

Political Representation and Perceived Ethnic Discrimination in Africa

Brice Romuald Gueyap Kounga*

Wilfried Youmbi Fotso[†]

September 14, 2026

Abstract

African governments are widely believed to favour the ethnic group of the leader, and two in five respondents in the pooled sample say their own group is treated unfairly by the state. We ask what this belief responds to: what the state gives, or who holds it. Following a single question asked of 255,659 Afrobarometer respondents in 42 countries over two decades, and identifying from the 30 ethnic-group-by-country cells in eleven countries whose access to the presidency changes, we find that gaining the presidency lowers the share of a group's members reporting frequent unfair treatment by 3.8 percentage points relative to other groups interviewed in the same country and round, a fifth of the mean. Inference is at the country level, where the estimate survives a wild cluster bootstrap. Over the same transitions, the confidence intervals rule out a favourable lived-poverty response larger than half of the belief response and a favourable local-public-goods response larger than a third of it. In a stacked design the response appears by the first survey round in which the group holds the presidency, and it is precisely estimated only for competitive electoral turnovers; successions that change the president's ethnicity without an electoral victory are not detectably negative. We therefore cannot separate ethnic representation from partisan victory, nor a belief about state treatment from a broader change in perceived group standing. A co-ethnic group's acquisition of the presidency, particularly through a competitive electoral turnover, is associated with a substantial relative change in reported group treatment, while the two contemporaneous material outcomes we observe respond much less.

Keywords: ethnic favouritism, political representation, perceived discrimination, elections, Africa.

JEL: D72, D83, H41, O55, P16.

*Department of Economics, Western University. E-mail: bgueyapk@uwo.ca. We thank Afrobarometer and its national partners for the survey data, and the authors of PLAD and ACPED for making their data public. Replication code and the two hand-built ethnic crosswalks are available from the authors. All errors are ours.

[†]Department of Economics, Wilfrid Laurier University. E-mail: wyoumbifotso@wlu.ca.

1 Introduction

Nearly forty percent of adults surveyed in sub-Saharan Africa report that members of their ethnic group are treated unfairly by their government, and 17.0 percent report that this happens often or always. Such perceptions are widely believed to matter. In standard accounts of ethnic conflict, groups mobilise not over differences in income as such, but over perceived injustice in how the state treats them (Stewart, 2008; Cederman et al., 2013), and perceptions of unfair treatment have been linked to tax non-compliance, withdrawal from formal institutions, protest, and support for armed groups. Despite this, we know surprisingly little about where these beliefs come from. A sizeable empirical literature documents what African governments actually allocate to the ethnic group of the leader: roads (Burgess et al., 2015), night-time light (Hodler and Raschky, 2014; De Luca et al., 2018; Dickens, 2018), schooling and infant mortality (Franck and Rainer, 2012). Almost nothing is known, by contrast, about how what citizens *believe* responds to the same political changes.

In this paper, we examine how perceived ethnic discrimination responds to a group’s access to executive power. We follow a single survey question, “How often, if ever, are [members of the respondent’s ethnic group] treated unfairly by the government?”, across seven rounds of the Afrobarometer surveys, covering 255,659 respondents in 42 countries between 2005 and 2023. We link each respondent to two measures of their ethnic group’s access to power on the exact date of the interview: whether the group holds the presidency, using a hand-audited register of heads of state that we construct from PLAD (Bomprezzi et al., 2024), and the group’s share of cabinet posts in the month of the interview, using the ACPED minister-month data (Raleigh and Wigmore-Shepherd, 2020). In linking survey respondents to characteristics of their ethnic group, our design follows Nunn and Wantchekon (2011); the difference is that our group-level treatment varies over time, which allows us to hold the group itself fixed.

