

Agreement, Diversity, and Polarization Indices for Approval Elections

Piotr Faliszewski¹, Jitka Mertlová², Krzysztof Sornat¹, Stanisław Szufa³ and Tomasz Wąs⁴

¹AGH University of Kraków

²Czech Technical University in Prague

³University of Geneva

⁴University of Oxford

Abstract

An index is a function that measures the extent to which an election has a particular feature. We seek indices that capture agreement, diversity, and polarization among voters in approval elections, normalized with respect to saturation. The latter means that if two elections differ by the fraction of candidates approved by an average voter, but otherwise are of similar nature, then they should have similar index values. We propose several indices, analyze their properties, and use them to derive a new map of approval elections, and compare various real-life elections from Pabulib, Preflib and other sources.

1 Introduction

Our main goal is to design three functions that given an approval election—i.e., an election where voters give yes/no answers regarding their support for particular candidates—indicate the extent to which the votes are in agreement, are diverse, or are polarized. Following Hashemi and Endriss [2014], we refer to functions that quantify election features (by giving values between 0 and 1) as *election indices*.

Such indices would be valuable because they would allow us to compare approval elections coming from different sources. For example, it is not uncommon to use participatory budgeting (PB) elections from Pabulib [Faliszewski *et al.*, 2023a] outside of the PB context and it is important to validate if this data remains relevant.¹ Further, as recently shown by Faliszewski *et al.* [2025b], for ordinal elections (where the voters rank the candidates) having good agreement, diversity, and polarization indices leads to appealing maps of elections (i.e., visualizations of election datasets where each election is a dot and distances between these dots give an idea of their elections’ similarity). While a map of approval elections exists [Szufa *et al.*, 2022], it seems to be less useful than those for ordinal ones (see the argument below).

One characteristic issue when dealing with approval elections is that they may largely differ in terms of their *saturation*, i.e., the average number of candidates approved by

a voter, while being very similar in nature. As an extreme example, consider one election where all voters approve the same 10% of candidates, and a second one, where all voters approve the same 50% of candidates. They both represent perfect agreement, but the map of Szufa *et al.* [2022] puts them far apart. Indeed, its main dimensions are election saturation and the position on the agreement-disagreement spectrum. We seek saturation-independent agreement, diversity, and polarization indices, so a map based on them would be capable of discovering more nuanced relations between data.

Our general strategy is to focus on designing a saturation-independent agreement index. Given such an index, we derive a diversity one by clustering the voters into as like-minded groups as possible and claiming that the election is diverse if the agreement within these groups is low. For polarization—which we define as the presence of two groups with conflicting preferences—we take the difference between the overall agreement and the agreements in two voter groups after clustering (the overall agreement in a polarized election would be low, but it would be high in the like-minded groups). This approach was proposed by Faliszewski *et al.* [2023b] in the ordinal case, for the special case of measuring agreement using the Kemeny rule [Kemeny, 1959]. We extend their idea.

Altogether, our contributions are as follows. First, we propose a number of agreement indices and analyze them theoretically and experimentally, paying careful attention to saturation independence. In particular, we verify if their results are intuitive on a number of selected election instances. Then, using these indices, we derive and test diversity and polarization ones (using both the general scheme and some other ideas). Finally, we identify three indices—capturing agreement, diversity, and polarization—that work best together in a certain measurable way, and use them to form a map of approval elections. Our map contains both synthetic and real-life elections, e.g., coming from Preflib [Mattei and Walsh, 2013] and Pabulib [Faliszewski *et al.*, 2023a], but also from other sources, including some new ones. Omitted proofs and discussions are in the full version of the paper [Faliszewski *et al.*, 2026a]; for code appendix, see: github.com/Project-PRAGMA/Approval-DAP-IJCAI-2026.

2 Preliminaries

For an integer t , we write $[t]$ to denote $\{1, \dots, t\}$. For a vector $x \in \mathbb{R}^t$, we write $x[i]$ to denote its i -th coordinate, as well as

¹E.g., for year 2025 Google Scholar points to 18 papers that cite the Pabulib database and perform numerical experiments. Among these, only 12 are indeed focused on participatory budgeting.

$\bar{x} = \frac{1}{t}(x[1] + \dots + x[t])$ to denote its mean. For $p \in [0, 1]$, $x \sim \text{Bernoulli}(p)$ indicates that x gets value 1 with probability p and value 0 with probability $1 - p$.

Elections. An (approval) election is a pair $E = (C, V)$, where $C = \{c_1, \dots, c_m\}$ is a set of candidates and $V = (v_1, \dots, v_n)$ is a collection of voters. Each voter has an approval ballot (or, a vote) where he or she indicates which candidates he or she approves (and he or she disapproves the other ones). Formally, we represent a vote as a binary vector, whose entries correspond to the candidates and have value 1 if a given candidate is approved and 0 otherwise. To simplify notation, we write v_i to refer both to the respective voter and his or her vote; the exact meaning will always be clear from the context. Additionally, for a voter v_i we write $A(v_i)$ to denote the set of candidates that this voter approves, and for a candidate c_j we write $A(c_j)$ to mean the set of voters approving c_j . Given an election $E = (C, V)$, by E^{-1} we mean an election obtained from E by reversing all entries of all votes (i.e., for each voter v_i and candidate c_j , in E^{-1} , for c_j this voter reports $1 - v_i[j]$). We refer to E^{-1} as a reverse of E .

Pearson Correlation Coefficient. Given two t -dimensional vectors, $x, y \in \mathbb{R}^t$, their Pearson correlation coefficient (PCC) is:

$$\text{pcc}(x, y) = \frac{\sum_{i=1}^t (x[i] - \bar{x})(y[i] - \bar{y})}{\sqrt{\sum_{i=1}^t (x[i] - \bar{x})^2} \sqrt{\sum_{i=1}^t (y[i] - \bar{y})^2}}. \quad (1)$$

$\text{pcc}(x, y)$ takes values between 1 and -1 , where 1 means perfect correlation (positive linear dependence between the entries of x and y), -1 means perfect anti-correlation (negative linear dependence between the entries of x and y), and 0 means a complete lack of linear correlation. In rare cases when x or y contains only 1s or only 0s, the denominator would be equal to 0, so we fix the value of $\text{pcc}(x, y)$ as 1.

