

Who Reviews The Reviewers? A Multi-Level Jury Problem

Ben Abramowitz
Tulane University
New Orleans, USA
babramow@tulane.edu

Omer Lev
Ben-Gurion University of the Negev
Beersheba, Israel
omerlev@bgu.ac.il

Nicholas Mattei
Tulane University
New Orleans, USA
nsmattei@tulane.edu

ABSTRACT

We consider the problem of determining a binary ground truth using advice from a group of independent reviewers (experts) who express their guess about a ground truth correctly with some independent probability (competence) p_i . In this setting, when all reviewers are competent with $p \geq 0.5$, the Condorcet Jury Theorem tells us that adding more reviewers increases the overall accuracy, and if all p_i 's are known, then there exists an optimal weighting of the reviewers. However, in practical settings, reviewers may be noisy or incompetent, i.e., $p_i \leq 0.5$, and the number of experts may be small, so the asymptotic Condorcet Jury Theorem is not practically relevant. In such cases we explore appointing one or more chairs (judges) who determine the weight of each reviewer for aggregation, creating multiple levels. However, these chairs may be unable to correctly identify the competence of the reviewers they oversee, and therefore unable to compute the optimal weighting. We give conditions on when a set of chairs is able to weight the reviewers optimally, and depending on the competence distribution of the agents, give results about when it is better to have more chairs or more reviewers. Through simulations we show that in some cases it is better to have more chairs, but in many cases it is better to have more reviewers.

KEYWORDS

Peer Review, Peer Selection, Jury Theorem

ACM Reference Format:

Ben Abramowitz, Omer Lev, and Nicholas Mattei. 2025. Who Reviews The Reviewers? A Multi-Level Jury Problem. In *Proc. of the 24th International Conference on Autonomous Agents and Multiagent Systems (AAMAS 2025)*, Detroit, Michigan, USA, May 19 – 23, 2025, IFAAMAS, 9 pages.

1 INTRODUCTION

People have been struggling with finding the *correct* answer for millennia.¹ In ancient times, when faced with a problem that required discovering a ground truth, two main approaches dominated. The first, less common today, was to approach deities and either ask them to intervene on the randomness of the world (as in the Book of Joshua, Chapter 7), which is a bit akin to sortition [20]; or to ask the deity's wisdom directly (e.g., the Oracle at Delphi). The second approach, still in widespread use today, is to try to assess the known information and draw a conclusion. This can either be

done by laymen, the basic premise of the jury system as established by Magna Carta, or by people with expertise. In both cases, groups of people are used (instead of single individuals) to increase the reliability and accuracy of the answers, building on the “wisdom of the crowds”.

Mathematical analysis of using a group of agents – a *jury*, or a set of *experts* – to assess information and make a decision has been done since at least Condorcet's time, in the late 18th century, when he established the Condorcet Jury Theorem (CJT) [16, 17]. In the standard jury setting, agents vote on a binary ground truth and the objective is to aggregate their votes, using a voting rule, to maximize the probability of the outcome being correct. It is typical to assume that agents guess the ground truth correctly with some independent probability (competence) p_i . We call agents competent when $p > 0.5$ and incompetent when $p \leq 0.5$.² According to the CJT, when the agents are competent, the collective accuracy of their majority vote tends to correctness as the number of agents increases. Even with a relatively small number of moderately competent agents, accuracy can be very high. However, this result, which basically tells us that groups are less prone to mistake than individuals, rests on a knife's edge. By symmetry, if the agents are even minimally incompetent then, as the population grows, their collective accuracy under majority voting tends to 0, and a small group of highly incompetent agents stands no chance.

In the world around us, this idea is used everywhere – in judicial settings (juries), in academic conferences (peer evaluation), in voting for political leaders or in referendums, and even in settings with inanimate agents, such as aggregating sensor outputs into a single reading or indicator.

The precariousness of the CJT stems from the underlying aggregation procedure, majority voting, being anonymous, thus treating all agents equally, regardless of their competence. When agent competences can be different, majority rule is generally sub-optimal, and if one knows the agents' level of expertise exactly, the optimal aggregation method for maximizing accuracy with any number of independent experts and any competences is to use a weighted majority rule in which each agent's weight is the log-odds of their competence [34, 41]. In this result, somewhat surprisingly, the optimal weight of each agent does not depend on the competences of the other agents or even on the total number of agents. That is, if agents are added or removed, we do not need to update the weights in order for the weighted majority rule to remain optimal. However, the assumption that each agent's competence is known exactly is highly unrealistic.

In this work, we consider a variant of the classic jury setting, inspired by the domain of academic peer review [40], which attempts to address these unrealistic assumptions. Since the quality of reviewers may not be known by the conference Program Chairs,

¹Fans of *The Hitchhiker's Guide to the Galaxy* know it is 42.



This work is licensed under a Creative Commons Attribution International 4.0 License.

²The term competent is not meant to express a value judgment.

many conferences (e.g., IJCAI) appoint more senior researchers as SPCs (or chairs) to evaluate the reviewers and decide how to aggregate their views. This multi-level process inspired our model: we have both reviewers, who we will call *experts*, but also chairs, who we will call *judges*, that evaluate the experts and assign them weights. While we recognize that it is unlikely there is a pure, objective ground truth for research, we build on a long line of research in this model [4, 33, 40], one could use other settings, such as the grading of MOOC exercises (which usually have a fixed, correct, answer) [48].

Analyzing this setting is particularly interesting when the number of agents is relatively small, and therefore we cannot rely on the asymptotic guarantees of the CJT; as well as when there is a potential for significant deviation among agent competency, even when the particular competence values are unknown; or when agents can be incompetent, i.e., they will make the wrong decision most of the time.³ We examine when such a two-level system works well, under what conditions it might be worthwhile to implement it, and when is it better to have an expert become a judge.

Contribution. We propose and investigate a novel model of multi-level jury problems for use when we have a small number of possibly unreliable agents. We show that when we know the agent competences exactly (or even approximately), we can find an optimal aggregation procedure, as long as the judges are competent. When the agents' (experts and judges) competences are unknown, we provide a set of numerical experiments demonstrating that adding more than a single judge is rarely helpful, and, indeed, in some cases, the potential damage of a less competent judge is enough to prefer to avoid judges completely.

2 RELATED WORK

There is a long history of studying the Condorcet jury model and its extensions (e.g., Ben-Yashar and Paroush [9], Berend and Paroush [11], Feld and Grofman [19], Grofman [22]), including work across computer science, philosophy, and economics. Our work is primarily based off the literature on weighting experts in both the offline [34, 41] and online settings [10, 15, 21, 47], though in this work we restrict our focus to a single decision. The overall CJT model can be seen in portfolio solver techniques where slower, more reliable algorithms evaluate ensembles of faster, less accurate algorithms [46] and in boosting techniques in ML, where one aggregates weakly competent classifiers into a better overall classifier [39]. These methods are particularly useful when each expert will use limited effort or energy [45].

In settings with repeated decisions, the competences of the experts can be estimated based on their performance history. Their competence might be estimated by their similarity to other agents, or how often they agreed with the decision outcomes in the past [5, 25, 38]. In our setting, however, we do not have access to this history and cannot use it to estimate competences. Indeed, in peer review, one may have a notion of other reviewers' competency, but rarely does one co-review with another to form a precise estimate.