Our empirical strategy compares a group to itself as its access to power changes. All specifications include ethnic-group-by-country fixed effects, which absorb the permanent level of each group’s reported grievance (including its historical position, any time-invariant discrimination, and the group’s use of the response scale), and country-by-round fixed effects, which absorb all national shocks, including the round’s translation of the question and the national political climate. Identification therefore comes from changes in a group’s access to power, relative to other groups interviewed in the same country in the same round. The within-group design also addresses a well-known measurement problem: because “often” may mean different things to a dominant group and a marginalised one (King et al., 2004), comparisons of levels across groups are difficult to interpret, but changes within a group are informative.

We find that representation has a large effect on the belief. When a respondent’s ethnic group gains the presidency, the probability that the respondent reports frequent unfair treatment falls by 3.8 percentage points (standard error 0.9), which is 21 percent of the estimation-sample mean and 0.10 standard deviations of the outcome. The probability of reporting any unfair treatment falls by 6.1 percentage points, and the shift runs across the whole scale rather than at one end: the share answering “never” rises by 6.1 points while each of the three positive categories falls by about two. Cabinet representation has a smaller effect in the same direction: a ten percentage-point increase in a group’s share of cabinet posts reduces reported

discrimination by 0.8 percentage points.

Because treatment switches on and off rather than being absorbed once taken, we examine the estimator itself. The two-way fixed-effects estimand here carries 0.0 percent negative weights, so it is a proper convex combination of the underlying effects, and the estimator of [de Chaisemartin and D’Haultfoeuille \(2020\)](#), built for treatments that switch in both directions, returns -0.0778 : the same sign, a larger magnitude. Inference does not rest on the analytic variance formula either. Clustering on country leaves only eleven of twenty-eight clusters with any treatment variation, which is the range in which cluster-robust standard errors are known to be too small, so we report a wild cluster bootstrap: $p = 0.015$.

We then examine whether anything material changes alongside the belief, and we are careful about what the answer establishes. Using the same respondents and specification, an index of ten amenities that the enumerator physically observed in the respondent’s enumeration area (schools, clinics, police stations, markets, electricity, piped water, sewage, telephone service, post offices and paved roads) changes by -0.03 standard deviations (95 percent confidence interval $[-0.09, +0.03]$, clustered on country), and the lived-poverty index, which records how often the household went without basic needs, changes by -0.01 standard deviations ($[-0.05, +0.03]$). These are bounds rather than zeros, and the intervals carry the claim: they rule out a favourable lived-poverty response larger than about half of the 0.098 belief response, and a favourable local-public-goods response larger than about a third of it. Two further limits apply to both. The measure describes the respondent’s locality rather than the ethnic group, and it is contemporaneous: the favouritism literature estimates allocation over five to twenty years, so a null within a survey round bounds the speed of the material response, not its existence.

A natural concern is that the fairness question simply measures approval of the government, and that co-ethnics of a sitting president are more favourable about everything. Three findings suggest this is not the whole story. First, although the same transitions raise presidential approval by 0.42 standard deviations (four times the grievance response), a regression that conditions on approval leaves 45 percent of the grievance coefficient. Approval is itself moved by co-ethnicity, so that decomposition is descriptive rather than a mechanism test: it shows the two responses are not collinear, not that 45 percent of the belief response is independent of approval. Second, the belief response survives adjustments that generic favourability should not survive: it is unchanged by controls for proximity to a national election (-0.0396 , s.e. 0.0091), and it is larger, not smaller, when country-by-round effects are dropped (-0.0531 , s.e. 0.0175). That last specification is not an absolute causal effect either, since dropping the country-by-round effects readmits national shocks; it establishes only that the pooled response is not mechanically zero-sum. Third, the response is asymmetric in a way generic approval is not: groups that lose the presidency show an increase in grievance that disappears within one round, while groups that gain it show a decline that persists.

