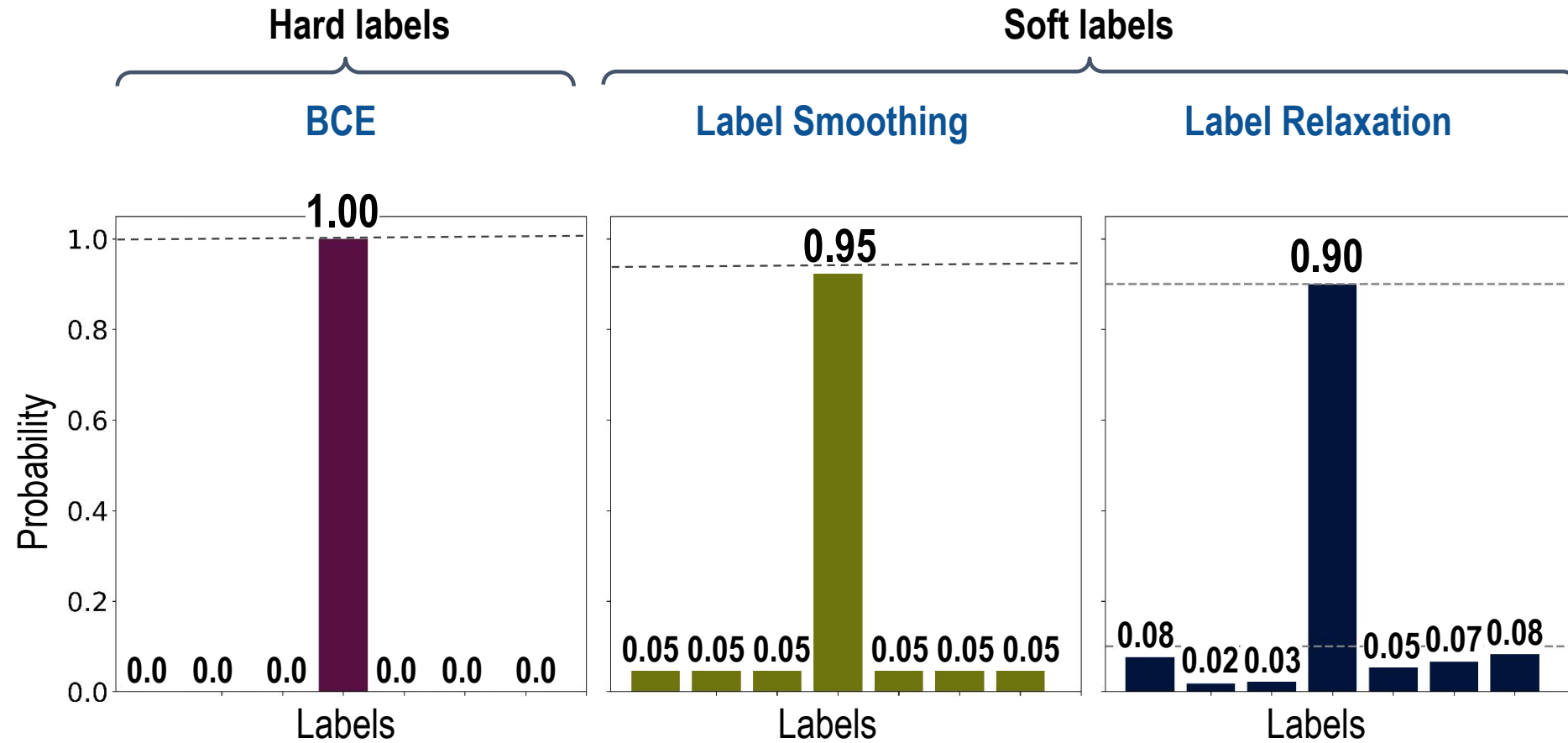
Predicted

Label

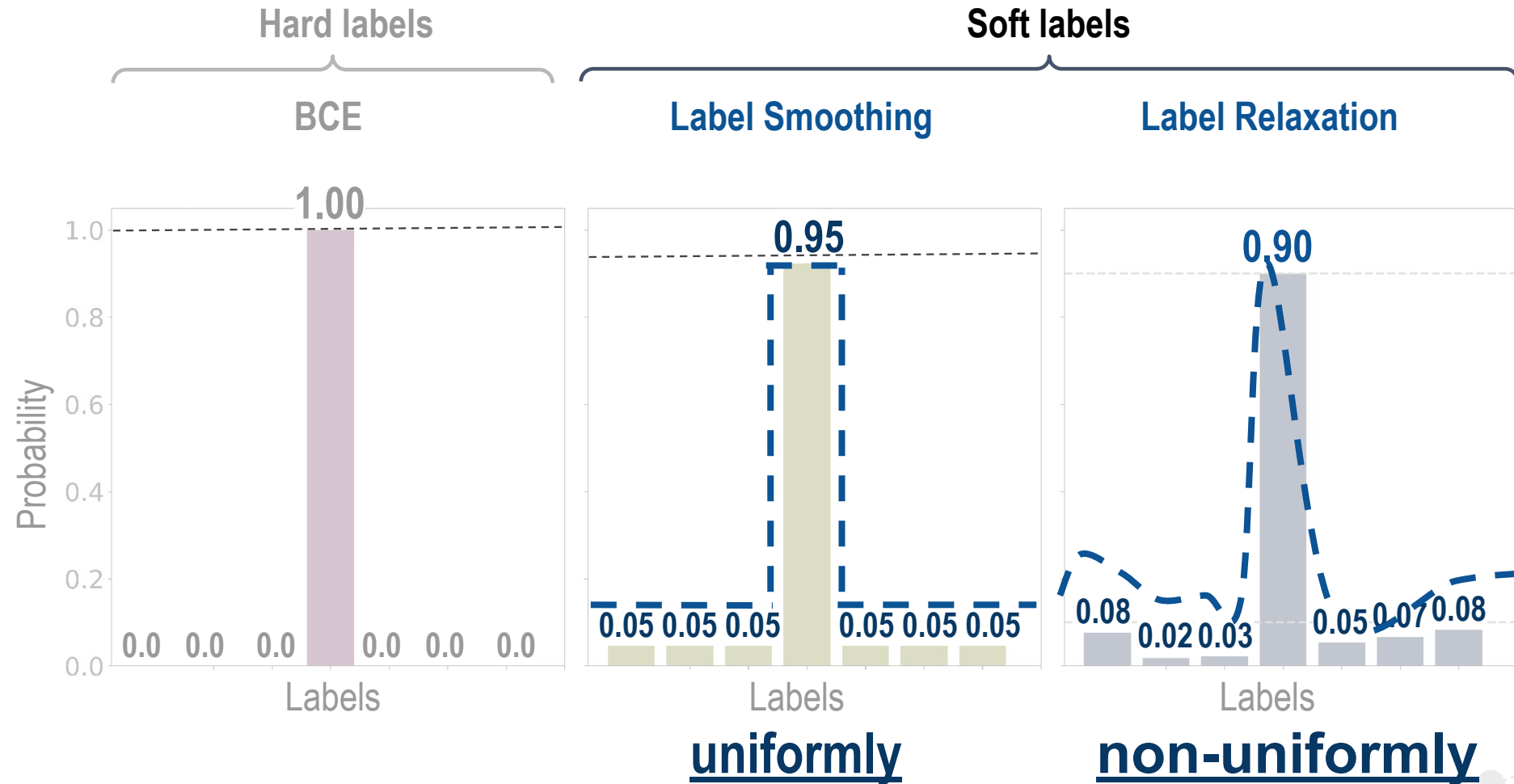
Smoothed: 0.95, 0.05, 0.95, 0.05 ...

Tail prediction: $p_s(t \mid h, r) = \begin{cases} 1 - \alpha & \text{if } t' = t \\ \frac{\alpha}{(|\mathcal{E}| - 1)} & \text{otherwise} \end{cases}$

BCE and soft labels



BCE and soft labels



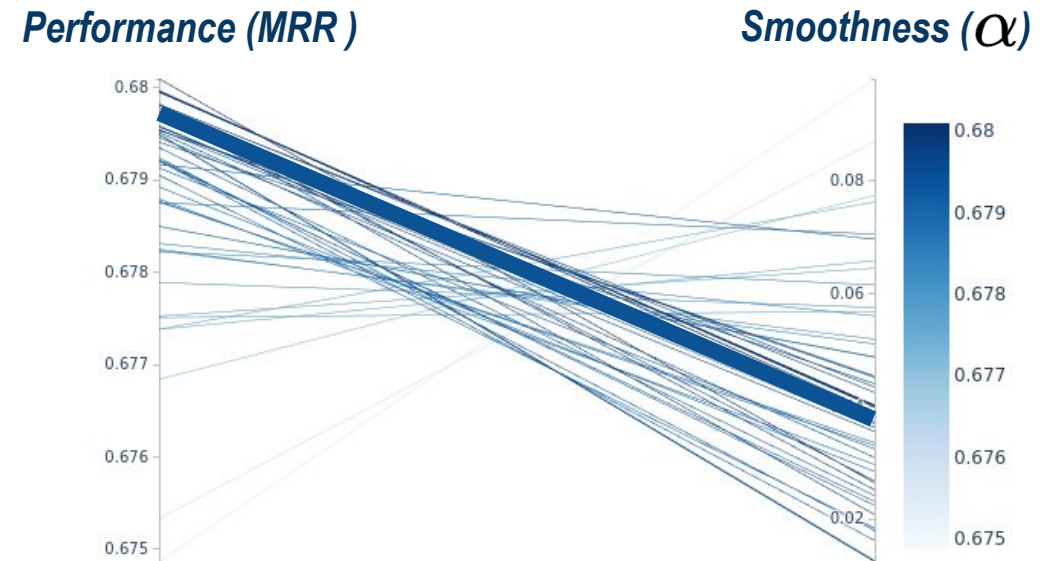
So far,
 α as a smoothing factor

How to choose the optimal α value?

$$\alpha \in (0, 1]$$

LS with Bayesian Optimization

- Smoothing ratio may depend on model architecture and training dataset
- We performed **50 runs** per model–dataset combination
 - Hyper-parameter optimization!