³In academic peer evaluation it is uncommon for reviewers to be given negative weights, which we will see is required for the log-odds rule to be optimal. But it is common in other settings, e.g., sensors or proxy voting scenarios where one might want to always do the opposite of a political rival.

Our emphasis on imperfect judges is also inspired by work on "wisdom of the crowds" and crowdsourcing [13, 14, 42], proxy voting [1, 2, 36] and truth-tracking in Liquid Democracy [8, 50]. For example, in proxy/delegative voting the voters are like judges assigning weights (voting units) to their proxies/delegates. But, as noted above, a major inspiration has been the work on academic peer evaluation [40], in which experts assess each other's competences. Some treat this matrix as a Markov chain, and use its eigenvector values as the experts' weights [23] in a manner reminiscent of some peer-evaluation models [29, 30, 35, 48]. In contrast, within our setting the agents who vote and the of agents who weight the voters are disjoint.

The problem of partitioning agents into judges and experts is also related to the problem of computing optimal committee sizes [32, 37]. There is also work on how group accuracy depends on the size of the group and mean competence [22, 24], which is reflected in our simulations. Lastly, we mention work extending the Condorcet jury problem to more than two outcomes [18, 26, 31].

3 MODEL AND NOTATION

Our model has two types of agents; judges and experts. Let E be a set of m experts and J be a set of n judges. The experts vote on a single binary issue where there is only one correct (ground truth) outcome. Without loss of generality, let the options be $\{1, 0\}$ where 1 is correct and 0 is incorrect. Each expert $e \in E$ has a *competence*, or probability p_e of voting correctly, independent of all other experts. We associate each expert's index with their vote, so expert $e \in E$ casts a vote $v_e \in \{1, 0\}$ with competence $p_e = P(v_e = 1)$. If an agent's competence is above $1/2$ we will say that they are competent, and call them incompetent otherwise. We assume no one is always correct or always incorrect, and so $0 < p_e < 1$.

Weighted Majority Rules. We refer to the probability of producing the correct outcome as the *accuracy*, and reserve the term *competence* to refer to individual agents' probabilities of voting correctly; i.e. accuracy is collective.

Definition 3.1 (Weighted Majority Rule). A weighted majority rule gives each expert $e \in E$ a weight $w_e \in \mathbb{R}$ and selects 1 as the winner if $\sum_{v_e=1} w_e > \sum_{v_e=0} w_e$, selects 0 as the winner if $\sum_{v_e=1} w_e < \sum_{v_e=0} w_e$, and uses a tie-breaking rule (e.g. coin flip) for the edge case where these sums are equal.⁴

Definition 3.2 (Simple Majority Rule). Simple majority rule refers to the weighted majority rule where all weights are equal and positive and ties are broken randomly.

The Condorcet Jury Theorem tells us that if $p_e \geq 0.5 + \epsilon$ for some $\epsilon > 0$ for all experts, then with simple majority accuracy tends to 1 asymptotically as the number of experts tends to infinity. A weighting function maps vectors of values in $(0, 1)$ (i.e. competences) to equal length vectors of real values. For any set of experts, including incompetent ones, the optimal aggregation method of experts' votes, is to apply the log-odds weighting function to the experts' competences and use the corresponding weighted majority rule [34, 41].

⁴Note that ties never occurred in our simulations as the need for tie-breaking, with real valued weights, is highly unlikely.

Definition 3.3 (Log-Odds Weighting Function). Given a vector of values in the open unit interval $\vec{p} = (p_1, \dots, p_m)$, the log-odds weighting function returns the vector $\vec{w} = (w_1, \dots, w_m)$ where $w_e = \log(\frac{p_e}{1-p_e})$ for all $1 \leq e \leq m$.

Any weighting of the agents implies a collection of winning coalitions – subsets of agents who, if they all vote the same way, determine the outcome regardless of the other votes [44]. Different weightings may yield the same rule because they imply the exact same winning coalitions. For example, with 5 agents there are exactly 7 distinct weighted majority rules [27, 28]. Multiplying the weights of all agents by a constant does not change the winning coalitions and therefore does not change the rule. Similarly, perturbing agent weights by small amounts may not change the winning coalitions. Therefore, while the weights may vary continuously, the accuracy under various weightings will change in discrete steps. In practice, weights may be finite precision rather than true real numbers, and this is also the case in our simulations that use floating point arithmetic, but as long as the rounding tends to be too small to change winning coalitions for most instances its effect will be negligible.

The log-odds weighting rule assigns a positive weight to competent experts when $p_e > 0.5$, weight of zero if $p_e = 0.5$, and negative weight to any incompetent expert with $p_e < 0.5$. In some settings it may be inappropriate to allow negative weights and better to assume any such weights are rounded up to zero. Bounding weights below by zero has the effect of ignoring the incompetent experts and is therefore qualitatively similar to assuming all experts are competent, though with a smaller number of experts. We therefore focus on the more informative setting where weights can be negative. Negative weights also have real-world motivation. A remote sensor may have drifted so far off its initial calibration to be *reliably* wrong, as has happened with many spacecraft [7]. However, we would like to believe that peer reviewers, jurors, and the like are not so reliably wrong that negative weights would be needed.

Multi-Level Jury Problems. Each judge $j \in J$ estimates the competence of each expert $e \in E$ as $p_{je} \in (0, 1)$ and assigns them a weight w_{je} based on this estimate. When assigning weights to the experts, our judges always use the log-odds weighting function on their competence estimates. Intuitively, our judges are trying to implement the optimal rule using their estimates of the experts' competences. Formally, judge j assigns expert e a weight $w_{je} = \log(\frac{p_{je}}{1-p_{je}})$. When there are multiple judges, the weight of an expert will be the average weight assigned to them by the judges $w_e = \frac{1}{n} \sum_j w_{je}$.

Perceived Competences. Our theoretical results depend on judges using the log-odds weighting but do not depend on how the judges form their competence estimates p_{je} . To perform empirical analysis we must make assumptions about where these estimates come from. Rather than drawing the estimates p_{je} from some named distribution with mean p_e , we take an approach inspired by peer review, and assume that judges and experts are fundamentally similar. Each judge j has competence p_j just like the experts and estimates the competence of expert e as $p_{je} = (p_j \cdot p_e) + (1 - p_j)(1 - p_e)$, i.e., the probability that expert e agrees with them. When p_{je} is derived in this way, we will refer to it as the judge's *perceived competence* of

the expert. As with peer review, a judge may estimate an expert's competency from knowing them professionally, but may not have observed many past reviews. A judge could also reach this estimate of competency if they observe enough votes from the expert but the ground truth is never revealed, as is the case in some peer prediction settings [49]. While there are many models one could select for the formulation of perceived competencies, this formulation does not rely on complex models of estimation or opinion formation. Indeed, any additional assumptions that improves accuracy, i.e., that judges are better at estimation than simple observation, would improve the limit results we obtain. So while this assumption is simplistic, it is in some ways a worst-case assumption for our theoretical results, and leads to interesting empirical results.

Example 3.4. Suppose we have 5 experts with competences $\vec{p}_E = (0.6, 0.6, 0.6, 0.7, 0.9)$. The optimal log-odds weighting is approximately $\vec{w}_E^* = (0.41, 0.41, 0.41, 0.85, 2.2)$. With these weights the most competent expert ($p_e = 0.9$) receives a weight ($w_e = 2.2$) that makes them a dictator in a weighted majority vote, since their weight is greater than all other experts combined. Hence, the accuracy under the log-odds weighting is exactly 0.9. If instead we use simple majority, the accuracy drops to 0.82.

