

Selective Prediction in Molecular Structure Retrieval

Workshop on Uncertainty Estimation in Machine Learning

Mira Jürgens

PhD student
Department of Data Science and Mathematical Modeling
Ghent University
Promotor: Willem Waegeman

Contributors



**Mira
Jürgens**
Ghent University



**Gaetan
De Waele**
University of Antwerp

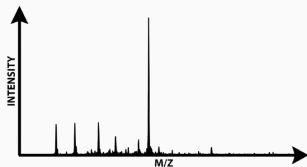


**Morteza
Rakhshaninejad**
Ghent University



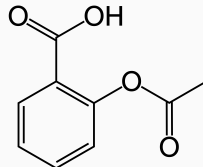
**Willem
Waegeman**
Ghent University

Molecular Identification from Mass Spectrometry



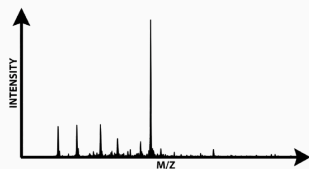
Spectrum x

identify?



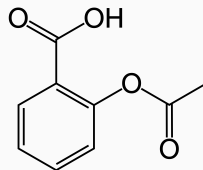
Molecule y

Molecular Identification from Mass Spectrometry



Spectrum x

identify?

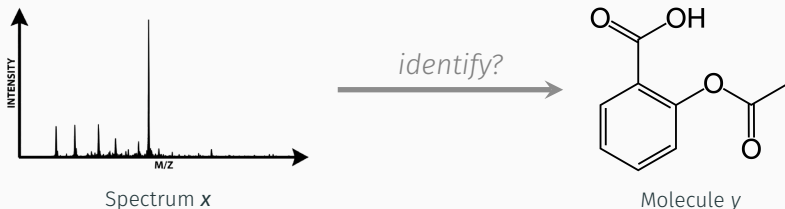


Molecule y

Retrieval task: Given spectrum x , rank a candidate set $\mathcal{C}(x)$ of molecules

$$\underbrace{\geq 10^{60}}_{\text{chemical space}} \xrightarrow{\text{filter by mass/formula}} \underbrace{|\mathcal{C}(x)| \leq 256}_{\text{candidates}}$$

Molecular Identification from Mass Spectrometry

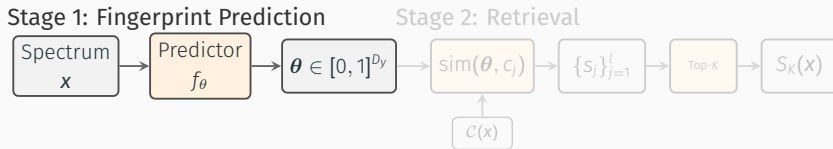


Retrieval task: Given spectrum x , rank a candidate set $\mathcal{C}(x)$ of molecules

$$\underbrace{\geq 10^{60}}_{\text{chemical space}} \xrightarrow{\text{filter by mass/formula}} \underbrace{|\mathcal{C}(x)| \leq 256}_{\text{candidates}}$$

Challenges: Sparse training data (< 1% coverage), ambiguous mapping (isomers, noise)

Two-Stage Retrieval Pipeline



$\mathcal{X} \subseteq \mathbb{R}^{D_x}$: MS/MS spectra

$\mathcal{Y} \subseteq \{0, 1\}^{D_y}$: fingerprints

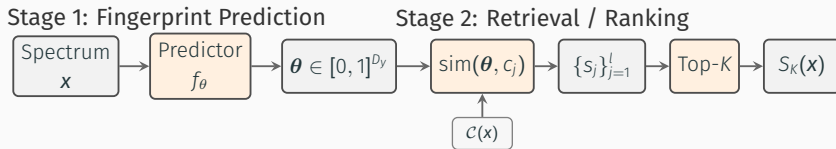
$\theta \in [0, 1]^{D_y}$: bit probabilities

$f_\theta : \mathcal{X} \rightarrow [0, 1]^{D_y}$: fingerprint predictor

Model from De Waele et al.¹

¹De Waele et al., "Small Molecule Retrieval from Tandem Mass Spectrometry: What are we optimizing for?"

