

Calibration as Decision Alignment

Aaron Roth

University of Pennsylvania

Fall 2026

Overview

- ▶ A probability forecast is an interface between prediction and decision making: a user uses it to inform a choice of action.

Overview

- ▶ A probability forecast is an interface between prediction and decision making: a user uses it to inform a choice of action.
- ▶ Calibration is a self-consistency requirement: forecasts must be unbiased conditional on their own value.

Overview

- ▶ A probability forecast is an interface between prediction and decision making: a user uses it to inform a choice of action.
- ▶ Calibration is a self-consistency requirement: forecasts must be unbiased conditional on their own value.
- ▶ It makes taking the forecast at face value at least as good as any alternative rule using only the forecast.

Forecasts Are Decision Interfaces

In the first part of the course, an alignment procedure produced the policy whose welfare we evaluated:

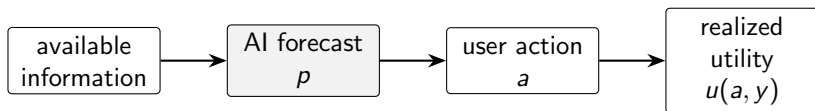
human comparisons $\longrightarrow$ policy.

Forecasts Are Decision Interfaces

In the first part of the course, an alignment procedure produced the policy whose welfare we evaluated:

human comparisons $\longrightarrow$ policy.

We now study systems whose output informs a later decision:

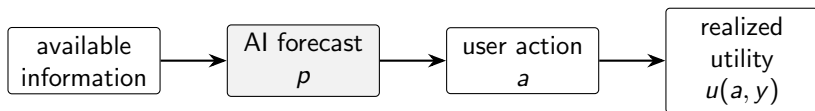


Forecasts Are Decision Interfaces

In the first part of the course, an alignment procedure produced the policy whose welfare we evaluated:

human comparisons $\longrightarrow$ policy.

We now study systems whose output informs a later decision:



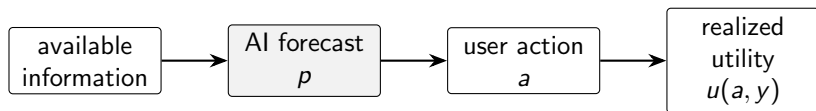
The user's utility and action set are now taken as given. We ask whether the AI's report supports that decision.

Forecasts Are Decision Interfaces

In the first part of the course, an alignment procedure produced the policy whose welfare we evaluated:

human comparisons $\longrightarrow$ policy.

We now study systems whose output informs a later decision:



The user's utility and action set are now taken as given. We ask whether the AI's report supports that decision.

What property lets downstream decision makers safely treat forecasts as if they are correct?

The User Maximizes Expected Utility Under the Forecast

Let $\mathcal{Y} = \{1, \dots, d\}$ be a finite outcome set and let $p \in \Delta(\mathcal{Y})$ be a forecast. A user has a finite action set $\mathcal{A}$ and utility $u(a, y)$.

The User Maximizes Expected Utility Under the Forecast

Let $\mathcal{Y} = \{1, \dots, d\}$ be a finite outcome set and let $p \in \Delta(\mathcal{Y})$ be a forecast. A user has a finite action set $\mathcal{A}$ and utility $u(a, y)$.

Represent action a by the vector

$$u_a = (u(a, 1), \dots, u(a, d)) \in \mathbb{R}^d.$$

Its expected utility under outcome distribution p is

$$\mathbb{E}_{Y \sim p}[u(a, Y)] = \sum_{y \in \mathcal{Y}} p_y u(a, y) = \langle u_a, p \rangle.$$

The User Maximizes Expected Utility Under the Forecast

Let $\mathcal{Y} = \{1, \dots, d\}$ be a finite outcome set and let $p \in \Delta(\mathcal{Y})$ be a forecast. A user has a finite action set $\mathcal{A}$ and utility $u(a, y)$.

Represent action a by the vector

$$u_a = (u(a, 1), \dots, u(a, d)) \in \mathbb{R}^d.$$

Its expected utility under outcome distribution p is

$$\mathbb{E}_{Y \sim p}[u(a, Y)] = \sum_{y \in \mathcal{Y}} p_y u(a, y) = \langle u_a, p \rangle.$$

Fixing a tie-breaking rule, the user's best response action is

$$\text{BR}_u(p) \in \arg \max_{a \in \mathcal{A}} \langle u_a, p \rangle.$$

Running Example: Deploy, Hold, or Audit

A proposal is either safe or harmful. Deploying a safe proposal gives utility 1, while deploying a harmful proposal gives utility -4 .