(Dis)similarity Measures Between Votes. Given two votes, u and v , over m candidates, their Hamming distance is $\text{ham}(u, v) = \sum_{i \in [m]} |u[i] - v[i]|$ and their Jaccard distances is $\text{jacc}(u, v) = 1 - \frac{\sum_{i \in [m]} \min(u[i], v[i])}{\sum_{i \in [m]} \max(u[i], v[i])}$. Hamming distance takes values between 0 and m , whereas Jaccard takes values between 0 and 1. We also sometimes measure similarity between two votes using Pearson correlation coefficient, which in the case of binary vectors is known as the phi coefficient [Yule, 1912], or Matthews coefficient (MCC) [Baldi *et al.*, 2000]. There are many other (dis)similarity measures between binary vectors (see, e.g., the discussion of Duan *et al.* [2014]) but we restrict our attention to these three as they are common in voting literature.

3 Election Data

We test our election indices on three kinds of data: Special, fixed elections of particular, well-understood structure, synthetically generated random elections, and elections obtained from various real-life data (ranging from actual approval elections to those derived from other types of preferences). In the discussion below we consider elections with candidate set $C = \{c_1, \dots, c_m\}$ and voter collection $V = (v_1, \dots, v_n)$.

Special Elections. Our special elections are:

p -Identity (p -ID). In a p -Identity election, all voters approve the same set of $\lfloor pm \rfloor$ candidates (e.g., $c_1, \dots, c_{\lfloor pm \rfloor}$) and disapprove the remaining ones.

k -Party. In a k -Party election there are k disjoint, equal-sized groups of candidates (± 1 , depending on divisibility) and k disjoint, equal-sized groups of voters (± 1 , depending on divisibility). For each $i \in [k]$, the voters in the i -th group approve exactly the candidates in the i -th group. For $x, y \in [0, 1]$, by (x, y) -2-Party we mean a 2-Party election where the first group includes an x fraction of candidates and y fraction of voters, and the second group includes the remaining ones.

For the next three elections we require $n = m$:

Diagonal. This is an m -Party election.

Triangle. For each $i \in [n]$, the i -th voter approves the first i candidates, i.e., $c_1, \dots, c_i$.

Cyclic. The first voter approves the first candidate and last $\lfloor m/2 \rfloor - 1$ candidates. For each $i \in \{2, \dots, n\}$, we obtain the i -th vote by taking the $(i - 1)$ -st one and shifting it cyclically by one position to the right, i.e., replacing each candidate c_j by c_{j+1} , with c_m mapped to c_1 .

Synthetic Elections. Next let us describe our models of generating random elections. If we only describe a process of sampling a single vote, it means that all the votes in the election are generated this way, independently from each other:

p -Impartial Culture (p -IC). For $p \in [0, 1]$, we sample vote v so that for each $i \in [m]$, we have $v[i] \sim \text{Bernoulli}(p)$.

$(p_1, \dots, p_m)$ -Independent Approval Model (IAM). For $(p_1, \dots, p_m) \in [0, 1]^m$, we sample vote v so that for each $i \in [m]$, $v[i] \sim \text{Bernoulli}(p_i)$. IAMs were discussed, e.g., by Faliszewski *et al.* [2025a], Lackner and Maly [2025], and Xia [2025].

(p, ϕ) -Resampling. Given $p, \phi \in [0, 1]$ and a fixed central vote u with $|A(u)| = \lfloor pm \rfloor$, we sample vote v as follows: For each $i \in [m]$, with probability $1 - \phi$ we let $v[i] = u[i]$ and with probability ϕ we let $v[i] \sim \text{Bernoulli}(p)$. This model is due to Szufa *et al.* [2022]. Some authors, such as Faliszewski *et al.* [2025a], allow arbitrary central votes, but we use the original variant as it gives pm approved candidates in expectation.

Euclidean. Here, we assume that each voter and candidate is associated with a point in a d -dimensional Euclidean space, and each voter v_i approves a candidate c if the distance between their corresponding points is smaller than some threshold value r_i (which can be different for each voter). Such a model was studied, e.g., by Godziszewski *et al.* [2021].

We also consider a few special distributions, to generate elections with particular features:

p -ID/IC. The first half of the voters form a p -ID election and the second half is generated using the p -IC model.

Lin-IC. In this model, v_1 is generated using 1-IC, v_2 using $(1 - 1/n)$ -IC, v_3 using $(1 - 2/n)$ -IC, and so on, until v_n , generated using $1/n$ -IC.

$N(\mathcal{D}, \phi)$. $\mathcal{D}$ is a distribution over elections. We first generate an election using $\mathcal{D}$ and then we replace each vote v with a new one as follows: We let p be the fraction of candidates approved in v and we generate a vote using (p, ϕ) -resampling model with central vote equal to v . In other words, $N(\mathcal{D}, \phi)$ introduces noise on top of elections generated using $\mathcal{D}$ (N stands for *noisy*).

4 Indices and Initial Experiments

By an election index, we mean a function f that given an election E outputs a value between 0 and 1, quantifying the extent to which E has a feature captured by f . In the following sections we define agreement, diversity, and polarization indices. For now, we discuss indices on a high level and describe our approach to evaluating them.

Normalization. Given an election $E = (C, V)$, we consider its average vote length (avl; the average number of candidates approved by a voter), reverse avl (the average number of candidates disapproved by a voter), and its saturation (satr; the fraction of candidates that an average voter approves):

$$\text{avl}(E) = \frac{1}{|V|} \sum_{v \in V} |A(v)|, \quad \text{rev-avl}(E) = \text{avl}(E^{-1}),$$

$$\text{satr}(E) = \text{avl}(E)/|C|.$$

While saturation is a natural election index by itself, for the other indices we are mostly interested in normalizing them in a way that would be independent of saturation. In other words, if two elections are similar in nature except that their saturations differ (as would be the case, e.g., for 0.5-IC and 0.25-IC elections), we would prefer an election index to give them similar values. Naturally, this is not a precise mathematical statement, but rather an intuitive desideratum. Yet, in some cases it is easy to argue that a particular index is not saturation-independent, and we often view it as an argument against this index.

