

Candidate Selection by Parties: Crime and Politics in India

Arvind Magesan

Economics Department, University of Calgary

Andrea Szabó

Economics Department, University of Houston

Gergely Ujhelyi

Economics Department, University of Houston

June 4, 2024

Abstract

We study how parties choose candidates, a key issue to understand political selection and ultimately policy choices. Do parties select candidates that voters like, or are their choices shaped by other considerations? What is the impact of policies that limit parties' choice sets, such as restrictions on candidates with a criminal history? To study these questions, we combine rich candidate level data from India with a model in which parties trade-off the electoral appeal of candidates against internal party preferences in a strategic game of candidate selection. We find that parties' preferences systematically deviate from voters'. While parties select candidates who are likely to win, all else equal they prefer those who are not overly popular. Selection decisions are also driven by strategic considerations, as well as payoffs that are independent of voter preferences, such as the ease of recruiting certain candidates. Our estimates provide a nuanced explanation for parties' motivation to run criminal candidates, and, through counterfactual simulations, shed light on the potential impacts of banning criminals from contesting elections.

1 Introduction

Do parties select candidates that voters like, or are their choices shaped by other considerations? In a representative democracy, elections aggregate voter preferences over candidates *who appear on the ballot*. Therefore the question of whom parties choose to run is key to understand political selection, and ultimately policy choices.

Although a sizeable literature now explores individuals’ own decision to enter politics (Besley 2005; Dal Bó and Finan 2018), we know much less about how parties select their candidates. As Dal Bó and Finan (2018, p.566) write: “political parties likely play a major role in who becomes a politician, and yet we have a very limited understanding of how political parties recruit and screen their candidates.”

Clearly, parties like to win, which requires running candidates that voters will support. At the same time, there are several reasons why a party may not want to run the most popular candidate.

Party elites may value candidate traits like loyalty, influence or wealth even if these do not lead to more votes. For example, wealthy candidates can bring resources to the party or self-finance their campaign. Similarly, candidates from influential groups can provide valuable connections and access for party leaders. Such benefits have been proposed as an explanation for the widespread nomination of (wealthy and well-connected) criminal politicians in India (Vaishnav 2017).

Sometimes a candidate who enjoys “too much” voter support may in fact threaten the position of party elites, or steer the party’s policies in a direction that is inconsistent with the elite’s preferences. In the US, these concerns are well illustrated by changes in the Republican party since the 2016 electoral campaign of Donald Trump. In India, where parties are highly centralized, candidates who become too popular can also give rise to internal power struggles and weaken the party leader (Chandra 2016).

In this paper, we study candidate selection in Indian national elections, combining a discrete-choice model of voter preferences over candidates with a strategic game of candidate selection between the two main party alliances. Estimating the model quantifies several factors that cause party choices to systematically deviate from the maximization of winning probabilities (or vote shares). We use our framework to model the impact of policies that affect the set of candidates available to parties, in particular the banning of candidates with criminal backgrounds.

Our primary source of data is the Election Commission of India. We focus on the 2009 and 2014 national elections, using additional information from state elections held between 2008-2017. The dataset contains vote shares and the number of eligible voters, as well

as several candidate characteristics, such as gender and caste. We include information on candidate wealth, education, and criminal backgrounds from affidavits that candidates are required to file with the Election Commission, and create an indicator for candidates with Muslim names.¹ We match to this data constituency characteristics from the Indian Census. This village level information comes from the SHRUG database (Asher et al. 2020) and allows us capture heterogeneity in voter preferences.

We begin by estimating a comprehensive BLP discrete-choice demand system describing voter preferences among candidates in Indian national elections (Berry, Levinsohn and Pakes 1995). Voters have preferences over specific candidate characteristics, such as education, wealth, whether the candidate has a criminal background, and whether the candidate is Muslim. These preferences are also shaped by constituency characteristics and unobserved candidate characteristics.

Our specification of the supply side focuses on candidate choices by the two main party alliances in Indian national politics, the NDA and the UPA, led by the two largest parties, the BJP and the INC, respectively. We use a simultaneous game of incomplete information to model the strategic interaction between these two players, and estimate their preferences over candidates using the Nested Pseudo-Likelihood procedure of Aguirregabiria and Mira (2007).

To make the estimation of this game conceptually and practically feasible, we assume that parties’ choice sets are comprised of clusters of observed candidate characteristics, or candidate “types.” We approximate the pool of potential candidate types that parties select from in the national election using the set of candidates contesting elections to Indian state legislatures.² We adopt a machine learning algorithm to identify candidate types in this state election data - essentially, these are combinations of candidate characteristics that tend to occur together in the data.³ This clustering algorithm identifies four candidate types, all of which turn out to have clear interpretations: an “educated type” (educated non-Muslim

¹Muslims are a particularly salient group in Indian politics, but this characteristic is rarely used in academic research due to a lack of data. We create an indicator for Muslim candidates based on their names using methods from text analysis, assigning candidates to different groups based on the “distance” of their name from libraries of Muslim and non-Muslim names. Throughout we use “Muslim” to describe ethnicity, proxied by name, rather than religion.

²A large fraction (almost 25% in our sample years) of candidates for state elections are independent candidates - many of these wanted to run under the banner of a major party but were not selected. In addition, many national level candidates begin their political careers at more local levels of politics (Dar 2019) further supporting the idea that candidates for state election approximate the pool from which national candidates are drawn.

³Our approach here is similar in spirit to Bandiera et al. (2020) who reduce high dimensional data on CEO activities to a small set of CEO “types” in order to study how CEOs behavior affects firm performance. Similarly, Hamilton et al. (2021) use a clustering approach to reduce the choice set of patients choosing between different medial treatments.

with no criminal history), an “uneducated type” (uneducated non-Muslim with no criminal history), a “Muslim type”, and a “criminal type” (non-Muslim with a criminal history who is also relatively wealthy).

Using our demand estimates, we construct counterfactual vote shares corresponding to all combinations of candidate types potentially chosen by the competing parties. This allows us to compute expected vote shares and win probabilities. Parties’ objective functions nest these expected vote shares and win probabilities, as well as a set of heterogeneous costs of running different candidates.

Estimating parties’ objective functions reveals that while parties prefer candidate types with better chances of winning, all else equal, they prefer candidates with a lower expected vote share. In other words, parties prefer to win with less popular candidates. This finding is consistent with the idea that party elites trade off anticipated electoral performance against the threat of a candidate becoming too powerful and undermining the leadership or its policies, as suggested by the literature on Indian electoral politics (see Section 3.1).

We also find that considerations that are independent of voter preferences (and hence win probabilities or vote shares) matter in parties’ objective functions. All else equal, candidate types that are more common in the relevant local candidate pool are less costly for parties to recruit and run. Parties also have direct preferences over candidate types: for example, the NDA has a particularly large direct cost of running a Muslim type, which is in line with its declared Hindu nationalist profile. Interestingly, we find that both parties obtain a positive direct payoff from running criminal types relative to others. This is consistent with Vaishnav (2017), where criminals bring benefits to the party in the form of (organized crime) networks and an ability to finance their own campaigns.

The pattern of estimated payoff parameters has a striking equilibrium implication: according to our results, criminal candidates are often strategic complements. In other words, parties are often compelled to run a criminal candidate because the other party is running a criminal as well. Intuitively, this happens when, faced with an opposition criminal candidate, the party’s only hope for a win is to also select a criminal candidate. From the parties’ joint perspective, running criminals may thus be inefficient.

Motivated by this last observation as well as recent policy proposals in India and around the world, we use our framework to model the impact of a ban on criminal candidates.⁴

⁴Restrictions on candidates with a criminal history is also a salient issue outside of India. In Brazil, where 40% of candidates running for governor in 2014 had pending court cases against them, the 2014 Clean Record Law banned convicted criminals from running for a period of 8 years (Paiva, Sakai and Schoenster 2014). In the US, states vary in whether they allow convicted criminals to run for office. In Louisiana and Maine, convicted felons are eligible to run; in Massachusetts, they are ineligible to run while incarcerated, while in Texas they are ineligible even after the completion of their sentence. See <https://restoration.ccresourcecenter.org>.

One implication of the ban is to change the distribution of candidate types contesting elections, leading to higher fractions of educated, uneducated, and Muslim types. According to our results, the ban also has implications for parties’ expected winning probabilities. We find that, with a criminal ban in effect, the vote share of third party candidates rises, lowering the winning probability of both major parties. It appears that some voters’ preference for the major parties is conditional on these parties’ ability to run criminal types.

In spite of this, we find that in 34% of constituencies *both* parties benefit from a criminal ban. The reason for this is the strategic complementarity of criminal candidates. Once their opponent is banned from running a criminal candidate, each party finds other candidate types more attractive. Thus, banning criminals allows parties to compete with candidates who deliver higher payoffs in equilibrium. This provides an explanation for why political parties could support a ban on criminal candidates while at the same time running such candidates in elections.

2 Related literature

Our paper contributes to a growing literature that studies parties’ role in political selection. We consider an important setting, India, where the previous literature on this issue has been largely descriptive. Our structural approach provides a methodological contribution, and allows us to ask several new questions.

In the political selection literature, several papers analyze the tradeoffs that parties make between electability and other considerations, and how this impacts the “quality” of candidates. Galasso and Nannicini (2011) present a theoretical model where high quality candidates are valued by swing voters but are expensive for parties to recruit. This tradeoff leads to parties running high quality candidates in more competitive districts, which is consistent with data from Italy. Mattozzi and Merlo (2015) analyze a similar tradeoff, between the quality of candidates and the effort they are willing to exert to build the party organization. In their model, parties can sometimes increase effort by not running the candidate preferred by voters. Besley et al. (2017) study a model where running the best candidate would jeopardize party leaders’ survival, and find evidence consistent with their predictions in Sweden.⁵

In this set of papers, candidate “quality” is measured in terms of education or residuals from a Mincerian wage regression. By contrast, our approach makes it possible to study selection on multiple dimensions simultaneously. We do not take an a priori stance on what

⁵See also Casas-Arce and Saiz (2015), who show evidence that, prior to the introduction of a gender quota, Spanish parties failed to maximize votes by not running more female candidates.

constitutes a high quality candidate, and we also let the data tell us what the relevant types of candidates are in parties' choice sets.⁶

In the developing country context, two recent papers conduct field experiments that create changes in the candidate nomination process to study party leaders' behavior. [Gulzar, Hai and Paudel \(2021\)](#) provide party leaders in Nepal with information sheets on potential candidates, including information on their party service, competence, and popularity among voters. [Casey, Kamara and Meriggi \(2021\)](#) survey voters in Sierra Leone about their preferred candidate, and then randomize an intervention where this information is shared with party leaders, and potential candidates present their qualifications and debate each-other in a public forum. Both papers find that providing information resulted in parties fielding more candidates preferred by voters, which is consistent with a lack of information about voter preferences at baseline.

Our approach is complementary to this strand of the literature. First, we quantify several different mechanisms behind candidate selection, including parties' payoff from winning, vote shares, direct costs, and recruiting costs. Second, neither experiment was designed to study parties' strategic behavior in the candidate selection process,⁷ something we explicitly account for. Third, while both experimental studies identify the average treatment effect of the intervention, neither of them was designed to measure the weights of different objectives in parties' candidate selection procedures. Our structural approach makes it possible to answer questions regarding *counterfactual* policy interventions, including in principle interventions in different settings.⁸ Fourth, creating experimental changes in real-world nomination procedures must necessarily be limited due to ethical considerations (see Appendix 3 of [Gulzar, Hai and Paudel \(2021\)](#) for a detailed discussion), which creates issues for the replicability of these studies in other time periods or other countries. Our study does not face this limitation, as our methodology can easily be applied to other contexts.

Conceptually, the paper closest to ours is [Iaryczower, Kim and Montero \(2023\)](#), who present a structural analysis of parties' policy positions (ideology) in Brazil. In their framework, candidates are exogenously given, and parties choose ideological positions constrained

⁶Like us, [Dal Bó et al. \(2017\)](#) also considers candidate selection on multiple dimensions in their comprehensive study of Swedish politicians. While their focus is on *self*-selection, they show some evidence that parties are more willing to promote individuals who are competent (as measured by cognitive and leadership abilities) independently of their background. This is consistent with the idea that in Sweden socioeconomic and ethnic background likely play a considerably smaller role in politics than in India or in developing democracies more generally.

⁷In the [Casey, Kamara and Meriggi \(2021\)](#) experiment parties were able to choose which districts were included in the experiment, and they mostly chose safe districts with little competition.

⁸For example, one could consider which component of party preferences might take different values in another setting, and use our model to draw conclusions there as well. The possibility of this type of inference adds to the external validity of our findings, and is one of the advantages of a model-based approach.

by their candidates’ ideology. Instead, we focus on how parties select their candidates. In doing so we take party ideology as exogenously given, which is appropriate in a setting where the central party organization sets the policy platform and the elected candidates follow it, as is the case for the two major Indian parties we study.