Adaptive Label Smoothing/Relaxation

- Smoothing ratio adjusts dynamically during the training, instead of relying on a fixed value.

α

39

Adaptive Label Smoothing/Relaxation

- Smoothing ratio adjusts dynamically during the training, instead of relying on a fixed ratio.
- Tracks the change in Kullback-Leibler (KL) divergence between predicted and target label distributions.

α

40

Adaptive Label Smoothing/Relaxation

- Smoothing ratio adjusts **dynamically** during the training, instead of relying on a fixed ratio.
- Tracks the change in Kullback-Leibler (KL) divergence between predicted and target label **distributions**.

$$\alpha_{t+1} = \begin{cases} \min(\alpha_t + \Delta\alpha, \alpha_{\max}), & \text{if } \mathcal{K} \, KL_{t+1} > KL_t \\ \max(\alpha_t - \Delta\alpha, \alpha_{\min}), & \text{otherwise} \end{cases}$$

KL: Kullback–Leibler divergence

$\Delta\alpha$: the step size (hyper-parameter)

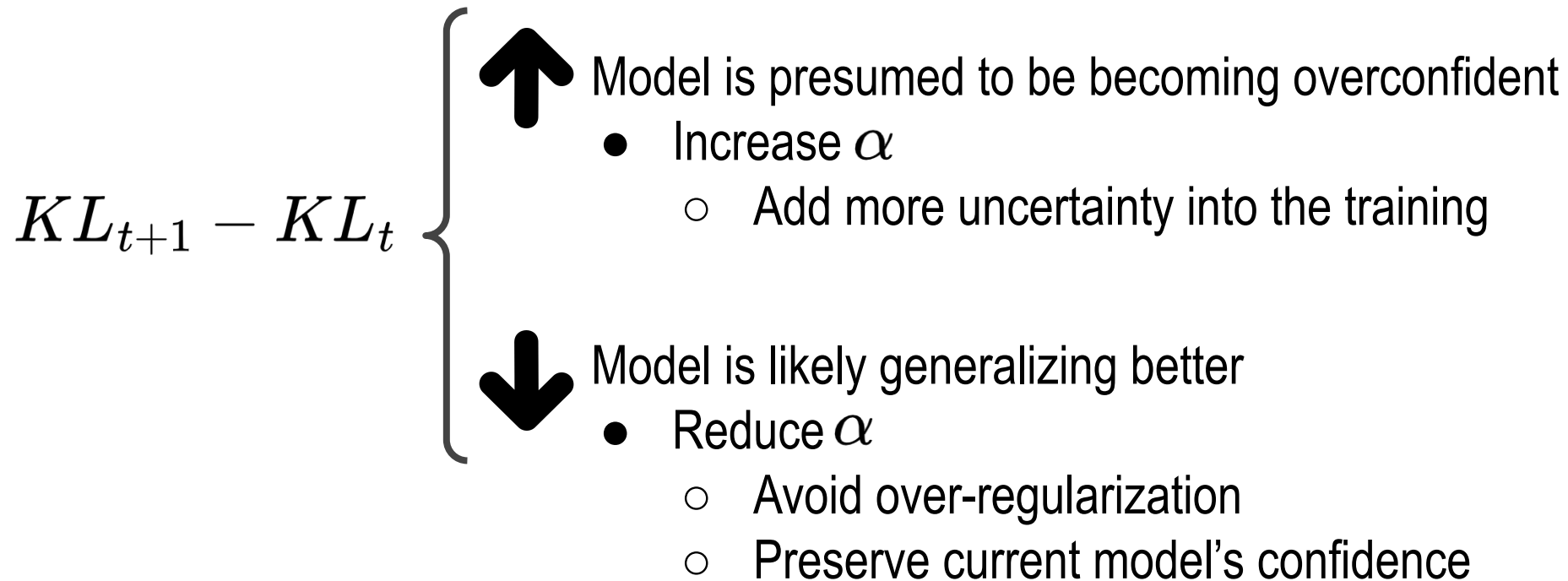
$\alpha_{\min}, \alpha_{\max}$: hyper-parameters

t: training step at time t

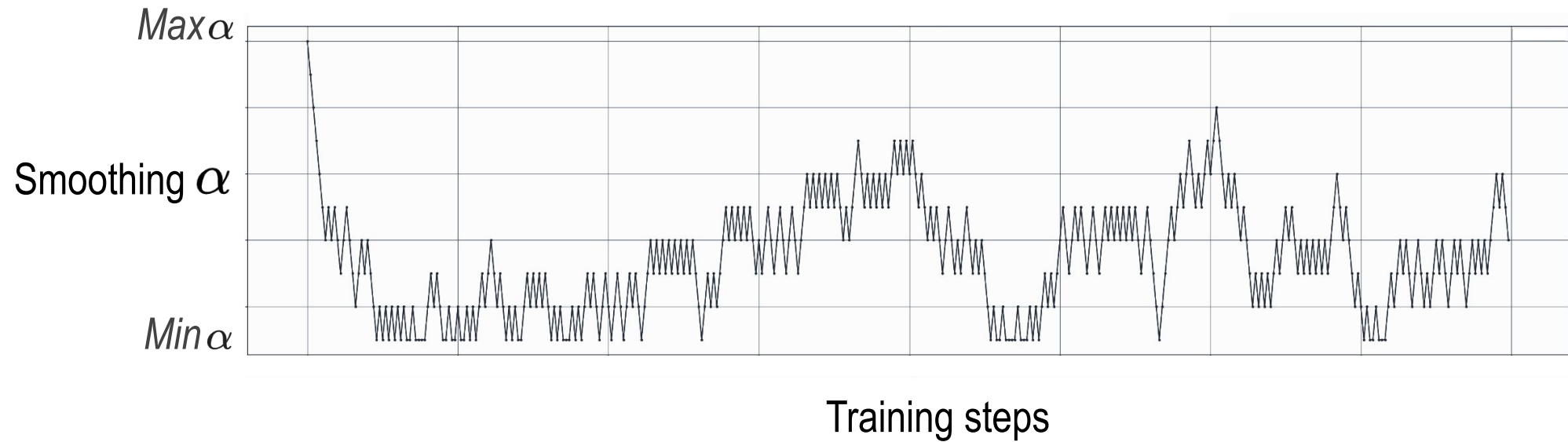
Adaptive Label Smoothing/Relaxation

$$KL_{t+1} - KL_t \left\{ \begin{array}{l} \uparrow \text{Model is presumed to be becoming overconfident} \\ \bullet \text{ Increase } \alpha \\ \circ \text{ Add more uncertainty into the training} \end{array} \right.$$

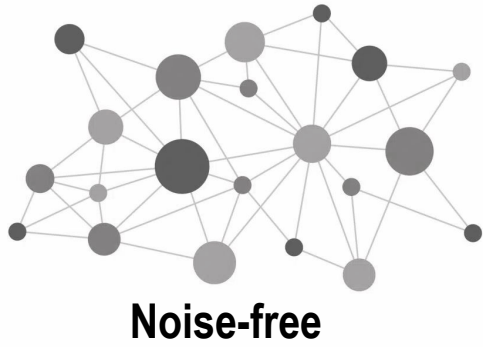
Adaptive Label Smoothing/Relaxation



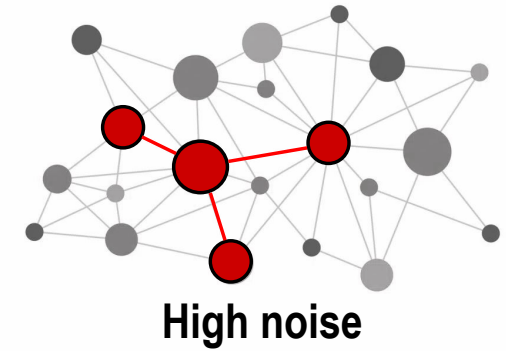
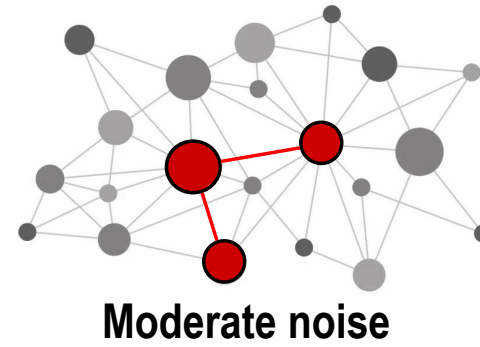
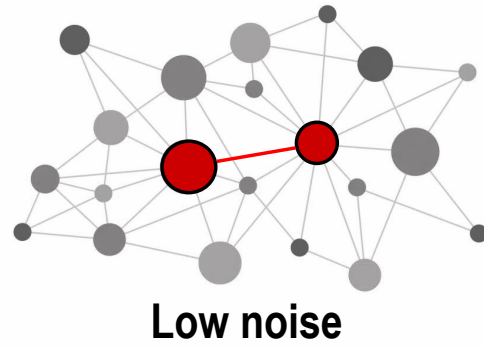
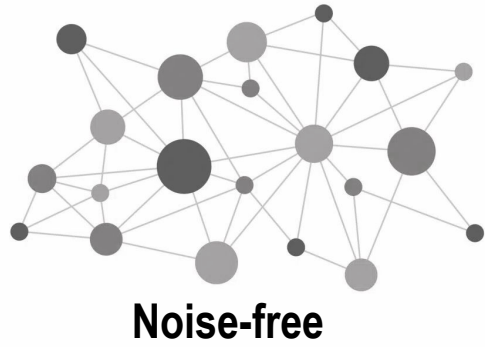
Adaptive Label Smoothing/Relaxation