Resampling Experiment. For each of our indices, we perform the following *resampling experiment*: For each $p \in \{0.1, \dots, 0.9\}$ and each $\phi \in \{0, 0.1, \dots, 1\}$, we generate 10 (p, ϕ) -resampling elections with 60 candidates and 60 voters, and compute for them the average value of the index in question. We arrange the results as a matrix, where each row corresponds to a fixed value of p and each column corresponds to a fixed value of ϕ . Color brightness represents the average value of the index. We present such matrices in the first row of Table 1, and we provide an enlarged example (including

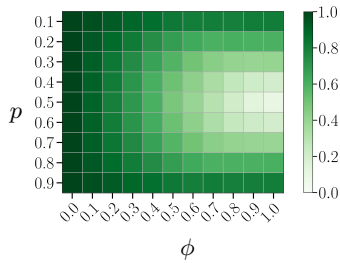


Figure 1: Results of the resampling experiment for Approval Agreement index (fixed p values in rows, and fixed ϕ values in columns).

axis labels and the color-bar scale) in Figure 1. The important feature of these plots is that each column corresponds to a fixed value of ϕ , while saturation (i.e., p) varies across rows. Hence, if the columns of the matrix are not (approximately) uniformly colored, then we claim that a given index is not saturation-independent (as, arguably, resampling elections with different values of p but the same value of ϕ are of the same nature). For such an index, we say that it *fails the resampling test of saturation independence*.

Index Values on Example Elections. In Table 1 we also give values of our indices on a number of characteristic (distributions over) elections. Each row gives index values for a given family of elections. All elections have 60 candidates and 60 voters and the results for each family are averaged over sampling 10 elections (we report standard deviation in smaller font). For each family we give its name and an icon presenting an example election; in this icon, columns are the candidates, rows are the voters, and darker spots indicate that a given voter approves a given candidate. We discuss this table throughout the following sections.

5 Agreement Indices

An agreement index should identify whether the voters have similar views, i.e., tend to approve and disapprove the same candidates. We consider two types of such indices: Global ones look at the extent to which all voters agree with each other, whereas local ones also detect large groups of voters that are internally in agreement, but that may disagree with each other. We propose six agreement indices that we describe below. All our indices give value 1 exactly for p -ID elections—which, indeed, represent perfect agreements—and for most of them we also characterize when they give value 0.

5.1 Approval Agreement

Our first index, Approval Agreement, adapts an idea from the world of ordinal elections (where the voters rank the candidates). There, for every pair of candidates a and b we compute the absolute difference between the fractions of voters that rank them one way or the other, and output the average of these values [Alcalde-Unzu and Vorsatz, 2013, Hashemi and Endriss, 2014, Can *et al.*, 2015, Faliszewski *et al.*, 2023b]. In the approval setting, we compare the numbers of voters approving and disapproving each given candidate.

Definition 5.1. Given an election $E = (C, V)$, its *Approval Agreement* is $\text{av-agr}(E) = \frac{1}{|C|} \sum_{c \in C} \left| \frac{|A(c)|}{|V|} - \left(1 - \frac{|A(c)|}{|V|}\right) \right| = \frac{1}{|C|} \sum_{c \in C} \left| 1 - 2 \frac{|A(c)|}{|V|} \right|$.

Unfortunately, this index has several drawbacks. Foremost, as we see in the first row of Table 1, it fails the resampling test of saturation independence (and, indeed, we see that it gives different values for $1/2$ -IC and $1/4$ -IC). This is also visible in the following proposition, where we show that it assumes value 0 if and only if each candidate is approved by exactly half of the voters and, hence, it is not possible to assume this value for saturation different from 0.5.

Proposition 5.2. For an election E , $\text{av-agr}(E) = 1$ exactly if E is an identity election, and $\text{av-agr}(E) = 0$ exactly if each candidate is approved by exactly half of the voters.

		av-agr	cntr-agr	pabi-agr	pcc-agr	jacc-agr	pcc ⁺ -agr	cntr-div	pcc-div	out-div	cntr-pol	pcc-pol	pabi-pol
resampling experiment:													
saturation independence:		✗	✗	✓	✓	✗	✓	✗	✓	✓	—	—	—
1.	1/3-ID		1.00 ±0.00	1.00 ±0.00	1.00 ±0.00	1.00 ±0.00	1.00 ±0.00	0.00 ±0.00	0.00 ±0.00	0.00 ±0.00	0.00 ±0.00	0.00 ±0.00	0.00 ±0.00
2.	2-Party		0.00 ±0.00	0.00 ±0.00	0.00 ±0.00	0.00 ±0.00	0.50 ±0.00	0.20 ±0.00	0.23 ±0.01	0.10 ±0.00	1.00 ±0.00	1.00 ±0.00	1.00 ±0.00
3.	N(2-Party, 0.6)		0.10 ±0.01	0.08 ±0.01	0.01 ±0.00	0.01 ±0.00	0.35 ±0.00	0.66 ±0.02	0.82 ±0.01	0.29 ±0.00	0.27 ±0.07	0.17 ±0.01	0.24 ±0.01
4.	3-Party		0.33 ±0.00	0.00 ±0.00	0.00 ±0.00	0.00 ±0.00	0.33 ±0.00	0.33 ±0.00	0.31 ±0.01	0.14 ±0.00	0.33 ±0.00	0.50 ±0.00	0.63 ±0.00
5.	4-Party		0.50 ±0.00	0.00 ±0.00	0.00 ±0.00	0.00 ±0.00	0.25 ±0.00	0.45 ±0.00	0.40 ±0.00	0.16 ±0.00	0.25 ±0.00	0.33 ±0.00	0.43 ±0.00
6.	2-Party(1/3)		0.33 ±0.00	0.24 ±0.00	0.10 ±0.00	0.11 ±0.00	0.56 ±0.00	0.15 ±0.00	0.18 ±0.00	0.10 ±0.00	0.76 ±0.00	0.89 ±0.00	0.99 ±0.00
7.	Cyclic		0.00 ±0.00	0.00 ±0.00	0.00 ±0.00	0.00 ±0.00	0.39 ±0.00	0.46 ±0.00	0.57 ±0.02	0.22 ±0.00	0.50 ±0.00	0.33 ±0.00	0.58 ±0.00
8.	Diagonal		0.97 ±0.00	0.00 ±0.00	0.00 ±0.00	0.00 ±0.00	0.02 ±0.00	0.97 ±0.00	0.97 ±0.00	0.64 ±0.01	0.02 ±0.00	0.02 ±0.00	0.01 ±0.00
9.	Triangle		0.50 ±0.00	0.49 ±0.00	0.33 ±0.00	0.50 ±0.00	0.51 ±0.00	0.41 ±0.00	0.32 ±0.00	0.18 ±0.00	0.01 ±0.00	0.14 ±0.00	0.47 ±0.00
10.	N(Tri,0.6)		0.22 ±0.01	0.20 ±0.01	0.06 ±0.00	0.15 ±0.02	0.33 ±0.01	0.89 ±0.01	0.77 ±0.03	0.26 ±0.00	0.00 ±0.02	0.05 ±0.00	0.38 ±0.01
11.	1/2-ID/IC		0.49 ±0.01	0.50 ±0.02	0.26 ±0.01	0.26 ±0.01	0.51 ±0.01	0.40 ±0.00	0.50 ±0.00	0.19 ±0.00	0.08 ±0.02	0.26 ±0.01	0.45 ±0.01
12.	1/2-IC		0.11 ±0.01	0.10 ±0.01	0.02 ±0.00	0.02 ±0.00	0.35 ±0.01	0.80 ±0.01	0.91 ±0.00	0.29 ±0.00	0.07 ±0.01	0.05 ±0.00	0.18 ±0.00
13.	1/4-IC		0.50 ±0.01	0.00 ±0.00	0.02 ±0.00	0.02 ±0.00	0.16 ±0.01	0.98 ±0.00	0.91 ±0.00	0.31 ±0.00	0.01 ±0.01	0.05 ±0.00	0.16 ±0.00
14.	Lin-IC		0.08 ±0.01	0.07 ±0.01	0.01 ±0.00	0.09 ±0.03	0.31 ±0.00	0.96 ±0.00	0.84 ±0.03	0.27 ±0.00	0.00 ±0.01	0.04 ±0.00	0.37 ±0.01