A judge with $p_j = 0.6$ would assign the experts weights of approximately $\vec{w}_E^{0.6} = (0.08, 0.08, 0.08, 0.16, 0.323)$ using the log-odds weighting of perceived competences. Note that the fifth expert is no longer a dictator. How high of a competence would the judge need to have to assign weights that results in the optimal weighting? The judge's competence would have to be greater than 0.962; a far cry from 0.6 and higher than all the experts. And yet, the judge's sub-optimal weighting yields an accuracy of 0.898, which is a great improvement over simple majority, and extremely close to optimal at 0.9! Example 3.4 is illustrated in Figure 1 where we plot the accuracy, sweeping p_j from 0.0 to 1.0.

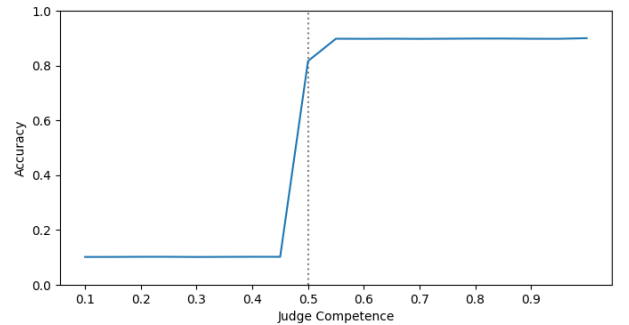


Figure 1: Accuracy of perceived optimal weightings from a single judge with the expert competences as in Example 3.4.

In Example 3.4 we rounded the values of the weights to two decimal places, which did not change the rule implemented. Similarly, whether judges are human, sensors, or algorithms, they do not (always) need to provide high precision of weights. More specifically, the smaller the number of agents, the less chance there is for small perturbations or rounding of the weights to change the rule. This

is also why Figure 1 will be piecewise linear regardless of the step size.

4 OPTIMALITY AND ROBUSTNESS

With a single judge, if $p_{je} = p_e$ for all $e \in E$, then all experts receive their optimal weight. As noted, small perturbations to the weights do not change the rule because the winning coalitions determined by the weights do not change. Thus, if p_{je} is close enough to p_e for all experts, they will still produce the optimal weighting. We now establish sufficient conditions for an ensemble of judges to produce the optimal weighting, and provide a condition under which the difference between an expert's weight and their optimal weight tends to be small.

Proposition 4.1 shows that when the geometric mean of the judges' perceived competence of experts odds is their true competence odds, all experts are assigned their optimal weights, $w_e = w_e^*$. This does not require individual judges to know the experts' true competences, and does not depend on the number of experts nor the number of judges.

PROPOSITION 4.1. *If each judge uses the log-odds weighting function on their estimates of expert competences, and the geometric mean of the judges' estimates of each expert's competence odds is the expert's true odds, then the weighted majority rule using the judges' average weights to weight each expert is exactly the optimal weighted majority rule.*

PROOF. Since judge j gives each expert a weight of $w_{je} = \log(\frac{p_{je}}{1-p_{je}})$,

$$\begin{aligned} w_e &= \frac{1}{n} \sum_j w_{je} = \frac{1}{n} \sum_j \log\left(\frac{p_{je}}{1-p_{je}}\right) = \\ &= \frac{1}{n} \log\left(\prod_j \frac{p_{je}}{1-p_{je}}\right) = \log\left(\left(\prod_j \frac{p_{je}}{1-p_{je}}\right)^{\frac{1}{n}}\right) \end{aligned}$$

We assume the geometric mean of judge's estimates of the experts' competence odds is correct, i.e., $(\frac{p_e}{1-p_e}) = (\prod_j \frac{p_{je}}{1-p_{je}})^{\frac{1}{n}}$. Thus, $w_e = \log(\frac{p_e}{1-p_e}) = w_e^*$. $\square$

This result requires assumptions on the judges' competence estimates that must hold for *all* experts. However, suppose there are errors in these collective competence estimates. We want to know how sensitive the weight of a single expert is to such errors.

COROLLARY 4.2. *If the geometric mean of judge estimates of competence is off by some multiplicative factor α for some expert, then the error of that expert's weight is only $\log(\alpha)$.*

PROOF. In the proof above, assume instead that $\alpha(\frac{p_e}{1-p_e}) = (\prod_j \frac{p_{je}}{1-p_{je}})^{\frac{1}{n}}$. Then $w_e = \log(\alpha \cdot \frac{p_e}{1-p_e}) = w_e^* + \log(\alpha)$. $\square$

Admittedly, it is not clear in what settings the conditions of Proposition 4.1 should be expected to hold. Neither can we make claims about what multiplicative factors are realistic in Corollary 4.2. But ultimately, what we care about most is the sensitivity of the accuracy to errors in competence estimates, which has to do with the set of winning coalitions induced by the weights. We care less about the sensitivity of the weights themselves, although the sensitivity of the weights gives some intuition.

In Example 3.4, we see that for all $p_j > 0.55$ the accuracy rivals that of the optimal rule, with almost no difference. The effect that dominates Figure 1 is when we move from $p_j < 0.5$ to $p_j = 0.5$ (when the rule becomes simple majority), with another slight bump with a move to $p_j > 0.55$.

Recent work [6] shows that when expert competences are drawn from certain distributions over the range $(1/2, 1)$, simple majority achieves an accuracy close to optimal. However, as one might expect, when experts can be incompetent the majority rule is no longer a good approximation to the optimal weighted majority rule. Thus, if judges can at least differentiate the competent from incompetent experts, the weighting they produce should be expected to outperform simple majority rule when there are incompetent agents. In our model, any minimally competent judge with $p_j > 0.5$ is able to achieve this. We will discuss this more in Section 5 with an array of experiments, including a suite of experiments with a single judge. For now we introduce some basic theoretical observations that help understand the phenomena observed in our experiments. Namely, the improvement in accuracy as the judge(s) competence(s) reach $p_j = 0.5$ appears to be largely but not entirely explained by the judge's ability to (1) assign experts weights of the correct sign, and (2) assign experts weights according to the weak order of their competences.

Example 4.3 (Two Experts). Suppose there are two experts with competences (p_1, p_2) such that $p_1 > p_2$. If $p_2 > 1/2$, i.e., both experts are competent, the optimal aggregation rule is to make p_1 dictator. However, if $p_1 > 1/2 > p_2$ then the optimal rule is either to make p_1 a dictator if $p_1 \geq 1 - p_2$, or else make p_2 an anti-dictator using negative weight such that the outcome is the opposite of however p_2 votes. If $1/2 > p_1 > p_2$, then the optimal rule makes p_2 the anti-dictator symmetrically with the first case.

From Example 4.3, we see that even with only two experts, if a judge can determine which experts are competent, and order their competences correctly, this is enough information to produce the optimal rule. With more experts, the situation is more complicated, but our experiments reveal that merely separating the competent experts from incompetent ones creates a large improvement in overall accuracy. Any chair with $p_j > 1/2$ using log-odds weightings of the perceived competences can achieve this improvement.