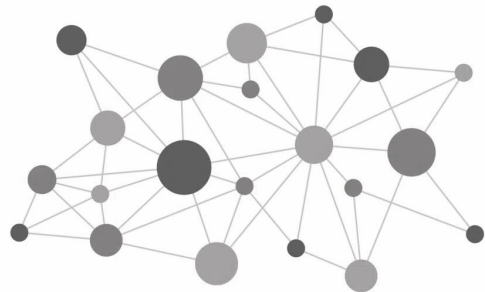
Experimental setup



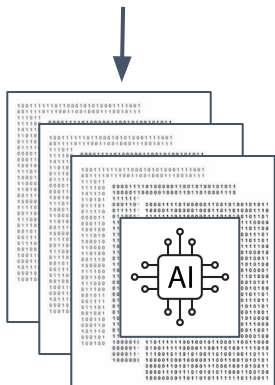
Experimental setup



Experimental setup

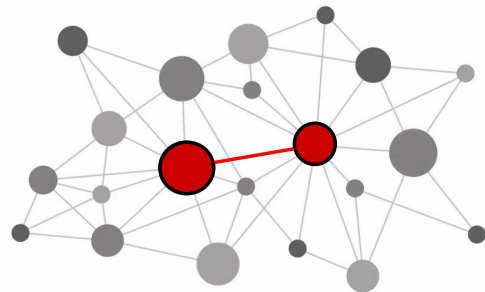


Noise-free

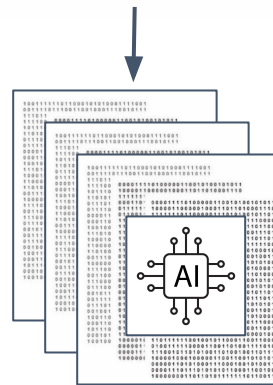


KGE Models

Hard vs. Soft labels?

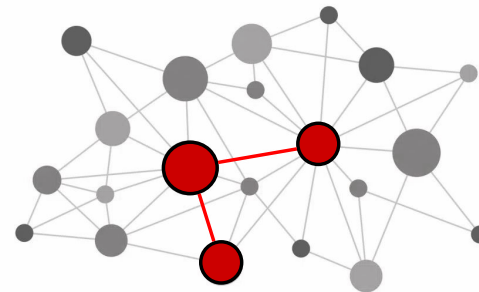


Low noise

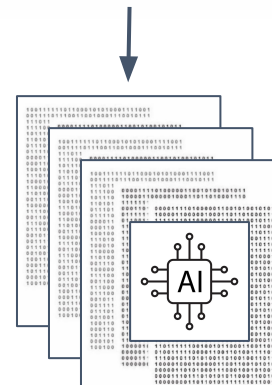


KGE Models

Hard vs. Soft labels?

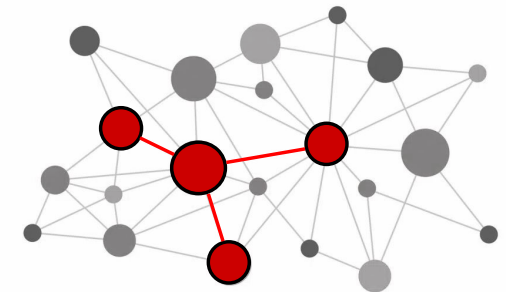


Moderate noise

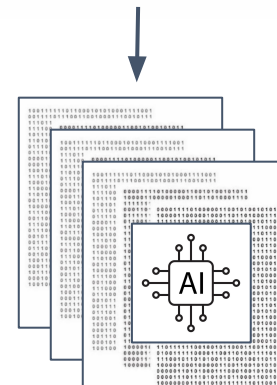


KGE Models

Hard vs. Soft labels?



High noise



KGE Models

Hard vs. Soft labels?

Experimental setup

Models: *DistMult*, *ComplEx*, *QMult*, *RotatE*, *MuRE*, *Keci*

Experimental setup

Models: *DistMult, ComplEx, QMult, RotatE, MuRE, Keci*

Datasets: *UMLS, KINSHIP, WN18RR, NELL-995-h100, FB15K-237*

Experimental setup

Models: *DistMult, ComplEx, QMult, RotatE, MuRE, Keci*

Datasets: *UMLS, KINSHIP, WN18RR, NELL-995-h100, FB15K-237*

Random Noise ratios: *0.0 .. 0.64*

Experimental setup

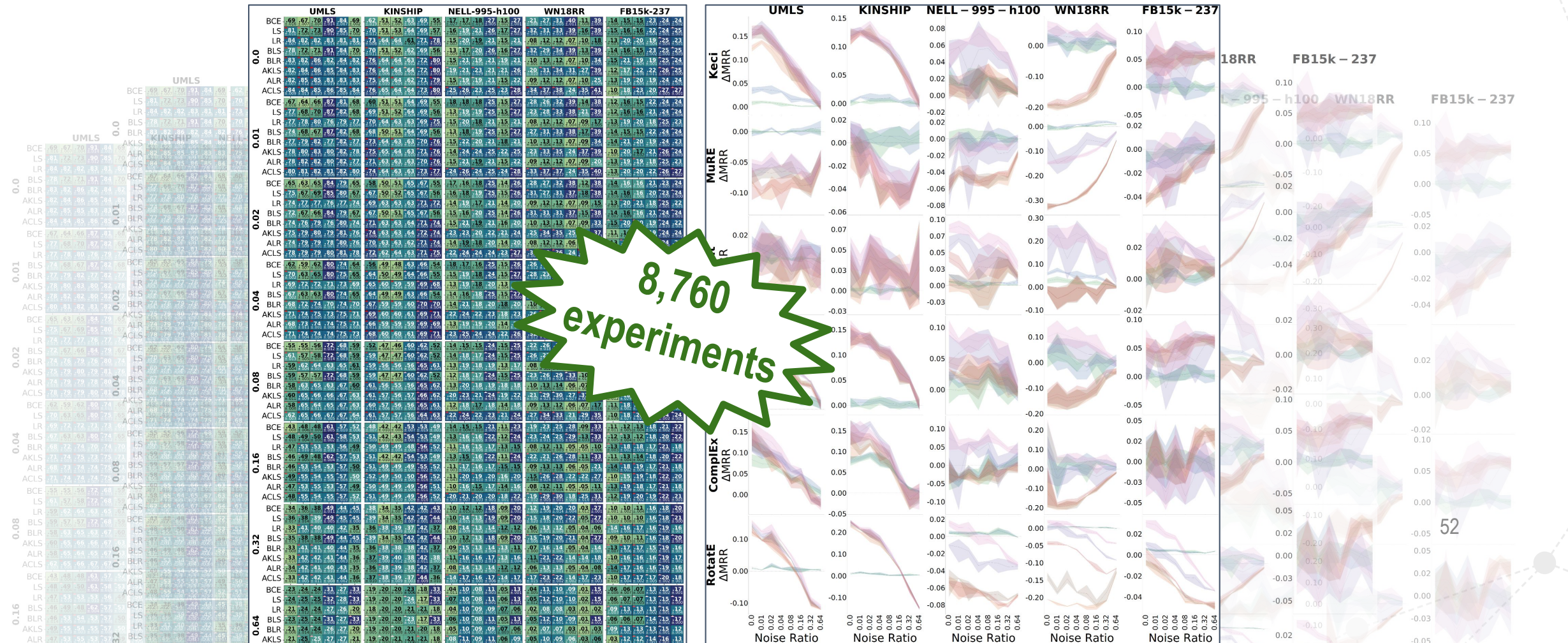
Models: *DistMult, ComplEx, QMult, RotatE, MuRE, Keci*

Datasets: *UMLS, KINSHIP, WN18RR, NELL-995-h100, FB15K-237*

Random Noise ratios: *0.0 .. 0.64*

Loss functions: $\left\{ \begin{array}{l} \textbf{Hard: } BCE \\ \textbf{Soft: } Smoothing, Relaxation, Bayesian/Adaptive versions \end{array} \right.$