3 Background and data

3.1 Background

3.1.1 Parties and elections

We study general elections to India’s national legislature (the *Lok Sabha*), a near ideal setting for analyzing strategic candidate selection by political parties.

The Lok Sabha is comprised by a large number of single member districts (called “constituencies”). In each election, voters in a given constituency elect one and only one representative from the available choices on the ballot. In contrast to a setting with proportional representation, each competing party selects a single candidate to run in the constituency.

There are two main competing (pre-election) alliances, the United Progressive Alliance (UPA) and the National Democratic Alliance (NDA), led by the two main national parties, the INC and the BJP, respectively. These alliances run candidates in almost all constituencies in each election and together win the majority of seats. In every constituency, the alliance contains a group of parties (and occasionally some independent candidates) that enter into a pre-election agreement about which candidate will run to represent the alliance, without competition from other members of the group. Because our model will treat alliances as the players, we will refer to the two alliances throughout simply as *parties*.

Although we study parties’ choices in national elections, we will also make use of data from state elections (elections to states’ legislative assemblies). These are separate elections, and in most cases are held in different years from the national election. The constituencies in the two elections are different, as each national constituency is subdivided into several state constituencies. The set of parties competing in national and state elections can also be different (state elections tend to be contested by a large number of regional parties), but the UPA and the NDA are major forces in state elections as well.

3.1.2 Party goals and procedures for candidate selection

Indian parties are known for their centralized organizations in which a central committee, or in some cases a charismatic leader, dictates all major decisions, including candidate selection.

Farooqui and Sridharan (2014) review nomination procedures used in different countries, noting that “the USA represents the decentralised extreme, that of party primaries” while “India lies near the other extreme in that most of its major parties are at the completely or near-completely top-down of the six types of party nomination processes, with the national party leadership having the final say.” (p.80) Although both the INC and the BJP have formal consultation procedures that involve local party organizations in the candidate selection process, in practice decisions are ultimately made by each party’s central committee (Roy 1966; Farooqui and Sridharan 2014).

Apart from electability, important factors in the candidate selection process include loyalty to the party leadership and service to the party organization. These considerations are often explicit in parties’ written procedures on candidate selection (Roy 1966). For example, the INC has declared: “Winnability alone should not be the benchmark for deciding nominees of the party during elections. There should be a balance required between loyalty and winnability.” (Jaipur Declaration of the AICC, 2013, quoted in Chandra (2016, p.41)). The main source of this tradeoff is threats to the party from factions that undermine its leadership, or from candidates who defect to other parties (Chhibber, Jensenius and Suryanarayan 2014). Chandra (2016) summarizes this tradeoff as follows:

“The central leadership’s decisions are influenced not only by the anticipated electoral performance of the candidate, but by two intra-party considerations: (1) to ensure the compliance of powerful factions within the party so that they work for the party candidate or at a minimum do not work against the candidate and do not defect from the party and (2) to undercut factions that are becoming too powerful so that they do not threaten the leadership or its loyalists. Expectations of the electoral performance of the candidate are often subordinated to these two criteria.” (p.224)

The leadership’s desire to undercut internal rivals has also been proposed as an explanation for why Indian parties often run newcomers instead of incumbent candidates who may be perceived as becoming too powerful (Lee 2020).

Another important factor influencing parties’ candidate selection is financial considerations: financial contributions to the party and a candidate’s ability to finance their own campaign. Farooqui and Sridharan (2014, p.87) describe, in the case of the BSP party, the process through which candidates effectively bid to receive the nomination. Similarly, Vaishnav (2017) argues that the main appeal of criminal politicians to Indian parties stems from the fact that these individuals can finance their own campaigns, including by breaking campaign finance laws if necessary.

3.1.3 Rules on criminal politicians

Based on the 1951 The Representation of the People Act, criminally convicted politicians are ineligible to run for election for a period of six years after the completion of their sentence. In practice, many indicted politicians run, and win, while undergoing lengthy trials. In 2019, 43% of candidates elected to the national legislature had been indicted on criminal charges at some point.⁹

Over the last several decades multiple commissions tasked with electoral reform in India have recommended tightening the restrictions on criminal candidates.¹⁰ Perhaps surprisingly, the two main parties themselves have at different times expressed a desire to ban criminal candidates - even while actively running such candidates in elections.¹¹ Our results below will help rationalize this phenomenon.

3.2 Data and sample

We study candidate selection in India’s 2009 and 2014 national elections using a dataset of official election returns combined with candidate and constituency characteristics.

3.2.1 Elections

Election returns come from the Election Commission of India (ECI) and for each constituency they contain turnout, each candidate’s name, party, and number of votes. Constituency boundaries were set in April 2008 and are unchanged throughout our sample period.

Our specification requires information on local (state legislative) elections matched to “corresponding” national elections.¹² Each national election constituency contains a subset

⁹<https://www.hindustantimes.com/india-news/mps-with-criminal-cases-increased-in-last-decade-report-101628621064962.html>

¹⁰For example, in 2004 the Election Commission unsuccessfully recommended that indicted politicians be banned from contesting elections (https://prsindia.org/files/bills_acts/bills_parliament/2008/bill200_20081202200_Election_Commission_Proposed_Electoral_Reforms.pdf). In 2013, in *Lily Thomas v Union of India*, the Indian Supreme Court closed a loophole allowing some politicians to run while appealing their convictions.

¹¹In a speech delivered in 2010, Sonia Gandhi (then chief executive of the INC) discussed the “need to build a consensus on how to prevent individuals with a criminal record from contesting elections.” This sentiment was echoed by members of the rival BJP, including then leader of the opposition Sushma Swaraj. See “Bar Criminals from Fighting Polls - Sonia Gandhi” <https://economictimes.indiatimes.com/news/politics-and-nation/bar-criminals-from-fighting-polls-sonia-gandhi/articleshow/5500935.cms?from=mdr>, accessed on March 25, 2024. Other politicians, most prominently then BJP MP Ashini Upadhyay have brought forward petitions to the Supreme Court seeking a ban of criminal politicians https://www.livelaw.in/pdf_upload/pdf_upload-367143.pdf, accessed on March 25, 2024.

¹²India’s state elections are notorious for the sheer number of candidates who contest them, which provides us with rich variation in candidate characteristics. The Indian Electoral Commission had to sharply increase the deposit that candidates pay to contest elections due primarily to the number of candidates contesting

of the state election constituencies (this assignment is also constant after April 2008). In most states, state elections are held in different years from the national election, every 5 years. For each state, we assign the first state election held after 2008 to the 2009 national election and the second state election to the 2014 national election. In practice this means that state elections held between 2008-2012 are assigned to the 2009 national election and state elections held between 2013-2017 are assigned to the 2014 national election.

3.2.2 Candidate characteristics

The ECI data contains information on candidates’ gender, age, and caste (Scheduled Caste, Scheduled Tribe, or General). We supplement this with information on candidates’ education, wealth, and criminal history collected and published by the civil group ADR at www.myneta.info. The latter is based on affidavits that all Indian candidates for national and state elections are required to file with the ECI.¹³

The criminal histories that candidates are required to report include previous criminal convictions as well as pending cases, i.e., cases where a judge decided to proceed with a criminal charge beyond the initial police investigation and prosecutorial action (roughly equivalent to an “indictment” in the US system). [Vaishnav \(2017, p.318\)](#) provides a detailed description.

Although this dataset on candidates is already quite rich, it does not contain information on one of the most important characteristics in Indian politics: whether the candidate is an ethnic Muslim. We construct a Muslim indicator based on candidate names using tools from text analysis. Specifically, we categorize candidates as Muslim or non-Muslim based on the “distance” between their name and text fragments commonly found in Muslim names. The details are in [Appendix 1](#).

3.2.3 Constituency characteristics

In the Indian electoral system, some constituencies are reserved and can only be contested by Scheduled Caste or Scheduled Tribe candidates and the ECI data contains indicators for these reserved constituencies.

For demographic and other characteristics of each constituency, we use the SHRUG dataset ([Asher et al. 2020](#)). Specifically, we use village-level information from the 2011 Indian Census, which the SHRUG allows to be matched to constituencies. We use the following

state elections, as the cost of administering these elections was becoming prohibitive in some constituencies. See [Kapoor and Magesan \(2018\)](#) for details.

¹³Previous studies using this dataset include [Fisman, Schulz and Vig \(2016\)](#), [Prakash, Rockmore and Upal \(2019\)](#), [Ujhelyi, Chatterjee and Szabó \(2021\)](#), and [Asher and Novosad \(2023\)](#), among others.

characteristics: literacy rate, share of working population, share of Scheduled Caste and Scheduled Tribe population, whether the village has access to paved roads, and whether the village is located in a rural or urban area.

We use both the village level information, and also aggregate it up to the constituency level. This creates some missing constituencies when villages could not be uniquely matched to constituencies (see [Asher et al. \(2020\)](#)).

3.2.4 Sample construction

Details of our sample construction are in [Appendix 1](#). We drop constituencies with missing demographic information, and states with very few constituencies. Our final sample includes the 15 largest Indian states, and contains 232 national constituencies and 1629 state constituencies (about half of the constituencies in these states). The national constituencies are contested by a total of 3208 candidates in 2009 and 3373 candidates in 2014. The state constituencies are contested by 17,965 candidates in the 2009 election period and 18,801 in the 2014 election period. Summary statistics of our data are in [Table 1](#).

4 Model

We consider a simultaneous move Bayesian game of candidate selection between competing parties. Candidates are described by a set of characteristics. In selecting a candidate, each party weighs its own internal preference over candidates against the preferences of voters, and thus the probability of winning. We discuss the decision problem of parties and voters in turn.

4.1 Parties

Consider an electoral constituency where competing parties choose which candidate to run. Each party p chooses one candidate out of a set of potential candidates $\mathcal{A}_p$. Let the choice of party p be given by a_p , and denote the vector of choices of p 's opponents by $\mathbf{a}_{-p}$. Voters cast their votes based on the candidates that parties choose to run: let $s_p(a_p, \mathbf{a}_{-p})$ represent the votes cast for party p given a candidate selection profile $(a_p, \mathbf{a}_{-p})$, and let $w_p(a_p, \mathbf{a}_{-p})$ represent its corresponding winning probability (both of these functions will be derived endogenously below).¹⁴

¹⁴Throughout, s_p will correspond to the share of eligible voters who voted for party p (rather than the share among those who turn out to vote). We will refer to s_p as “vote share” for simplicity.

Table 1: Summary statistics

	N	Mean	Std. Dev.	Median	10%	90%
<i>A. Candidate characteristics - state elections</i>						
UPA (0/1)	36766	0.10				
NDA (0/1)	36766	0.11				
Education (0/1)	29137	0.62				
Muslim (0/1)	36766	0.11				
Criminal history (0/1)	30465	0.20				
Assets (log)	29903	13.90	2.36	14.00	10.84	16.83
Male (0/1)	36766	0.93				
Age	36766	44.76	11.27	44	31	61
<i>B. Candidate characteristics - national elections</i>						
UPA (0/1)	6581	0.07				
NDA (0/1)	6581	0.07				
Education (0/1)	2900	0.75				
Muslim (0/1)	6581	0.12				
Criminal history (0/1)	3068	0.25				
Assets (log)	3012	14.78	2.52	14.97	11.55	17.71
Male (0/1)	6581	0.93				
Age	6581	46.25	11.96	45	31	63
SC or ST (0/1)	6581	0.35				
<i>C. Constituency characteristics - national elections</i>						
Eligible voters (1000)	464	1421.50	195.59	1426.24	1173.14	1685.34
Turnout (%)	464	64.97	12.39	65.99	47.57	81.05
N. of candidates before aggregation	464	14.18	6.06	14	7	22
N. of candidates after aggregation	464	5.71	1.46	5	4	8
Reserved constituency (0/1)	464	0.28				
Literate population (share)	464	0.61	0.10	0.61	0.49	0.74
ST and SC population (share)	464	0.27	0.14	0.24	0.13	0.46
Rural population (share)	464	0.82	0.11	0.83	0.68	0.94
Population with paved roads (share)	464	0.83	0.19	0.90	0.58	1.00
Working population (share)	464	0.42	0.07	0.43	0.32	0.50

Notes: Education: 1 if completed high school. Criminal history: 1 if has at least one criminal case. N. of candidates after aggregation: number of candidates after small-party candidates are aggregated as described in the text. Assets (log): log of real assets in Rp.

A key innovation of our approach is to allow for the fact that parties may care about the candidate they run beyond its effect on votes. A party may experience *direct* costs or benefits from running specific candidates. This could reflect considerations such as the availability of certain types of candidates (e.g., a party with few Muslim members may find it more costly to run a Muslim candidate), internal politics (e.g., some candidates may be loyal to the party leadership, while others may come from a competing faction within the party) or party finances (e.g., some candidates may be able to finance their own campaigns, making them a “cheaper” choice for the party).