Running Example: Deploy, Hold, or Audit

A proposal is either safe or harmful. Deploying a safe proposal gives utility 1, while deploying a harmful proposal gives utility -4 . The user may instead hold, or pay $1/5$ for a perfect audit and then deploy exactly when the proposal is safe:

	safe	harmful
deploy	1	-4
hold	0	0
audit	$4/5$	$-1/5$

Running Example: Deploy, Hold, or Audit

A proposal is either safe or harmful. Deploying a safe proposal gives utility 1, while deploying a harmful proposal gives utility -4 . The user may instead hold, or pay $1/5$ for a perfect audit and then deploy exactly when the proposal is safe:

	safe	harmful
deploy	1	-4
hold	0	0
audit	$4/5$	$-1/5$

Let r denote the reported probability that the proposal is safe, so the forecast vector is $p = (r, 1 - r)$.

Each Forecast Induces One Best Action

Write $BR_u(r)$ as shorthand for $BR_u((r, 1 - r))$.

Each Forecast Induces One Best Action

Write $BR_u(r)$ as shorthand for $BR_u((r, 1 - r))$.

Under safe probability r , the three expected utilities are

$$U(\text{deploy}; r) = r(1) + (1 - r)(-4) = 5r - 4,$$

$$U(\text{hold}; r) = 0,$$

$$U(\text{audit}; r) = r\left(\frac{4}{5}\right) + (1 - r)\left(-\frac{1}{5}\right) = r - \frac{1}{5}.$$

Each Forecast Induces One Best Action

Write $BR_u(r)$ as shorthand for $BR_u((r, 1 - r))$.

Under safe probability r , the three expected utilities are

$$U(\text{deploy}; r) = r(1) + (1 - r)(-4) = 5r - 4,$$

$$U(\text{hold}; r) = 0,$$

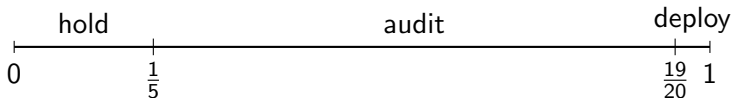
$$U(\text{audit}; r) = r\left(\frac{4}{5}\right) + (1 - r)\left(-\frac{1}{5}\right) = r - \frac{1}{5}.$$

Therefore, away from the two ties,

$$BR_u(r) = \begin{cases} \text{hold}, & r < 1/5, \\ \text{audit}, & 1/5 < r < 19/20, \\ \text{deploy}, & r > 19/20. \end{cases}$$

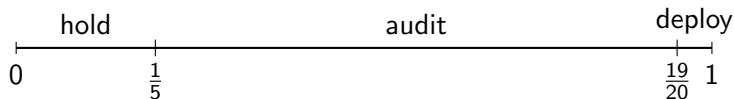
For This User, Forecasts Matter Through Decision Regions

The numerical report partitions the probability interval into three regions:



For This User, Forecasts Matter Through Decision Regions

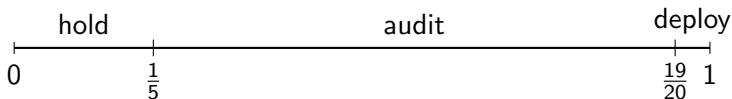
The numerical report partitions the probability interval into three regions:



For this user, changing a forecast from 0.30 to 0.50 does not change the action: both reports lead to an audit.

For This User, Forecasts Matter Through Decision Regions

The numerical report partitions the probability interval into three regions:



For this user, changing a forecast from 0.30 to 0.50 does not change the action: both reports lead to an audit.

Changing a forecast from 0.19 to 0.21 does change the action.

Errors near a decision boundary can therefore matter more than equally large errors elsewhere.

Unbiased Forecasting On Average Is Not Enough

Consider equally many rounds of two kinds:

reported safe probability	outcome	chosen action	utility
0.99	harmful	deploy	-4
0.01	safe	hold	0

Unbiased Forecasting On Average Is Not Enough

Consider equally many rounds of two kinds:

reported safe probability	outcome	chosen action	utility
0.99	harmful	deploy	-4
0.01	safe	hold	0

Both the average forecast and the realized safe frequency equal $1/2$:

$$\frac{0.99 + 0.01}{2} = \frac{0 + 1}{2} = \frac{1}{2}.$$

Unbiased Forecasting On Average Is Not Enough

Consider equally many rounds of two kinds:

reported safe probability	outcome	chosen action	utility
0.99	harmful	deploy	-4
0.01	safe	hold	0

Both the average forecast and the realized safe frequency equal $1/2$:

$$\frac{0.99 + 0.01}{2} = \frac{0 + 1}{2} = \frac{1}{2}.$$

Yet the user deploys every harmful proposal and holds every safe proposal. The errors cancel only because we averaged across reports that induce different decisions.