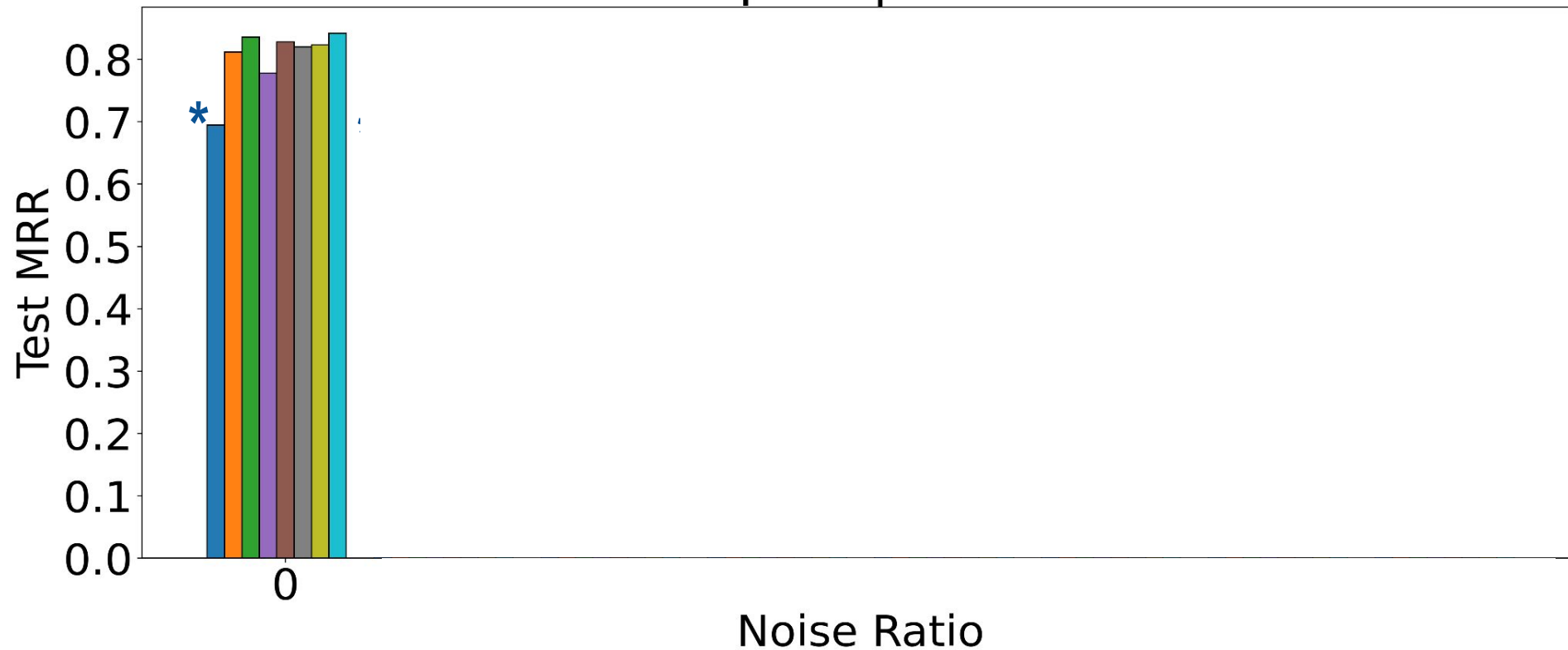
Results



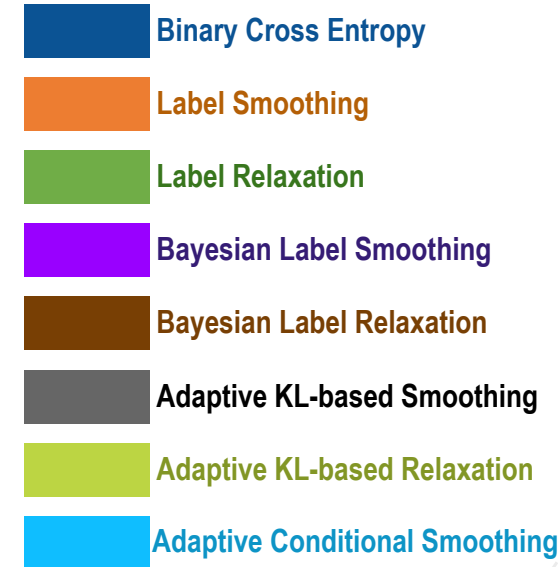
A. Memariani (DICE): Link Prediction Under Non-targeted Attacks: Do Soft Labels Always Help?