Table 1: Values of our indices for elections from various distributions. The elections have 60 candidates and 60 voters. Index values are averaged over 10 samples and the standard deviation is reported in smaller font.

The index also behaves poorly on Party elections. As we see in Table 1, for 2-Party its value is 0, but it increases as we consider 3-Party and 4-Party, to become close to 1 for Diagonal. Indeed, the index treats the fact that many voters disapprove the same candidates as a sign of agreement among them (while in some cases—such as elections within parliaments—this indeed might be a reasonable view, in others—such as in participatory budgeting—it is ungrounded).

5.2 Central Agreement(s)

Our second index, Central Agreement, is designed to resemble the Kemeny voting rule for ordinal elections [Kemeny, 1959]. In that setting, a Kemeny ranking is a ranking that minimizes the sum of swap distances to all the votes in a given election. The larger is this value, the bigger is the disagreement among the voters. We implement the same idea for approval elections using the Hamming distance (while computing a Kemeny ranking is well-known to be intractable [Bartholdi III *et al.*, 1989, Hemaspaandra *et al.*, 2005], in our case the analogous task is very simple).

Let $E = (C, V)$ be an election. A vote u over candidate set C is central for E if for each candidate $c \in C$, u agrees

regarding c with at least half of the voters. So, if exactly half of the voters approve some candidate and half disapprove it, then there is more than one central vote for the given election. We write $\text{cen}(E)$ to denote the set of all central votes for E .

Definition 5.3. Let $E = (C, V)$ be an election and let u be an arbitrary member of $\text{cen}(E)$. We define $\text{cntr-agr}(E) = 1 - \frac{\sum_{v_i \in V} \text{ham}(v_i, u)}{|V| \cdot \min(\text{avl}(E), \text{rev-avl}(E))}$.

The denominator in the definition of $\text{cntr-agr}(E)$ is chosen to be (nearly) equal to the largest possible value of the numerator for elections of the same size and saturation as E .

Proposition 5.4. For an election E , $\text{cntr-agr}(E) = 1$ exactly if E is an identity election, and $\text{cntr-agr}(E) = 0$ exactly if either $\text{cen}(E)$ contains a vote approving all candidates or $\text{cen}(E)$ contains a vote disapproving all candidates.

Central Agreement has two major drawbacks. First, as indicated by Proposition 5.4, it is quite non-discriminative and often assumes value 0 (e.g., on the majority of Pabulib elections). Second, it fails the resampling test of saturation independence (see Table 1). Interestingly, Central Agreement and Approval Agreement differ mostly by their normalization

(replacing the denominator with m in Proposition 5.5 below, we would get Approval Agreement).

Proposition 5.5. *Let $E = (C, V)$, then $\text{cntr-agr}(E) = 1 - \sum_{c \in C} (1 - |1 - \frac{2|A(c)|}{|V|}|) / (2 \cdot \min(\text{avl}(E), \text{rev-avl}(E)))$.*

We do not consider central agreement based on the Jaccard distance as it would be intractable [Chierichetti *et al.*, 2010] and we find sufficiently many other agreement indices.

5.3 Pairwise Agreements

The final four indices are based on computing expected (dis)similarity between two randomly selected votes.

Definition 5.6. *For an election $E = (C, V)$, we define Hamming, Jaccard, PCC, and PCC^+ Pairwise Agreement Indices:*

$$\begin{aligned} \text{pair-agr}(E) &= 1 - \frac{\sum_{u \in V} \sum_{v \in V} \text{ham}(u, v)}{2|V|^2|C| \cdot \text{satr}(E)(1 - \text{satr}(E))}, \\ \text{jacc-agr}(E) &= \frac{1}{|V|^2} \sum_{u \in V} \sum_{v \in V} (1 - \text{jacc}(u, v)), \\ \text{pcc-agr}(E) &= \frac{1}{|V|^2} \sum_{u \in V} \sum_{v \in V} \text{pcc}(u, v), \\ \text{pcc}^+\text{-agr}(E) &= \frac{1}{|V|^2} \sum_{u \in V} \sum_{v \in V} \max(0, \text{pcc}(u, v)). \end{aligned}$$

The normalization for Hamming Pairwise Agreement follows that for Central Agreement: the denominator is equal to the maximum possible value of the numerator among elections with the same size and saturation (assuming that it is possible to ensure that every candidate has the same approval score). The normalizations for Jaccard and PCC^+ use the fact that we average over values that are between 0 and 1. For PCC Agreement, we need an additional argument.

