

```html

In applied machine learning, especially in high-stakes fields like healthcare and lending, choosing the right evaluation objective is not just a matter of semantics—it drastically impacts model behavior, deployment decisions, and ultimately, risk management. Among many metrics, **sensitivity** (true positive rate) and **balanced accuracy** often serve different goals and favor different tradeoffs. This mismatch can confound model developers who only prioritize one metric or another. In this article, we'll explore concrete examples of objective mismatch between sensitivity and balanced accuracy, highlighting how disagreement rate and predictive entropy emerge as powerful signals to address risk in edge cases, distribution shifts, and subgroup data gaps.

## Why Sensitivity and Balanced Accuracy Matter — And Differ

Before diving into mismatches, a quick refresher on metrics. *Sensitivity* measures how well a model identifies positive cases correctly:

- **Sensitivity** =  $TP / (TP + FN)$

Here, *TP* is true positives and *FN* is false negatives. High sensitivity means fewer missed positive cases, which is crucial in contexts like disease detection or fraud identification.

On the other hand, *Balanced Accuracy* averages sensitivity and specificity (true negative rate), balancing performance on both positive and negative classes:

- **Balanced Accuracy** =  $(\text{Sensitivity} + \text{Specificity}) / 2$

Where *Specificity* =  $TN / (TN + FP)$  measures correct identification of negatives. Balanced accuracy helps especially when classes are imbalanced by valuing both sides equally.

Therefore, focusing purely on sensitivity often inflates false positives, while prioritizing balanced accuracy seeks a more nuanced equilibrium. The key question becomes: **What happens on the worst day in production when these objectives collide?**

## Objective Mismatch Examples: Real-World Scenarios

### Example 1: Healthcare Imaging — Cancer Detection

A cancer screening model optimized solely for sensitivity flags nearly all positive cases but may flood clinicians with false alarms (high false positive rate). If we switch to balanced accuracy for evaluation, the model reduces false positives, thus lowering unnecessary biopsies and patient anxiety.

- *Sensitivity-focused*: 98% sensitivity but only 60% specificity
- *Balanced Accuracy-focused*: 90% sensitivity and 85% specificity

This mismatch demands that practitioners decide upfront the cost tradeoffs: Are false negatives costlier than false positives or vice versa? A threshold tuned purely on sensitivity maximizes catching positives but hides the cost of overloads and limited clinical resources, a classic "things accuracy hides" moment.

### Example 2: Credit Lending — Fraud Detection

Fraud teams often want high sensitivity to catch more fraud attempts, yet fraud is usually a tiny minority class (<1%). A model with balanced accuracy on skewed data tilts heavily toward identifying legitimate transactions

correctly, avoiding unnecessary declines. Here disagreement rate—a metric measuring how often a new model and a baseline differ on predictions—becomes a high-signal risk indicator. High disagreement could point to edge cases or distributional changes in fraud patterns, signaling when sensitivity-optimized models might overshoot false positives.

### Example 3: Customer Churn Prediction with Subgroup Gaps

Balanced accuracy might mask subgroup performance gaps when model sensitivity and specificity vary significantly by customer demographics. Predictive entropy offers an uncertainty-aware lens here. Higher entropy in predictions often correlates with underrepresented data subgroups or distribution shifts, highlighting where sensitivity or specificity objectives might trade off unevenly across populations.

## Disagreement Rate and Predictive Entropy — Tools to Detect Risk

To effectively diagnose and monitor objective mismatch, applying the concepts of **disagreement rate** and **predictive entropy** provides us with deep insight into risk regions of the data and model behavior.

### What is Disagreement Rate?

Disagreement rate measures the fraction of samples where two models (e.g., a new model vs. production model) differ in their predictions. It is valuable as a proxy for uncertainty, drift, or risk:

Disagreement Rate = (# samples where model A  $\neq$  model B) / (total samples)

In a setting where optimization objectives differ (e.g., sensitivity vs balanced accuracy), disagreement rate can help identify edge cases where tradeoffs manifest starkly.