Example Result

Model: ComplEx | Dataset: UMLS



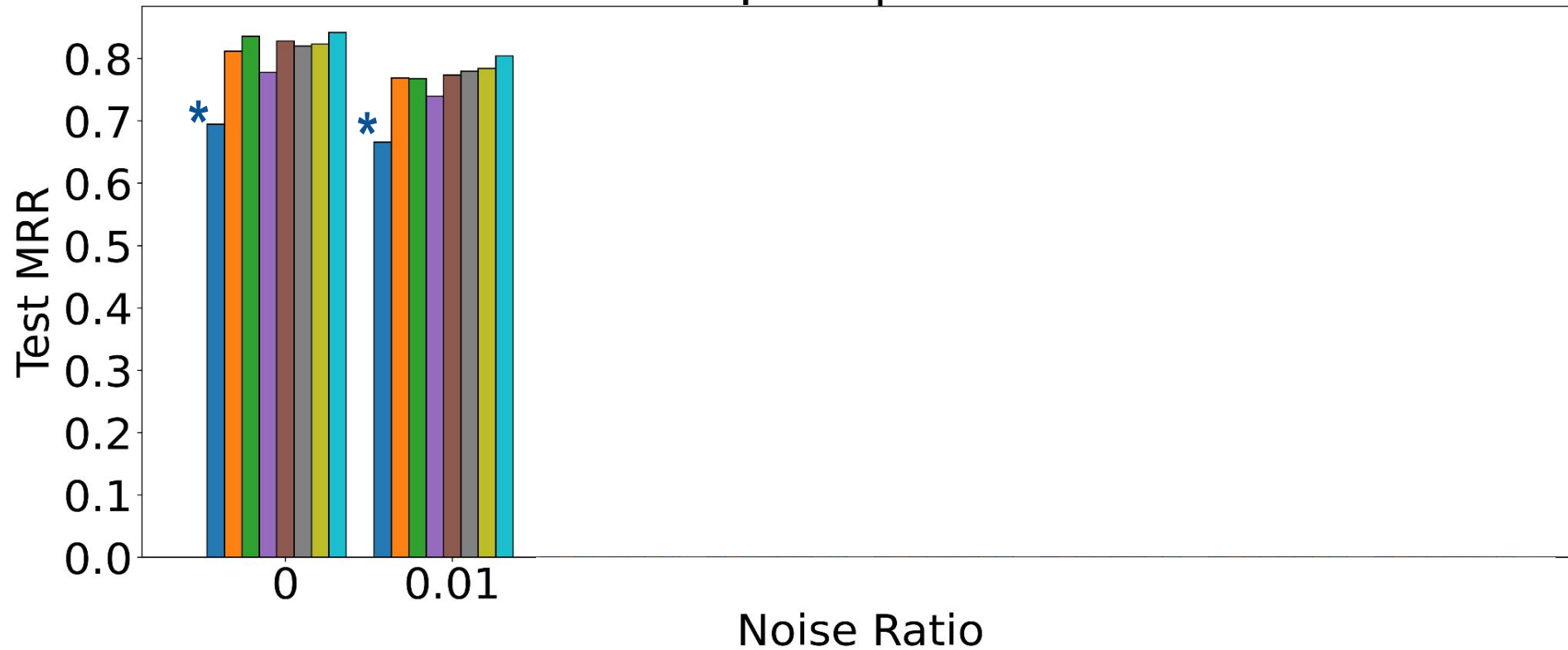
Loss



* BCE as the baseline

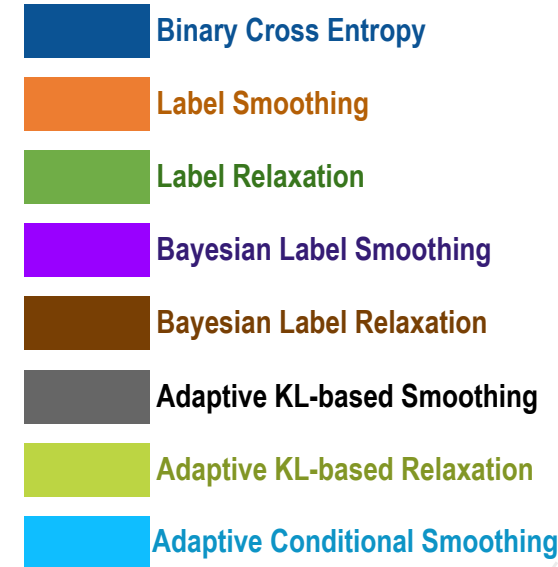
Example Result

Model: ComplEx | Dataset: UMLS



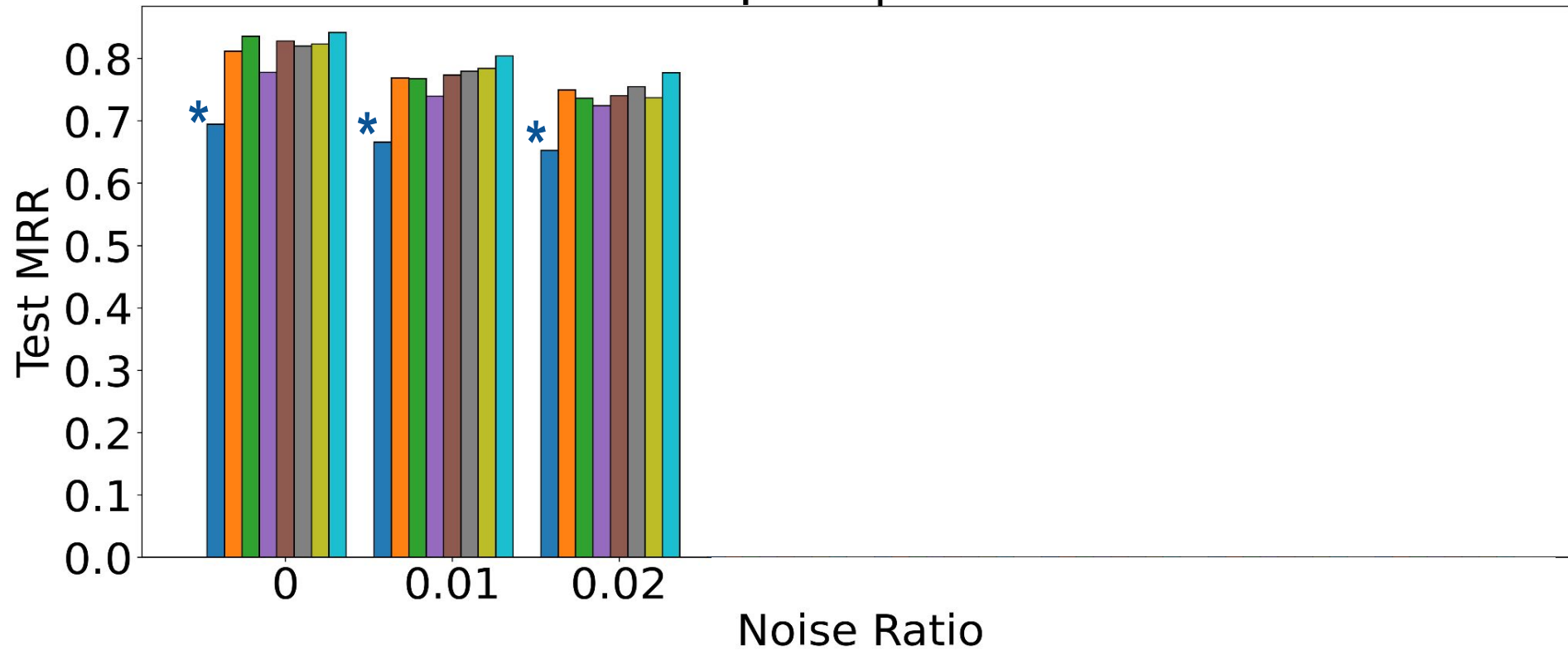
* BCE as the baseline

Loss



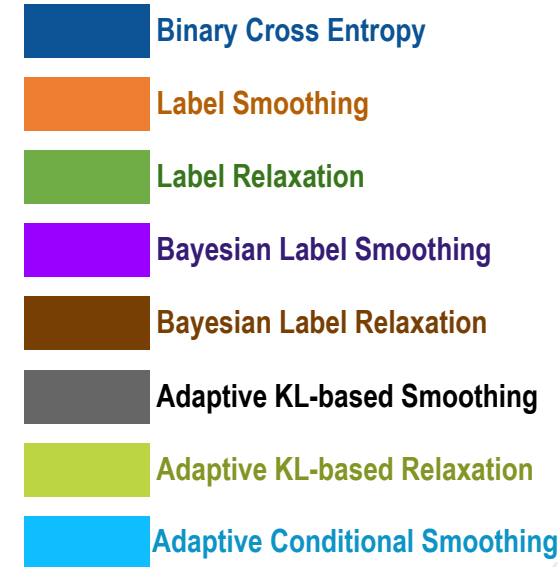
Example Result

Model: ComplEx | Dataset: UMLS



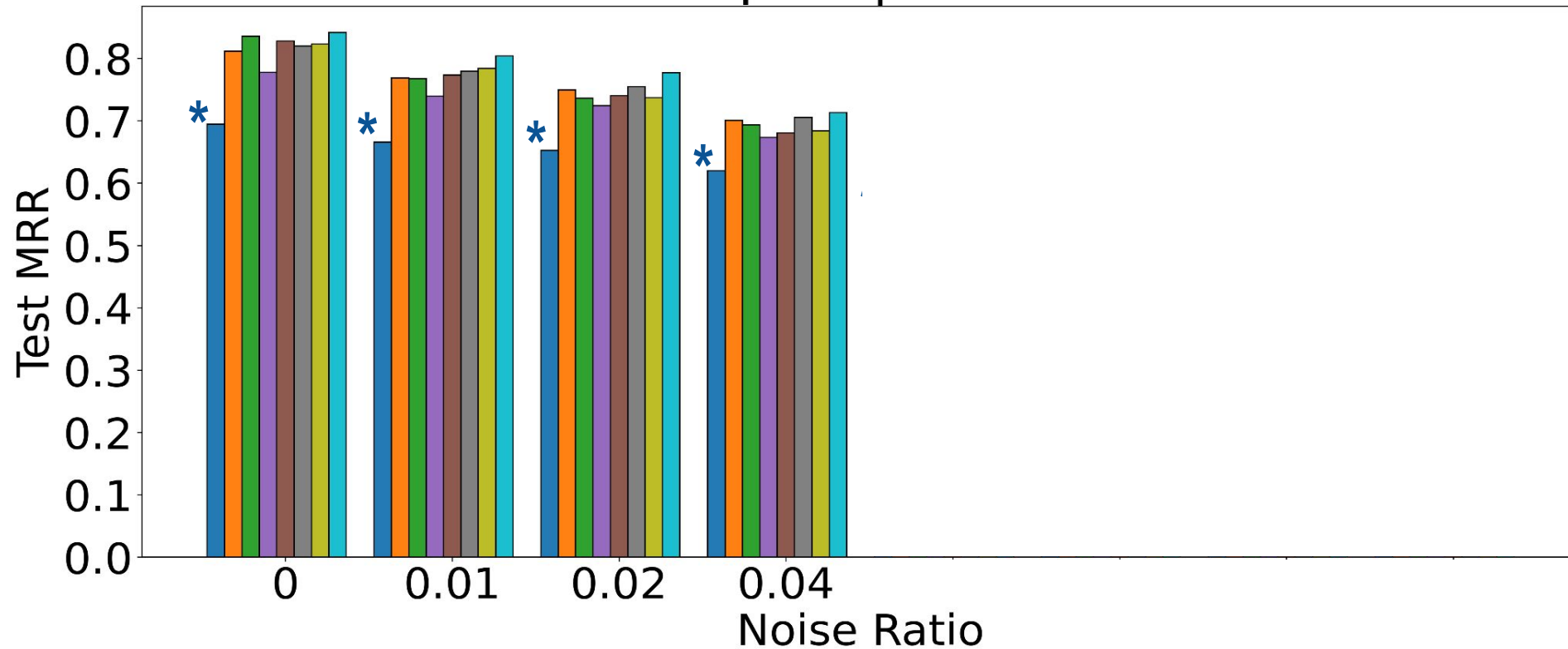
* BCE as the baseline

Loss

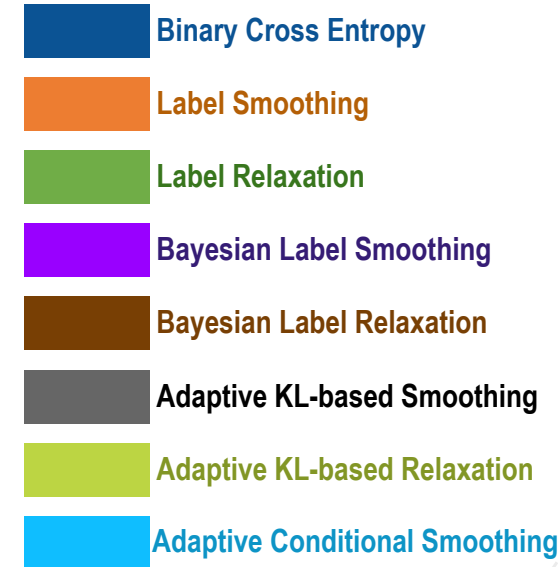


Example Result

Model: Complex | Dataset: UMLS



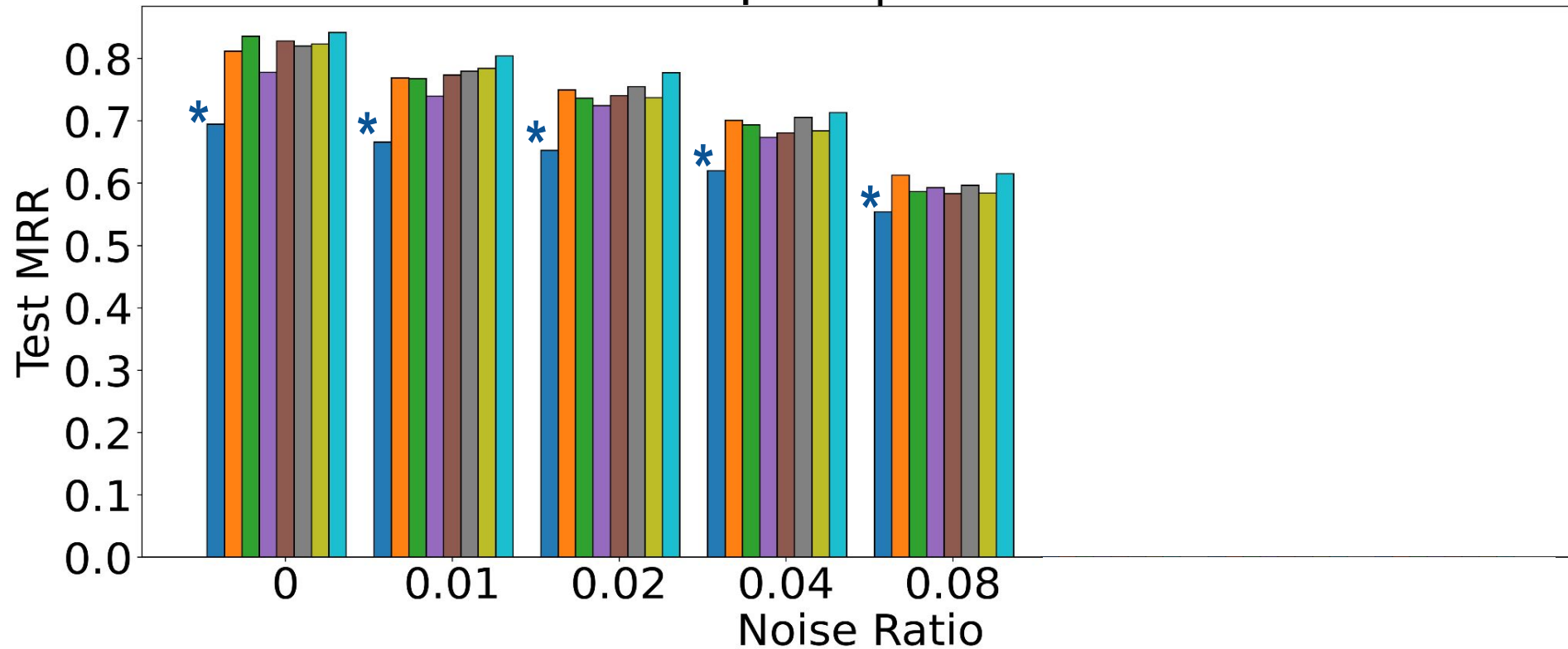
Loss



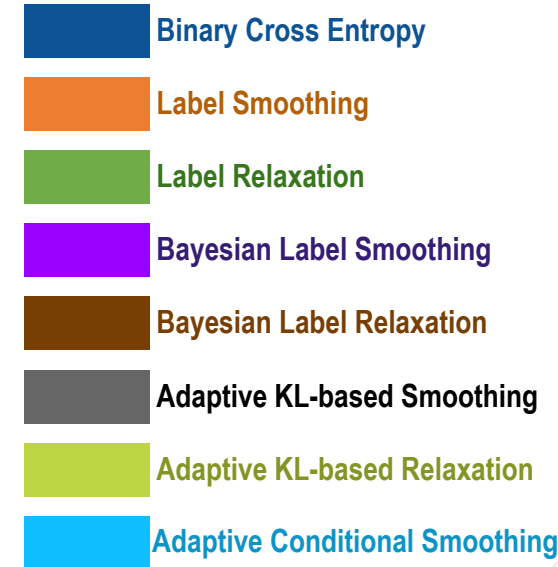
* BCE as the baseline

Example Result