PROPOSITION 4.4 (CORRECT SIGN). *If $\text{sign}(p_{je} - 0.5) = \text{sign}(p_e - 0.5)$, then $\text{sign}(w_{je}) = \text{sign}(w_e^*)$.*

PROOF. Proposition 4.4 follows directly from the fact that $\frac{p}{1-p} > 1$ if and only if $p > 1/2$, and therefore $\log(\frac{p}{1-p}) > 0$ if and only if $p > 1/2$. Symmetrically for $p < 1/2$. $\square$

PROPOSITION 4.5 (CORRECT ORDER). *If p_{je} is a strictly monotonic increasing function of p_e , then the weak order of expert weights given by judge j is the weak order of the experts' competences.*

Proposition 4.5 follows from the monotonicity of the log-odds weighting. When a single judge applies the log-odds weighting to their perceived competences of a small set of experts then we can make the following observations.

OBSERVATION 1. *If $p_j = 1/2$ then all experts are equally weighed. If $p_j = 1$ then experts are optimally weighed.*

$n \backslash m$	1	2	3	4	5	6	7	8	9	10
0	0.750	0.749	0.843	0.843	0.897	0.897	0.929	0.929	0.950	0.950
1	0.750	0.831	0.880	0.913	0.936	0.951	0.964	0.973	0.980	0.984
2	0.751	0.833	0.881	0.914	0.937	0.953	0.964	0.974	0.980	0.984
3	0.748	0.833	0.880	0.914	0.937	0.952	0.964	0.972	0.980	0.984
4	0.748	0.832	0.882	0.914	0.937	0.952	0.964	0.974	0.980	0.984
5	0.747	0.833	0.882	0.915	0.937	0.953	0.965	0.972	0.980	0.984
6	0.748	0.832	0.881	0.914	0.937	0.952	0.964	0.974	0.980	0.985
7	0.750	0.833	0.881	0.913	0.937	0.954	0.964	0.974	0.980	0.984
8	0.748	0.834	0.881	0.915	0.937	0.952	0.964	0.974	0.980	0.985
9	0.751	0.832	0.881	0.914	0.937	0.953	0.965	0.974	0.980	0.984
10	0.748	0.833	0.881	0.915	0.937	0.954	0.966	0.974	0.979	0.984

Table 1: Accuracy with all agent competences drawn from the uniform distribution over (0.501,0.999)

If $p_j = 1/2$ then the judge will perceive all experts as having a competence of $1/2$, and therefore assign them the weight of 0, which we treat as giving them equal weight, i.e. majority vote. When $p_j = 1$, the judge knows exact competences of the experts and therefore assign them their optimal weights.

OBSERVATION 2. *If $p_j > 1/2$, then $p_{je} > p_{je'}$ iff $p_e > p_{e'}$.*

This means that a judge’s perceived competences of the experts preserves the order of their true competences if $p_j > 1/2$. When weights are based on perceived competences, this means whenever $p_j > 0.5$, the judge will assign all experts’ weights with the correct sign and in the correct order. This is because when $p_j > 0.5$, p_{je} is monotonically increasing in p_e and $p_{je} > 0.5$ iff $p_e > 0.5$. Thus, even a single barely competent judge might give us an edge over simple majority.

THEOREM 4.6 (MINIMAL COMPETENT SINGLE JUDGE). *If $p_j > 0.5$ and the judge assigns experts weights according to their perceived competences, the weights given by the judge will have the correct sign and the correct order.*

PROOF. Let $p_j = 1/2 + \epsilon_j$ and $p_e = 1/2 + \epsilon_e$.

$$\begin{aligned} p_{je} &= (1/2 + \epsilon_j)(1/2 + \epsilon_e) + (1/2 - \epsilon_j)(1/2 - \epsilon_e) \\ &= 1/2 + 2\epsilon_j\epsilon_e \end{aligned}$$

If $\epsilon_j > 0$ and $\epsilon_e > 0$, then this value is greater than $1/2$, if $\epsilon_j > 0$ and $\epsilon_e < 0$, then this value is less than $1/2$, and if $\epsilon_e = 0$ then this value is exactly $1/2$. The proposition then follows from Proposition 4.5 and Proposition 4.4. $\square$

We can generalize Proposition 4.4 by replacing the requirement that $p_j > 1/2$ with the requirement that the geometric mean of judges’ estimated competence odds is greater than 1. Notice that the requirements for this theorem are far weaker than for the optimality demanded by Proposition 4.1.

PROPOSITION 4.7 (CORRECT SIGN). *If the geometric mean of every expert’s estimated competence from the judges is greater than 1 whenever $p_j > 1/2$, less than 1 whenever $p_j < 1/2$, equal to 1 when $p_j = 1/2$, every expert will be assigned a weight with the same sign as their optimal weight.*

PROOF. Suppose $p_e > 1/2$. Their optimal weight is positive, and so we need the following to hold:

$$\begin{aligned} \sum_{j \in J} \log \left(\frac{p_{je}}{1 - p_{je}} \right) &> 0 \\ \prod_{j \in J} \frac{p_{je}}{1 - p_{je}} &> 1 \\ \left(\prod_j \frac{p_{je}}{1 - p_{je}} \right)^\gamma &> 1 \end{aligned}$$

for $\gamma > 0$. When $\gamma = \frac{1}{n}$ this is the geometric mean. The case for $p_e < 1/2$ is symmetric with flipped inequality, and for $p_e = 1/2$ is the same but with strict equality. $\square$

Proposition 4.4, Proposition 4.5, and Theorem 4.6 can be immediately generalized to multiple judges by requiring their respective assumptions to hold for all judges individually, because the sum of non-negative (strictly monotonic) functions remains non-negative (strictly monotonic). While our results on correct sign and order are interesting, they are not sufficient to always outperform equal weighing. As we will see in the next section, with our perceived competences model we see a discrete interpolation between equal and optimal weighting as a single judge’s competence increases. Sign and order preservation appear to be part of the explanation as to why, but this explanation is incomplete. We pose the conjecture that as single judge competency grows, the perceived competence monotonically changes from $\frac{1}{2}$ to the optimal value.

5 EMPIRICAL RESULTS

We gain a deeper understanding of the behavior we see in Example 3.4 and Section 4 with a set of numerical experiments. In our simulations, the agents’ competences are drawn i.i.d from various distributions. All experiments were run for 100,000 iterations for each parameterization of the problem instance so that the variances are negligible. When there is no judge ($n = 0$), unweighted majority vote is used with random tie-breaking.

We consider three distributions of expert competences: uniform, truncated normal, and truncated exponential. The uniform distribution reflects settings where the experts can equally have any competence; the exponential distribution models settings where the expertise tends to be rare [12]; and the normal distribution is appropriate for a common expertise [43].

We provide Tables 1-3, with one for each family of competence distribution over (0.5, 1.0) where no agents are incompetent and all agents, judges and experts, have their competences drawn from the same distribution. By looking across the rows, we see the improvement from incrementally adding experts, by looking down each column we see any change from adding judges, and by looking at the diagonals we see how best to choose n and m to divide the agents into roles given a fixed number of agents $n + m$.