Two-Stage Retrieval Pipeline



$\mathcal{X} \subseteq \mathbb{R}^{D_x}$: MS/MS spectra

$\mathcal{Y} \subseteq \{0, 1\}^{D_y}$: fingerprints

$\boldsymbol{\theta} \in [0, 1]^{D_y}$: bit probabilities

$f_{\theta} : \mathcal{X} \rightarrow [0, 1]^{D_y}$: fingerprint predictor

$\mathcal{C}(\mathbf{x}) \subset \mathcal{Y}$: candidates

$s_j = \text{sim}(\boldsymbol{\theta}, c_j)$: similarity score

$S_K(\mathbf{x}) \subset \mathcal{C}(\mathbf{x})$: top-K candidates

Model from De Waele et al.¹

¹De Waele et al., "Small Molecule Retrieval from Tandem Mass Spectrometry: What are we optimizing for?"

MassSpecGym: A benchmark for the discovery and identification of molecules

Roman Bushuiev^{1,2}, Anton Bushuiev², Niek F. de Jonge³, Adamo Young⁴,
Fleming Kretschmer⁵, Raman Samusevich^{1,2}, Janne Heirman⁶, Fei Wang^{7,8},
Luke Zhang⁹, Kai Dührkop⁵, Marcus Ludwig¹⁰, Nils A. Haupt⁵, Apurva Kalia¹¹,
Corinna Brungs¹, Robin Schmid¹, Russell Greiner^{7,8}, Bo Wang⁴, David S. Wishart^{7,12},
Li-Ping Liu¹¹, Juho Rousu¹³, Wout Bittremieux⁶, Hannes Rost⁹, Tytus D. Mak¹⁴,
Soha Hassoun^{11,15}, Florian Huber¹⁶, Justin J.J. van der Hooft^{3,17}, Michael A. Stravs¹⁸,
Sebastian Böcker⁵, Josef Sivic², Tomáš Pluskal¹

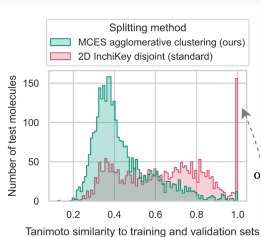
	Hit rate @ 1 ↑	Hit rate @ 5 ↑	Hit rate @ 20 ↑
Random	0.37 (0.24-0.54)	2.01 (1.68-2.39)	8.22 (7.53-8.89)
DeepSets	1.47 (1.18-1.77)	6.21 (5.64-6.82)	19.23 (18.24-20.26)
Fingerprint FFN	2.54 (2.17-2.99)	7.59 (6.96-8.28)	20.00 (19.01-20.98)
DeepSets + Fourier features	5.24 (4.71-5.83)	12.58 (11.80-13.42)	28.21 (27.10-29.36)
MIST	14.64 (13.82-15.54)	34.87 (33.69-36.10)	59.15 (57.89-60.39)

Figure 1: Retrieval results from Massspecgym²

²R. Bushuiev et al., “MassSpecGym: A benchmark for the discovery and identification of molecules”.