Because the outcome is a survey response, we devote particular attention to measurement threats. The first is composition: [Green \(2021\)](#) shows that citizens in autocracies switch their stated ethnic identity toward a new president’s group, which could move group averages mechanically. We find that the president’s group’s share of the survey sample rises by +0.63 percentage points ($p = 0.07$) around

a transition, the direction Green’s results predict. Bounding its consequence requires care about the denominator, because presidents come from large groups: across the cells that actually switch, whose mean share is 15.6 percent, an increase of +0.63 points makes entrants 3.9 percent of the post-transition group. Even under the extreme assumption that every entrant reports no unfair treatment at all, composition accounts for at most 0.71 points of the 3.8-point estimate. That is a real contribution rather than a negligible one, and it is a bound rather than an estimate; the stronger defence is that our identifying transitions sit in electoral regimes, where Green finds no switching. The second is selection: item non-response does not change detectably with co-ethnic status, though the interval admits a shift of up to 1.4 points. The third is the enumerator. [Adida et al. \(2016\)](#) show that responses to ethnically sensitive questions in Africa depend on the ethnicity of the interviewer, and [Nunn and Wantchekon \(2011\)](#) control for interviewer characteristics for the same reason. We provide what is to our knowledge the first estimate of this effect for the discrimination question: being interviewed by a co-ethnic raises reported grievance by 2.1 percentage points. Controlling for interviewer co-ethnicity leaves our estimate unchanged. Adding interviewer fixed effects reduces it sharply, from -0.0569 to -0.0179 on that subsample, and because interviewers are assigned to enumeration areas this could reflect either the enumerator or the locality. The two are partially separable under a proxy construction. Afrobarometer does not publish a primary-sampling-unit identifier, but the sample design places eight respondents in each enumeration area and the enumerator’s amenity observations are recorded once per area, so a local cluster can be reconstructed from the locality name and the amenity vector. Doing so recovers 12,646 reconstructed local clusters with a median of eight respondents, and the decomposition is informative: comparing co-ethnics to their neighbours in the same reconstructed cluster gives -0.0365 , statistically indistinguishable from our headline estimate, while region effects, which are far coarser, change nothing. Locality at the scale at which people actually live accounts for most of the collapse; the residual attributable to the enumerator leaves the estimate at -0.0246 when both are absorbed. We report the full ladder rather than a single number.

Our findings contribute to three literatures, and we are explicit about what in each is already known. The first is the literature on ethnic favouritism in Africa ([Franck and Rainer, 2012](#); [Hodler and Raschky, 2014](#); [Burgess et al., 2015](#); [Kramon and Posner, 2013](#); [De Luca et al., 2018](#); [Dickens, 2018](#)), which measures what leaders deliver to their co-ethnics. [De Luca et al. \(2018\)](#) identify from leader turnover within country-and-group cells, which is our design applied to the material side of the comparison. Closest to our result is [Bandyopadhyay and Green \(2025\)](#), who use Afrobarometer with individual co-ethnicity with the president and find no material co-ethnic benefit alongside a co-ethnic premium in subjective evaluations that survives conditioning on the services being evaluated. We are therefore not the first to document a wedge of this kind. What we add is the fairness item itself rather than evaluations of services, within-group identification from turnover rather than cross-sectional co-ethnicity, and bounds on the material response rather than a null. The implication we draw, that perceived and measured favouritism are distinct objects and that using one as a proxy for the other is a substantive assumption, has an established analogue for vertical income inequality ([Gimpelson and Treisman, 2018](#); [Cruces et al., 2013](#)), and our contribution is to show it holds for the horizontal, ethnic case where the conflict literature relies on the

proxy most heavily (Cederman et al., 2011; Baldwin and Huber, 2010). The second is the literature on ethnic identity in Africa (Eifert et al., 2010; Green, 2021; Robinson, 2014; Müller-Crepon, 2025; Depetris-Chauvin et al., 2020; Ali et al., 2019; McNamee, 2019; Letsa and Wilfahrt, 2020; Chlouba, 2026), which studies who claims an identity and when it becomes salient. Robinson (2014) is the canonical treatment of the ethnic-versus-national item we use as our comparison outcome, and reports no relationship between head-of-state ethnicity and national identification; Depetris-Chauvin et al. (2020) show identification moving causally in these data in response to a national shared experience. Perceived discrimination is a different outcome from either: in our data, transitions that move it by 0.10 standard deviations move ethnic-versus-national identification by a statistically insignificant 0.01, and the two move in opposite directions in Cameroon after 2018. The third is the literature linking grievances to conflict (Stewart, 2008; Cederman et al., 2013; Mamdani, 1996), which has largely relied on measured inequality as a proxy for grievance; our estimates suggest the proxy misses much of the variation that politics generates. Among the handful of papers that use this survey question at all (Tiku and Sylwester, 2024; Villamil, 2023; Gadjanova, 2022; Tuki, 2025), none exploits within-group variation in representation, and none addresses the enumerator problem.

