

```html

In applied machine learning, it's common practice to compare the outputs of multiple models making predictions on the same data inputs. Yet, anyone working with ensembles, model rollouts, or multiple candidate models quickly runs into a persistent question: **why do my models disagree on the same input?** This disagreement isn't just noise or a minor inconvenience. It can be a high-signal risk indicator, highlighting edge cases, distributional shifts, data gaps, or fundamental tradeoffs embedded in objectives and losses.

In this post, I'll break down four core causes of model disagreement and explain why disagreement metrics like the **disagreement rate** and **predictive entropy** are incredibly useful tools for you. Along the way, you'll learn how to transform model disagreements from frustrating quirks into actionable signals for better model reliability and risk management.



## Overview

1. Disagreement as a High-Signal Risk Indicator
2. Edge Cases and Distribution Shift
3. Data Gaps and Subgroup Coverage
4. Objective Mismatch and Loss Function Tradeoffs
5. Measuring Disagreement: Disagreement Rate and Predictive Entropy

## Disagreement as a High-Signal Risk Indicator

When models trained on similar data disagree meaningfully on the same input, it signals uncertainty, risk, or brittleness in your system's decision-making. Intuitively, if two or more models—ideally complementary but competent—produce different predictions, the input instance likely sits in a gray zone where your system is less confident and more prone to error.

This disagreement can emerge in multiple forms:

- **Class label clashes:** disagreement on the predicted class or outcome.

- **Probability mismatches:** prediction probabilities are calibrated to different values.
- **Ranking or scoring differences:** models order options differently.

Far from being a nuisance, disagreements have actionable value:

- **Flagging high-risk or ambiguous cases for review:** humans or expert systems can focus on inputs where disagreement occurs.
- **Identifying model weaknesses:** disagreement highlights gaps in data, feature representation, or model capacity.
- **Detecting distributional shifts or emerging edge cases:** novel inputs often cause models to diverge.

In production systems, disagreement can be part of a risk scoring or alerting system that boosts operational safety and highlights when intervention is needed.

## Edge Cases and Distribution Shift

One of the most common reasons for model disagreement is that the input represents an edge case or falls outside the training distribution. Models excel at interpolating within distributions they have seen during training, but when an example radically differs—due to data drift, temporal changes, or rare scenarios—models have different ways of extrapolating and are prone to diverge.

Consider an example from lending risk scoring: a suddenly applied-for loan from a previously unseen occupation or employment type might cause one model built on historical credit data to flag high risk, while another model trained with additional alternative data disagrees. This disagreement is a clear sign that your data and models don't yet account adequately for this new population segment.

### Distribution Shift Examples

- **Temporal drift:** changing customer behavior or market conditions over time.
- **Geographical drift:** models trained on one region making errors on submissions from a new region.
- **Feature distribution shift:** changes in input feature ranges or correlations unseen during training.

Disagreement helps surface these shifts early — a critical advantage for monitoring and maintaining trust in your models. On the worst day in production, you want your system to shout, "Warning: input looks suspicious or novel!" rather than silently failing.

## Data Gaps and Subgroup Coverage

Another deep cause of disagreement lies in **data coverage gaps**—certain subgroups or regions of your feature space are underrepresented or missing in your training data. As a result, models trained independently may learn different generalizations or rely on different signals, creating divergence.

For example:

- **Minority demographic groups:** models trained on skewed samples may disagree about risk for these groups.
- **Rare event types:** fraud detection or healthcare diagnosis models may disagree when the observed pattern is infrequent.
- **Data quality issues:** missing or noisy features affect models differently depending on robustness.

Ensuring balanced and representative [reportz](#) data coverage is a foundational step towards harmonizing model predictions. In practice, disagreement analysis can highlight which subgroups systematically cause divergent predictions, providing actionable insights for data collection priorities.