To capture these considerations, we specify party p ’s payoff from choosing candidate $a_p \in \mathcal{A}_p$ as

$$b^w w_p(a_p, \mathbf{a}_{-p}) + b^s s_p(a_p, \mathbf{a}_{-p}) + c_p(a_p) + \varepsilon_p(a_p). \quad (1)$$

The party cares about its winning probability, with weight b^w , as well as its vote share. The weight b^s on the latter could be positive if, e.g., the party leader’s status is enhanced by a large vote share. It could also be negative if popular candidates may challenge the leader’s authority or defect and form new parties, as discussed in Section 3.1. More generally, a winning candidate with a high vote share could take the party in directions that are disliked by the elite. Similarly, a losing candidate with a high vote share may be more difficult to replace in the next election.¹⁵

The term $c_p(a_p) + \varepsilon_p(a_p)$ captures payoffs from a_p that are independent of w_p and s_p . For clarity, we will refer to these direct payoffs as “costs.” The difference between $c_p(a_p)$ and $\varepsilon_p(a_p)$ is that the former is observable to all competing parties, while the latter is party p ’s private information. The private component $\varepsilon_p(a_p)$ is i.i.d. across parties and candidates, with cdf $G(\cdot)$.

An important factor affecting $c_p(a_p)$ is the pool of potential candidates available to the party. This will be determined by who the party’s members are, and who among its members has both the motivation and ability to run for office. For example, finding a Muslim candidate could be easier if the pool includes more Muslims. Or, running a candidate with a criminal history could be more socially acceptable within the party if the party has many such candidates in the pool.

As explained below, we will measure a party’s pool of potential candidates using information on the party’s candidates in state elections. This is motivated by the fact that (i) many Indian parties competing in state elections have clear affiliations to a national party (either the party is the same, or they belong to the same electoral alliance), and (ii) it is common

¹⁵Instead of vote shares, these considerations could alternatively be captured by vote margins (difference in vote share relative to the winner) or by the number of votes (implying higher values for larger constituencies). We find that using any of these alternatives makes little difference to the results - see Appendix 8.

for national politicians to begin their political careers in state elections. To highlight this, write $c_p(a_p) = c_p(a_p, L_p)$, where L_p denotes the pool of party p 's candidates in relevant state elections.

Given the presence of private information, this setup gives rise to a simultaneous game of incomplete information between parties competing in the constituency. The solution concept is Bayesian Nash Equilibrium (BNE) in pure strategies.¹⁶ For a realization of the private costs $\varepsilon_p \equiv \{\varepsilon_p(a)\}_{a \in \mathcal{A}_p}$, a party chooses candidate $a_p(\varepsilon_p)$. Let $P(\mathbf{a})$ denote the ex ante probability of a profile of choices $\mathbf{a}$. Then given ε_p and $P(\mathbf{a}_{-p})$, in a BNE party p chooses a_p to maximize its expected payoff given all other parties' strategies:

$$a_p \in \arg \max_a U_p(a, P(\mathbf{a}_{-p})),$$

where

$$U_p(a, P(\mathbf{a}_{-p})) \equiv E_P[b^w w_p(a, \mathbf{a}_{-p}) + b^s s_p(a, \mathbf{a}_{-p}) | a] + c_p(a, L_p) + \varepsilon_p(a) \quad (2)$$

is the expected value of (1) over the possible realizations of opponents' choices $\mathbf{a}_{-p}$.

For the purposes of estimation it is convenient to express strategies as *choice probabilities* (CPs). In particular, define payoffs net of the private cost as

$$\tilde{U}_p(a, P(\mathbf{a}_{-p})) \equiv U_p(a, P(\mathbf{a}_{-p})) - \varepsilon_p(a) \quad (3)$$

so that a_p maximizes party p 's expected payoffs iff:

$$\tilde{U}_p(a_p, P(\mathbf{a}_{-p})) + \varepsilon_p(a_p) \geq \tilde{U}_p(a, P(\mathbf{a}_{-p})) + \varepsilon_p(a) \quad \forall a \in \mathcal{A}_p.$$

The probability of party p choosing action a_p given the opponent's strategy $P(\mathbf{a}_{-p})$ is then:

$$\begin{aligned} P(a_p) &= \int_{\varepsilon_p} \mathbf{1} \left\{ \varepsilon_p(a) - \varepsilon_p(a_p) \leq \tilde{U}_p(a_p, P(\mathbf{a}_{-p})) - \tilde{U}_p(a, P(\mathbf{a}_{-p})), \quad \forall a \in \mathcal{A}_p \right\} dG(\varepsilon_p) \\ &\equiv \Lambda_p(a_p; \mathbf{P}_{-p}) \end{aligned} \quad (4)$$

Equilibrium in the game is fully characterized by a fixed point in $P(\mathbf{a})$ of the system of equations defined by (4) for all p . Stacking equations by actions and parties, an equilibrium vector of CPs $\mathbf{P}^*$ satisfies:

$$\mathbf{P}^* = \mathbf{\Lambda}(\mathbf{P}^*)$$

Under the assumption that $\varepsilon_p(a)$ follows the Type-I Extreme Value Distribution, the

¹⁶As long as $G(\cdot)$ is atomless, the existence of a pure-strategy BNE is guaranteed: see Fudenberg and Tirole (1991, Ch 6.8).

equilibrium CPs satisfy

$$\Lambda_p(a_p; \mathbf{P}_{-p}) = \frac{\exp \left\{ \tilde{U}_p(a_p, \mathbf{P}(\mathbf{a}_{-p})) \right\}}{\sum_a \exp \left\{ \tilde{U}_p(a, \mathbf{P}(\mathbf{a}_{-p})) \right\}} \quad \forall p.$$

An important assumption, and limitation, of the above model is that it considers each constituency in isolation. That is, conditional on observables, choosing a candidate in one constituency has no bearing on the candidate chosen in another constituency. In practice it is possible that even after conditioning on observables, party decisions in one constituency affect decisions in another. Allowing for a party to jointly decide on candidates across constituencies (a “Colonel Blotto” type game) would render the model inestimable, as this would leave us with as many markets as we have national elections (two).

4.2 Voters

To model parties’ winning probabilities and vote shares as a function of the set of candidates running, we consider the individual decisions made by a continuum of voters. We assume expressive voting with a flexible specification of voter preferences over candidates’ characteristics (Ujhelyi, Chatterjee and Szabó (2021) - USC (2021) from now on).¹⁷

Specifically, each candidate a_p can be described by a vector of characteristics $\mathbf{x}_p = \mathbf{x}(a_p)$, such as their education level or criminal history. Given a set of candidates that parties have chosen to run, voter i ’s utility from voting for the candidate of party p is

$$V_{ip} = \beta_i \mathbf{x}_p + \xi_p + \eta_{ip}. \quad (5)$$

The first term represents voters’ (potentially heterogenous) preferences over the characteristics of p ’s candidate. The second term, ξ_p , allows for unobserved (to the researcher) candidate characteristics valued by voters or, equivalently, shocks to parties’ popularity in the given constituency. The distribution of ξ_p is left unspecified, and it can be correlated with $\mathbf{x}_p$. Finally, η_{ip} are individual preference shocks drawn from a Type-I Extreme Value distribution. To model the sources of preference heterogeneity among voters, write

$$\beta_i = \beta + \mathbf{\Pi} \mathbf{d}_i, \quad (6)$$

¹⁷The assumption of expressive voting is supported by extensive survey evidence on Indian voters’ motivations. Banerjee (2017) provides a book-length discussion of the meaning that voters attach to the act of voting. Based on a recent survey, Heath and Ziegfeld (2022) estimate that at most 1.1% of individuals vote strategically.

where $\mathbf{d}_i$ is a vector of voter demographics, while $\boldsymbol{\beta}$ and $\boldsymbol{\Pi}$ contain the parameters.

To complete the voter's choice set, let $p = 0$ indicate the option to abstain and

$$V_{i0} = \boldsymbol{\pi}_0 \mathbf{d}_i + \eta_{i0}$$

the voter's associated utility. This allows for the utility of abstention (hence the cost of voting) to vary across voters.

Voter i chooses option p (vote for one of the parties or abstain) if $V_{ip} > V_{ip'}$ for all $p' \neq p$. Thus, voters choose between their options based on the observed and unobserved candidate characteristics, the benefit of abstention, and their idiosyncratic shocks. This implicitly defines the set for which voter i will choose option p , $\{(\mathbf{d}_i, \boldsymbol{\eta}_i) | V_{ip} > V_{ip'} \text{ for all } p' \neq p\}$. Given a distribution of $\mathbf{d}_i$ and $\boldsymbol{\eta}_i$, integrating over this set yields parties' vote shares as a function of their candidate choices. Under the assumed Type-I EV distribution for η_{ip} and given a distribution $F(\mathbf{d}_i)$, these vote shares can be written as

$$s_p(\mathbf{x}_p, \mathbf{x}_{-p}) = \int \frac{\exp[\boldsymbol{\beta}_i \mathbf{x}_p + \xi_p - \boldsymbol{\pi}_0 \mathbf{d}_i]}{1 + \sum_{q>0} \exp[\boldsymbol{\beta}_i \mathbf{x}_q + \xi_q - \boldsymbol{\pi}_0 \mathbf{d}_i]} dF(\mathbf{d}_i). \quad (7)$$

This setup allows the domain of voter preferences to be different from the parties' choice sets (voter preferences are defined over characteristics $\mathbf{x}$ while the parties' choice set is $\mathcal{A}_p$). This is realistic because not every party may have access to candidates with all possible combinations of characteristics.¹⁸

To model the relationship between party choices and candidate characteristics, we simply assume that, after choices are made, candidate characteristics are drawn from a distribution $H(\cdot|a)$ for each party. Thus, party p 's vote share is the expected vote share over the realizations of these characteristics:

$$s_p(a_p, \mathbf{a}_{-p}) = \int \int_{\mathbf{x}_p \mathbf{x}_{-p}} s_p(\mathbf{x}_p, \mathbf{x}_{-p}) dH_p(\mathbf{x}_p|a_p) dH_{-p}(\mathbf{x}_{-p}|\mathbf{a}_{-p}) \quad (8)$$

where $s_p(\mathbf{x}_p, \mathbf{x}_{-p})$ is given in (7). Similarly, party p 's winning probability is given by

$$w_p(a_p, \mathbf{a}_{-p}) = \int \int_{\mathbf{x}_p \mathbf{x}_{-p}} \mathbf{1}\{s_p(\mathbf{x}_p, \mathbf{x}_{-p}) > s_{-p}(\mathbf{x}_p, \mathbf{x}_{-p})\} dH_p(\mathbf{x}_p|a_p) dH_{-p}(\mathbf{x}_{-p}|\mathbf{a}_{-p}) \quad (9)$$

¹⁸This feature also allows our model to potentially be extended to situations where a party leader picking candidates may not have full control over their characteristics (for example, the leader might delegate candidate choice to subordinates). Modeling this explicitly could be an interesting avenue for future work.

In equilibrium, a party's expected vote share and winning probability takes into account its opponent's strategy, captured by $P(\mathbf{a}_{-p})$ from equation (4):

$$E_P[s_p(a_p, \mathbf{a}_{-p})|a_p] = \sum_{\mathbf{a}_{-p}} s_p(a_p, \mathbf{a}_{-p})P(\mathbf{a}_{-p}) \quad (10)$$

$$E_P[w_p(a_p, \mathbf{a}_{-p})|a_p] = \sum_{\mathbf{a}_{-p}} w_p(a_p, \mathbf{a}_{-p})P(\mathbf{a}_{-p}). \quad (11)$$

5 Specification and estimation

5.1 Overview

Our ultimate goal is to estimate parameters in the parties' objective function (2). To do this, we proceed in two stages. In the first stage, we estimate voters' utility functions (5) with a BLP procedure, using as instruments variables that enter parties objective function but do not directly enter voter utilities. This yields estimates of the voter preference parameters β and Π , as well as the popularity shocks ξ . By construction, for the candidates observed in the data, the vote shares predicted with these estimates perfectly match the observed vote shares.

With these estimates, we can use our model of voters to predict parties' vote shares given *any* combination of candidates, based on (7). In the second stage, we use these estimated vote share functions to estimate benefit and cost parameters in (2) using a Pseudo-Maximum-Likelihood procedure.

5.2 Estimating voter preferences

Estimation of voter preferences follows the Generalized Method of Moments (GMM) algorithm proposed by [Berry, Levinsohn and Pakes \(1995\)](#). Detailed treatments of the procedure can be found in [Berry, Levinsohn and Pakes \(1995\)](#), [Nevo \(2000\)](#) and [Nevo \(2001\)](#). Here we modify a previous application of this procedure to Indian *state* elections in USC (2021).