Proposition 5.7. *For each election E , $\text{pcc-agr}(E) \in [0, 1]$.*

At first sight, computing each of these indices requires $O(|C||V|^2)$ time: Computing similarity between a pair of votes requires $O(|C|)$ time and there are $O(|V|^2)$ pairs to consider. Yet, by changing the order of summation, we can compute Hamming Pairwise Agreement in $O(|C||V|)$ time.

Theorem 5.8. *For an election $E = (C, V)$, we have that $\text{pair-agr}(E) = 1 - \frac{\sum_{c \in C} |A(c)| \cdot (|V| - |A(c)|)}{|V|^2|C| \cdot \text{satr}(E)(1 - \text{satr}(E))}$.*

Next we characterize when Hamming and Jaccard Pairwise indices assume values 1 and 0. In case of PCC-based indices we only provide some examples and observations.

Theorem 5.9. *Each of Pairwise, Jaccard, PCC, and PCC^+ Agreement Indices assumes value 1 exactly for identity elections. Hamming Pairwise Agreement has value 0 exactly if all candidates have equal approval score, and Jaccard Pairwise Agreement has value 0 exactly if all the voters approve disjoint sets of candidates.*

Proposition 5.10. *For each k -Party election $E = (C, V)$ such that both $|C|$ and $|V|$ are divisible by k , we have $\text{pcc-agr}(E) = 0$ and $\text{pcc}^+\text{-agr}(E) = 1/k$.*

Remark 5.11. *Let $p \in [0, 1]$ be a number and let u and v be two votes such that a p^2 fraction of candidates is approved in both votes, a $(1 - p)^2$ fraction is disapproved in both, a $p(1 - p)$ fraction is only approved in u , and a $p(1 - p)$ fraction is only approved in v . Then $\text{pcc}(u, v) = 0$. As votes in p -IC elections are close to satisfying these conditions, we expect PCC and PCC^+ Agreements to be close to 0 for IC elections.*

We note that Hamming Pairwise Agreement and PCC Agreement tend to give very similar results for the examples in Table 1. In fact, if all voters approve the same number of candidates, then the two indices are equal (but this is not the case when vote lengths differ; see the Triangle election).

Proposition 5.12. *Let $E = (C, V)$ be an election such that $0 < |A(u)| = |A(v)| < |C|$ for all $u, v \in V$. Then we have that $\text{pair-agr}(E) = \text{pcc-agr}(E)$.*

In the resampling experiment both Hamming and PCC Agreements give results independent of saturation. Overall, we view both of them as good global agreement indices.

Jaccard and PCC^+ Pairwise Agreement indices capture local agreement, by detecting large agreeing groups, even if they disagree with other such groups (indeed, in PCC^+ we only sum positive PCC values to account for like-minded voters, disregarding negative correlations that could cancel them out). For example, they give high—and nearly identical—values for party elections (the fewer parties, the higher agreement, as there are larger agreeing groups) and detect agreeing groups in $1/2$ -ID/IC, and Lin-IC. Yet, we prefer PCC^+ Agreement as it passes the resampling test of saturation independence and assigns close-to-zero values to IC elections. The fact that Jaccard Pairwise Agreement fails the resampling test of saturation is expected, as it treats approvals and disapprovals asymmetrically. If such asymmetry is natural in a given setting it still might be a valuable index.

6 Diversity and Polarization Indices

Intuitively, an election is diverse if (a) its votes tend to be different from each other, and (b) they cover the space of all possible votes (modulo saturation). For example, we view Diagonal as a maximally diverse election, because for its saturation (one approval per voter) it contains exactly one copy of each possible vote. Yet, we view Cyclic, which is similar in spirit to Diagonal, as much less diverse than $1/2$ -IC; these two types of election have the same saturation, but the votes in the latter are notably more different from each other and, intuitively, cover the preference space more evenly.

On the other hand, we say that an election is polarized if it consists of two groups of voters that agree internally but strongly disagree with each other. Consequently, 2-Party is an example of a maximally polarized election. We acknowledge that elections with three, four, or more internally consistent but disagreeing groups could also be seen as maximally polarized—and seeking indices capturing this intuition would be valuable—but we leave this for future work.

For a broader view of diversity and polarization, see the special issue edited by Levin *et al.* [2021]. In voting literature, see, e.g., the works of Nehring and Puppe [2002], Hashemi and Endriss [2014], Can *et al.* [2015], Karpov *et al.* [2024], Ammann and Puppe [2025], and Faliszewski *et al.* [2023b, 2025b, 2026b, 2026c].

6.1 Clustering-Based Diversity and Polarization

Our first approach to designing diversity indices follows the view of Faliszewski *et al.* [2023b, 2026b], that an election is diverse if even after clustering its voters into like-minded

groups, these groups disagree internally. As opposed to them, we parameterize our definition with an agreement index.

For a given integer k and an election $E = (C, V)$, let $\mathcal{P}_k(V)$ be a set of all complete partitions of voters into k disjoint subsets $V_1, \dots, V_k$, and let a be some agreement index. We aim to find a partition that maximizes average agreement within the subsets, weighted by the size of each cluster, i.e., $\sum_{i=1}^k \frac{|V_i|}{|V|} a(C, V_i)$. In principle, there are many ways in which the values for different k can be aggregated to give a single index. We follow a simple approach of Faliszewski *et al.* [2025b] and take the average of values for k from 1 to 5. This results in the following definition.