Calibration Conditions on What Was Reported

We work with a fixed realized sequence

$$(p_1, y_1), \dots, (p_T, y_T), \quad p_t \in \mathcal{P} \subset \Delta(\mathcal{Y}), \quad y_t \in \mathcal{Y},$$

where the finite set $\mathcal{P}$ contains the possible reports.

Calibration Conditions on What Was Reported

We work with a fixed realized sequence

$$(p_1, y_1), \dots, (p_T, y_T), \quad p_t \in \mathcal{P} \subset \Delta(\mathcal{Y}), \quad y_t \in \mathcal{Y},$$

where the finite set $\mathcal{P}$ contains the possible reports.

Let $N_p := |\{t : p_t = p\}|$. For a report with $N_p > 0$, let e_y denote the point mass on outcome y and define the empirical outcome distribution in that report cell:

$$q_p := \frac{1}{N_p} \sum_{t: p_t = p} e_{y_t}.$$

Calibration Conditions on What Was Reported

We work with a fixed realized sequence

$$(p_1, y_1), \dots, (p_T, y_T), \quad p_t \in \mathcal{P} \subset \Delta(\mathcal{Y}), \quad y_t \in \mathcal{Y},$$

where the finite set $\mathcal{P}$ contains the possible reports.

Let $N_p := |\{t : p_t = p\}|$. For a report with $N_p > 0$, let e_y denote the point mass on outcome y and define the empirical outcome distribution in that report cell:

$$q_p := \frac{1}{N_p} \sum_{t: p_t = p} e_{y_t}.$$

The sequence is exactly calibrated when $q_p = p$ for every report that occurs. For binary outcomes: among rounds assigned probability r , the positive outcome occurs with frequency r .

Report-Cell Bias Measures Approximate Calibration

It is convenient to avoid dividing by the size of a report cell.
Define its unnormalized bias

$$B_p := \sum_{t:p_t=p} (e_{y_t} - p).$$

When $N_p > 0$, this is equivalently $B_p = N_p(q_p - p)$. Exact calibration means $B_p = 0$ for every $p \in \mathcal{P}$.

Report-Cell Bias Measures Approximate Calibration

It is convenient to avoid dividing by the size of a report cell.
Define its unnormalized bias

$$B_p := \sum_{t:p_t=p} (e_{y_t} - p).$$

When $N_p > 0$, this is equivalently $B_p = N_p(q_p - p)$. Exact calibration means $B_p = 0$ for every $p \in \mathcal{P}$.

The average sequential ℓ_1 calibration error is

$$\text{Cal}_{1,T}^{\text{seq}} := \frac{1}{T} \sum_{p \in \mathcal{P}} \|B_p\|_1.$$

Report-Cell Bias Measures Approximate Calibration

It is convenient to avoid dividing by the size of a report cell.
Define its unnormalized bias

$$B_p := \sum_{t:p_t=p} (e_{y_t} - p).$$

When $N_p > 0$, this is equivalently $B_p = N_p(q_p - p)$. Exact calibration means $B_p = 0$ for every $p \in \mathcal{P}$.

The average sequential ℓ_1 calibration error is

$$\text{Cal}_{1,T}^{\text{seq}} := \frac{1}{T} \sum_{p \in \mathcal{P}} \|B_p\|_1.$$

Population analogue: if $\hat{P}$ is a random forecast, calibration requires

$$\mathbb{E}[e_Y \mid \hat{P} = p] = p \quad \text{for every } p \text{ in the support of } \hat{P}.$$

Could Another Rule Use the Same Report Better?

On round t , the best response policy chooses

$$a_t = \text{BR}_u(p_t).$$

Could Another Rule Use the Same Report Better?

On round t , the best response policy chooses

$$a_t = \text{BR}_u(p_t).$$

A *reinterpretation* is any rule

$$\pi : \mathcal{P} \rightarrow \mathcal{A}.$$

It sees the same forecast p_t but may choose a different action $\pi(p_t)$.

Could Another Rule Use the Same Report Better?

On round t , the best response policy chooses

$$a_t = \text{BR}_u(p_t).$$

A *reinterpretation* is any rule

$$\pi : \mathcal{P} \rightarrow \mathcal{A}.$$

It sees the same forecast p_t but may choose a different action $\pi(p_t)$.

Its average realized gain over the best response policy is

$$R_T^{\text{rep}}(\pi) := \frac{1}{T} \sum_{t=1}^T (u(\pi(p_t), y_t) - u(a_t, y_t)).$$

Is reinterpretation ever worthwhile?