We should also say what our interpretations are not. The reading in Section 7 on which representation changes a group’s perceived standing rather than its allocation is group-position theory (Bobo and Hutchings, 1996; Major et al., 2002) applied to a new setting, not a new mechanism; the finding that descriptive representation moves attitudes toward institutions has a long history in the representation literature, though the canonical result there is more equivocal than it is sometimes reported to be, since Gay (2002) finds the effect confined to the constituent-representative relationship and absent at the system level. Isaksson (2015) is the nearest precedent for comparing a subjective and an objective measure of state treatment on the same African respondents, using experienced corruption and group political standing.

The remainder of the paper is organised as follows. Section 2 provides historical background and a simple conceptual framework whose predictions organise the results. Section 3 describes the data, including the construction of the leader register and the two ethnic crosswalks. Section 4 presents the empirical strategy. Section 5 reports the main results, Section 6 examines identification concerns, and Section 7 discusses interpretation. Section 8 examines Cameroon, where the executive never changes but the salience of the ethnic cleavage does. Section 9 concludes.

2 Background and conceptual framework

2.1 Ethnicity and the state in Africa

The belief we study has a history. Colonial rule did not find fixed ethnic groups in Africa; to a substantial degree it made them, administering subjects through codified “customary” authorities and hardening fluid identities into census categories with differential access to the state (Mamdani, 1996). What decolonisation transferred was therefore not only a state but a particular relationship between ethnicity and the state: group membership had been, within living memory, an official determinant of what one could receive

and where one could stand. The post-independence concentration of power in presidencies sharpened the stakes. In systems where the executive controls appointments, procurement and the location of public investment, the question “who holds the presidency?” is a question about the terms on which one’s group deals with the state. A large empirical literature confirms that the answer matters materially: co-ethnic regions receive more roads (Burgess et al., 2015), emit more night-time light (Hodler and Raschky, 2014; De Luca et al., 2018), targeted at ethnic homelands (Dickens, 2018), and their children are more likely to be educated and to survive infancy (Franck and Rainer, 2012), with the effects concentrated where institutions constrain the executive least.

Citizens, then, are not wrong to believe that ethnic favouritism exists; the literature just cited is the proof that on average it does. The question this paper asks is a different one: when a group’s access to power changes, does the group’s *belief* about its treatment track what it measurably receives, or does it track the access itself?

2.2 A simple framework

A minimal framework clarifies what the empirical results can and cannot distinguish. Let θ_{gt} denote how unfairly the state actually treats group g at time t , the object the survey question asks about. Citizens do not observe θ_{gt} . A respondent forms the belief

$$g_{igt} = \mathbb{E}[\theta_{gt} \mid R_{gt}, s_{igt}], \quad (1)$$

where R_{gt} is the group’s representation in the executive (public, salient, and free to observe) and s_{igt} is the respondent’s noisy private signal of actual treatment: the state of the local clinic, the household’s material circumstances, what neighbours and media report. The weight each input receives depends on its precision. Representation is a useful input into (1) precisely because favouritism is real on average: a citizen who knows the literature’s average effect rationally revises θ downward when a co-ethnic takes power.