## Objective Mismatch and Loss Function Tradeoffs

Sometimes models disagree because they optimize different objectives or use different loss functions—even when trained on the same data. This is what I call objective mismatch. For example, you might have:

- A model trained to maximize accuracy (or cross-entropy loss)
- An ensemble component trained to minimize a cost-sensitive loss that heavily penalizes false negatives
- A model fine-tuned on a fairness or subgroup calibration criterion
- A model optimized for ranking rather than pure classification error

These objective differences can lead to different decision boundaries and scorings, and thus to disagreement on certain inputs — particularly those near classification thresholds or carrying different cost tradeoffs.

Understanding these tradeoffs upfront is key. Rather than expecting perfect alignment between models with different goals, embrace disagreement as a reflection of competing operational priorities. Designing thresholds and risk scoring tuned to actual costs instead of “vibes” is crucial for operational success.

## Measuring Disagreement: Disagreement Rate and Predictive Entropy

Quantifying disagreement between models is the first step to leveraging it. Two common metrics are **disagreement rate** and **predictive entropy**.

### Disagreement Rate

Disagreement rate measures the proportion of inputs on which two or more models produce different predictions (typically class labels). Formally, for two models  $m_1$  and  $m_2$  and a dataset  $D$ :



Definition  $\text{DisagreementRate} = \frac{1}{|D|} \sum_{x \in D} \mathbb{1}(m_1(x) \neq m_2(x))$

This straightforward metric gives a clear, interpretable readout of how often models conflict. In practice, you can extend disagreement rate to multiple models by calculating pairwise disagreement averages or majority vote discordance.

## Predictive Entropy

Whereas disagreement rate measures differences on discrete labels, **predictive entropy** assesses the uncertainty within a probabilistic model's output distribution. For a model's predicted probability distribution  $p(y|x)$  over classes  $y$ :

Definition  $H[p(y|x)] = - \sum_y p(y|x) \log p(y|x)$

Predictive entropy reflects how "spread out" or uncertain a model's predictions are. High entropy means the model is unsure and assigns moderate probabilities across multiple classes, while low entropy implies confident predictions.

When comparing multiple models, you can examine:

- Each model's predictive entropy individually — high entropy indicates uncertainty even if models agree on labels.
- The ensemble predictive entropy — aggregating predicted probabilities across models and computing entropy over the ensemble average.
- Mutual information or difference metrics that isolate epistemic uncertainty (model disagreement) from aleatoric uncertainty (intrinsic noise).

## Things Accuracy Hides: Why Just Reporting Accuracy Isn't Enough

My pet peeve? Teams that show you a shiny accuracy number and claim their models fully capture the problem. Accuracy aggregates performance and hides the key insights disagreement exposes:

- **Which inputs are uncertain?** Accuracy doesn't tell you where models diverge.
- **Are there subgroups with systematic disagreement?** Accuracy is often dominated by majority classes, hiding rare but critical gaps.
- **How do models behave under distribution shift or noise?** Without disagreement or uncertainty metrics, you're flying blind.

Massive test set accuracy improvements can mask increasing disagreement and brittleness on the "worst day in prod." Always complement accuracy with disagreement and uncertainty analysis for a true risk-aware ML system.

## Summary and Recommendations

Model disagreement reveals crucial dimensions of risk and blind spots in your ML pipelines. To harness this signal effectively:

1. **Track disagreement rate regularly:** Include it as a core monitoring metric alongside accuracy and loss.
2. **Analyze disagreement by subgroup:** Detect data coverage gaps and representation biases.
3. **Leverage predictive entropy:** Quantify uncertainty even when models agree on labels.
4. **Consider objective mismatches:** Align or explicitly manage models with different operational priorities.
5. **Use disagreement to flag risky cases for review:** Implement alerts or human-in-the-loop staging for disagreements.

Ever notice how by embracing model disagreement thoughtfully, you transform a frustrating “why do my models disagree?” question into a powerful diagnostic and risk-management asset that keeps your applied ml system robust and aligned with real-world needs.

...