First Group the Realized Gain by Report

Write $a(p) = \text{BR}_u(p)$. Since $u(a, y) = \langle u_a, e_y \rangle$,

$$\begin{aligned} TR_T^{\text{rep}}(\pi) &= \sum_{t=1}^T \langle u_{\pi(p_t)} - u_{a(p_t)}, e_{y_t} \rangle \\ &= \sum_{p \in \mathcal{P}} \left\langle u_{\pi(p)} - u_{a(p)}, \sum_{t: p_t=p} e_{y_t} \right\rangle. \end{aligned}$$

The definition of B_p gives

$$\sum_{t: p_t=p} e_{y_t} = N_p p + B_p.$$

First Group the Realized Gain by Report

Write $a(p) = \text{BR}_u(p)$. Since $u(a, y) = \langle u_a, e_y \rangle$,

$$\begin{aligned} TR_T^{\text{rep}}(\pi) &= \sum_{t=1}^T \langle u_{\pi(p_t)} - u_{a(p_t)}, e_{y_t} \rangle \\ &= \sum_{p \in \mathcal{P}} \left\langle u_{\pi(p)} - u_{a(p)}, \sum_{t: p_t=p} e_{y_t} \right\rangle. \end{aligned}$$

First Group the Realized Gain by Report

Write $a(p) = \text{BR}_u(p)$. Since $u(a, y) = \langle u_a, e_y \rangle$,

$$\begin{aligned} TR_T^{\text{rep}}(\pi) &= \sum_{t=1}^T \langle u_{\pi(p_t)} - u_{a(p_t)}, e_{y_t} \rangle \\ &= \sum_{p \in \mathcal{P}} \left\langle u_{\pi(p)} - u_{a(p)}, \sum_{t: p_t=p} e_{y_t} \right\rangle. \end{aligned}$$

The definition of B_p gives

$$\sum_{t: p_t=p} e_{y_t} = N_p p + B_p.$$

Optimality Handles Forecasts; Calibration Handles Errors

Substituting $\sum_{t:p_t=p} e_{y_t} = N_p p + B_p$ gives

$$\begin{aligned} TR_T^{\text{rep}}(\pi) = & \underbrace{\sum_{p \in \mathcal{P}} N_p \langle u_{\pi(p)} - u_{a(p)}, p \rangle}_{\text{comparison in the forecasted world}} \\ & + \underbrace{\sum_{p \in \mathcal{P}} \langle u_{\pi(p)} - u_{a(p)}, B_p \rangle}_{\text{cost of replacing forecasts by outcomes}} . \end{aligned}$$

Optimality Handles Forecasts; Calibration Handles Errors

Substituting $\sum_{t:p_t=p} e_{y_t} = N_p p + B_p$ gives

$$\begin{aligned} TR_T^{\text{rep}}(\pi) = & \underbrace{\sum_{p \in \mathcal{P}} N_p \langle u_{\pi(p)} - u_{a(p)}, p \rangle}_{\text{comparison in the forecasted world}} \\ & + \underbrace{\sum_{p \in \mathcal{P}} \langle u_{\pi(p)} - u_{a(p)}, B_p \rangle}_{\text{cost of replacing forecasts by outcomes}} . \end{aligned}$$

The first sum is nonpositive because $a(p) = \text{BR}_u(p)$ maximizes $\langle u_a, p \rangle$.

Calibration Controls the Remaining Error

Define the largest coordinatewise utility difference between two actions:

$$D_u := \max_{a,b \in \mathcal{A}} \|u_a - u_b\|_\infty.$$

Calibration Controls the Remaining Error

Define the largest coordinatewise utility difference between two actions:

$$D_u := \max_{a,b \in \mathcal{A}} \|u_a - u_b\|_\infty.$$

For each report cell, Hölder's inequality gives

$$\begin{aligned} \langle u_{\pi(p)} - u_{a(p)}, B_p \rangle &\leq |\langle u_{\pi(p)} - u_{a(p)}, B_p \rangle| \\ &\leq \|u_{\pi(p)} - u_{a(p)}\|_\infty \|B_p\|_1 \\ &\leq D_u \|B_p\|_1. \end{aligned}$$

Calibration Controls the Remaining Error

Define the largest coordinatewise utility difference between two actions:

$$D_u := \max_{a,b \in \mathcal{A}} \|u_a - u_b\|_\infty.$$

For each report cell, Hölder's inequality gives

$$\begin{aligned} \langle u_{\pi(p)} - u_{a(p)}, B_p \rangle &\leq |\langle u_{\pi(p)} - u_{a(p)}, B_p \rangle| \\ &\leq \|u_{\pi(p)} - u_{a(p)}\|_\infty \|B_p\|_1 \\ &\leq D_u \|B_p\|_1. \end{aligned}$$