This framework yields two predictions and two further implications that organise the results. First, beliefs should respond to representation even if nothing local changes, because R enters the expectation directly (Section 5.1). Second, if actual allocation responds little or slowly within a survey round, the private signal s should not move, and measured local conditions should show no response (Section 5.2). Third, an information mechanism operating through political attentiveness would predict a larger response among urban and more educated respondents; the fully interacted specifications in Section 6 do not support that prediction. Fourth, if cabinet representation provides information about group treatment independently of presidential approval, cabinet inclusion should move the fairness item more than approval; the data do not support this implication once presidential co-ethnicity is controlled for and the president’s own group is excluded (Section 5.4). That the auxiliary implications fail is itself informative: it narrows the mechanism toward public status, electoral victory, or broad perceived group standing.

The framework is deliberately agnostic about whether responding to R is a mistake. If representation

genuinely improves a group’s treatment in ways our amenity index cannot see (police behaviour, bureaucratic access, dignity in encounters with officials), then the belief response is informative updating on an imperfect but relevant signal. If it does not, the response is a heuristic that overweights the visible. Section 7 returns to this distinction; the data in hand bound the material channel but cannot close it entirely.

3 Data

3.1 Data sources and description

We combine four sources. The first is the Afrobarometer, a nationally representative repeated cross-section fielded face-to-face by national partners under a common protocol, with roughly 1,200 to 2,400 respondents per country per round. We pool the seven rounds that carry the discrimination question, Round 3 (2005–06, 18 countries) through Round 9 (2021–23, 39 countries), for a total of 255,659 respondents with valid responses in 42 countries. Beyond the attitudinal items, the surveys record each respondent’s self-reported ethnic group from country-specific code lists, the exact date of the interview, the enumerator’s observations of the enumeration area, and (in three rounds) the interviewer’s own ethnic group; these last two features are, as will become clear, essential to the design.

The second source is a register of every head of state in office in the 38 sample countries between 2005 and 2023, with the leader’s ethnic group and region of birth, which we construct starting from PLAD ([Bomprezzi et al., 2024](#)) and audit leader by leader against multiple codings (Section 3.3). The third is ACPED ([Raleigh and Wigmore-Shepherd, 2020](#)), a panel of 316,202 minister-months covering 37 countries from 1996 to mid-2021, from which we build each ethnic group’s monthly share of cabinet posts. The fourth is a set of hand-built linking files: a crosswalk from each leader’s group to the Afrobarometer labels of that country, a crosswalk from Afrobarometer’s 669 ethnic labels to ACPED’s politically relevant groups, and a day-precise register of 150 national elections and 14 coups used for the timing analyses. All four are constructed for this paper and are part of the replication package; because the headline estimates rest on the coding judgements inside them, Section 3.3, Appendix A and the replication documentation accompanying the package record every correction, exclusion and sensitivity.

The unit of observation throughout is the individual respondent. Each respondent is assigned the leader in office and the cabinet composition prevailing on their own interview date, so treatment varies within country-round wherever fieldwork spans a transition, and across rounds otherwise. Table 1 reports summary statistics for the pooled sample.

3.2 Outcome: the grievance item

The dependent variable comes from the question “How often, if ever, are _____s [respondent’s ethnic group] treated unfairly by the government?”, answered on a four-point scale: never (0), sometimes (1), often (2), always (3). It is renumbered every round (q81 in Round 3, Q82 in Round 4, Q85a in Rounds 5

and 7, Q88a in Round 6, Q82a in Round 8 and Q84b in Round 9), which is one reason it is under-used. We build three coded outcomes: an indicator for “often” or “always” (the headline measure), an indicator for any unfair treatment, and the raw 0–3 scale. Pooling Rounds 3–9 gives 255,659 valid responses in 42 countries; 17.0% choose often or always and 39.9% report any unfair treatment.¹

Round 2 is excluded. Its nearest item asks how often “your identity group” is treated unfairly, with no ethnicity variable to fix the referent, so it is not the same question. Rounds 3 and 4 ask it of the respondent’s ethnic group in the same words and on the same four-point scale as Rounds 5–9, and their value labels are identical. Two variables the later rounds carry are absent before Round 5 and are simply missing there: the lived-poverty index, and the interviewer’s ethnic group. Neither is a headline outcome. The enumerator’s amenity checklist is shorter in Round 3 (eight of the ten items; no mobile-phone service and no road), which country-by-round fixed effects absorb; Table 7 is unchanged when the index is restricted to the eight items observed in every round.