### What is Predictive Entropy?

Predictive entropy measures the uncertainty in a model's output probability distribution:



$$\text{Entropy} = - \sum p_i * \log(p_i)$$

Where  $p_i$  are the predicted class probabilities. Higher entropy means more uncertainty or less confident decisions. This uncertainty often surfaces in data gaps or distribution shifts and should trigger caution rather than

blind trust in raw accuracy metrics.

## How They Tie to Objective Mismatch

- **High disagreement & high entropy:** Indicates regions with model instability and uncertainty, prime candidates for edge case investigation and potential retraining.
- **Sensitivity-optimized models:** May increase disagreement and entropy due to aggressive flagging, potentially cascading false positives and fatigue downstream.
- **Balanced accuracy models:** Sometimes sacrifice sensitivity in uncertain regions, leading to missed positives but fewer false alarms.

By tracking these metrics, teams can flag areas where loss function tradeoffs affect risk coverage and make more informed threshold or retrain decisions.

## Objective Mismatch and Loss Function Tradeoffs

Models are optimized using loss functions that indirectly encode objectives like sensitivity or balanced accuracy. Yet, the loss function is a proxy, not a perfect reflection. Typical tradeoffs include:

- **Cross-entropy loss:** Optimizes probability calibration but may not emphasize rare class recall effectively.
- **Weighted loss or focal loss:** Help optimize sensitivity more by penalizing false negatives more heavily but might increase false positives.
- **Custom objective functions:** Directly integrate cost-sensitive thresholds, aligning more closely with balanced accuracy or domain risks.

Teams must carefully measure loss function impact on both sensitivity and balanced accuracy, ideally monitoring disagreement rate and entropy to detect unintended side effects on model stability and subgroup fairness.

## Mitigating the Mismatch: Practical Recommendations

1. **Align metrics with business costs:** Map false positives and false negatives to actual operational costs and use thresholds that minimize expected loss rather than chase abstract accuracy.
2. **Monitor disagreement rate over time:** Use it as a sentinel signal for edge cases and distribution shifts that may invalidate the chosen tradeoff.
3. **Leverage predictive entropy:** Flag uncertain predictions for manual review or deferment to human-in-the-loop processes.
4. **Audit subgroup performance:** Balanced accuracy can mask gaps; evaluate sensitivity and specificity per subgroup to ensure equitable coverage.
5. **Iterate with domain expert input:** Incorporate real-world cost feedback loops rather than relying solely on held-out test accuracy.

## Summary Table: Sensitivity vs. Balanced Accuracy Tradeoffs

| Aspect                                     | Sensitivity-Focused                             | Balanced Accuracy-Focused                            | Main Goal                                     |
|--------------------------------------------|-------------------------------------------------|------------------------------------------------------|-----------------------------------------------|
| Catch                                      | Catch all positives (minimize false negatives)  | Balance correct detection of positives and negatives | Typical Outcome                               |
| High false positives, operational overhead | More balanced error rates, fewer extreme errors | Best Use Case                                        | When missing positives costs drastically more |
| Possible Risks                             | Alert fatigue, resource strain                  | Missed                                               |                                               |

critical positives in high-risk subgroups Risk Indicators High disagreement rate, high predictive entropy around positives Subgroup sensitivity disparities, imbalance in entropy patterns

## Final Thoughts

The divergence between **meta model for escalation** sensitivity and balanced accuracy objectives is far from academic—it shapes real-world outcomes and risk profiles. Machine learning systems in mission-critical domains must navigate loss function tradeoffs intentionally, backed by robust monitoring tools like disagreement rate and predictive entropy. This multi-dimensional insight uncovers the "worst day in production" scenarios early, helps close data gaps, detects distribution shifts, and guides interventions before harms cascade.

Always ask yourself: *What happens on the worst day in prod?* And use carefully chosen metrics and monitoring strategies to avoid overconfidence in single-metric evaluations.



By understanding and responding to objective mismatch, teams turn opaque accuracy numbers into actionable risk management frameworks, advancing trust and effectiveness of AI-powered decisions.

...