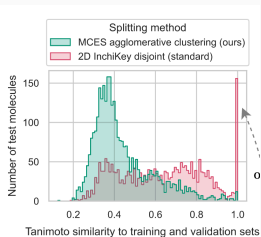
Structural Challenges

MCES-based split

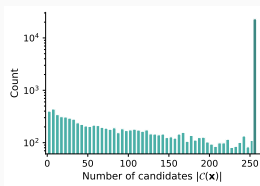


Structural Challenges

MCES-based split

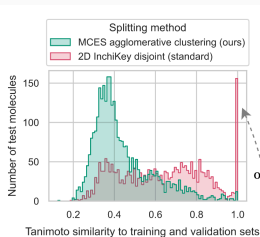


Varying number of candidates

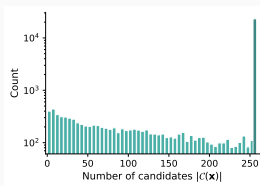


Structural Challenges

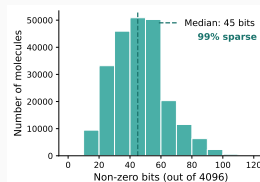
MCES-based split



Varying number of candidates

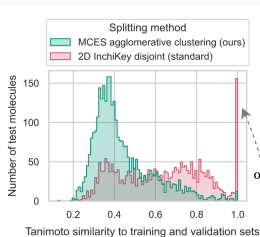


Extreme sparsity

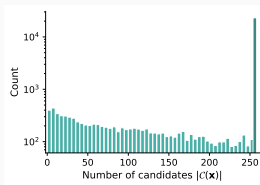


Structural Challenges

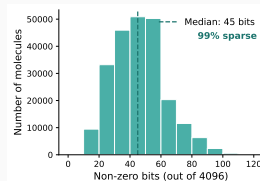
MCES-based split



Varying number of candidates



Extreme sparsity

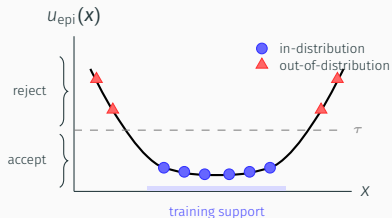


How can we use uncertainty to get guarantees on our predictions?

Selective Prediction in Molecular Retrieval

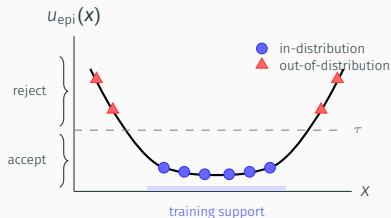
"Downstream Tasks" as evaluation in uncertainty quantification

Out-of-Distribution Detection

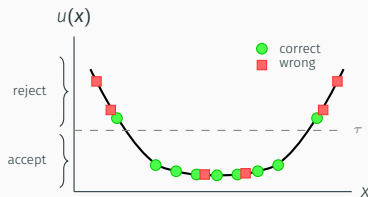


"Downstream Tasks" as evaluation in uncertainty quantification

Out-of-Distribution Detection

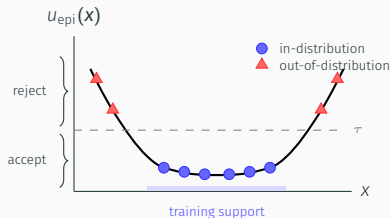


Selective Prediction

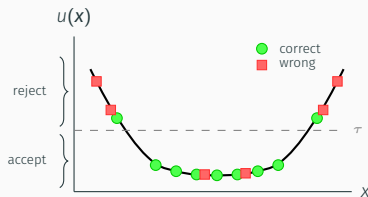


"Downstream Tasks" as evaluation in uncertainty quantification

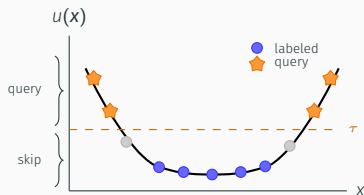
Out-of-Distribution Detection



Selective Prediction



Active Learning



Selective prediction setting

Following Geifman et al.³, we define

- $f: \mathcal{X} \rightarrow \mathcal{Y}$ as our **predictor**,
- $g_\tau: \mathcal{X} \rightarrow \{0, 1\}$ as our **selector**, and

$$(f, g_\tau)(x) = \begin{cases} f(x), & g_\tau(x) = 1 \\ \text{reject}, & g_\tau(x) = 0 \end{cases}$$

as the **selective predictor**.

³Geifman and El-Yaniv, “Selective Classification for Deep Neural Networks”.

Selective prediction setting

Following Geifman et al.³, we define

- $f: \mathcal{X} \rightarrow \mathcal{Y}$ as our **predictor**,
- $g_\tau: \mathcal{X} \rightarrow \{0, 1\}$ as our **selector**, and

$$(f, g_\tau)(x) = \begin{cases} f(x), & g_\tau(x) = 1 \\ \text{reject}, & g_\tau(x) = 0 \end{cases}$$

as the **selective predictor**. Having an *uncertainty score* $u_f: \mathcal{X} \rightarrow \mathbb{R}$, g_τ can be defined as

$$g_\tau(x) := g_\tau(x, u_f) = \begin{cases} 1, & u_f(x) < \tau \\ 0, & u_f(x) \geq \tau. \end{cases}$$

³Geifman and El-Yaniv, “Selective Classification for Deep Neural Networks”.