3.3 Treatment 1: the head of state’s ethnic group

We built a register of every head of state in office in each sample country between 2005 and 2023, with the leader’s ethnic group and first-level administrative unit of birth. The starting point is PLAD v7.0 (Bomprezzi et al., 2024), which codes leader birthplace, GADM identifiers and up to two ethnicity fields with a precision flag. PLAD is not usable unaudited: in our 38-country subsample we found and made thirteen changes, listed in Appendix Table A3: eight are corrections to a date or an attribution, four are spells we drop because no defensible attribution exists, and one records a finer attribution that Afrobarometer’s response categories cannot represent. We drop a leader spell outright whenever the leader’s group is not reliably documented in any source we could reach (Masisi, Damiba, Guebuza, Magufuli, Samia Suluhu Hassan and Motlanthe) or is contested between sources we regard as equally credible (Barrow, Koroma, Mogae, Mkapka and Amadou Toumani Touré). Eleven spells are excluded in total. The pre-2011 spells were checked against three independent codings, PLAD, ACPED’s politically relevant ethnicity field, and the leader-to-survey-category crosswalk in Green (2021), with priority given to the category a respondent could actually have selected; the five exclusions above are the cases those three disagree on. Two are worth naming because the alternative would have been consequential: no source attributes any ethnic group to Motlanthe, and Mogae is coded Tswana by one source and Kalanga by another, which is the difference between a plurality and a seven-percent minority. Following the convention that ethnicity is not a meaningful political cleavage there, we drop Eswatini, Lesotho, Cabo Verde, São Tomé and Príncipe, Mauritius, Seychelles, Sudan, Tunisia, Algeria, Egypt, Morocco and Burundi.

Matching the leader’s group to the respondent’s self-reported label required a hand-built country-by-country crosswalk with explicit include and exclude rules, since Afrobarometer spellings drift across rounds (Cameroon “Bamiléké”/“Bameleké”; Sierra Leone “Temne”/“Themne”; Tanzania “Sukuma”/

¹Recent commentary assumes the item was dropped after Round 8. It was not: we verified directly in the Round 9 merged file that it is asked in 36 of 39 countries (all except Seychelles, Sudan and Tunisia), yielding 45,768 valid responses.

“Wasukuma”/“Msukuma”). Naive string matching produces false positives (Angola’s “Mbundu” matches “Ovimbundu”, a different group), which is why the exclude rules matter. The resulting co-ethnicity sample is 196,769 respondents in 28 countries, of whom 196,728 enter the estimation once singleton cells are dropped; 23.3% are co-ethnic with the sitting head of state, and the country-level shares track the leaders’ groups’ population shares closely (Zimbabwe 77%, Burkina Faso 56%, Ghana 46%, Madagascar 29%, Mozambique 2%, Tanzania 1%), which is a first validation of the coding.

The same drift has a second, less obvious cost. The group-by-country fixed effects are formed on the respondent’s label, so a group written two ways across rounds becomes two cells and its treatment variation is discarded rather than used. Most of the drift is an accent artefact (“Malinké” and “Soninké” lose their final vowel under normalisation), but two cases cost real variation: Namibia’s Ovambo, whose president Pohamba held office across four of our rounds, is “Oshiwambo” in Round 3 and “Wambo” thereafter, and Senegal’s Haalpulaar is “Pular” in Round 3 and “Pulaar/Toucouleur” thereafter. We alias these explicitly. We do *not* merge labels that bundle distinct groups, since folding one into a treated cell would treat the others as co-ethnic; Round 4, for example, gives Nigeria a single code reading “Ijaw/Kalabari/Andoni/Nembe/Ogoni/Okirika”. Those remain their own cells and contribute nothing to the estimate.