When competences are bounded below by 0.5, we see that there is always a benefit to increasing the number of experts, but no benefit to increasing the number of judges beyond $n = 1$. Notably, without any judges, adding an expert only increases accuracy when incrementing the number of experts from even to odd. When adding

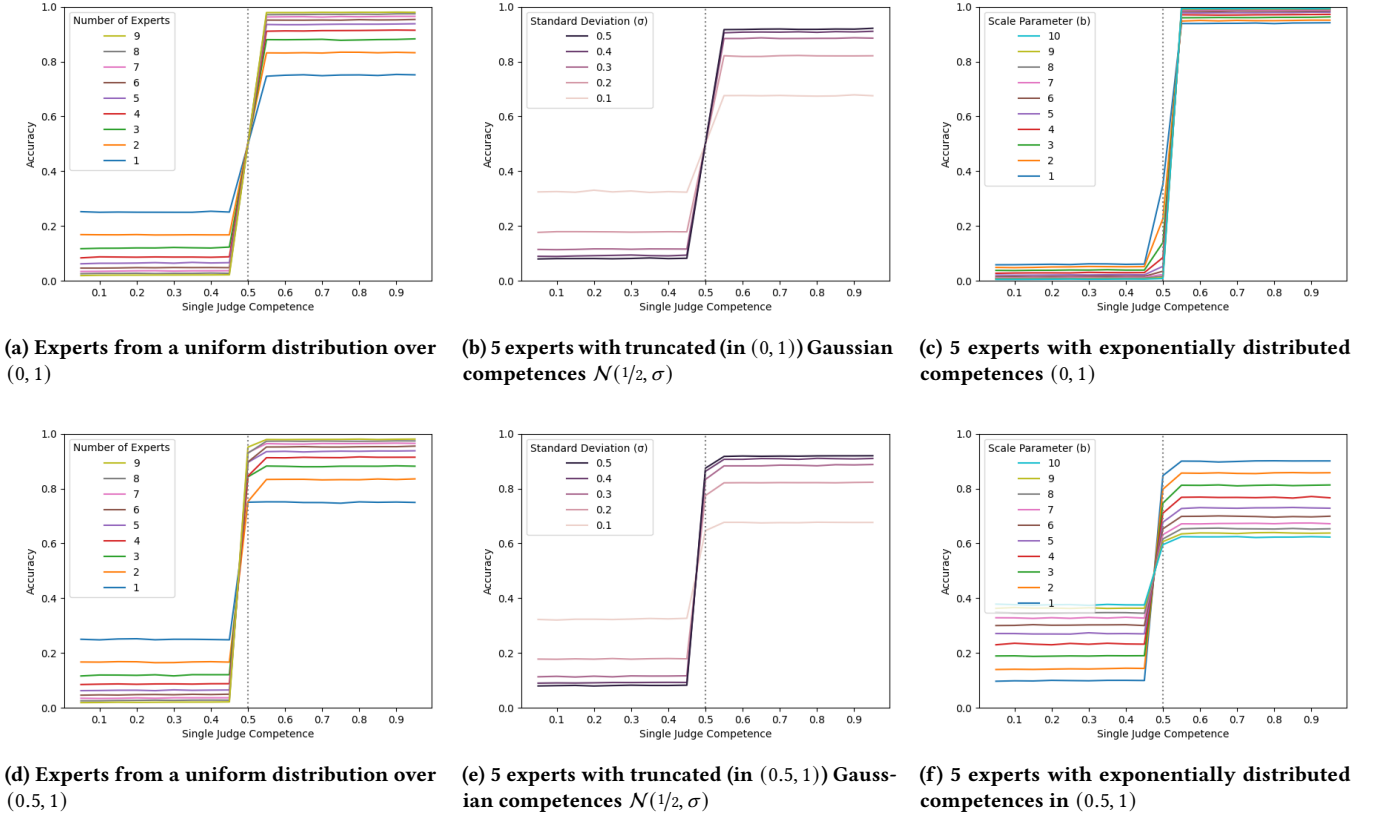


Figure 2: Accuracy with a single judge and expert competences drawn i.i.d. from a distribution with support $[0.001, 0.999]$ (top row) or support $[0.501, 0.999]$ (bottom row).

a single expert makes m even, there is no observable increase in accuracy if there is no judge. With one or more competent judges, increasing the number of experts always increases accuracy. By contrast, adding a single judge shows increasing accuracy whenever there is more than 1 expert, but increasing the number of judges further shows no benefit. See Appendix for a suite of similar experiments varying the support and parameters of the various distributions.

In all three of these tables, there is a noticeable benefit to adding a single judge, but no benefit to adding judges beyond that. We therefore take a closer look in Section 5.1 at the case with a single judge, and examine how the accuracy varies if the judge’s competence differs from the experts, even allowing the judge to be incompetent. Lastly, we will look in Section 5.2 at how to determine the number of agents to set as judges versus experts when all agent competences are drawn from the same distribution and agents can be incompetent. Further tables can be found in the Appendix.

5.1 Single Judge

The top row of Figure 2 shows accuracy as a function of the single judge’s competence when expert competences are distributed over the interval $[0.001, 0.999]$ according to the uniform, truncated normal ($\mathcal{N}(1/2, \text{varying } \sigma)$), and truncated exponential distributions

$n \backslash m$	1	2	3	4	5	6	7	8	9	10
0	0.581	0.578	0.621	0.618	0.647	0.648	0.673	0.669	0.693	0.692
1	0.580	0.616	0.639	0.658	0.673	0.691	0.705	0.720	0.732	0.741
2	0.580	0.612	0.640	0.660	0.680	0.690	0.705	0.719	0.729	0.743
3	0.581	0.612	0.640	0.659	0.676	0.691	0.707	0.718	0.730	0.742
4	0.581	0.614	0.638	0.658	0.673	0.692	0.707	0.719	0.727	0.741
5	0.581	0.616	0.635	0.658	0.674	0.692	0.706	0.718	0.729	0.743
6	0.580	0.612	0.635	0.660	0.675	0.692	0.706	0.717	0.731	0.742
7	0.579	0.613	0.639	0.659	0.678	0.692	0.707	0.718	0.730	0.741
8	0.581	0.615	0.639	0.659	0.676	0.689	0.708	0.718	0.732	0.741
9	0.582	0.613	0.639	0.659	0.675	0.692	0.707	0.719	0.733	0.741
10	0.579	0.612	0.641	0.662	0.676	0.692	0.705	0.718	0.731	0.743

Table 2: Accuracy with all agent competences drawn from the Gaussian distribution $\mathcal{N}(0.5, 0.1)$ truncated over (0.501, 0.999)

(using the density function $e^{-x}/(1-e^{-b})$ for varying values of b) respectively. Only Figures 2a and 2b exhibit true symmetry because competences are drawn from a symmetric distribution with mean 0.5, and Figures 2d and 2e show behavior most similar to Figure 1.