Selective prediction - Evaluation

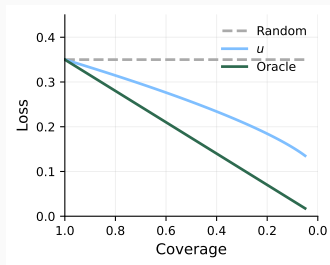
Risk-coverage-curve

Define the *coverage* of the selector g_τ as

$$\phi(g_\tau) = \mathbb{E}[g_\tau(X)]$$

and the *selective risk* as

$$R(f, g) = \frac{\mathbb{E}[\ell(f(x), y)g(x)]}{\phi(g)}.$$



Selective prediction - Evaluation

Risk-coverage-curve

Define the *coverage* of the selector g_τ as

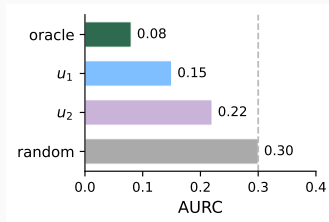
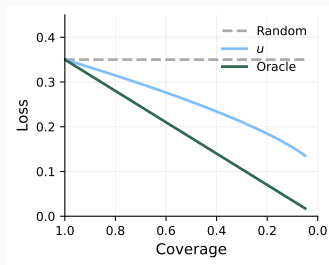
$$\phi(g_\tau) = \mathbb{E}[g_\tau(X)]$$

and the *selective risk* as

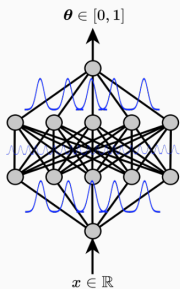
$$R(f, g) = \frac{\mathbb{E}[\ell(f(x), y)g(x)]}{\phi(g)}.$$

Area-under-risk-coverage-curve

$$\text{AURC} = \int_0^1 R(f, g_\tau) d\tau$$

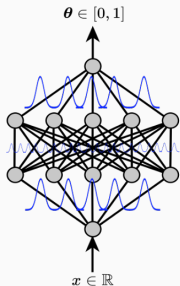


Levels of uncertainty



Use a *second-order distribution* to represent uncertainty:

$$q(\boldsymbol{\theta}|\mathbf{x}) : \mathcal{X} \rightarrow \mathbb{P}(\Theta).$$



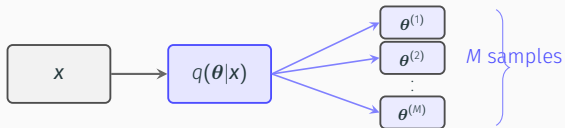
Use a *second-order distribution* to represent uncertainty:

$$q(\theta|x) : \mathcal{X} \rightarrow \mathbb{P}(\Theta).$$

Standard:



Second-order model:



Here: use Deep Ensembles, MC Dropout, Laplace approximation to sample from $q(\theta|x)$.

On which *levels* can we model uncertainty?

On the fingerprint level

Assuming **independence between bits**, we model $\theta_d \sim q_d(\theta_d|\mathbf{x})$, $d = 1, \dots, D_y$.
Calculate *uncertainty per bit*.

⁴Depeweg et al., “Decomposition of uncertainty in Bayesian deep learning for efficient and risk-sensitive learning”.

On the fingerprint level

Assuming **independence between bits**, we model $\theta_d \sim q_d(\theta_d|\mathbf{x})$, $d = 1, \dots, D_y$.
Calculate *uncertainty per bit*.