5.2.1 Specification and endogenous characteristics

We focus on four candidate characteristics $\mathbf{x}$: education, Muslim, crime, and assets.¹⁹ To be consistent with the estimation of party objectives below, we standardize all these variables

¹⁹We also considered three other characteristics observed in the data: caste, gender, and age. Gender has very little variation (almost all candidates are male). Caste has very little variation once we control for constituency reservation. Age does not seem to be an important characteristic for either voters or parties in these elections.

to have 0 mean and unit standard deviation. We also include in this vector an indicator for candidates where one or more characteristics were imputed. In addition, we include in (5) the following control variables: party and alliance fixed effects (to control for a portion of ξ_p that is common across constituencies), state and year fixed effects and an indicator for reserved constituencies.

To deal with the presence of many small parties and independent candidates, we follow USC (2021) and aggregate these candidates in each constituency. Specifically, in each constituency we aggregate into one “small party” category parties that are not part of either the UPA or the NDA alliance and only run a few candidates in the data.

The BLP procedure requires the use of instrumental variables (IV) for two reasons: first, to identify the “nonlinear” parameters Π , and second, to identify the parameters on any variables in $\mathbf{x}$ that parties can adjust in response to the popularity shocks ξ_p (the typical source of endogeneity in the identification of demand). Because the focus of our study is parties’ choice of their candidates, we treat all four candidate characteristics as endogenous.

Variables that enter parties’ objective function (2) but do not directly enter voter utility are valid instruments. In our specification, the characteristics of a party’s candidates in state elections, L_p , satisfy this condition. Recall the idea behind the presence of these variables in (2): a party’s available pool of candidates affects its cost of choosing candidates with specific characteristics. For example, a party has a lower cost of finding a highly educated candidate when most candidates in its pool are highly educated. The prevalence of high education (for example) in the pool of candidates is measured by the prevalence of this characteristic among the candidates a party runs in state elections in the same geographic area.²⁰

The characteristics of a party’s candidates in state elections are valid instruments as long as they are uncorrelated with voter valuations ξ_p in the national election. One potential threat to the identifying assumption is if L_p reflects party choices in the state election that are also affected by popularity shocks, and shocks at the national and state level are correlated. For example, a scandal involving a criminal politician could reduce the popularity of parties running criminal candidates at both the state and national level, and parties could respond by reducing the number of criminal politicians in both types of elections. Note, however, that (i) most state elections are held in different years than the corresponding national election, and (ii) the competitive environment (e.g., the number of competing parties, or voter priorities) can be quite different between national and state elections. For these reasons, it is not clear that popularity shocks are generally correlated, nor that parties would respond to similar

²⁰The idea corresponds to Industrial Organization applications that use cost shifters that affect firm profits but not consumer utilities as instruments for endogenous variables (typically, prices). USC (2021), which studied state elections, used candidate characteristics in neighboring constituencies based on a similar logic.

shocks in the same way at the national and state levels.

We create our instruments as a function of the party alliance (UPA/NDA/neither) and the state assembly constituencies overlapping with the national election constituency. For example, we instrument the assets of a UPA candidate with the average assets of all UPA candidates running in the state constituencies contained in the given national constituency. We create these instruments both for the same election and for the other election in the data.

To get a preliminary sense of the relevance and strength of these instruments (which we also explore in more detail in the next subsection), we regress each characteristic on the corresponding instruments and control variables. The results are shown in Table 2. In columns 1 and 2, the state election averages of the education characteristic have large and significant association with the education level of a party’s candidates in the national election. Columns 3-8 show similar patterns for the other characteristics as well. These correlations, which are interesting in their own right, provide support for the idea that state candidate averages proxy for the pool of characteristics that a party is able to draw from at the national level.

Table 2: Characteristics regressed on instruments

Dep. var.:	Education		Muslim		Crime		Assets	
	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)
IV same election	0.19 (0.05)	0.16 (0.05)	0.54 (0.07)	0.34 (0.08)	0.29 (0.06)	0.28 (0.06)	0.32 (0.05)	0.31 (0.05)
IV other election		0.34 (0.08)		0.44 (0.10)		0.10 (0.09)		0.10 (0.08)
Adj. R^2	0.10	0.11	0.09	0.09	0.14	0.14	0.37	0.37
F	13.23	15.22	60.73	42.73	22.94	13.34	42.73	20.90
Mean dep. var.	0.14	0.14	-0.03	-0.03	0.05	0.05	-0.10	-0.10
Std. dev.	0.84	0.84	0.88	0.88	0.96	0.96	1.00	1.00

Notes: *IV same election* is the average of the given characteristic among an alliance’s candidates in state election constituencies in the given year. *IV other election* is the same variable for the other election in the data (2009 for 2014 and vice versa). Regressions control for state, year, party and alliance fixed effects, and indicators for imputed characteristics and reserved constituencies. The F statistic is the [Olea and Pflueger \(2013\)](#) effective F statistic. Robust standard errors in parentheses. N = 2,649.

5.2.2 Differentiation IVs and specification choice

In order to identify nonlinear parameters, we use the “differentiation IVs” proposed by [Gandhi and Houde \(2019\)](#). The idea behind these instruments is to use the menu of choices

available to each decision-maker to identify preference heterogeneity among decision-makers. In our application, preference heterogeneity among voters for a candidate’s education level (say) is identified based on how many candidates in a voter’s choice set have similar education levels.

To construct the differentiation IVs, we first predict each endogenous candidate characteristic using the above instruments (and all exogenous variables). We then use these predicted characteristics to form instruments: for each (predicted) characteristic of a candidate that enters the nonlinear part of voter utility, we compute the number of candidates in the constituency whose corresponding (predicted) characteristic is within one standard deviation.

To guide our specification choice and evaluate the strength of our instruments, we use the methods proposed by [Gandhi and Houde \(2019\)](#). The details are in Appendix 2. The idea is to evaluate whether the instruments are “strong enough” to reject the linear (Logit) specification, i.e., $\Pi = \mathbf{0}$. First, we enter the differentiation IVs as controls in a Logit specification. We find that the differentiation IVs for Muslim and assets are statistically significant while the differentiation IVs for education and crime are not. This suggests that the former two are capable of capturing departures from the Logit model. As an alternative diagnostic, we also run a specification that includes the differentiation IVs as instruments instead of controls. The overidentification J-test clearly rejects this specification, which also provides support for focusing on the nonlinear specifications ([Gandhi and Houde 2019](#)).

To further evaluate the nature of preference heterogeneity, we estimate a random coefficients specification where voter demographics $\mathbf{d}_i$ in (6) are replaced with random variables drawn from a standard normal distribution. This specification indicates the presence of significant heterogeneity in voters’ preference for candidate assets, but not for the other three characteristics.

5.2.3 Adding voter demographics

Based on the specification checks described in the previous section, we first focus on identifying the relevant sources of preference heterogeneity in voters’ valuation of candidate assets. Our main demographic variables are literacy, rural population, presence of paved roads, working population, and lower caste population. We interact the differentiation IV for assets with the average value of each of these demographic variables in the constituency.²¹

We again evaluate these instruments using a Logit specification. This supports using the

²¹Following [Gandhi and Houde \(2019\)](#), the idea is to identify a parameter π_m^k on the interaction of the demographic d_m with the characteristic x^k using the instrument $\bar{d}_m \hat{x}^k$, where $\bar{d}_m$ is the average value of the demographic in the constituency and $\hat{x}^k$ is the differentiation IV for the candidate characteristic.

interaction of the asset differentiation IV with literacy, rural population, and presence of paved roads (Table A.2). Estimating nonlinear specifications using different combinations of these instruments and corresponding nonlinear parameters yields a clear favorite, shown in column 1 of Table 3. This specification passes the overidentification J test, and results in nonlinear coefficients that are jointly statistically significant based on the Newey-West test.

Table 3: Voter preference parameter estimates

	(1)	(2)	(3)	(4)
<i>Linear parameters (β)</i>				
Education	-0.19 (0.69)	0.94 (0.82)	-0.59 (0.73)	-0.57 (0.80)
Muslim	-0.50 (0.28)	-0.37 (0.27)	-0.54 (0.29)	-0.54 (0.30)
Crime	0.66 (0.51)	0.73 (0.57)	0.56 (0.50)	0.59 (0.50)
Assets	2.40 (3.48)	-2.00 (1.14)	2.81 (2.72)	4.84 (6.05)
<i>Nonlinear parameters (π)</i>				
Assets $\times$ Literacy	-3.45 (3.49)	4.15 (1.85)		-2.42 (5.95)
Assets $\times$ Road		2.59 (1.50)	-3.78 (2.85)	-3.60 (3.37)
Assets $\times$ Rural	3.84 (1.27)		3.97 (1.21)	3.26 (2.47)
J	1.481	9.715	8.880	9.407
df	4	4	4	4
p-value	0.83	0.05	0.06	0.05
Newey-West p-value	0.001	0.005	0.007	0.001

Notes: BLP estimates. Specifications include state, year, party and alliance fixed effects, indicators for imputed characteristics, and reserved constituencies. J is the overidentification J-statistic with its degree of freedom (df) and p-value. The bottom row shows the the p-value of the Newey-West D-test for the null that all nonlinear parameters are jointly 0. Robust standard errors clustered by constituency in parentheses.

According to the estimates in column 1 of Table 3, all else equal voters dislike Muslim candidates and like candidates with a criminal history (consistent with Vaishnav (2017) and USC (2021)) though the latter is not statistically significant. Rural voters have a preference for candidates with more assets. A possible explanation is that wealthier candidates have higher social status and thus easier access to other government officials to advance local interests and “get things done.” According to Vaishnav (2017), providing these connections between the local community and the state apparatus is very important in voters’ evaluation

of the candidates.

5.3 Specifying the parties' choice sets

In our specification of voter preferences, we conceptualized parties' choice of candidates a_p as a choice of a bundle of candidate characteristics $\mathbf{x}_p$. Applied directly to parties' problem, this would imply very large choice sets $\mathcal{A}_p$, containing all the possible combinations of characteristics. This is neither practical for estimation, nor realistic as a model of party choices. For example, it is unlikely that a party would view two candidates who are identical in all dimensions but whose assets are slightly different as substantively different options.

For a more conceptually appealing (and computationally feasible) model of parties' problem, we assume the existence of a smaller set of candidate "types" that parties consider when choosing who to run. For example, a type could be an "educated non-Muslim with some criminal history in the second quartile of the asset distribution." Rather than making ad hoc assumptions about these types, we use machine learning tools to let the data tell us what they should be.

5.3.1 Data and variables for constructing candidate types

Our goal is to describe the pool of potential national candidates, as opposed to the set of candidates actually selected by the parties. As argued above, the candidates running in state elections provide a good proxy for the pool of candidates that national parties can select from. Thus, we define candidate types based on the characteristics of the candidates running in state elections.

There are many state candidates (36,766 in our data), which yields rich variation in candidate characteristics in the pool of potential candidates. This data also has many (13,787) independent candidates. These candidates are often individuals who wanted to run with a major party but were not selected and are thus highly informative about the type of candidates parties choose from.

As above, we use the candidate characteristics education, assets, criminal history and Muslim. Of these, education, assets, and criminal history have missing values. In order to assign each candidate to a type, we impute these missing characteristics as described in Appendix 1.²² We standardize these variables to have 0 mean and unit standard deviation.

²²An alternative approach we considered is to use a Missing indicator as an additional characteristic. However, there are differences between missing values in the state and national election data (for example, most independent candidates' education, assets, and criminal history is missing in the former). Thus, using Missing as an additional characteristic would mechanically make candidate types in the two datasets less comparable.

5.3.2 Clustering algorithm

We use k-means clustering to create the types. This iterative procedure partitions the data into K clusters based on the 4 variables described above (Muslim, education, assets and crimes). The algorithm begins by specifying K initial centroids and forming clusters by assigning each candidate to the closest centroid. Throughout, we use Euclidean distance to compute candidates' distance from a centroid. Next, new centroids are computed based on the average characteristics of the candidates assigned to each cluster. Using these centroids, candidates are reassigned to the closest cluster, and the process continues until no candidate is reassigned from their current cluster.

To use k-means clustering, one must first choose the number of clusters K . There is currently no cross-validation method to assess the relative performance of different values, and one option is to choose the number K based on substantive considerations (Athey and Imbens 2019).²³ Alternatively, we can use a set of common measures from the machine learning literature to select K . The first measure relies on the Within Cluster Sum of Squares (WCSS), also referred to as the Inertia score, which measures the similarity of units within the same cluster. The second measure is the Silhouette Coefficient (SC), which measures how far apart the clusters are from one another. The formula for each measure is given in Appendix 6. Ideally, the clusters should result in a low WCSS and a large SC.