Definition 6.1. For an election $E = (C, V)$, and agreement index a , we define the a -Diversity index as:

$$a\text{-div}(E) = 1 - \frac{1}{5} \sum_{k=1}^5 \max_{\{V_1, \dots, V_k\} \in \mathcal{P}_k(V)} \sum_{i=1}^k \frac{|V_i|}{|V|} a(C, V_i).$$

Taking Central Agreement as the underlying index, we obtain Central Diversity, denoted cntr-div , and taking PCC Pairwise Agreement we obtain PCC Diversity, denoted pcc-div . We do not consider Hamming Pairwise Agreement for the sake of focus and due to its similarity to pcc-agr . In both cases, obtaining the optimal partition of votes is not computationally feasible, thus we rely on clustering heuristics. We approximate cntr-div with an algorithm based on the k -means clustering: We employ the k -medoids version using Hamming distance, where cluster centers are restricted to observed votes and are iteratively updated to minimize the sum of intra-cluster distances. For pcc-div we use spectral clustering [Ng *et al.*, 2001]; our input for spectral clustering is the number k of clusters and the affinity matrix, where for each pair of votes u, v we provide value $1 + \frac{1}{2} \text{pcc}(u, v)$ (value 0 means total dissimilarity, and 1 means the votes are equal).

To quantify polarization, we measure the increase of agreement arising from clustering the election into two groups instead of viewing it as a single group (again, the general idea is due to Faliszewski *et al.* [2023b, 2026b]).

Definition 6.2. For an election $E = (C, V)$, and agreement index a , we define the a -Polarization index as:

$$a\text{-pol}(E) = \max_{\{V_1, V_2\} \in \mathcal{P}_2(V)} \sum_{i=1}^2 \frac{|V_i|}{|V|} a(C, V_i) - a(E).$$

We consider Central Polarization (cntr-pol) and PCC Polarization (pcc-pol), based on the Central and PCC Pairwise Agreement indices, and we use the same clustering heuristics as in the case of diversity indices.

Additionally, we consider the Pairwise Polarization index, which outputs the standard deviation of Hamming distances between all pairs of votes. The intuition is that in a polarized electorate each two voters either have very similar or close-to-opposite preferences (factor of $2/|C|$ ensures values in $[0, 1]$).

Definition 6.3. For an election $E = (C, V)$, the Pairwise Polarization index is defined as $\text{pair-pol}(E) = 2\sqrt{\frac{\sum_{u,v \in V} (\text{ham}(u,v) - \sum_{x,y \in V} \text{ham}(x,y)/|V|^2)^2}{|V|^2|C|^2}}$.

6.2 Outer Diversity

Recently, Faliszewski *et al.* [2026c] proposed an alternative diversity index, based on the average distance between the

votes in the election and all possible votes. If the average distance is small, the votes are well spread in the whole space, i.e., diverse. It was originally proposed in the ordinal setting and for domains (sets of votes without specified multiplicity) rather than elections, but we adapt it to our setting as follows.

For two collections of votes over the same set of candidates, $U = (u_1, \dots, u_k)$ and $V = (v_1, \dots, v_n)$, by $\text{ham}(U, V)$ we mean the average Hamming distance between n copies of votes in U and k copies of votes in V , matched so as to minimize this distance, i.e.:

$$\text{ham}(U, V) = \frac{1}{kn} \min_{\pi: [kn] \rightarrow [kn]} \sum_{i=1}^{kn} \text{ham}(u_i, v_{\pi(i)}),$$

where π is a bijection and indices from U are taken modulo k and from V modulo n .

For a candidate set C , let $\mathcal{U}_C$ consist of a single copy of each vote from $\{0, 1\}^{|C|}$. Then, in principle, we could measure diversity of an election $E = (C, V)$ based on $\text{ham}(V, \mathcal{U}_C)$, i.e., the average distance between votes in V and all possible votes. However, if the saturation of E is very small or large, the value of $\text{ham}(V, \mathcal{U}_C)$ is always substantial (as most votes in $\mathcal{U}_C$ have moderate saturation). Thus, no election with small/large saturation could be considered diverse according to such a measure. To circumvent this issue, instead of $\mathcal{U}_C$ we use $p\text{-}\mathcal{U}_C$, defined as a collection in which the number of copies of each vote is proportional to the probability of drawing such a vote from $p\text{-IC}$ (so, $1/2\text{-}\mathcal{U}_C = \mathcal{U}_C$).

We would like the index to be maximum and equal to 1 if $V = p\text{-}\mathcal{U}_C$, and minimum and equal to 0 if V is a singleton. From the following lemma we get that for a single vote v with saturation p , it holds that $\text{ham}(\{v\}, p\text{-}\mathcal{U}_C) = 2p(1 - p)$.

Lemma 6.4. For candidate set C , $q \in \{0, \frac{1}{|C|}, \dots, 1\}$ and $v \in \{0, 1\}^{|C|}$ with $q|C|$ ones it holds that $\text{ham}(\{v\}, p\text{-}\mathcal{U}_C) = p(1 - q) + q(1 - p)$.

Definition 6.5. For an election $E = (C, V)$ with $\text{satr}(E) = p$, the Outer Diversity index is defined as $\text{out-div}(E) = 1 - \frac{1}{2p(1-p)} \text{ham}(V, p\text{-}\mathcal{U}_C)$.

As discussed, the value 1 obtained for $V = p\text{-}\mathcal{U}_C$, and a single vote v always yields the value of 0. It remains to prove that this is indeed the minimum and maximum value.

Proposition 6.6. For each election E , $\text{out-div}(E) \in [0, 1]$.

6.3 Comparison of the Indices

As we see in Table 1, Central Diversity fails the resampling test of saturation independence, but the other diversity indices pass it. For polarization indices the test is not meaningful, as we expect low polarization for all values of p and ϕ parameters. Clustering-based diversity indices give close-to-1 values for Diagonal and IC elections (including Lin-IC), while outer diversity performs poorly and gives low values for these elections. We tested that its values are much closer to 1 for analogous elections with a notably larger number of voters, but overall this index is not satisfying (even if outer diversity for ordinal preference domains works well [Faliszewski *et al.*, 2026c]). Further, all diversity indices give higher values for the noisy variants of respective elections (see, e.g., the values for 2-Party and N(2-Party, 0.6)).

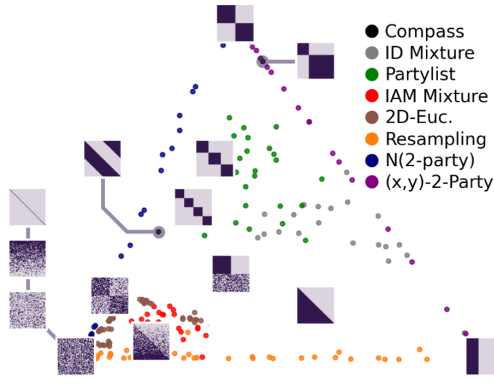


Figure 2: Map of synthetic elections.