Information theoretic decomposition⁴ (bit-level)

Assume $q(\boldsymbol{\theta}|\mathbf{x}) = \prod_{d=1}^D q_d(\theta_d|\mathbf{x})$ and $y_d|\theta_d \sim \text{Bernoulli}(\theta_d)$.

$$\mathbb{H}\left[\mathbb{E}_{\theta_i \sim q_i}[y_d|\theta_d]\right] = \underbrace{\mathbb{E}_{\theta_d \sim q_d}\left[\mathbb{H}[y_d|\theta_d]\right]}_{\text{aleatoric}} + \underbrace{\mathbb{E}_{\theta_d \sim q_d}\left[D_{\text{KL}}(\text{Bern}(\theta_d) \parallel \text{Bern}(\bar{\theta}_d))\right]}_{\text{epistemic}},$$

where $\bar{\theta}_d = \mathbb{E}_{\theta_d}[\theta_d]$ is the **aggregated bit probability**.

⁴Depeweg et al., “Decomposition of uncertainty in Bayesian deep learning for efficient and risk-sensitive learning”.

On the fingerprint level

Assuming **independence between bits**, we model $\theta_d \sim q_d(\theta_d|\mathbf{x})$, $d = 1, \dots, D_y$.
Calculate *uncertainty per bit*.

Information theoretic decomposition⁴ (bit-level)

Assume $q(\boldsymbol{\theta}|\mathbf{x}) = \prod_{d=1}^D q_d(\theta_d|\mathbf{x})$ and $y_d|\theta_d \sim \text{Bernoulli}(\theta_d)$.

$$\mathbb{H}\left[\mathbb{E}_{\theta_i \sim q_i}[y_d|\theta_d]\right] = \underbrace{\mathbb{E}_{\theta_d \sim q_d}\left[\mathbb{H}[y_d|\theta_d]\right]}_{\text{aleatoric}} + \underbrace{\mathbb{E}_{\theta_d \sim q_d}\left[D_{\text{KL}}(\text{Bern}(\theta_d) \parallel \text{Bern}(\bar{\theta}_d))\right]}_{\text{epistemic}},$$

where $\bar{\theta}_d = \mathbb{E}_{\theta_d}[\theta_d]$ is the **aggregated bit probability**.

Fingerprint level aggregation: Sum over bit, i.e.

$$u_{\text{tot}}(\mathbf{x}) = \sum_{d=1}^D u_{\text{tot},d}(\mathbf{x}), \quad u_{\text{al}}(\mathbf{x}) = \sum_{d=1}^D u_{\text{al},d}(\mathbf{x}),$$
$$u_{\text{ep}}(\mathbf{x}) = \sum_{d=1}^D u_{\text{ep},d}(\mathbf{x}).$$

⁴Depeweg et al., “Decomposition of uncertainty in Bayesian deep learning for efficient and risk-sensitive learning”.

On the retrieval level

For spectrum $\mathbf{x}$ with candidate set $\mathcal{C}(\mathbf{x}) = \{c_1, \dots, c_M\}$, each sample $\boldsymbol{\theta}^{(s)} \sim q(\boldsymbol{\theta}|\mathbf{x})$ induces a *distribution over candidates*:

$$p_c(\mathbf{x}, \boldsymbol{\theta}) = \frac{\exp(\text{sim}(\boldsymbol{\theta}, c)/T)}{\sum_{c' \in \mathcal{C}(\mathbf{x})} \exp(\text{sim}(\boldsymbol{\theta}, c')/T)}, \quad c \in \mathcal{C}(\mathbf{x})$$

On the retrieval level

For spectrum $\mathbf{x}$ with candidate set $\mathcal{C}(\mathbf{x}) = \{c_1, \dots, c_M\}$, each sample $\boldsymbol{\theta}^{(s)} \sim q(\boldsymbol{\theta}|\mathbf{x})$ induces a *distribution over candidates*:

$$p_c(\mathbf{x}, \boldsymbol{\theta}) = \frac{\exp(\text{sim}(\boldsymbol{\theta}, c)/T)}{\sum_{c' \in \mathcal{C}(\mathbf{x})} \exp(\text{sim}(\boldsymbol{\theta}, c')/T)}, \quad c \in \mathcal{C}(\mathbf{x})$$

Information-theoretic decomposition (retrieval)