The WCSS is shown on the left panel of Figure 1 for values of K ranging from $K = 1, \dots, 20$. Notice that this measure is decreasing in K by construction - the more clusters, the more similar the members of the cluster will be. As can be seen in the figure, up to $K = 4$ each additional cluster creates large drops in the WCSS. By comparison, the gains from additional clusters beyond 4 are much smaller.²⁴

On the right panel of Figure 1 we display the Silhouette Coefficient over the same range of values of K . The largest gain in SC , by far, occurs when moving from $K = 3$ to $K = 4$, consistent with the results for $WCSS$ above. When moving from $K = 4$ to $K = 5$, SC actually decreases slightly, and increases in SC are relatively small for $K > 4$.

Each of these validation checks supports using $K = 4$.²⁵ In addition, as discussed in the

²³Alternative unsupervised methods such as Density Based Clustering (DBSCAN) do not require the researcher to input the number of types K . The drawback of this type of method is that the researcher must select other hyperparameters, and the resulting clusters can be highly sensitive to these choices. Moreover, DBSCAN does not classify all points in the sample. Trebbi and Weese (2019) develop an interesting method to identify the number of organized insurgent groups using correlations in the timing of attacks over space. It is not obvious how to apply their method in our setting however, as identification of the number of clusters would require covariates that vary independently across our candidate characteristics, which we do not have.

²⁴This approach for validating the choice of K is often referred to as the “elbow method” (Thorndike (1953)).

²⁵We also found that with $K = 4$, the clusters resulting from the algorithm were invariant to the starting values of the centroids. This was not the case with $K = 5$ and $K = 6$.

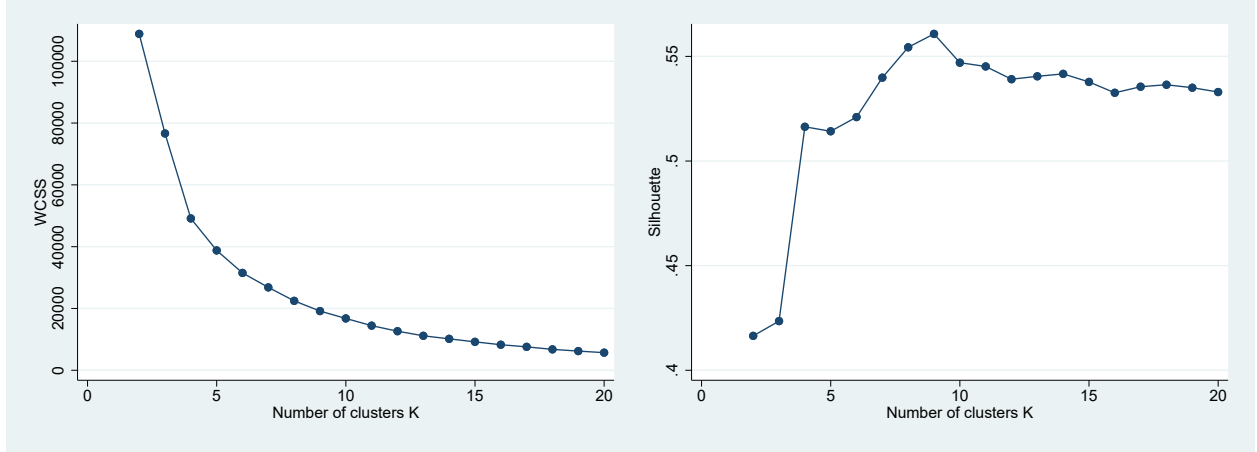


Figure 1: WCSS and Silhouette scores for different values of K

next section, $K = 4$ also yields candidate types that are easy to interpret and match well some of the important categories of politicians mentioned in the literature.

5.3.3 Candidate types

Table 4 shows the centroids of each of the four candidate types resulting from the k-means clustering algorithm. For ease of interpretation, the table shows values on the original scale of each variable (rather than the standardized scale on which the clustering algorithm is run).

The types resulting from the algorithm turn out to have very clear interpretations. Types 1 and 2 are both non-Muslim candidates with no criminal history, but Type 1 has high education while Type 2 has low education. Type 3 contains all the Muslim candidates. Type 4 contains all the non-Muslim candidates with criminal history, and these candidates are also richer than the other three types. For simplicity, we will refer to the four types as *educated*, *uneducated*, *Muslim*, and *criminal* types, respectively.

These candidate types obtained from the clustering algorithm appear quite sensible. Given the salience of the Muslim characteristic in Indian politics, it is not surprising that one of the types we obtain is defined by this characteristic. The positive correlation between criminality and wealth, reflected in our criminal type, is consistent with [Vaishnav \(2017\)](#) and [Asher and Novosad \(2023\)](#). Finally, the importance of education among non-Muslim non-criminal candidates is consistent with a large literature that uses education as the main measure of candidate valence (e.g., [Galasso and Nannicini \(2011\)](#); [Besley and Reynal-Querol \(2011\)](#)).

In sum, the types resulting from the clustering reflect meaningful choices for the parties, and not simply some artificial combination of candidate characteristics. To further support

this, in Appendix 6 we ask how robust the types are to using a different clustering algorithm (agglomerative hierarchical clustering). Remarkably, we find that the two algorithms perfectly agree on the Muslim and criminal types. They differ somewhat in the allocation of the remaining candidates, but we show that k-means clustering dominates the hierarchical algorithm in terms of validation.

Table 4: Centroids of candidate types

	Assets	Crimes	Education	Muslim
Type 1	14.35	0.00	1.00	0.00
Type 2	14.27	0.00	0.00	0.00
Type 3	14.30	0.17	0.44	1.00
Type 4	15.06	1.00	0.68	0.00

Notes: Centroids resulting from the k-means clustering algorithm. The algorithm is run on standardized variables; the table shows the centroids transformed back to the original scale for ease of interpretation. Assets: real Rp in logs; Crimes: 1 if has at least one criminal case; Education: 1 if completed high school; Muslim: 1 if Muslim name.

The top panel of Table 5 shows the distribution of candidate types among all candidates, as well as the candidates of the UPA and the NDA. Compared to the UPA, the NDA has a somewhat higher share of uneducated types and (not surprisingly) a smaller share of Muslim types.

5.3.4 Types in the national elections

Next we assign national candidates to the candidate types created from the state elections data. To do this, recall that the types in Table 4 have clear definitions in terms of characteristics: Type 1 contains all the educated, non-Muslim, non-criminal candidates; Type 2 all the uneducated, non-Muslim, non-criminal candidates; Type 3 all the Muslim candidates; and Type 4 all the non-Muslim criminal candidates. We assign national candidates to types using these definitions.²⁶

The resulting distribution of candidate types is shown in the lower panel of Table 5. Relative to state candidates, there are relatively more candidates of the educated type (and fewer of the uneducated type). There are also more of the criminal type, but this difference between state and national elections is less pronounced for candidates of the UPA and the NDA. The share of the Muslim type is similar between state and national elections. As in

²⁶We also tried an alternative assignment where we assign each candidate to the type with the closest centroid. This creates a very similar assignment: in particular, only 3 UPA or NDA candidates are assigned to a different type.

Table 5: Distribution of candidate types

	All candidates	UPA	NDA
<i>State elections</i>			
Type 1	34.74	46.78	44.15
Type 2	39.63	16.37	20.76
Type 3	11.20	11.27	6.55
Type 4	14.43	25.58	28.55
Total	100	100	100
N	36,764	3,788	4,109
<i>National elections</i>			
Type 1	45.16	50.98	52.95
Type 2	20.63	5.47	9.09
Type 3	12.48	10.94	5.23
Type 4	21.73	32.6	32.73
Total	100	100	100
N	6,577	457	440

Notes: Candidates in state and national elections assigned to each type by the clustering algorithm. Types 1-4 are the Educated, Uneducated, Muslim, and Criminal types, respectively.

the state elections, the most pronounced difference between the NDA and the UPA is the former’s lower share of the Muslim type and, to a lesser extent, its higher share of candidates of the uneducated type.

Table 6 presents information on the average electoral performance of each candidate type in the raw data. The first two columns are measures of winning probability: “winner” is an indicator equal to 1 if the candidate won, and “closeness to winner” is the candidate’s vote share divided by the winner’s vote share. These values suggest that different candidate types have different winning probabilities. The criminal type has the highest average on both measures. In spite of these differences, the vote share of different types is very similar on average. A possible explanation is that parties do not attempt to increase their vote shares beyond the minimum necessary to win, in line with the discussion in Section 3.1. The choice frequencies of the different types, shown in the last column, indicate a very different pattern from either win probabilities or vote shares. This is suggestive of the fact that parties’ choices are driven by other considerations.

Table 6: Electoral performance and frequency of different candidate types in the raw data

	Winner	Closeness to winner	Vote share	Choice share
Type 1	0.31	0.65	0.19	0.52
Type 2	0.35	0.67	0.19	0.07
Type 3	0.23	0.62	0.19	0.08
Type 4	0.35	0.70	0.20	0.33

Notes: Average values from the data for UPA and NDA candidates, by candidate type. Types 1-4 are the Educated, Uneducated, Muslim, and Criminal types, respectively. *Winner* is an indicator for candidates who won. *Closeness to winner* is the candidate’s vote share divided by the winner’s vote share. *Vote share* is the vote share observed in the data. Choice share is relative frequency in the data. $N = 897$.

5.4 Specification, identification, and estimation of party objectives

We focus on candidate selection by the two major party alliances in Indian politics, the UPA and the NDA. Each party’s choice set contains the $K = 4$ candidate types obtained from the clustering algorithm. Based on (2), we specify party p ’s objective function in constituency c as

$$\begin{aligned}
 U_p(a_{pc}, P(\mathbf{a}_{-p,c})) = & \sum_{\mathbf{a}_{-p,c}} [b^w w_{pc}(a_{pc}, \mathbf{a}_{-p,c}) + b^s s_{pc}(a_{pc}, \mathbf{a}_{-p,c})] P(\mathbf{a}_{-p,c}) \\
 & + \sum_{k=2}^4 (c_{kp}^0 + \mathbf{c}_k \mathbf{L}_{pc}) \mathbf{1}\{a_{pc} = k\} + \varepsilon_{pc}(a_{pc}),
 \end{aligned} \tag{12}$$

where $\mathbf{L}_{pc}$ contains the proxies for the pool of candidate characteristics (the average of candidate education, assets, Muslim and criminal history in the assembly constituencies corresponding to constituency c), and c_{kp}^0 and $\mathbf{c}_k$ are type-specific cost parameters. As will become clear below, the cost parameters are only identified for three of the four types; we use $k = 1$ as the excluded category in (12). We will refer to the costs that depend on the candidate pool, $\mathbf{c}_k$, as “recruitment costs” and the costs c_{kp}^0 as “direct costs.”

Our goal is to estimate the parameters $\boldsymbol{\theta} = (b^w, b^s, \{c_{kp}^0, \mathbf{c}_k\}_{k=2}^4)$, which include parties’ payoffs from increasing their probability of winning and their vote share, and their costs of running different candidate types.

To compute the vote shares and winning probabilities corresponding to different hypothetical candidate choices by the parties, we use our demand estimates. We compute the vote shares and associated winning probabilities for all possible action profiles $(a_{pc}, \mathbf{a}_{-p,c})$

of the players in a given constituency (16 profiles), holding fixed all exogenous variables (including non-strategic parties' choices) as well as the popularity shocks ξ_{pc} .

To use our demand estimates, we also need to assign a vector of characteristics to each type choice in the data. As described in Section 4.2, when parties choose types, there is some uncertainty over the actual realization of characteristics. We capture this uncertainty using the distribution of candidate characteristics for each type in the state election data. For each combination of types, we draw 1000 values of their characteristics from the data, and compute vote shares and winning probabilities as expected values across these draws (equations (8) and (9), respectively).

We provide a formal discussion of the identification of the party objective function parameters in Appendix 3, but it is worth briefly discussing the sources of identification here. As is typically the case in models of discrete choice, the cost parameters are identified only up to differences with respect to a reference alternative. That is, we identify $c_k - c_1$ for $k = 2, 3, 4$ where Type 1 is the reference type. The identification of these differences, for given parameters (b^w, b^s) is standard (see Appendix 3) and depends on the magnitude of the observed probability of selecting type k , $P(a_{pc} = k)$ relative to the probability of selecting the reference type.