$n \backslash m$	1	2	3	4	5	6	7	8	9	10
0	0.672	0.672	0.750	0.747	0.797	0.799	0.834	0.835	0.862	0.865
1	0.674	0.745	0.795	0.831	0.857	0.881	0.897	0.913	0.925	0.935
2	0.672	0.744	0.795	0.829	0.858	0.882	0.898	0.912	0.926	0.936
3	0.675	0.744	0.796	0.829	0.858	0.879	0.899	0.913	0.925	0.936
4	0.671	0.747	0.791	0.830	0.858	0.880	0.897	0.913	0.924	0.938
5	0.672	0.746	0.794	0.831	0.857	0.879	0.897	0.912	0.925	0.936
6	0.672	0.746	0.791	0.829	0.858	0.880	0.898	0.913	0.926	0.935
7	0.674	0.742	0.793	0.831	0.857	0.880	0.896	0.913	0.925	0.936
8	0.672	0.745	0.792	0.829	0.858	0.880	0.897	0.913	0.925	0.937
9	0.671	0.747	0.796	0.829	0.858	0.881	0.898	0.913	0.925	0.936
10	0.674	0.745	0.795	0.832	0.857	0.878	0.897	0.912	0.927	0.936

Table 3: Accuracy with all agent competences drawn from the exponential distribution with scale parameter $b = 2$ truncated over (0.501, 0.999)

In the top row of Figure 2 we can have highly incompetent experts, but even in this setting whenever the judge has competence $p_j > 1/2$, high overall accuracy is achieved. This is because the ability of the judges to differentiate competent experts from incompetent ones is of primary importance, and Proposition 4.7 shows that a judge using perceived competences is able to do this. In Figures 2a-2c, once the judge passes a minimum threshold of competence, little is gained from increasing p_j . Interestingly, in Figure 2b we see that when expert competences are distributed normally with mean $1/2$, higher variance leads to higher collective accuracy. This appears to be because a judge with sufficiently high competence can differentiate between highly competent and minimally competent experts, and then leverage the benefits of having highly competent experts when they are present.

In the bottom row of Figure 2 we show accuracy as a function of the single judge’s competence when expert competences are distributed over the interval $[0.501, 0.999]$ according to our three distributions. This is closer to prior work on the Condorcet jury theorem where all experts are assumed to be competent [16]. Unlike in Figure 2a and 2b, which are symmetrical distributions, we now see an asymmetry around $p_j = 0.5$. When the judge’s competence is $1/2$, the judge gives all experts the same weight, so when experts competence is symmetrical around $1/2$ (as in the upper row of Figure 2), resulting in an accuracy of $1/2$, but when they cannot be incompetent, the resulting accuracy is higher – almost optimal [6]. In contrast to the top row of Figure 2, when all experts are competent, there is a large difference in accuracy for truncated exponential distributions with different scale parameters.

5.2 Should We Add a Judge or an Expert?

Empirically, with a single judge the accuracy improves as the judge’s competence grows, and we know from the Condorcet Jury Theorems that as the number of experts increases, if the experts are competent, accuracy will increase. Hence, if the conditions of Proposition 4.7 hold, then accuracy will increase as the number of experts increases whether they are competent or incompetent, as long as they have competences that are not equal to $1/2$. We examine the balance between the benefits of increasing the number of judges

and increasing the number of experts. That is, with a fixed set of agents of unknown competences, how should they be partitioned between experts and judges? This problem is faced by any scientific conference with a hierarchical structure: how to divide its Program Committee between reviewers and SPCs, ACs, etc.

We first draw agents’ competences from uniform distributions with varying lower bounds and examine the optimal number of agents to set aside as judges rather than experts when we have 5 and 11 agents, respectively (Figures 3a and 3d). For both number of agents, we see that *setting aside a single agent as judge diminishes the accuracy compared to the simple majority rule in almost all cases*. This is more pronounced when there is a possibility the judge will have competence below $1/2$, i.e., when a lone judge is incompetent they give all competent experts negative weights and incompetent experts positive weights. The only case where a judge is helpful is when the minimum competence of agents is $1/2$, perhaps because there is high enough chance that the judge will be helpful, and the agents’ competence is not guaranteed to be high enough that losing the judge as an expert is too big a hit. Even adding more judges, at best, returns the accuracy to the level of a simple majority rule, though most commonly it does not. In the 11 agent case, Figure 3d, this effect is even more pronounced than for 5 agents. With 11 agents, adding enough judges can eventually bring peak accuracy to slightly above the simple majority, though it requires roughly an even split between judges and experts (or even slightly more judges). Further experiments show that if we simply add a judge, increasing the number of agents by one without reducing the number of experts, accuracy drops, so adding a judge is harmful, up until the point where the lower bound on competences guarantees the judge will be helpful (see Appendix).

In contrast to the uniform distribution, when drawing competences from the normal and the exponential distributions, things are a bit different. They show that a single judge can be productive. With normal distributions, when the mean is high but not extremely so (Figures 3b and 3e), adding a judge helps. When the mean is very high (0.8 and above), aggregating all agents as experts seems to be better than having a judge, for whom there is still a probability of being bad. But when agents are with a lower mean, having a judge seems to help, and this is true even for a mean of 0.5, in which there is a probability of 0.5 that the judge will be incompetent. This pattern appears for 5 agents, but, as in the uniform case, it is more pronounced for 11 agents.

For the exponential distributions (Figures 3c and 3f), this property is stronger – it is *always* beneficial, for our parameters, for agents to have one judge, and that improves over a simple majority. This is likely due to the fact that the small loss of accuracy from having one fewer experts is made up for by the ability of even a minimally competent judge (and all agents are competent in this distribution) to distinguish highly competent experts from less competent ones. Unlike in the uniform case, adding more judges (after a single one) is never helpful compared to a single judge (though sometimes two judges are better than simple majority).

These results imply that the division of labor in scientific conference is counter-intuitive: multiple layers above regular experts (e.g., SPC, AC) does not seem to be helpful. It seems better, given a finite set of agents to assign, to have a flat hierarchy (i.e., more experts, fewer judges,) and use simple majority, despite this idea