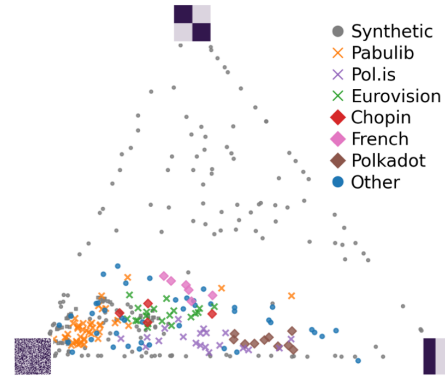


Figure 3: Map of real-life elections.

All polarization indices give value 1 for 2-Party and lower values for its noisy variant, as well as for 3- and 4-Party, as expected. Pairwise Polarization still gives nearly-1 for $(1/3, 1/3)$ -2-Party whereas Central and PCC Polarization give lower ones, but both approaches could be justified. Overall, Pairwise Polarization tends to give highest values. For example, for Triangle it gives nearly $1/2$, while the other indices give values much closer to 0. Both approaches could be justified, either by Triangle’s similarity to 2-Party, or by saying that in Triangle all voters agree on best-to-worst ranking of the candidates, but have different approval thresholds.

7 Experiments

To present experimental results for our indices, we use the *map of elections* framework [Szufa *et al.*, 2025]. We construct the map as follows. First, we assemble a collection of elections. Next, we compute pairwise distances between elections using the feature distance, as in the work of Faliszewski *et al.* [2025b]. Specifically, let $f(E) = (f_1(E), \dots, f_k(E))$ denote the feature vector of election E , where $f_1, \dots, f_k$ are our indices. For two elections E and F , we define their distance as the (standard) Euclidean distance between their feature vectors, $d_f(E, F) = \|f(E) - f(F)\|_2$. Finally, we embed the resulting distance matrix using multidimensional scaling (MDS [Kruskal, 1964]).

For the map, we use the following three indices: pcc-agr, pcc-div, and pcc-pol.² In Figure 2, we show a map based on synthetic elections, while Figure 3 presents a map based on real-life ones (all our datasets will be publicly available). Each point corresponds to a single election. We describe the selected/generated elections in the full version of the paper.

Interestingly, the synthetic map forms a triangle whose vertices correspond to three extreme cultures: ID (maximizing agreement), IC (maximizing diversity), and 2-party (maximizing polarization). The edge between IC and ID is populated by Resampling elections; the edge between 2-party and IC is filled by $N(2\text{-Party})$ elections; and the edge between

2-party and ID is filled by (x, y) -2-Party elections. Both 2D-Euclidean (brown) and IAM (red) elections lie close to IC, reflecting high diversity together with low polarization and agreement. Finally, the partylist (green) elections occupy the upper region of the triangle, indicating high polarization.

Turning to the real-life map (Figure 3), we observe that most elections exhibit very low polarization and, on average, substantially higher diversity than agreement. Compared with other real-world datasets, the Polkadot blockchain data (stakeholders voting over validators) shows relatively high agreement. Both Polkadot and Pol.is (users approving comments on an online deliberation platform) exhibit particularly low polarization. Intuitively, this may reflect a shared “search for ground truth” dynamic in both settings. Notably, the Eurovision (song contest) data looks similar to the Chopin (piano competition) one, which is interesting given that both involve people evaluating musical performances. In comparison to other real-life data, the French (presidential) elections seem relatively polarized. Regarding Pabulib (participatory budgeting voting), with few exceptions, the data shows high diversity, with low agreement and low polarization. (This may stem from the fact that, in most cities, people can approve only a very limited number of projects, so the saturation is extremely low). To conclude, it is worth noting that our map of real-life elections resembles the one from the work of Faliszewski *et al.* [2025b], in which (while focusing on ordinal elections), most real-life instances appeared in a similar region of the map, i.e., low polarization, mid to high diversity, and low to mid agreement.

Finally, we note that MDS embeddings of maps in Figures 2 and 3 achieve mean multiplicative distortions of 1.014 and 1.010, respectively, indicating extremely good representation. In particular, this is far more accurate than the original maps presented in [Szufa *et al.*, 2025, Section 4.3.2], where overall distortion is between 1.21 and 1.26.

8 Future Work

Three main directions for future work include (a) a deeper study of agreement indices in ordinal elections, (b) finding more principled diversity/polarization indices (also for polarization with multiple disagreeing groups), and (c) looking for further features of (approval) elections and their indices.

²The sum of these indices is the most consistently close to 1 among those studied. This means they are complementary and, together, provide a good high-level summary of the nature of a given election. For elections with more than 200 candidates or more than 1000 voters, we computed the indices using their sampled subsets.

Acknowledgments

This project has received funding from the European Research Council (ERC) under the European Union’s Horizon 2020 research and innovation programme (grant agreement No 101002854), from the European Union under the project Robotics and advanced industrial production (reg. no. CZ.02.01.01/00/22_008/0004590), and from the Foundation for the University of Geneva. T. Wąs was supported by the UK Engineering and Physical Sciences Research Council (EPSRC) under grant EP/X038548/1. The research presented in this paper has been partially supported by the funds of Polish Ministry of Science and Higher Education assigned to AGH University of Science and Technology. Further, it was inspired by research proposed for NCN project UMO-2025/58/A/ST6/00371.