By contrast, the benefit parameters are pinned down in levels, not differences. To see this, note that given the Type-I Extreme Value assumption we can express the relationship between choice probabilities and party payoffs as:

$$\begin{aligned} \ln(P(a_{pc} = k)) - \ln(P(a_{pc} = 1)) &= b^w \times (E[w_{pc}(k, \mathbf{a}_{-p,c})] - E[w_{pc}(1, \mathbf{a}_{-p,c})]) \quad (13) \\ &+ b^s \times (E[s_{pc}(k, \mathbf{a}_{-p,c})] - E[s_{pc}(1, \mathbf{a}_{-p,c})]) \\ &+ c_k - c_1 + \eta_{pc}(k) \end{aligned}$$

where we have assumed a single type-specific cost term c_k for simplicity. As we can treat choice probabilities and win probabilities as known, this can be viewed as a regression of $\ln(P(a_{pc} = k)) - \ln(P(a_{pc} = 1))$ on $(E[w_{pc}(k, \mathbf{a}_{-p,c})] - E[w_{pc}(1, \mathbf{a}_{-p,c})])$ and $(E[s_{pc}(k, \mathbf{a}_{-p,c})] - E[s_{pc}(1, \mathbf{a}_{-p,c})])$ where the intercept is the cost difference $c_k - c_1$. Then b^w is, loosely, identified as the covariance across constituencies between the probability of selecting type k relative to the reference option, $\ln(P(a_{pc} = k)) - \ln(P(a_{pc} = 1))$, and the difference in expected win probabilities, $E[w_{pc}(k, \mathbf{a}_{-p,c})] - E[w_{pc}(1, \mathbf{a}_{-p,c})]$. If the party tends to select candidate k in the constituencies where they are likely to win, b^w will be positive. See Appendix 3 for a more formal and detailed discussion.

Estimation proceeds by recursively updating the choice probabilities using (pseudo) maximum likelihood estimates of the parameter vector θ up to convergence as in Aguirregabiria

and Mira (2007). Specifically, consider an initial choice probability estimate $\hat{P}^0(\mathbf{a}_{-p,c})$. In constituency c , party p chooses a_{pc} to maximize (12). Again defining utility net of the unobservable as $\tilde{U}_p(a, P(\mathbf{a}_{-p,c}); \boldsymbol{\theta})$, the implied probability that $a_{pc} = a$ is:

$$P(a|\hat{\mathbf{P}}^0, \boldsymbol{\theta}) = \frac{\exp \left\{ \tilde{U}_p(a, \hat{\mathbf{P}}^0(\mathbf{a}_{-p,c}), \boldsymbol{\theta}) \right\}}{\sum_{a'} \exp \left\{ \tilde{U}_p(a', \hat{\mathbf{P}}^0(\mathbf{a}_{-p,c}), \boldsymbol{\theta}) \right\}} \quad (14)$$

where we emphasize the fact that the choice probabilities are a function of the estimates $\hat{\mathbf{P}}^0$ as well as the parameters $\boldsymbol{\theta}$.

Denoting parties' choices observed in the data with a_{pc}^* , the log likelihood is

$$\ell(\boldsymbol{\theta}, \hat{\mathbf{P}}^0) = \sum_{pc} \sum_a \mathbf{1}\{a = a_{pc}^*\} \ln P(a|\hat{\mathbf{P}}^0, \boldsymbol{\theta}).$$

The estimates $\hat{\boldsymbol{\theta}}^0$ solve

$$\max_{\boldsymbol{\theta}} \ell(\boldsymbol{\theta}, \hat{\mathbf{P}}^0)$$

With these estimates, we can construct new estimates of the choice probabilities as

$$\hat{P}^1(a|\hat{\mathbf{P}}^0) = \frac{\exp \left\{ \tilde{U}_p(a, \hat{P}^0(\mathbf{a}_{-p,c}), \hat{\boldsymbol{\theta}}^0) \right\}}{\sum_{a'} \exp \left\{ \tilde{U}_p(a', \hat{P}^0(\mathbf{a}_{-p,c}), \hat{\boldsymbol{\theta}}^0) \right\}}.$$

Given these new choice probabilities, the equilibrium probability that $a_{pc} = a$ in equation (14) becomes $P(a|\hat{\mathbf{P}}^1, \boldsymbol{\theta})$. In turn, this yields an updated log likelihood $\ell(\boldsymbol{\theta}, \hat{\mathbf{P}}^1)$, which is maximized to obtain a new estimate $\hat{\boldsymbol{\theta}}^1$. We iterate in this way until convergence. The resulting estimator, $\hat{\boldsymbol{\theta}}_{NPL}$ is asymptotically equivalent to the Maximum Likelihood estimator. See Aguirregabiria and Mira (2007) for details. Standard errors are obtained from the inverse Hessian of the likelihood function evaluated at $\hat{\boldsymbol{\theta}}_{NPL}$.

6 Estimation results

6.1 What do parties maximize?

In Table 7, we first estimate a version of the model where parties care only about their expected winning probability and vote share. We then introduce the recruitment costs $\mathbf{c}_k \mathbf{L}_{pc}$ in column 2, and also add the direct costs c_{kp}^0 in column 3. Allowing for these costs substantially improves the fit of the model: as we move from column 1 to 3, the Log Likelihood increases

by 29.7%.²⁷ Chi-square tests at the bottom of columns 2 and 3 indicate that the cost parameters are highly jointly significant. Costs, which are independent of voter preferences, are important in explaining parties' choices of which candidates to run.

According to our estimates, parties like to win ($b^w > 0$), but, all else equal, choosing candidates with higher vote shares has a disutility ($b^s < 0$). This is consistent with the idea that party elites balance anticipated electoral performance of a popular candidate against the possibility that the candidate becomes too powerful and threatens the elite - as discussed in the literature on Indian parties (see Section 3.1). To interpret the magnitude of the estimates, we can use the standard deviation of parties' equilibrium payoffs, which is 1.49. All else equal, winning raises parties' payoff by 2.2 standard deviations ($3.24/1.49$). Winning with a one standard deviation higher vote share (7 percentage points in the data) lowers payoffs by two thirds of a standard deviation ($14.15 \times 0.07/1.49 = 0.66$).

The recruitment cost parameters $\mathbf{c}_k$ show that a higher prevalence of some candidate characteristic in a party's candidate pool lowers the party's cost of selecting a candidate type with that characteristic. For example, the positive estimate for c_4^{crime} shows that a higher prevalence of criminal candidates in the relevant pool increases a party's payoff from choosing the criminal Type 4. The negative estimates of c_k^{educ} for $k = 2, 3, 4$ indicate that the payoff from choosing the educated Type 1 (the excluded category) increases with a higher prevalence of educated candidates in the pool. In this way, the supply of candidate characteristics available to parties affects who they choose to run, mirroring the patterns seen in Table 2.

Column 3 shows that parties' costs of running specific candidates are not restricted to the recruitment costs. According to these estimates, the NDA has the largest direct cost from running a Muslim type (Type 3): running a Muslim candidate relative to an educated candidate lowers its payoff by 1.7 standard deviations, all else equal ($2.58/1.49$). This is in line with the Hindu nationalist profile of the BJP, the NDA's leading party, and indicates that the NDA's aversion to running Muslim types is not due simply to voter preferences.

The UPA has the largest direct cost from uneducated types (Type 2), and both parties incur the lowest direct cost (or highest benefit) from running a criminal type (Type 4), followed by the (excluded) educated type. The idea that parties have relatively low costs of running criminal candidates is in line with Vaishnav (2017)'s argument that these candidates are willing to fund their own campaigns (perhaps circumventing campaign finance laws), which makes them cheaper for the parties.

Figure 2 plots the recruitment costs for each type, and compares them to the direct costs

²⁷See Section 6.3 below for a more detailed evaluation of model fit. In Appendix 5 we provide more results illustrating how the fit improves as we introduce costs.

(indicated with vertical lines). Both recruitment costs and direct costs matter; for Type 2 and 3, direct costs are larger in magnitude compared to recruitment costs. Interestingly, the recruitment costs for Type 4 are mostly positive (i.e., a benefit). This means that, in most constituencies, parties incur relatively low costs from selecting criminal types compared to other types due to the ample supply of criminality in the candidate pool.

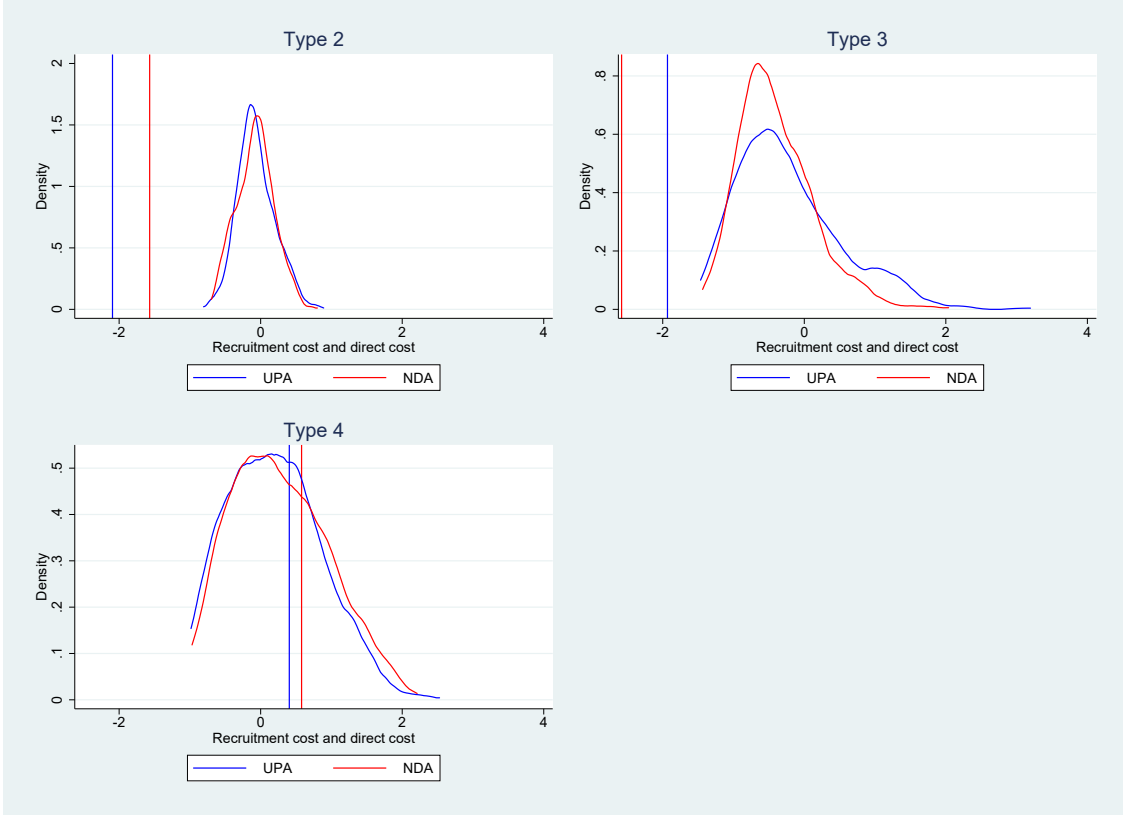


Figure 2: Distribution of each party’s costs of different types across constituencies

Based on column 3 of Table 7. Kernel density plots of recruitment costs $\mathbf{c}_k \mathbf{L}_{pc}$; direct costs c_{kp}^0 indicated with vertical lines. Types 2-4 are the Uneducated, Muslim, and Criminal types, respectively (Type 1 is the excluded category).

In Appendix 8 we present a series of robustness checks on these estimates. We show that controlling for several sources of observable constituency level heterogeneity, including election year, reserved constituencies, or the share of rural population, does not change our estimates in a material way. We also apply the method of Bonhomme, Lamadon and Manresa (2022) to control for *unobservable* constituency level heterogeneity, and show that this also leaves our main estimates unaffected.

Table 7: Party objective function estimates

	(1)	(2)	(3)
b^w	4.26 (0.84)	3.08 (0.95)	3.24 (1.02)
b^s	-5.35 (1.22)	-9.20 (1.45)	-14.15 (1.73)
$c_{2,NDA}^0$			-1.57 (0.27)
$c_{3,NDA}^0$			-2.58 (0.30)
$c_{4,NDA}^0$			0.58 (0.23)
$c_{2,UPA}^0$			-2.09 (0.32)
$c_{3,UPA}^0$			-1.93 (0.29)
$c_{4,UPA}^0$			0.40 (0.23)
c_2^{educ}		-2.01 (0.26)	-0.53 (0.34)
c_3^{educ}		-2.30 (0.26)	-0.63 (0.34)
c_4^{educ}		-0.64 (0.20)	-0.41 (0.23)
c_2^{crime}		-0.47 (0.22)	-0.03 (0.24)
c_3^{crime}		-0.31 (0.21)	0.36 (0.23)
c_4^{crime}		1.06 (0.14)	1.10 (0.15)
c_2^{asset}		-0.73 (0.19)	0.32 (0.26)
c_3^{asset}		-1.08 (0.20)	-0.12 (0.26)
c_4^{asset}		0.06 (0.14)	-0.06 (0.16)
c_2^{Muslim}		0.44 (0.32)	-0.07 (0.41)
c_3^{Muslim}		1.74 (0.27)	1.48 (0.29)
c_4^{Muslim}		0.12 (0.24)	0.03 (0.24)
Log likelihood	-1189.00	-919.41	-836.21
Joint significance of c 's	-	355.02	482.45

Notes: Estimates of party objective functions in (12). The cost parameters c are measured relative to the excluded category, Type 1. $c_{k,p}^0$ is party p 's direct cost of choosing a type k candidate. c_k^l is the impact of characteristic l in the pool of candidates on the recruiting cost of a type k candidate. Types 1-4 are the Educated, Uneducated, Muslim, and Criminal types, respectively. Standard errors in parentheses. Number of markets: 434.