Model: Complex | Dataset: UMLS



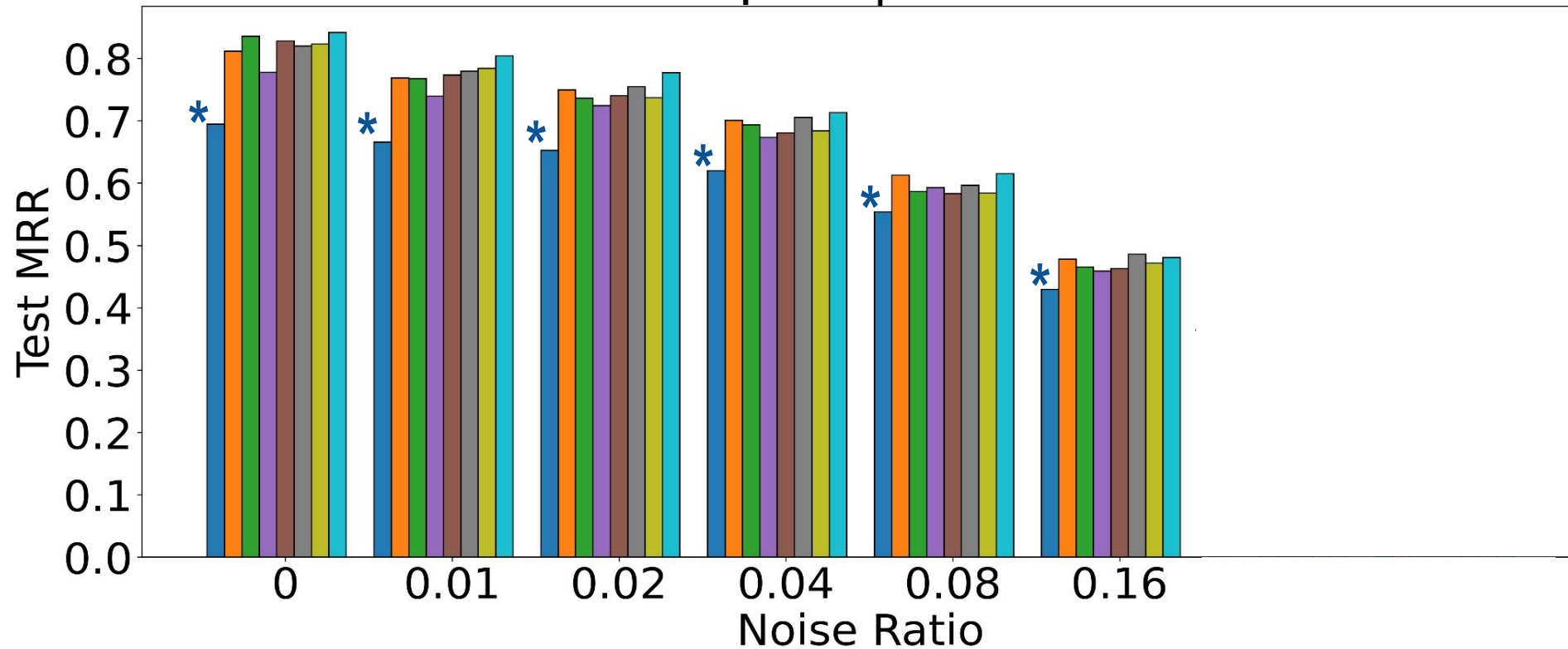
Loss



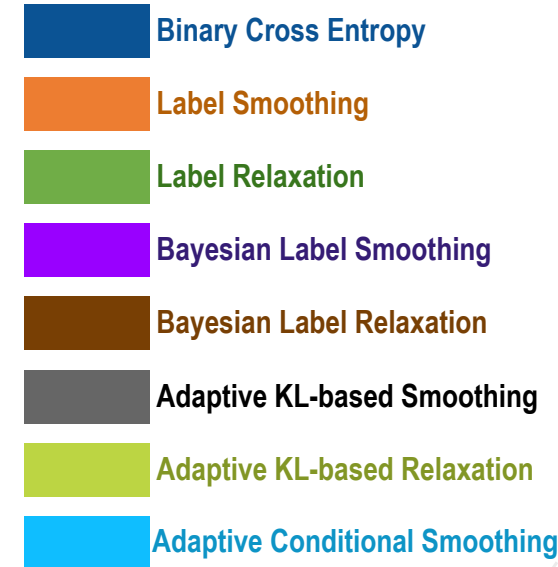
* BCE as the baseline

Example Result

Model: ComplEx | Dataset: UMLS



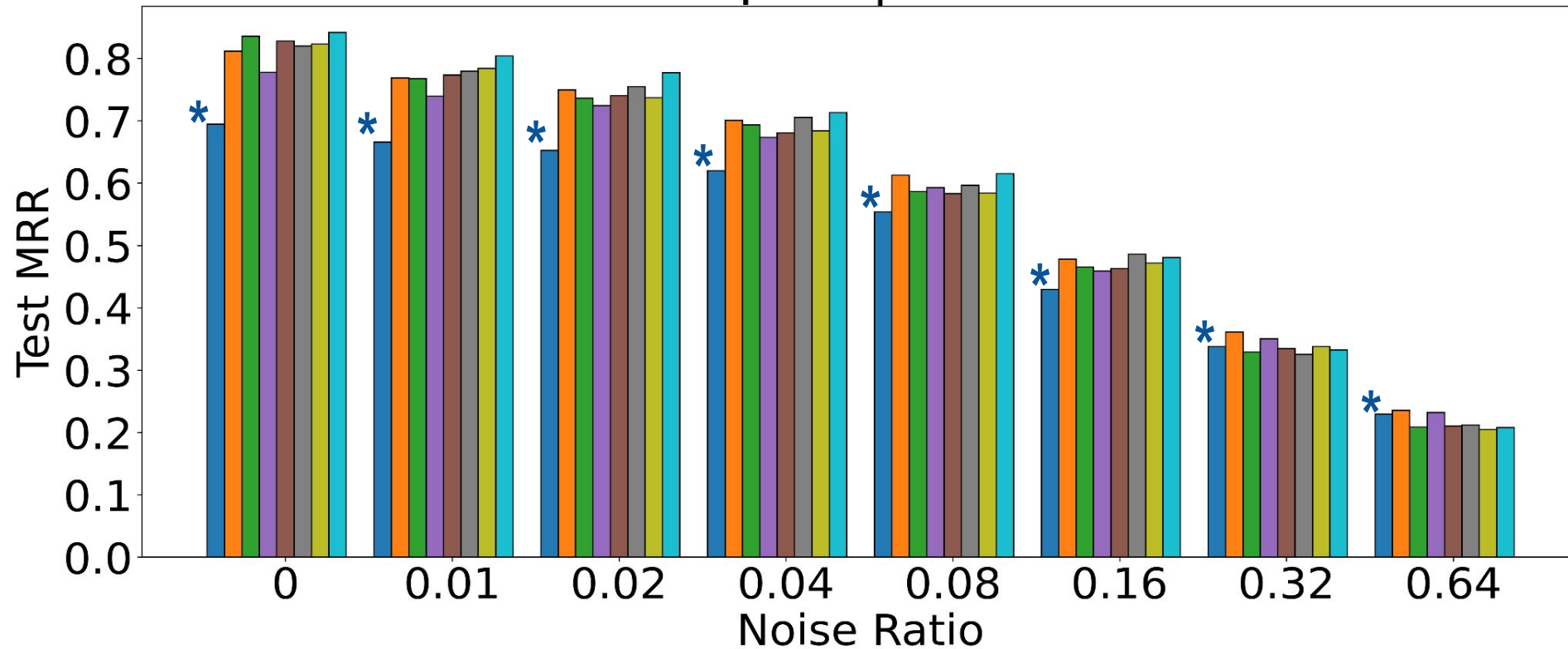
Loss



* BCE as the baseline

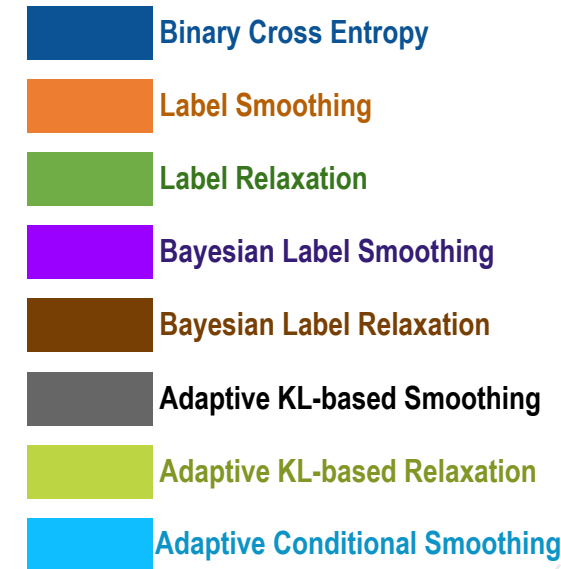
Example Result

Model: ComplEx | Dataset: UMLS



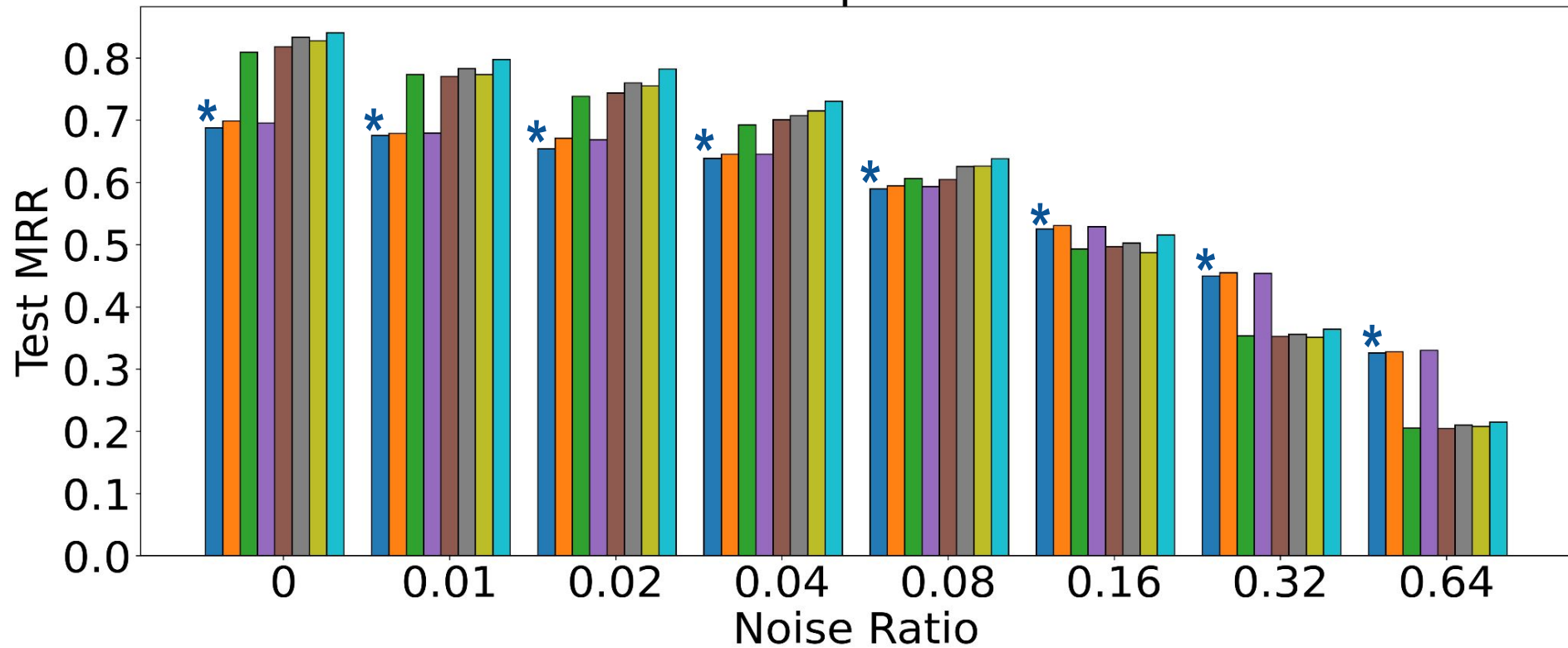
* BCE as the baseline

Loss



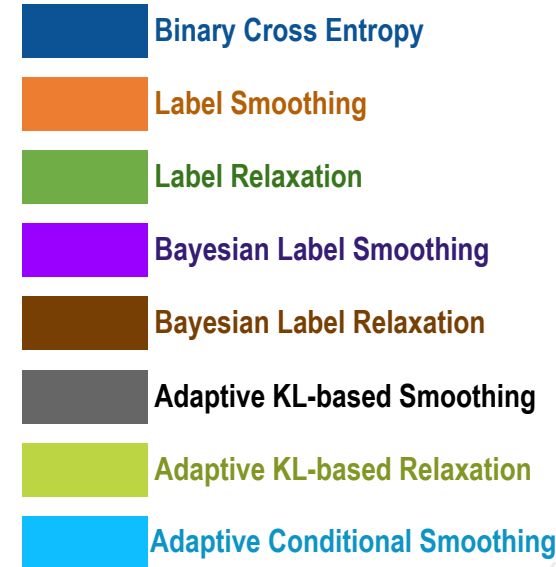
Example Result

Model: RotatE | Dataset: UMLS



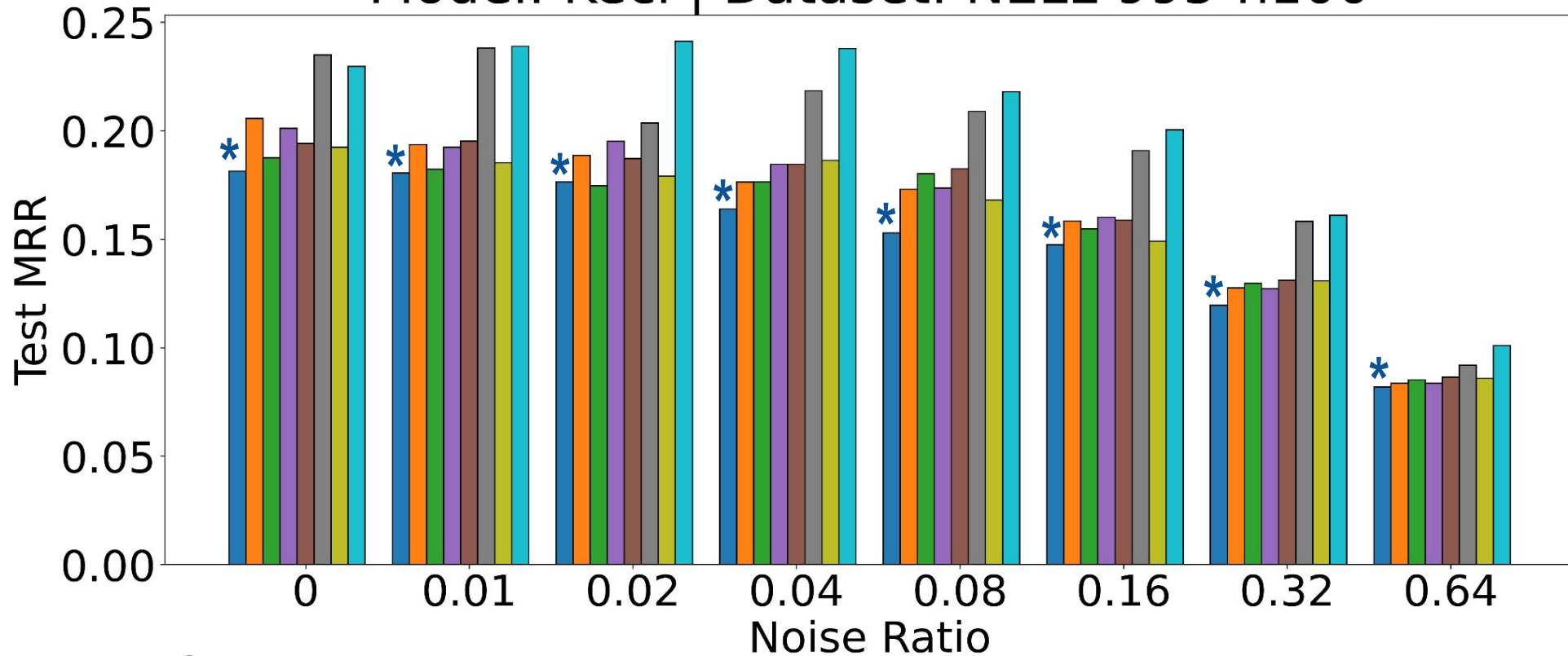
* BCE as the baseline

Loss



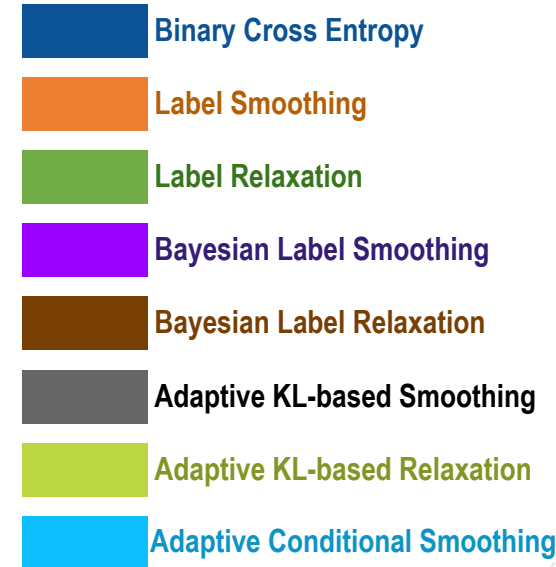
Example Result

Model: Keci | Dataset: NELL-995-h100



* BCE as the baseline

Loss



Statistical Testing

- Beyond improvements, we assessed their significance
 - Soft-label loss function that are better than BCE:

Statistical Testing

- Beyond improvements, we assessed their significance
 - Soft-label loss function that are better than BCE
- 8 noise ratios across 5 datasets, resulting in 48 groups
 - Significant improvements across 41 out of 48 groups

Conclusions

- Soft labels enable more robust link prediction under noisy data.
- Suitability of KL divergence for label smoothing.
- Bayesian optimization helps, but is expensive;
 - Adaptive smoothing is better with less computational cost.