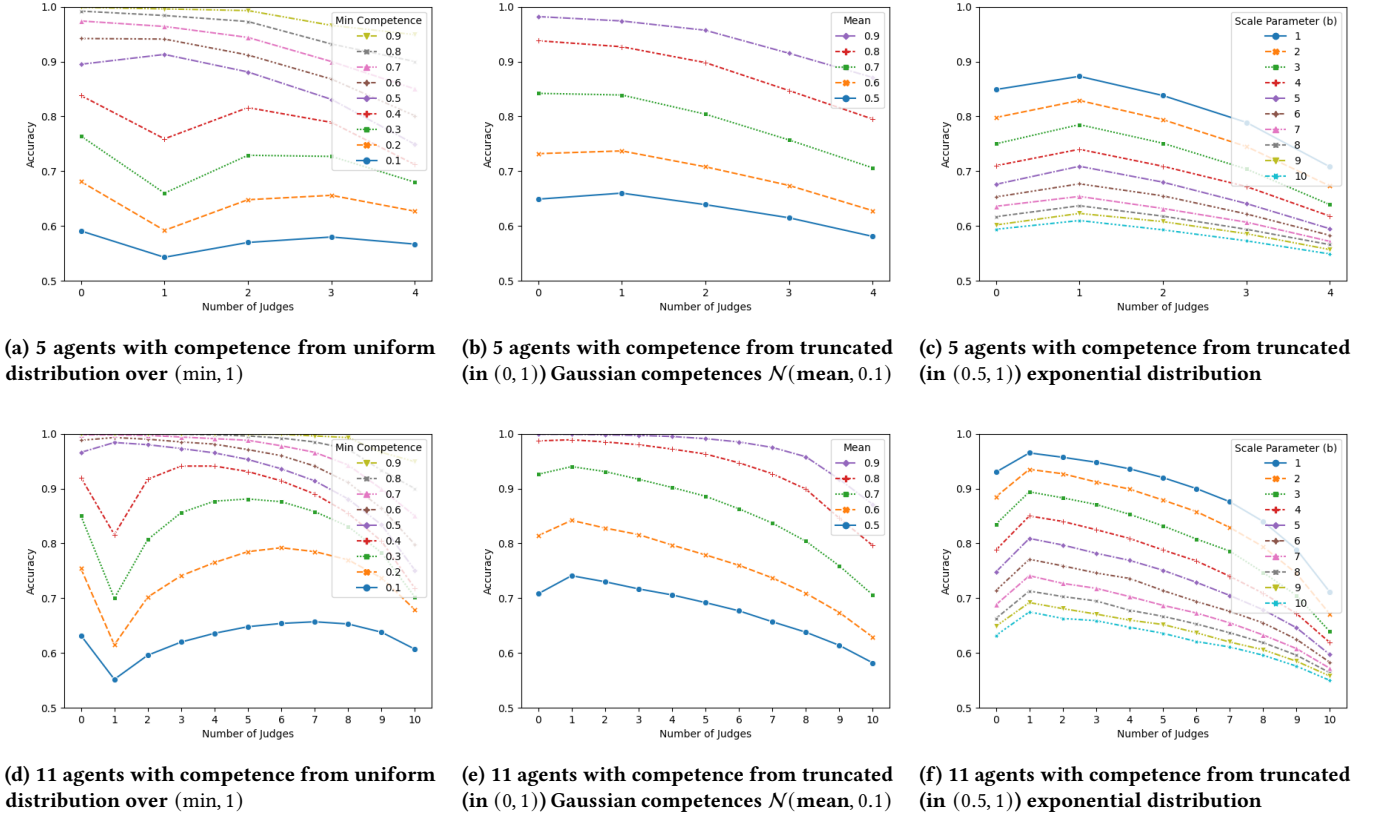


Figure 3: Accuracy partitioning a set of agents randomly into judges and experts with agent competences i.i.d. from the same distribution.

being often frowned upon. We did not investigate the case where a judge is more competent than experts, but even then it is not clear more judges are better, as losing a top expert incurs a cost. Indeed, it is not at all clear that the best judge is the one with the highest competence, and we leave this open question to future research.

6 DISCUSSION AND FUTURE WORK

We consider a multi-level jury problem in which experts are given weights according to estimates from judges of their competence. We focus on settings where there is a small number of agents, so the classic asymptotic results from the literature do not apply, as well as cases where it is possible for agents to be incompetent (i.e., their chance of being correct might be less than $1/2$).

We show that when judges are even minimally competent, we can guarantee that the weak order and sign of the weights assigned to experts will be correct. Additionally, we showed that if judges use the log-odds weighting and are reasonably accurate as a group, we will recover the optimal weighting function. Moreover, we show some cases where judges bring a meaningful benefit to the process. However, our results regarding how to divide a group of agents – a particularly relevant issue for scientific conferences – indicate that multiple judges may be unhelpful, and there are cases (e.g., uniform distributions) in which an additional expert is more valuable than a judge.

There are several interesting future directions. One is to reconsider the problem we have presented here when the weights given by experts must all be non-negative, or when it is required for each judge j that $\sum_{e \in E} w_{je} = 1$ (as required in Aziz et al. [3, 4] for the setting of peer evaluation). Another is to examine what happens when the hierarchy level is increased by adding an additional layers (as in large conferences, which have Area Chairs in charge of SPCs, in charge of PC members). At what point does it no longer become helpful (or begin to be helpful)? Can a guarantee of minimal quality of judges change the value proposition of having them? Further, it would be interesting to explore how this framework extends to non-binary information, as when PCs provide more detailed information to SPCs than just acceptance or rejection.

ACKNOWLEDGMENTS

Nicholas Mattei was supported by NSF Awards IIS-RI-2007955, IIS-III-2107505, IIS-RI-2134857, IIS-RI-2339880 and CNS-SCC-2427237 as well as the Tulane University Jurist Center for Artificial Intelligence and the Tulane University Center for Community-Engaged Artificial Intelligence. Ben Abramowitz was supported by the NSF under Grant #2127309 to the Computing Research Association for the CIFellows Project. Omer Lev was supported, in part, by NSF-BSF grant #2021659, and by Israel Science Fund (ISF) grants #1965/20 and #3152/20.