6.2 Why do parties choose criminal candidates?

Our estimates reveal a nuanced set of reasons that guide parties' choice of criminal candidates. This formalizes previous explanations proposed in the literature, and adds several new considerations.

First, according to our estimates, voters like the criminal type, particularly because these candidates are wealthy. Since parties like to win, they have an incentive to choose criminal types to increase their winning probability. At the same time, criminal types can generate particularly large vote shares, and our results show that this creates a disutility for party leaders for a given winning probability.²⁸ As discussed above, the disutility from vote shares that are “too large” can reflect threats from factionalism and defections from the party. Therefore, the finding that this disutility is particularly important for criminal types links together two influential literatures on candidate selection in India, one on criminal candidates (Vaishnav 2017), and one on party factions (Chandra 2016).

Second, criminal types also affect party payoffs independently of votes. We found that in most constituencies recruiting costs favor this type: the large supply of criminal candidates makes it relatively cheap to run them (Figure 2). We also found that criminal candidates have the lowest direct costs among the four types. This is consistent with the descriptive literature on how criminals' private wealth and connections can benefit the party. In Appendix 4 we further illustrate the importance of wealth as a determinant of candidate choice by sorting parties' equilibrium payoff by the wealth of their chosen candidate.

Importantly, this pattern of costs and benefits raises the possibility that parties' decision to run criminals is a *strategic* response to the other party's criminal candidate. A party could run a criminal simply because its opponent is doing that as well.

Intuitively, there are two reasons why a party could find it relatively advantageous to run a criminal when its opponent is running a criminal. First, because voters like the criminal type, running a different type against the opponent's criminal would yield a low probability of winning. Second, when the opponent is running a criminal, choosing a criminal increases the probability of winning without generating vote shares that are “too large” and therefore costly for the party. This suggests that, under some conditions, criminals may be strategic complements for the parties.

This possibility is supported by the fact that the correlation between the two parties' choice probability of a criminal type is positive (0.30), while the correlation between one party's choice probability for a criminal type and the other party's choice probability for

²⁸To quantify this effect for the criminal type, in Table A.13 in the Appendix we show that without the distutility of large vote shares (setting $b^s = 0$), the share of criminal candidates chosen by each party would more than double.

any other type is always negative (ranging between -0.14 and -0.29). See Table A.9 in the Appendix.

To investigate this directly, we compute $\frac{\partial P_p(4)}{\partial P_{-p}(4)}$, the change in party p 's probability of choosing a criminal (Type 4) in response to a change in its opponent's probability of choosing a criminal. A positive derivative for both parties indicates strategic complementarity (see Appendix 7). Figure 3 shows the distribution of this derivative across constituencies. For the UPA (NDA), the derivative is positive in 85% (87%) of the constituencies. The derivative is positive for both parties in 77% of the constituencies (and negative for both in only 6%). In most cases, criminal candidates are strategic complements.

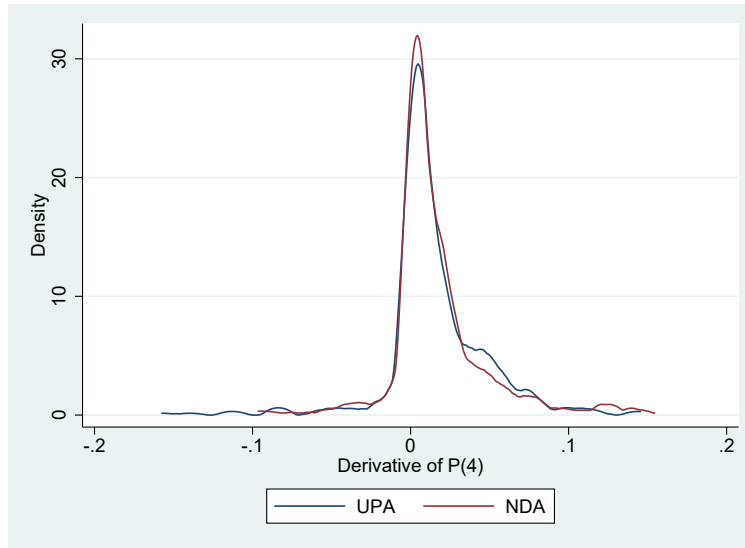


Figure 3: Strategic complementarity of criminals

Notes: Kernel density plots of $\partial P_p(4) \partial P_{-p}(4)$ based on equation (24) in Appendix 7.

Because running a criminal has costs for the party, strategic complementarity raises the possibility that in equilibrium parties run criminals more often than they would like to. As we show below, this has important policy implications for measures that restrict parties' ability to run criminal candidates.

6.3 Model validation and fit

To explore the model's ability to fit the data, we first use the estimated model to simulate party choices in every constituency. We repeat this 100 times and compare the average over simulations with the actual choices observed in the data for each type and each party. The result is in Panel A of Table 8. The model performs well.

Table 8: Model fit

	UPA actual	UPA predicted	NDA actual	NDA predicted	All actual	All predicted
<i>Panel A: Actual and predicted type choices</i>						
Type 1	217	221	229	234.80	446	455.80
Type 2	24	24	40	41	64	61.60
Type 3	49	46.20	22	19.80	71	66.00
Type 4	144	146.20	143	138.40	287	284.60
<i>Panel B: 5-fold cross validation</i>						
Type 1	41.4	40.68	43.4	42.00	84.80	82.68
Type 2	4.8	3.88	7.8	7.68	12.60	11.56
Type 3	9.2	10.32	4.4	5.64	13.60	15.96
Type 4	27.60	28.12	27.40	27.68	55.00	55.80

Notes: Panel A shows the number of candidates of each type observed in the data and predicted by the model (average across 100 simulations). Panel B uses a 5-fold cross validation procedure as described in the text (values shown are averages across 20% subsamples). Types 1-4 are the Educated, Uneducated, Muslim, and Criminal types, respectively.

To evaluate the model’s “out of sample” performance we use a cross-validation procedure. We repeat the following 5 times. First, we hold out the first 20% of the sample, and estimate the model using the remaining 80% of observations. We then solve for the equilibrium in each constituency in the 20% hold-out sample, and use these to simulate party choices on that sample. We repeat this process for the next 20-80 split, etc., and take the average across the 5 sets of predictions (this is essentially a k-fold cross validation, with $k = 5$).

The results are in Panel B of Table 8. While there is variation across the folds in predictive ability, on average over the k-folds the model does just as well in predicting outcomes as in Panel A.

7 Policy experiment: banning the criminal type

What is the impact of banning candidates with a criminal history from contesting the election? As discussed above, this is a relevant policy question in India and many other countries.

To model this kind of policy, we consider a counterfactual scenario where we make it prohibitively costly to choose the criminal type (Type 4) for both parties. Although Type 3 also contains some candidates with criminal history, Type 4 candidates *always* have a criminal history (Table 4). They are also wealthier, and thus more closely match the kind of criminal candidates who are often considered problematic in the Indian context (see [Vaishnav \(2017\)](#)).

Using our parameter estimates, we compute a new equilibrium. This involves computing

the expected vote share of candidate types once the criminal type has been banned. We do this in two ways: banning the criminal type only for the UPA and the NDA but leaving third party candidates as is, or also removing third parties that run criminal candidates in the demand model (we focus on the former in the main analysis). We study the change in parties' choices, winning probabilities, and payoffs in the counterfactual equilibrium compared to the baseline equilibrium.

7.1 Changes in candidate characteristics

A first observation is that, because candidates are bundles of correlated characteristics, removing a candidate type directly affects the distribution of characteristics among the candidates running for election. The criminal type is relatively wealthy, and comes from the non-Muslim majority. Removing these candidates may therefore raise the share of less wealthy and minority candidates. Such changes may be unintended side effects of a policy of banning candidates with a criminal history.

The first two columns of Table 9 compare the average choice probabilities (i.e., the predicted share of each candidate type among the candidates contesting the election) in the no-criminal counterfactual equilibrium and the baseline equilibrium. The full distribution of the changes in probabilities is shown in Figure A.14 in the Appendix.

We find that the choice probabilities of candidate types 1-3 all increase. Note that in single-agent models, reducing the choice set could never reduce the choice probability of a remaining option. In the context of our model where parties play best responses to each other's strategies, this is not necessarily the case. To illustrate with an example that ignores incomplete information, take a constituency where the NDA ran a Type 1 candidate and the UPA a Type 4 candidate. Losing the option of running a Type 4 candidate, the UPA might switch to Type 2. If the NDA's best response is to switch to Type 3, then all else equal the share of Type 2 and 3 candidates would increase, but the share of Type 1 would decrease.

Although choice probabilities for types 1-3 could decrease, we find that, for each type, this is the case in less than 2% of the constituencies. In absolute terms, the increase in average choice probabilities is largest for the educated type (Type 1), whose expected share increases by 25.4 percentage points, from 51.4 to 76.8 percent (Table 9). Eliminating criminal candidates increases the share of uneducated candidates (Type 2) by 55 percent (from 7.4 to 11.5), and the share of minority (Muslim) candidates (Type 3) by 43 percent (from 8.2 to 11.7). The latter effect is larger for the NDA, where the share of Muslim candidates was relatively low initially.

As shown in Table 9, these changes also impact the political viability of different candidate

Table 9: Different types’ probability of being chosen and winning, with and without the Criminal type

	Choice probability		Conditional win probability	
	Baseline	Counterfactual	Baseline	Counterfactual
<i>All candidates</i>				
Type 1	51.4	76.8	19.0	20.1
Type 2	7.4	11.5	20.8	22.0
Type 3	8.2	11.7	13.0	13.8
Type 4	33.1	0.0	38.2	-
<i>UPA</i>				
Type 1	50.0	75.4	16.7	17.8
Type 2	5.5	8.7	18.0	19.1
Type 3	11.3	16.0	11.3	12.0
Type 4	33.2	0.0	34.2	-
<i>NDA</i>				
Type 1	52.8	78.2	21.4	22.5
Type 2	9.2	14.4	23.6	25.0
Type 3	5.1	7.4	14.6	15.5
Type 4	33.0	0.0	42.2	-

Notes: Types 1-4 are the Educated, Uneducated, Muslim, and Criminal types, respectively. Values shown are the averages across all the constituencies in the data. Conditional win probability is the probability that a party would win conditional on choosing a given type.

types, captured by the conditional winning probabilities $E_P[w_p(a_p, \mathbf{a}_{-p})|a_p]$ (the probability that a type would win conditional on being selected). In each case, this conditional winning probability increases on average. Overall, we find that banning criminal candidates would raise the probability of running, as well as the probability of winning, of educated, uneducated, and Muslim candidates.

7.2 Changes in parties’ winning probabilities and payoffs

Figure 4 shows the changes in the two parties’ vote share and probability of winning following a ban on criminal types. The left panel shows that vote shares almost always decline, by an average of 3.1 percentage points. To put this in context, the average winning margin in the data is 8.1 percentage points. The reduction in vote shares also means a reduction in the probability of winning. The probability that the UPA wins falls in 93% of the markets, by an average of 3 percentage points. The probability that the NDA wins falls in 92% of the

markets, by an average of 5 percentage points.