References

- [Alcalde-Unzu and Vorsatz, 2013] Jorge Alcalde-Unzu and Marc Vorsatz. Measuring the cohesiveness of preferences: An axiomatic analysis. *Social Choice and Welfare*, 41(4):965–988, 2013.
- [Ammann and Puppe, 2025] Matthias Ammann and Clemens Puppe. Preference diversity. *Review of Economic Design*, 2025. Online First.
- [Baldi et al., 2000] Pierre Baldi, Søren Brunak, Yves Chauvin, Claus A. F. Andersen, and Henrik Nielsen. Assessing the accuracy of prediction algorithms for classification: An overview. *Bioinformatics*, 16(5):412–424, 2000.
- [Bartholdi III et al., 1989] John Bartholdi III, Craig Tovey, and Michael Trick. Voting schemes for which it can be difficult to tell who won the election. *Social Choice and Welfare*, 6(2):157–165, 1989.
- [Can et al., 2015] Burak Can, Ali Ihsan Ozkes, and Ton Storcken. Measuring polarization in preferences. *Mathematical Social Sciences*, 78:76–79, 2015.
- [Chierichetti et al., 2010] Flavio Chierichetti, Ravi Kumar, Sandeep Pandey, and Sergei Vassilvitskii. Finding the Jaccard median. In *Proceedings of SODA-2010*, pages 293–311, 2010.
- [Duan et al., 2014] Lian Duan, W. Nick Street, Yanchi Liu, Songhua Xu, and Brook Wu. Selecting the right correlation measure for binary data. *ACM Transactions on Knowledge Discovery from Data*, 9(2):13:1–13:28, 2014.
- [Faliszewski et al., 2023a] Piotr Faliszewski, Jarosław Flis, Dominik Peters, Grzegorz Pierczyński, Piotr Skowron, Dariusz Stolicki, Stanisław Szufa, and Nimrod Talmon. Participatory budgeting: Data, tools and analysis. In *Proceedings of IJCAI-2023*, pages 2667–2674, 8 2023.
- [Faliszewski et al., 2023b] Piotr Faliszewski, Andrzej Kaczmarczyk, Krzysztof Sornat, Stanisław Szufa, and Tomasz Wąs. Diversity, agreement, and polarization in elections. In *Proceedings of IJCAI-2023*, pages 2684–2692, 2023.
- [Faliszewski et al., 2025a] Piotr Faliszewski, Łukasz Janeczko, Andrzej Kaczmarczyk, Marcin Kurdziel, Grzegorz Pierczyński, and Stanisław Szufa. Learning real-life approval elections. In *Proceedings of AAMAS-2025*, pages 704–712, 2025.
- [Faliszewski et al., 2025b] Piotr Faliszewski, Jitka Mertlová, Pierre Nunn, Stanisław Szufa, and Tomasz Wąs. Distances between top-truncated elections of different sizes. In *Proceedings of AAAI-2025*, pages 13823–13830, 2025.
- [Faliszewski et al., 2026a] Piotr Faliszewski, Jitka Mertlová, Krzysztof Sornat, Stanisław Szufa, and Tomasz Wąs. Agreement, diversity, and polarization indices for approval elections. Technical Report arXiv:2605.14983 [cs.GT], arXiv.org, May 2026.
- [Faliszewski et al., 2026b] Piotr Faliszewski, Krzysztof Sornat, Stanisław Szufa, and Tomasz Wąs. Diversity of structured domains via k -Kemeny scores. In *Proceedings of AAAI-2026*, pages 16880–16888, 2026.
- [Faliszewski et al., 2026c] Piotr Faliszewski, Krzysztof Sornat, Stanisław Szufa, and Tomasz Wąs. Outer diversity of structured domains. In *Proceedings of AAMAS-2026*, 2026. To appear. See authors’ website.
- [Godziszewski et al., 2021] Michał Godziszewski, Paweł Batko, Piotr Skowron, and Piotr Faliszewski. An analysis of approval-based committee rules for 2D-Euclidean elections. In *Proceedings of AAAI-2021*, pages 5448–5455, 2021.
- [Hashemi and Endriss, 2014] Vahid Hashemi and Ulle Endriss. Measuring diversity of preferences in a group. In *Proceedings of ECAI-2014*, pages 423–428, 2014.
- [Hemaspaandra et al., 2005] Edith Hemaspaandra, Holger Spakowski, and Jörg Vogel. The complexity of Kemeny elections. *Theoretical Computer Science*, 349(3):382–391, 2005.
- [Karpov et al., 2024] Alexander Karpov, Klas Markström, Søren Riis, and Bei Zhou. Local diversity of Condorcet domains. Technical Report arXiv:2401.11912 [econ.TH], arXiv.org, January 2024.
- [Kemeny, 1959] John Kemeny. Mathematics without numbers. *Daedalus*, 88:577–591, 1959.
- [Kruskal, 1964] Joseph Kruskal. Multidimensional scaling by optimizing goodness of fit to a nonmetric hypothesis. *Psychometrika*, 29(1):1–27, 1964.
- [Lackner and Maly, 2025] Martin Lackner and Jan Maly. Approval-based shortlisting. *Social Choice and Welfare*, 64(1):97–142, 2025.
- [Levin et al., 2021] Simon A. Levin, Helen V. Milner, and Charles Perrings. The dynamics of political polarization. *Proceedings of the National Academy of Sciences*, 118(50):e2116950118, 2021.
- [Mattei and Walsh, 2013] Nicholas Mattei and Toby Walsh. Preflib: A library for preferences. In *Proceedings of ADT-2013*, pages 259–270, 2013.

- [Nehring and Puppe, 2002] Klaus Nehring and Clemens Puppe. A theory of diversity. *Econometrica*, 70(3):1155–1198, 2002.
- [Ng *et al.*, 2001] Andrew Ng, Michael Jordan, and Yair Weiss. On spectral clustering: Analysis and an algorithm. In *Proceedings of NeurIPS-2001*, pages 849–856, 2001.
- [Szufa *et al.*, 2022] Stanisław Szufa, Piotr Faliszewski, Łukasz Janeczko, Martin Lackner, Arkadii Slinko, Krzysztof Sornat, and Nimrod Talmon. How to sample approval elections? In *Proceedings of IJCAI-2022*, pages 496–502, 2022.
- [Szufa *et al.*, 2025] Stanisław Szufa, Niclas Boehmer, Robert Brederbeck, Piotr Faliszewski, Rolf Niedermeier, Piotr Skowron, Arkadii Slinko, and Nimrod Talmon. Drawing a map of elections. *Artificial Intelligence*, 343:104332, 2025.
- [Xia, 2025] Lirong Xia. A linear theory of multi-winner voting. Technical Report arXiv:2503.03082 [cs.GT], arXiv.org, March 2025.
- [Yule, 1912] G. Udny Yule. On the methods of measuring the association between two attributes. *Journal of the Royal Statistical Society*, 75:579–652, 1912.