Therefore

$$TR_T^{\text{rep}}(\pi) \leq D_u \sum_{p \in \mathcal{P}} \|B_p\|_1.$$

Calibration Rules Out Profitable Reinterpretation

Theorem

For every fixed forecast and outcome sequence and every reinterpretation $\pi : \mathcal{P} \rightarrow \mathcal{A}$,

$$R_T^{\text{rep}}(\pi) \leq D_u \text{Cal}_{1,T}^{\text{seq}}.$$

Calibration Rules Out Profitable Reinterpretation

Theorem

For every fixed forecast and outcome sequence and every reinterpretation $\pi : \mathcal{P} \rightarrow \mathcal{A}$,

$$R_T^{\text{rep}}(\pi) \leq D_u \text{Cal}_{1,T}^{\text{seq}}.$$

If the reports are exactly calibrated, then

$$R_T^{\text{rep}}(\pi) \leq 0 \quad \text{for every } \pi.$$

No alternative rule seeing only the same report can obtain higher realized average utility.

Calibration Rules Out Profitable Reinterpretation

Theorem

For every fixed forecast and outcome sequence and every reinterpretation $\pi : \mathcal{P} \rightarrow \mathcal{A}$,

$$R_T^{\text{rep}}(\pi) \leq D_u \text{Cal}_{1,T}^{\text{seq}}.$$

If the reports are exactly calibrated, then

$$R_T^{\text{rep}}(\pi) \leq 0 \quad \text{for every } \pi.$$

No alternative rule seeing only the same report can obtain higher realized average utility.

Because full calibration does not mention u , the conclusion holds for every downstream utility function and action set.

Calibration Promises Interpretation, Not Informativeness

Suppose half the proposals are safe and half harmful, and the forecast always reports $r = 1/2$. The forecast is exactly calibrated.

Calibration Promises Interpretation, Not Informativeness

Suppose half the proposals are safe and half harmful, and the forecast always reports $r = 1/2$. The forecast is exactly calibrated. The best response policy audits every proposal and obtains average realized utility

$$\frac{1}{2} - \frac{1}{5} = \frac{3}{10}.$$

Calibration Promises Interpretation, Not Informativeness

Suppose half the proposals are safe and half harmful, and the forecast always reports $r = 1/2$. The forecast is exactly calibrated. The best response policy audits every proposal and obtains average realized utility

$$\frac{1}{2} - \frac{1}{5} = \frac{3}{10}.$$

Now suppose an additional feature x_t reveals whether the proposal is safe. A contextual rule that sees x_t deploys exactly on safe rounds and holds on harmful rounds, obtaining average utility $1/2$.

Calibration Promises Interpretation, Not Informativeness

Suppose half the proposals are safe and half harmful, and the forecast always reports $r = 1/2$. The forecast is exactly calibrated. The best response policy audits every proposal and obtains average realized utility

$$\frac{1}{2} - \frac{1}{5} = \frac{3}{10}.$$

Now suppose an additional feature x_t reveals whether the proposal is safe. A contextual rule that sees x_t deploys exactly on safe rounds and holds on harmful rounds, obtaining average utility $1/2$. There is no contradiction: the theorem compares only with reinterpretations of the constant report. Calibration cannot recover information not present in the forecast.

Errors Can Cancel Within One Decision Region

Recall that every forecast in the interval

$$\frac{1}{5} < r < \frac{19}{20}$$

leads the best response policy to audit.

Errors Can Cancel Within One Decision Region

Recall that every forecast in the interval

$$\frac{1}{5} < r < \frac{19}{20}$$

leads the best response policy to audit.

Suppose equal numbers of rounds receive the following reports:

report r	empirical safe frequency	action	Prediction conditional bias
0.30	0.20	audit	-0.10 per round
0.50	0.60	audit	+0.10 per round

Errors Can Cancel Within One Decision Region

Recall that every forecast in the interval

$$\frac{1}{5} < r < \frac{19}{20}$$

leads the best response policy to audit.

Suppose equal numbers of rounds receive the following reports:

report r	empirical safe frequency	action	Prediction conditional bias
0.30	0.20	audit	-0.10 per round
0.50	0.60	audit	+0.10 per round

The prediction conditional bias cancels after the two report cells are combined. That cancellation is harmless for a user who takes the same action on both cells.

A Weaker Form of Hindsight Criticism

A report reinterpretation may remap every forecast value even if best response maps both to the same action. A weaker action remapping comparator only remaps chosen actions.

A Weaker Form of Hindsight Criticism

A report reinterpretation may remap every forecast value even if best response maps both to the same action. A weaker action remapping comparator only remaps chosen actions.