One reason for the decline in the main parties' vote share is that once criminal types are not available, some voters find the main parties less attractive relative to third parties: the probability that a third party wins increases in 97% of the markets. The pattern that yields this result is apparent already in the data. Across constituencies where both main parties run criminal types, the average total vote share of all third parties is 21.1%. Across constituencies where neither main party runs a criminal type, this vote share increases to 28.6%. The main parties' equilibrium response to a criminal ban can mitigate this decrease in popularity. However, the results on Figure 4 indicate that they cannot fully offset it, and hence their winning probability goes down.²⁹

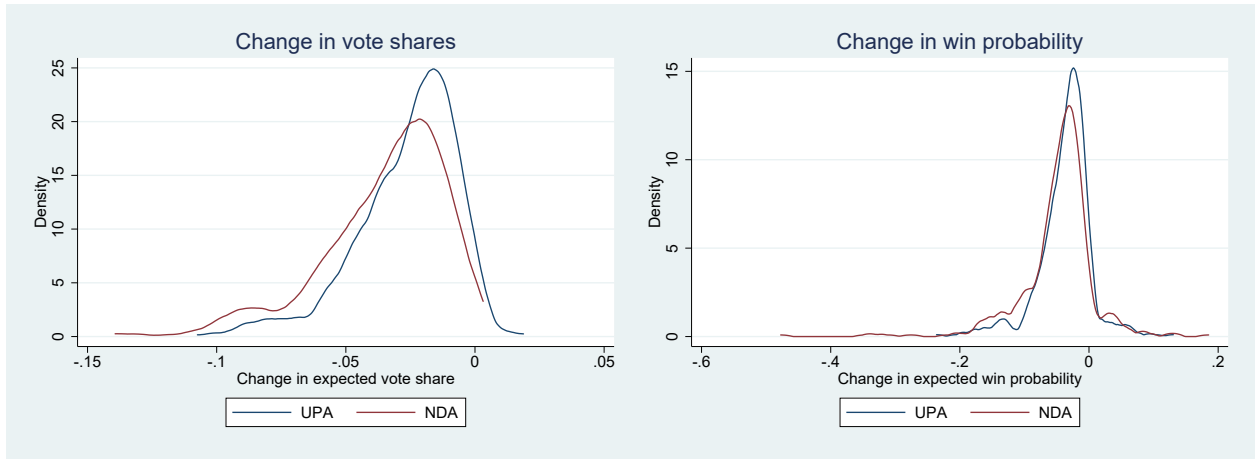


Figure 4: Distribution of changes in parties' vote shares and winning probabilities
Kernel density plots of changes in parties' unconditional expected vote shares and winning probabilities (counterfactual minus the baseline).

Computing the changes in parties' expected equilibrium payoffs yields the distributions shown in Figure 5. On average, both the UPA's and the NDA's payoffs decrease when the criminal type is banned. Both parties' payoffs decline in 28% of the constituencies. However, there are also many constituencies where payoffs rise: payoffs increase for both parties in 34% of the constituencies, for only the UPA in 17%, and for only the NDA in 21%.

²⁹It is important to note that the reduction in winning probabilities for the two major parties is *not* due to voters' switching to third parties who run criminals. First, we obtain similar patterns when we restrict attention to the 78 markets where none of the third parties runs a criminal: here the probability that the UPA wins falls in 91% of cases, the probability that the NDA wins falls in 88% of cases, and the probability that a third party wins increases in all cases but one (99%). Second, we also ran a set of counterfactuals where we removed all third parties running criminals, and measured the impact of the criminal ban relative to this modified baseline. We again found that the criminal ban lowered the winning probability for both the UPA (81% of markets) and the NDA (80% of markets), and increased the probability of third party winners (82% of markets).

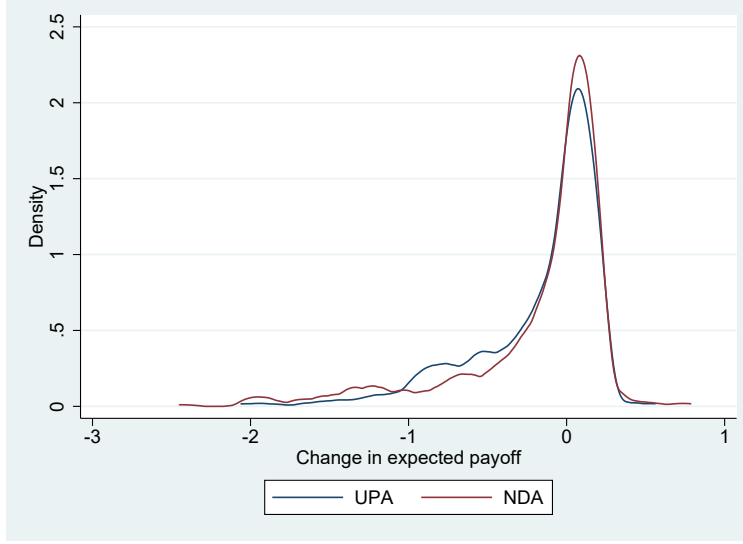


Figure 5: Distribution of the change in parties' payoffs
Kernel density plots of the change in parties' expected payoffs (counterfactual payoff minus baseline payoff).

The reason for this is related to the strategic complementarity discussed in Section 6.2. According to our estimates, a party can be incentivized to choose a criminal because its opponent chooses a criminal. Banning the criminal type allows parties to profitably choose different candidates. These candidates yield lower equilibrium vote shares, and in some cases lower recruiting costs, which raises parties' payoffs. Even though their probability of winning goes down, this decline is not large enough to undo the positive effect.

These findings suggest that, in some cases, the main Indian parties may be willing to collectively support a ban on criminal candidates. This helps explain instances when these parties express support for a ban on criminal candidates while at the same time running such candidates in elections.

8 Conclusion

We estimate a model of candidate selection by political parties to study why parties in a representative democracy select the candidates they do. Our setting is India, and we combine a rich demand side specification of voter preferences with a supply side game between parties that incorporates direct payoffs from candidate selection. To make the problem feasible, we use a machine learning algorithm to assign candidates to types based on detailed information on their characteristics.

We find that, while parties systematically select candidates with a higher probability

of winning the election, selection decisions are shaped by other considerations. All else equal, parties would prefer to win with candidates who are not overly popular with voters. This is consistent with a common perception in Indian politics that the major parties are reluctant to select candidates who might become too powerful and threaten the party and its elites. In addition, we find that selection decisions are shaped by considerations that are independent of voter preferences. Parties are more likely to select candidates whose characteristics are abundant in local candidate pools, and they also have distinct direct preferences over candidate types.

Our estimates provide a detailed explanation for the factors that cause parties to run criminal candidates. Apart from electoral considerations and party costs, we also find an important strategic complementarity. This creates an incentive for parties to run criminals because their opponent does so as well.

We use the estimated model to study the consequences of a ban on criminal candidates, a question of current policy relevance in India and elsewhere. We find that the ban reduces the chances that India's major national parties win the election, as the ban causes more voters to choose third party alternatives. In spite of this, banning criminal candidates can increase parties' payoffs by allowing them to compete with more desirable candidates.

Our paper provides insights for other countries where crime in politics is a salient problem, and more generally a method to identify and estimate the forces that guide parties' selection of candidates.

References

- Aguirregabiria, Victor, and Pedro Mira.** 2007. “Sequential estimation of dynamic discrete games.” *Econometrica*, 75(1): 1–53.
- Asher, Sam, and Paul Novosad.** 2023. “Rent-Seeking and Criminal Politicians: Evidence from Mining Booms.” *The Review of Economics and Statistics*, 105(1): 20–39.
- Asher, Sam, Tobias Lunt, Ryu Matsuura, and Paul Novosad.** 2020. “Development Research at High Geographic Resolution: An Analysis of Night Lights, Firms, and Poverty in India using the SHRUG Open Data Platform.” *working paper*.
- Athey, Susan, and Guido W Imbens.** 2019. “Machine learning methods that economists should know about.” *Annual Review of Economics*, 11: 685–725.
- Bandiera, Oriana, Andrea Prat, Stephen Hansen, and Raffaella Sadun.** 2020. “CEO behavior and firm performance.” *Journal of Political Economy*, 128(4): 1325–1369.
- Banerjee, Mukulika.** 2017. *Why India Votes?* Routledge.
- Berry, Steven, James Levinsohn, and Ariel Pakes.** 1995. “Automobile prices in market equilibrium.” *Econometrica*, 63(4): 841–890.
- Besley, Tim.** 2005. “Political Selection.” *Journal of Economic Perspectives*, 19(3): 43–60.
- Besley, Timothy, and Marta Reynal-Querol.** 2011. “Do democracies select more educated leaders?” *American political science review*, 105(3): 552–566.
- Besley, Timothy, Olle Folke, Torsten Persson, and Johanna Rickne.** 2017. “Gender quotas and the crisis of the mediocre man: Theory and evidence from Sweden.” *American Economic Review*, 107(8): 2204–2242.
- Bonhomme, Stéphane, Thibaut Lamadon, and Elena Manresa.** 2022. “Discretizing unobserved heterogeneity.” *Econometrica*, 90(2): 625–643.
- Casas-Arce, Pablo, and Albert Saiz.** 2015. “Women and power: unpopular, unwilling, or held back?” *Journal of political Economy*, 123(3): 641–669.
- Casey, Katherine, Abou Bakarr Kamara, and Niccoló F. Meriggi.** 2021. “An experiment in candidate selection.” *American Economic Review*, 111(5): 1575–1612.
- Chandra, Kanchan.** 2016. *Democratic dynasties: State, party, and family in contemporary Indian politics*. Cambridge University Press.
- Chhibber, Pradeep, Francesca Refsum Jensenius, and Pavithra Suryanarayan.** 2014. “Party organization and party proliferation in India.” *Party Politics*, 20(4): 489–505.
- Dal Bó, Ernesto, and Frederico Finan.** 2018. “Progress and perspectives in the study of political selection.” *Annual Review of Economics*, 10: 541–575.

- Dal Bó, Ernesto, Frederico Finan, Olle Folke, Torsten Persson, and Johanna Rickne.** 2017. “Who becomes a politician?” *The Quarterly Journal of Economics*, 132(4): 1877–1914.
- Dar, Aaditya.** 2019. “Parachuters vs. Climbers: Economic Consequences of Barriers to Political Entry in a Democracy.” *working paper*.
- Farooqui, Adnan, and Eswaran Sridharan.** 2014. “Incumbency, internal processes and renomination in Indian parties.” *Commonwealth & Comparative Politics*, 52(1): 78–108.
- Fisman, Raymond, Florian Schulz, and Vikrant Vig.** 2016. “Financial disclosure and political selection: Evidence from India.” *working paper*.
- Fudenberg, Drew, and Jean Tirole.** 1991. *Game theory*. MIT press.
- Galasso, Vincenzo, and Tommaso Nannicini.** 2011. “Competing on good politicians.” *American Political Science Review*, 105(1): 79–99.
- Gandhi, Amit, and Jean-François Houde.** 2019. “Measuring substitution patterns in differentiated-products industries.” *NBER Working paper w26375*.
- Gulzar, Saad, Zuhad Hai, and Binod Kumar Paudel.** 2021. “Information, candidate selection, and the quality of representation: Evidence from Nepal.” *The Journal of Politics*, 83(4): 1511–1528.
- Hamilton, Barton H., Andrés Hincapié, Robert A. Miller, and Nicholas W. Papageorge.** 2021. “Innovation and Diffusion of Medical Treatment.” *International Economic Review*, 62(3): 953–1009.
- Heath, Oliver, and Adam Ziegfeld.** 2022. “Why So Little Strategic Voting in India?” *American Political Science Review*, 116(4): 1523–1529.
- Iaryczower, Matias, Galileu Kim, and Sergio Montero.** 2023. “Representation Failure.” *working paper*.
- Kapoor, Sacha, and Arvind Magesan.** 2018. “Independent candidates and political representation in India.” *American Political Science Review*, 112(3): 678–697.
- Lee, Alexander.** 2020. “Incumbency, parties, and legislatures: Theory and evidence from India.” *Comparative Politics*, 52(2): 311–331.
- Mattozzi, Andrea, and Antonio Merlo.** 2015. “Mediocracy.” *Journal of Public Economics*, 130: 32–44.
- Nevo, Aviv.** 2000. “A practitioner’s guide to estimation of random-coefficients logit models of demand.” *Journal of Economics & Management Strategy*, 9(4): 513–548.
- Nevo, Aviv.** 2001. “Measuring market power in the ready-to-eat cereal industry.” *Econometrica*, 69(2): 307–342.

- Olea, José Luis Montiel, and Carolin Pflueger.** 2013. “A Robust Test for Weak Instruments.” *Journal of Business & Economic Statistics*, 31(3): 358–369.
- Paiva, Natália, Juliana Sakai, and Lauren Schoenster.** 2014. *Quatro em cada dez candidatos a governador têm ocorrências na Justiça ou nos Tribunais de Contas*. Transparencia Brasil, <https://www.transparencia.org.br/publicacoes>.
- Prakash, Nishith, Marc Rockmore, and Yogesh Uppal.** 2019. “Do criminally accused politicians affect economic outcomes? Evidence from India.” *Journal of Development Economics*, 141: 102370.
- Roy, Ramashray.** 1966. “Selection of Congress candidates. I: The formal criteria.” *Economic and Political Weekly*, 833–840.
- Thorndike, Robert.** 1953. “Who belongs in the family?” *Psychometrika*, 18(4): 267–276.
- Trebbi, Francesco, and Eric Weese.** 2019. “Insurgency and small wars: Estimation of unobserved coalition structures.” *Econometrica*, 87(2): 463–496.
- Ujhelyi, Gergely, Somdeep Chatterjee, and Andrea Szabó.** 2021. “None of the Above: Protest Voting in the World’s Largest Democracy.” *Journal of the European Economic Association*, 19(3): 1936–1979.
- Vaishnav, Milan.** 2017. *When crime pays: Money and muscle in Indian politics*. Yale University Press.