Let $\hat{y} \in \mathcal{C}(\mathbf{x})$ be the predicted candidate. Then:

$$\underbrace{\mathbb{H}[\bar{p}(\mathbf{x})]}_{\text{total}} = \underbrace{\mathbb{E}_{\boldsymbol{\theta} \sim q} [\mathbb{H}[p(\mathbf{x}, \boldsymbol{\theta})]]}_{\text{aleatoric}} + \underbrace{\mathbb{E}_{\boldsymbol{\theta} \sim q} [D_{KL}(p(\mathbf{x}, \boldsymbol{\theta}) \parallel \bar{p}(\mathbf{x}))]}_{\text{epistemic}}$$

where $\bar{p}_c(\mathbf{x}) = \mathbb{E}_{\boldsymbol{\theta}}[p_c(\mathbf{x}, \boldsymbol{\theta})]$ is the **aggregated candidate probability**.

Alternative: Rank-based uncertainty

Rank variance (retrieval level)

For each sample $\theta^{(s)}$, let $r^{(s)}(c)$ denote the rank of candidate c .

Define:

$$u_{\text{rank-var}}^{(k)}(\mathbf{x}) = \frac{1}{k} \sum_{c \in \text{Top-}k} \text{Var}_s[r^{(s)}(c)]$$

Alternative: Rank-based uncertainty

Rank variance (retrieval level)

For each sample $\theta^{(s)}$, let $r^{(s)}(c)$ denote the rank of candidate c .

Define:

$$u_{\text{rank-var}}^{(k)}(\mathbf{x}) = \frac{1}{k} \sum_{c \in \text{Top-}k} \text{Var}_s[r^{(s)}(c)]$$

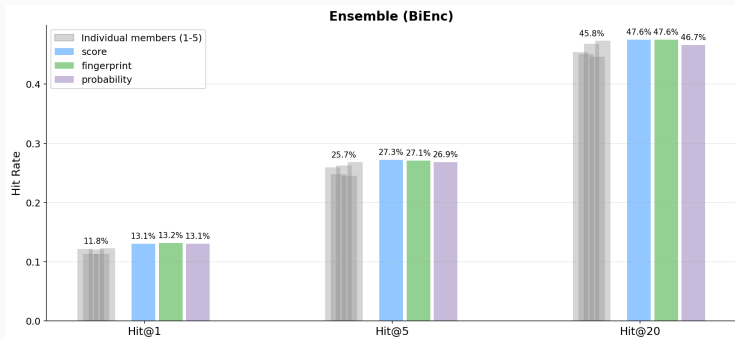
Why rank variance?

- Directly measures *ranking stability*
- Bypasses softmax compression

Experimental results

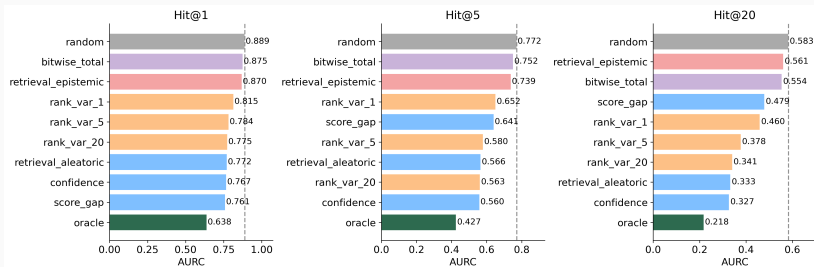
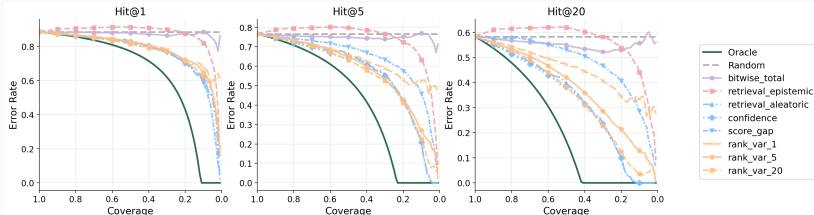
Aggregation of samples leads to improved hit rate