An *action remapping* is a function

$$\phi : \mathcal{A} \rightarrow \mathcal{A}.$$

It asks, for example:

Whenever I audited, would I have done better by deploying instead?

A Weaker Form of Hindsight Criticism

A report reinterpretation may remap every forecast value even if best response maps both to the same action. A weaker action remapping comparator only remaps chosen actions.

An *action remapping* is a function

$$\phi : \mathcal{A} \rightarrow \mathcal{A}.$$

It asks, for example:

Whenever I audited, would I have done better by deploying instead?

Its average gain is

$$R_T(\phi) := \frac{1}{T} \sum_{t=1}^T (u(\phi(a_t), y_t) - u(a_t, y_t)).$$

Swap Regret Tests Every Action Remapping

The user's average swap regret is

$$\text{SwapReg}_T := \max_{\phi: \mathcal{A} \rightarrow \mathcal{A}} R_T(\phi).$$

Swap Regret Tests Every Action Remapping

The user's average swap regret is

$$\text{SwapReg}_T := \max_{\phi: \mathcal{A} \rightarrow \mathcal{A}} R_T(\phi).$$

Because the identity remapping $\phi(a) = a$ is allowed,

$$\text{SwapReg}_T \geq 0.$$

Swap Regret Tests Every Action Remapping

The user's average swap regret is

$$\text{SwapReg}_T := \max_{\phi: \mathcal{A} \rightarrow \mathcal{A}} R_T(\phi).$$

Because the identity remapping $\phi(a) = a$ is allowed,

$$\text{SwapReg}_T \geq 0.$$

A small value means that hindsight reveals no systematic criticism of the form

Whenever the forecast led me to action a , I should instead have taken $\phi(a)$.

Constant remappings $\phi(a) \equiv b$ are included, so this benchmark is only stronger than the fixed-action comparisons measured by external regret in Lecture 5.

Decision Calibration Uses Action Cells

For every action $a \in \mathcal{A}$, define its action-cell bias

$$B_a := \sum_{t: a_t = a} (e_{y_t} - p_t).$$

Decision Calibration Uses Action Cells

For every action $a \in \mathcal{A}$, define its action-cell bias

$$B_a := \sum_{t: a_t = a} (e_{y_t} - p_t).$$

The sequence is *decision calibrated* for this user when

$$B_a = 0 \quad \text{for every } a \in \mathcal{A}.$$

Decision Calibration Uses Action Cells

For every action $a \in \mathcal{A}$, define its action-cell bias

$$B_a := \sum_{t: a_t = a} (e_{y_t} - p_t).$$

The sequence is *decision calibrated* for this user when

$$B_a = 0 \quad \text{for every } a \in \mathcal{A}.$$

This condition allows errors from distinct reports to cancel—but only when those reports induced the same action for this user. For any nonempty action cell, this says that the average forecast equals the empirical outcome distribution within the rounds inducing that action.

Decision Calibration Uses Action Cells

For every action $a \in \mathcal{A}$, define its action-cell bias

$$B_a := \sum_{t: a_t = a} (e_{y_t} - p_t).$$

The sequence is *decision calibrated* for this user when

$$B_a = 0 \quad \text{for every } a \in \mathcal{A}.$$

This condition allows errors from distinct reports to cancel—but only when those reports induced the same action for this user. For any nonempty action cell, this says that the average forecast equals the empirical outcome distribution within the rounds inducing that action. Exact calibration is norm-independent. For approximate guarantees we use $\|B_a\|_2$, matching the Euclidean approachability algorithm in the next lecture.

Full Calibration Implies Decision Calibration

Because the action is a deterministic function of the report,

$$\{t : a_t = a\} = \bigcup_{p: \text{BR}_u(p)=a} \{t : p_t = p\}.$$

The report cells in this union are disjoint, so

$$\begin{aligned} B_a &= \sum_{t: a_t=a} (e_{y_t} - p_t) \\ &= \sum_{p: \text{BR}_u(p)=a} \sum_{t: p_t=p} (e_{y_t} - p) \\ &= \sum_{p: \text{BR}_u(p)=a} B_p. \end{aligned}$$

Full Calibration Implies Decision Calibration

Because the action is a deterministic function of the report,

$$\{t : a_t = a\} = \bigcup_{p: \text{BR}_u(p)=a} \{t : p_t = p\}.$$

The report cells in this union are disjoint, so

$$\begin{aligned} B_a &= \sum_{t: a_t=a} (e_{y_t} - p_t) \\ &= \sum_{p: \text{BR}_u(p)=a} \sum_{t: p_t=p} (e_{y_t} - p) \\ &= \sum_{p: \text{BR}_u(p)=a} B_p. \end{aligned}$$

Thus $B_p = 0$ for every report implies $B_a = 0$ for every action. The 0.30/0.50 example shows that the converse fails.