REFERENCES

- [1] Ben Abramowitz and Nicholas Mattei. 2019. Flexible representative democracy: an introduction with binary issues. In *Proceedings of the 28th International Joint Conference on Artificial Intelligence*. 3–10.
- [2] Ben Abramowitz and Nicholas Mattei. 2025. Flexible representative democracy. *Social Choice and Welfare* 64, 1 (2025), 263–308.
- [3] Haris Aziz, Omer Lev, Nicholas Mattei, Jeffrey S. Rosenschein, and Toby Walsh. 2016. Strategyproof Peer Selection: Mechanisms, Analyses, and Experiments. In *Proceedings of the 30th Conference on Artificial Intelligence (AAAI)*. Phoenix, Arizona, 397–403.
- [4] Haris Aziz, Omer Lev, Nicholas Mattei, Jeffrey S. Rosenschein, and Toby Walsh. 2019. Strategyproof peer selection using randomization, partitioning, and apportionment. *Artificial Intelligence (AIJ)* 275 (October 2019), 295–309. <https://doi.org/10.1016/j.artint.2019.06.004>
- [5] Eyal Baharad, Jacob Goldberger, Moshe Koppel, and Shmuel Nitzan. 2012. Beyond Condorcet: Optimal aggregation rules using voting records. *Theory and decision* 72, 1 (2012), 113–130.
- [6] Roy Baharad, Shmuel Nitzan, and Erel Segal-Halevi. 2022. One person, one weight: when is weighted voting democratic? *Social Choice and Welfare* (2022), 1–27.
- [7] Itzhack Bar-Itzhack and Richard Harman. 2003. The effect of Sensor Failure on the Attitude and Rate Estimation of the MAP Spacecraft. In *AIAA Guidance, Navigation, and Control Conference and Exhibit*. 5485.
- [8] Ruben Becker, Gianlorenzo D’angelo, Esmail Delfaraz, and Hugo Gilbert. 2021. Unveiling the Truth in Liquid Democracy with Misinformed Voters. In *International Conference on Algorithmic Decision Theory*. Springer, 132–146.
- [9] Ruth Ben-Yashar and Jacob Paroush. 2000. A nonasymptotic Condorcet jury theorem. *Social Choice and Welfare* 17, 2 (2000), 189–199.
- [10] Daniel Berend and Aryeh Kontorovich. 2015. A finite sample analysis of the Naive Bayes classifier. *J. Mach. Learn. Res.* 16, 1 (2015), 1519–1545.
- [11] Daniel Berend and Jacob Paroush. 1998. When is Condorcet’s jury theorem valid? *Social Choice and Welfare* 15, 4 (1998), 481–488.
- [12] Daniel Berend and Luba Sapir. 2003. Between the expert and majority rules. *Advances in Applied Probability* 35, 4 (2003), 941–960.
- [13] Daren C Brabham. 2013. *Using crowdsourcing in government*. IBM Center for the Business of Government Washington, DC.
- [14] Daren C Brabham. 2015. *Crowdsourcing in the public sector*. Georgetown University Press.
- [15] Nicolo Cesa-Bianchi, Yoav Freund, David Haussler, David P Helmbold, Robert E Schapire, and Manfred K Warmuth. 1997. How to use expert advice. *Journal of the ACM (JACM)* 44, 3 (1997), 427–485.
- [16] Marquis De Condorcet. 1785. *Essai sur l’application de l’analyse à la probabilité des décisions rendues à la pluralité des voix*. De l’Imprimerie Royale (Paris).
- [17] Franz Dietrich and Kai Spiekermann. 2023. Jury Theorems. In *The Stanford Encyclopedia of Philosophy* (Spring 2023 ed.), Edward N. Zalta and Uri Nodelman (Eds.). Metaphysics Research Lab, Stanford University.
- [18] Patricia Everaere, Sébastien Konieczny, and Pierre Marquis. 2010. The epistemic view of belief merging: can we track the truth? In *ECAI 2010*. IOS Press, 621–626.
- [19] Scott L Feld and Bernard Grofman. 1984. The accuracy of group majority decisions in groups with added members. *Public Choice* 42, 3 (1984), 273–285.
- [20] Bailey Flanigan, Paul Gözl, Anupam Gupta, Brett Hennig, and Ariel D. Procaccia. 2021. Fair algorithms for selecting citizens’ assemblies. *Nature* 596 (2021), 548–552.
- [21] Rupert Freeman, David Pennock, Chara Podimata, and Jennifer Wortman Vaughan. 2020. No-regret and incentive-compatible online learning. In *International Conference on Machine Learning*. PMLR, 3270–3279.
- [22] Bernard Grofman. 1978. Judgmental competence of individuals and groups in a dichotomous choice situation: Is a majority of heads better than one? *Journal of Mathematical Sociology* 6, 1 (1978), 47–60.
- [23] Bernard Grofman and Scott L Feld. 1983. Determining optimal weights for expert judgment. In *Information Pooling and Group Decision Making: Proceedings of the Second University of California, Irvine, Conference on Political Economy*. JAI Press Greenwich, CT, 167–72.
- [24] Bernard Grofman, Scott L Feld, and Guillermo Owen. 1984. Group size and the performance of a composite group majority: Statistical truths and empirical results. *Organizational Behavior and Human Performance* 33, 3 (1984), 350–359.
- [25] Bernard Grofman, Guillermo Owen, and Scott L Feld. 1983. Thirteen theorems in search of the truth. *Theory and decision* 15, 3 (1983), 261–278.
- [26] Jonas Karge and Sebastian Rudolph. 2022. The More the Worst-Case-Merrier: A Generalized Condorcet Jury Theorem for Belief Fusion.. In *KR*.
- [27] Drora Karotkin, Samuel Nitzal, and Jacob Paroush. 1988. The essential ranking of decision rules in small panels of experts. *Theory and Decision* 24 (1988), 253–268.
- [28] Drora Karotkin and Jacob Paroush. 1994. Variability of decisional ability and the essential order of decision rules. *Journal of Economic Behavior & Organization* 23, 3 (1994), 343–354.
- [29] Omer Lev, Nicholas Mattei, Paolo Turrini, and Stanislav Zhydkov. 2021. Peer Selection with Noisy Assessments. *arXiv preprint arXiv:2107.10121* (2021).
- [30] Omer Lev, Nicholas Mattei, Paolo Turrini, and Stanislav Zhydkov. 2023. Peer-Nomination: A novel peer selection algorithm to handle strategic and noisy assessments. *Artificial Intelligence (AIJ)* 316 (March 2023), 103843.
- [31] Christian List and Robert E Goodin. 2001. Epistemic democracy: Generalizing the Condorcet jury theorem. (2001).
- [32] Malik Magdon-Ismael and Lirong Xia. 2018. A mathematical model for optimal decisions in a representative democracy. *Advances in Neural Information Processing Systems* 31 (2018).
- [33] Nicholas Mattei, Paolo Turrini, and Stanislav Zhydkov. 2021. PeerNomination: Relaxing exactness for increased accuracy in peer selection. In *Proceedings of the Twenty-Ninth International Conference on International Joint Conferences on Artificial Intelligence (IJCAI)*. 393–399.
- [34] Shmuel Nitzan and Jacob Paroush. 1982. Optimal decision rules in uncertain dichotomous choice situations. *International Economic Review* (1982), 289–297.
- [35] Lawrence Page, Sergey Brin, Rajeev Motwani, and Terry Winograd. 1999. *The PageRank citation ranking: Bringing order to the web*. Technical Report. Stanford InfoLab.
- [36] Marcus Pivato and Arnold Soh. 2020. Weighted representative democracy. *Journal of Mathematical Economics* 88 (2020), 52–63.
- [37] Manon Revel, Tao Lin, and Daniel Halpern. 2021. The Optimal Size of an Epistemic Congress. *arXiv preprint arXiv:2107.01042* (2021).
- [38] Jan-Willem Romeijn and David Atkinson. 2011. Learning juror competence: A generalized Condorcet jury theorem. *Politics, Philosophy & Economics* 10, 3 (2011), 237–262.
- [39] Robert E Schapire and Yoav Freund. 2013. Boosting: Foundations and algorithms. *Kybernetes* 42, 1 (2013), 164–166.
- [40] Nihar B. Shah. 2022. Challenges, experiments, and computational solutions in peer review. *Commun. ACM* 65, 6 (2022), 76–87.
- [41] Lloyd Shapley and Bernard Grofman. 1984. Optimizing group judgmental accuracy in the presence of interdependencies. *Public Choice* 43, 3 (1984), 329–343.
- [42] James Surowiecki. 2005. *The wisdom of crowds*. Anchor.
- [43] Nassim Nicholas Taleb. 2007. *The Black Swan: The Impact of the Highly Improbable*. Random House.
- [44] Alan Taylor and William Zwicker. 1992. A characterization of weighted voting. *Proceedings of the American mathematical society* 115, 4 (1992), 1089–1094.
- [45] Zoi Terzopoulou. 2023. Voting with limited energy: A study of plurality and Borda. In *Proceedings of the 2023 International Conference on Autonomous Agents and Multiagent Systems*. 2085–2093.
- [46] Chris Thornton, Frank Hutter, Holger H Hoos, and Kevin Leyton-Brown. 2013. Auto-WEKA: Combined selection and hyperparameter optimization of classification algorithms. In *Proceedings of the 19th ACM SIGKDD international conference on Knowledge discovery and data mining*. 847–855.
- [47] Volodimir G Vovk. 1990. Aggregating strategies. *Proc. of Computational Learning Theory, 1990* (1990).
- [48] Toby Walsh. 2014. The PeerRank Method for Peer Assessment. In *Proceedings of the 21st European Conference on Artificial Intelligence (ECAI)*. Prague, Czech Republic, 909–914.
- [49] Jens Witkowski and David C Parkes. 2012. Peer prediction without a common prior. In *Proceedings of the 13th ACM Conference on Electronic Commerce*. 964–981.
- [50] Yuzhe Zhang and Davide Grossi. 2022. Tracking Truth by Weighting Proxies in Liquid Democracy. In *Proceedings of the 21st International Conference on Autonomous Agents and Multiagent Systems*. 1482–1490.