Example: Deep Ensemble

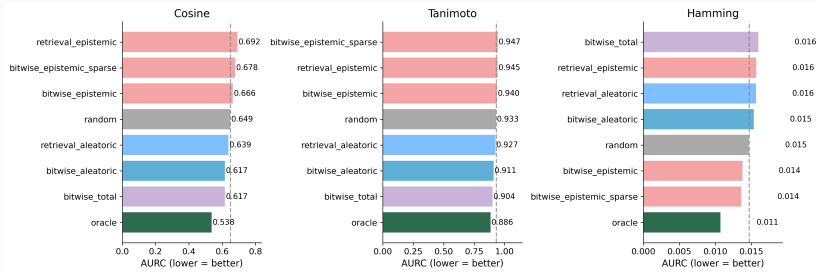
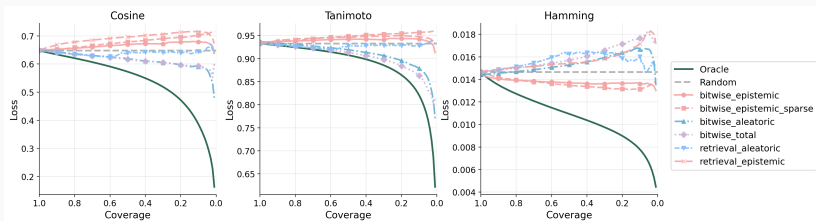


Different aggregation techniques do not seem to greatly influence performance

Retrieval: "First order" retrieval uncertainties beat "second-order" ones



Fingerprint similarity: Bitwise uncertainties most meaningful



From evaluation to guarantees

Can we guarantee a maximum error rate on the selected samples?

Selective risk control

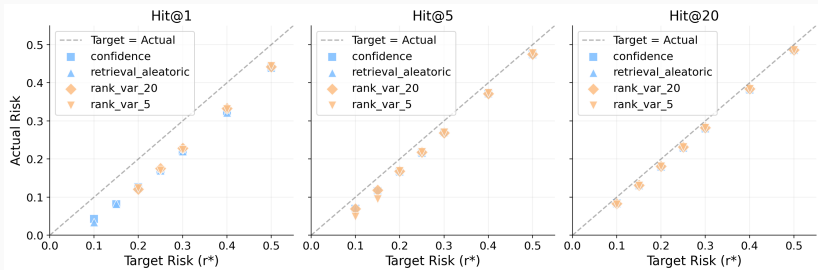
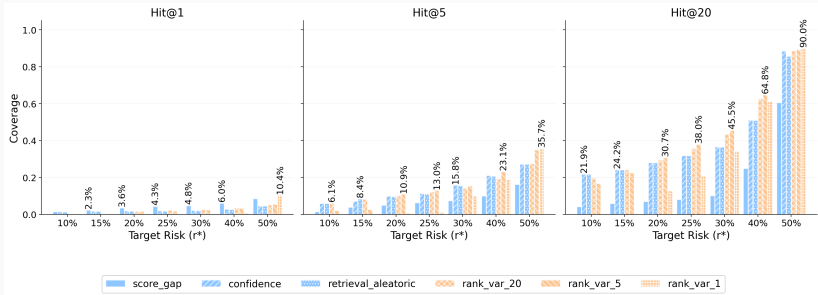
For a given target risk rate r^* and $\delta > 0$, find threshold τ^* such that

$$\mathbb{P}(R(f, g_{\tau^*}) \geq r^*) < \delta.$$

Method: Use calibration set to find τ^* via binary search and generalisation bounds⁵.

⁵Geifman and El-Yaniv, “Selective Classification for Deep Neural Networks”.

Selective risk control guarantees for hit rate



Open questions

- What are targeted uncertainty measures for the two-step retrieval task?
- How to properly integrate bitwise and retrieval uncertainty?

Conclusion

- Complex two-stage problem requires a careful modeling of uncertainty
- Similar to Hofman et al.⁶: *"One Size does not fit all"*.

⁶Hofman et al., "Uncertainty Quantification for Machine Learning: One Size Does Not Fit All".