The Same Transfer Argument Now Groups by Action

Fix an action remapping ϕ . Its total realized gain is

$$TR_T(\phi) = \sum_{t=1}^T \langle u_{\phi(a_t)} - u_{a_t}, e_{y_t} \rangle$$

The Same Transfer Argument Now Groups by Action

Fix an action remapping ϕ . Its total realized gain is

$$TR_T(\phi) = \sum_{t=1}^T \langle u_{\phi(a_t)} - u_{a_t}, e_{y_t} \rangle$$

Write $e_{y_t} = p_t + (e_{y_t} - p_t)$ to separate the forecasted comparison from the forecast error:

$$\begin{aligned} TR_T(\phi) = & \underbrace{\sum_{t=1}^T \langle u_{\phi(a_t)} - u_{a_t}, p_t \rangle}_{\text{comparison in the forecasted world}} \\ & + \underbrace{\sum_{t=1}^T \langle u_{\phi(a_t)} - u_{a_t}, e_{y_t} - p_t \rangle}_{\text{cost of replacing forecasts by outcomes}} . \end{aligned}$$

The Same Transfer Argument Now Groups by Action

Fix an action remapping ϕ . Its total realized gain is

$$TR_T(\phi) = \sum_{t=1}^T \langle u_{\phi(a_t)} - u_{a_t}, e_{y_t} \rangle$$

Write $e_{y_t} = p_t + (e_{y_t} - p_t)$ to separate the forecasted comparison from the forecast error:

$$\begin{aligned} TR_T(\phi) = & \underbrace{\sum_{t=1}^T \langle u_{\phi(a_t)} - u_{a_t}, p_t \rangle}_{\text{comparison in the forecasted world}} \\ & + \underbrace{\sum_{t=1}^T \langle u_{\phi(a_t)} - u_{a_t}, e_{y_t} - p_t \rangle}_{\text{cost of replacing forecasts by outcomes}} . \end{aligned}$$

Because $a_t = \text{BR}_u(p_t)$, every term in the first sum is nonpositive.

Action-Cell Bias Controls the Remaining Error

On every round in the action- a cell, the utility difference $u_{\phi(a_t)} - u_{a_t}$ equals the fixed vector $u_{\phi(a)} - u_a$. Therefore

$$\begin{aligned} & \sum_{t=1}^T \langle u_{\phi(a_t)} - u_{a_t}, e_{y_t} - p_t \rangle \\ &= \sum_{a \in \mathcal{A}} \left\langle u_{\phi(a)} - u_a, \sum_{t: a_t = a} (e_{y_t} - p_t) \right\rangle \\ &= \sum_{a \in \mathcal{A}} \langle u_{\phi(a)} - u_a, B_a \rangle. \end{aligned}$$

Action-Cell Bias Controls the Remaining Error

On every round in the action- a cell, the utility difference $u_{\phi(a_t)} - u_{a_t}$ equals the fixed vector $u_{\phi(a)} - u_a$. Therefore

$$\begin{aligned} & \sum_{t=1}^T \langle u_{\phi(a_t)} - u_{a_t}, e_{y_t} - p_t \rangle \\ &= \sum_{a \in \mathcal{A}} \left\langle u_{\phi(a)} - u_a, \sum_{t: a_t = a} (e_{y_t} - p_t) \right\rangle \\ &= \sum_{a \in \mathcal{A}} \langle u_{\phi(a)} - u_a, B_a \rangle. \end{aligned}$$

For approximate decision calibration, define the largest Euclidean utility difference between two actions:

$$D_{u,2} := \max_{a,b \in \mathcal{A}} \|u_a - u_b\|_2.$$

For each action cell, Cauchy–Schwarz gives

$$|\langle u_{\phi(a)} - u_a, B_a \rangle| \leq \|u_{\phi(a)} - u_a\|_2 \|B_a\|_2 \leq D_{u,2} \|B_a\|_2.$$

Decision Calibration Controls Swap Regret

Theorem

For every fixed forecast and outcome sequence,

$$\text{SwapReg}_T \leq \frac{D_{u,2}}{T} \sum_{a \in \mathcal{A}} \|B_a\|_2.$$

Decision Calibration Controls Swap Regret

Theorem

For every fixed forecast and outcome sequence,

$$\text{SwapReg}_T \leq \frac{D_{u,2}}{T} \sum_{a \in \mathcal{A}} \|B_a\|_2.$$

Exact decision calibration therefore gives

$$\text{SwapReg}_T = 0.$$

The forecast need not be calibrated separately at every reported probability; it needs its residual errors to cancel on the cells relevant to the user's decision function