Future works

- Focus on more structured inconsistencies:
 - Systematic errors in information extraction
 - Misalignments in ontological hierarchies
- Adaptive version of smoothing based on similarity among entities
- Hyper-parameter optimization based on other strategies

Thank you

Adel Memariani
adel.memariani@upb.de

Data Science Group (DICE)
Heinz Nixdorf Institute
Paderborn University

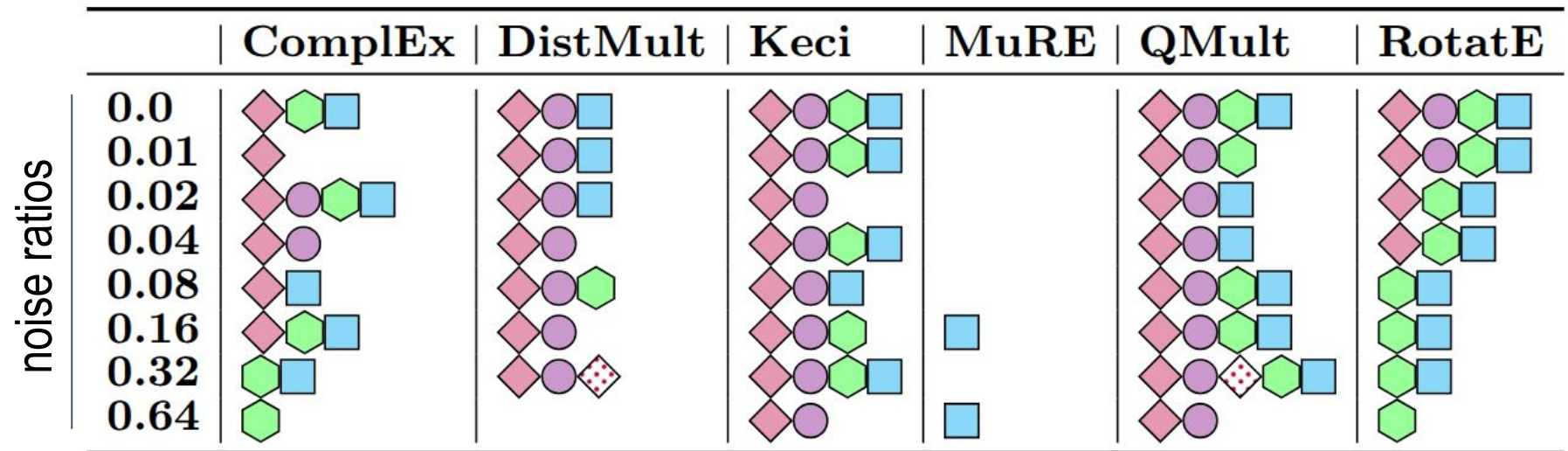
Link Prediction Under Non-targeted Attacks: Do Soft Labels Always Help?

- Soft labels enable more robust link prediction under noisy data
- Suitability of KL divergence for label smoothing
- Bayesian optimization helps, but is expensive
 - Adaptive smoothing is better with less computational cost

Appendix

Statistical Testing

- **Significant improvements** in link prediction performance across **41 out of 48** groups

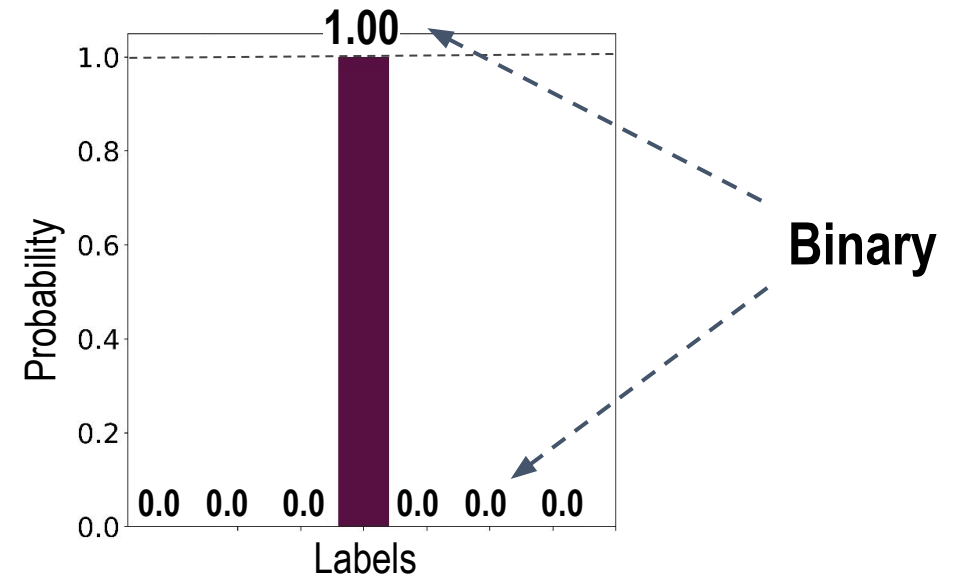


Wilcoxon signed-rank, $p < 0.05$, on MRR scores. Soft label-based loss functions that outperformed BCE and are statistically significant. ◆ Adaptive Conditional Smoothing ● AKLS Adaptive KL-based Smoothing

□ LS Label Smoothing ◆ BLR Bayesian Label Relaxation ● BLS Bayesian Label Smoothing

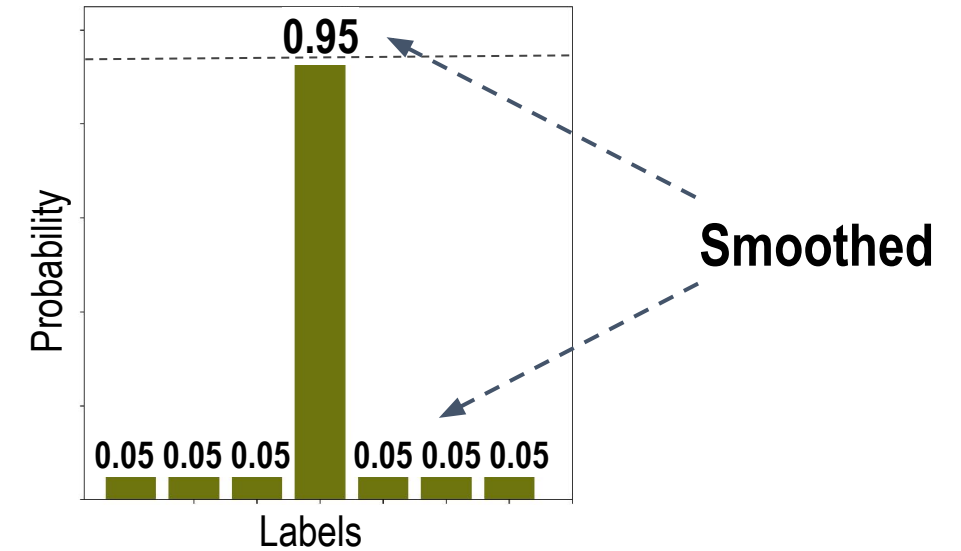
Binary Class Entropy (BCE) Loss

- Target as 100% true class, 0% everything else



Label Smoothing (LS)

- Redistributes a small portion of the probability mass uniformly across all other labels.

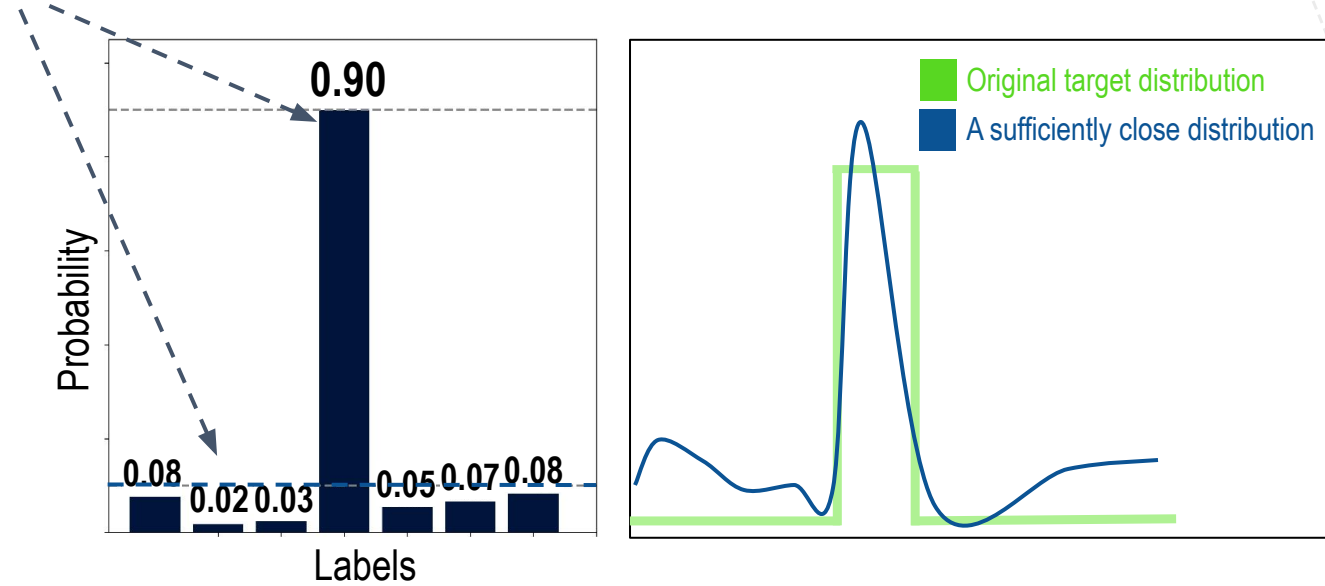


$$p_s(t \mid h, r) = \begin{cases} 1 - \alpha & \text{if } t' = t \\ \frac{\alpha}{(|\mathcal{E}| - 1)} & \text{otherwise} \end{cases}$$

Label Relaxation (LR)

- Relax the target into a **set of distributions**
- A **set** of acceptable targets.
 - Instead of a precise probability
 - Labels vary within a range.

Relaxed (not uniformly)



$$Q_{\alpha}^{(h,r,t)} = \{p' \in \mathcal{P}(Y) \mid p'(t \mid h, r) \geq 1 - \alpha\}$$

Credal set of distributions that are sufficiently close to the **original target**