Calibration Types and Their Benchmarks

condition	errors grouped by	Benchmark class controlled
full calibration	reported forecast p	every reinterpretation $\pi(p)$
decision calibration	induced action $a = \text{BR}_u(p)$	every action remapping $\phi(a)$

Calibration Types and Their Benchmarks

condition	errors grouped by	Benchmark class controlled
full calibration	reported forecast p	every reinterpretation $\pi(p)$
decision calibration	induced action $a = \text{BR}_u(p)$	every action remapping $\phi(a)$

Full calibration is utility-independent and serves arbitrary users. Decision calibration is weaker, and its cells depend on the specified users' utilities and action sets.

Calibration Types and Their Benchmarks

condition	errors grouped by	Benchmark class controlled
full calibration	reported forecast p	every reinterpretation $\pi(p)$
decision calibration	induced action $a = \text{BR}_u(p)$	every action remapping $\phi(a)$

Full calibration is utility-independent and serves arbitrary users. Decision calibration is weaker, and its cells depend on the specified users' utilities and action sets. We will see decision calibration is easier to obtain in high dimensional settings.

Decision-Calibrated Forecasting Appears Circular

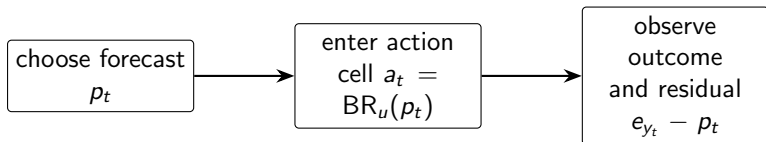
So far we have assumed small action-cell biases. Can a forecaster keep them small even in adversarial sequential forecasting settings?

Decision-Calibrated Forecasting Appears Circular

So far we have assumed small action-cell biases. Can a forecaster keep them small even in adversarial sequential forecasting settings? For action a , the active cell is the event

$$E_a(p_t) := \mathbf{1}\{\text{BR}_u(p_t) = a\}.$$

Choosing p_t determines which cell is active; only afterward does the outcome reveal the residual that updates its bias:



Next, we will see that Blackwell approachability extends Lecture 5's minimax argument to vector payoffs, choosing forecasts whose joint action-cell bias approaches a small ball around zero.

Summary

- ▶ A probability forecast provides a decision interface: the user chooses an action maximizing expected utility under the report.

Summary

- ▶ A probability forecast provides a decision interface: the user chooses an action maximizing expected utility under the report.
- ▶ Full calibration makes residual errors cancel within every report cell, ruling out profitable reinterpretations of the report.

Summary

- ▶ A probability forecast provides a decision interface: the user chooses an action maximizing expected utility under the report.
- ▶ Full calibration makes residual errors cancel within every report cell, ruling out profitable reinterpretations of the report.
- ▶ Calibration guarantees that it is optimal to “trust the forecast and act accordingly” amongst all forecast-based policies.

Summary

- ▶ A probability forecast provides a decision interface: the user chooses an action maximizing expected utility under the report.
- ▶ Full calibration makes residual errors cancel within every report cell, ruling out profitable reinterpretations of the report.
- ▶ Calibration guarantees that it is optimal to “trust the forecast and act accordingly” amongst all forecast-based policies.
- ▶ For a fixed user, action-cell calibration is enough to control swap regret and may be strictly weaker than full calibration. Hopefully it is easier to achieve.

Thanks!

See you next class!

References

- ▶ A. Philip Dawid. “The Well-Calibrated Bayesian.” *Journal of the American Statistical Association*, 1982.
- ▶ Dean P. Foster and Rakesh V. Vohra. “Calibrated Learning and Correlated Equilibrium.” *Games and Economic Behavior*, 1997.
- ▶ Shengjia Zhao, Michael P. Kim, Roshni Sahoo, Tengyu Ma, and Stefano Ermon. “Calibrating Predictions to Decisions: A Novel Approach to Multi-Class Calibration.” *NeurIPS*, 2021.
- ▶ Georgy Noarov, Ramya Ramalingam, Aaron Roth, and Stephan Xie. “High-Dimensional Prediction for Sequential Decision Making.” *ICML*, 2025.
- ▶ Aaron Roth and Mirah Shi. “Forecasting for Swap Regret for All Downstream Agents.” *ACM EC*, 2024.
- ▶ Lunjia Hu and Yifan Wu. “Calibration Error for Decision Making.” arXiv:2404.13503, 2024.
- ▶ David Blackwell. “An Analog of the Minimax Theorem for Vector Payoffs.” *Pacific Journal of Mathematics*, 1956.