3.4 Treatment 2: cabinet representation

The second treatment comes from ACPED v2 ([Raleigh and Wigmore-Shepherd, 2020](#)), a minister-month panel of 316,202 rows covering 37 African countries from December 1996 to June 2021, in which every minister carries a politically relevant ethnic group in the Ethnic Power Relations classification ([Vogt et al., 2015](#)). The raw file records only ministers who hold office, so we reconstruct the full (country $\times$ month $\times$ group) grid and fill zero cabinet share for group-months with no minister, a step the source file cannot express and one that matters, because complete exclusion is exactly the state we want to measure. This yields 97,820 group-months. We define *cab_share* as the group’s share of cabinet posts and *cab_gap* as that share minus the group’s population share.

Mapping Afrobarometer’s 669 country-specific ethnic labels onto ACPED’s coarser politically relevant groups required a second hand-built crosswalk, constructed country by country and documented in Appendix A; 614 of 669 labels are mapped. The merge succeeds for 75–80% of respondents in the covered countries. Because ACPED ends in June 2021, this treatment is available for Rounds 3 to 8 (166,889 respondents, 29 countries, 243 country-groups). Round 9 is analysed with the co-ethnicity treatment only.

3.5 Material and non-evaluative outcomes

The key comparison outcome is built from Afrobarometer’s own enumerator observations, recorded for every enumeration area independently of anything the respondent says: presence of a school, a health clinic, a police station, a market, an electricity grid, piped water, a sewage system, a mobile-phone service, a post office and a paved or tarred road. We use the share of the ten present as a 0–100 index. As a second,

individual-level material measure we use Afrobarometer’s lived-poverty index, the average frequency with which the household went without food, water, medical care, cooking fuel and cash income over the past year. Both are imperfect: the public-goods index is coarse and local rather than group-level, and the lived-poverty index is self-reported, though it asks about behaviour rather than attitude. Their virtue is that they are measured on exactly the same respondents, in the same interviews, as the belief. Table 1 reports the summary statistics.

Table 1: Summary statistics, Afrobarometer Rounds 3–9

	Mean	SD	N
<i>Beliefs about the state</i>			
Group treated unfairly, often or always	0.170	0.375	255,659
Group treated unfairly, ever	0.399	0.490	255,659
Unfair treatment, 0–3 scale	0.640	0.925	255,659
Approval of the president, 1–4	2.658	1.002	283,188
Trust in the president, 0–3	1.680	1.160	246,297
Satisfaction with democracy, 1–4	2.417	0.998	281,220
Ethnic identity before national	0.134	0.341	264,146
<i>Measured conditions</i>			
Local public goods index, 0–100	56.540	26.804	305,942
school in the enumeration area	0.855	0.352	304,804
health clinic	0.593	0.491	302,528
electricity grid	0.640	0.480	305,286
piped water	0.553	0.497	304,005
paved road	0.656	0.475	198,835
Lived poverty index, 0–4	1.304	0.941	196,918
<i>Respondent characteristics</i>			
Age	37.200	14.755	304,500
Female	0.501	0.500	305,976
Education, 0–9	3.400	2.202	304,971
Urban	0.423	0.494	302,601
<i>Treatments</i>			
Co-ethnic with the head of state	0.235	0.424	206,483
Group’s share of cabinet posts (%)	20.695	21.634	175,444
Group’s population share (%)	22.632	21.476	175,444

Pooled Rounds 3–9, unweighted. The public-goods components are enumerator observations of the respondent’s enumeration area; Round 3 observes eight of the ten items. The lived-poverty index is asked from Round 5 onward. Cabinet share and population share come from ACPED v2 and are available for Rounds 3–8.

4 Empirical strategy

For respondent i of ethnic group g in country c interviewed in round r on date $t(i)$, we estimate

$$y_{igcr} = \beta T_{gc,t(i)} + \alpha_{gc} + \delta_{cr} + X_i' \gamma + \varepsilon_{igcr}, \quad (2)$$

where T is either an indicator for the group holding the presidency on the interview date, or the group's cabinet share in the interview month. The group-by-country fixed effects α_{gc} absorb the permanent level of a group's reported grievance, including its scale use, its historical position and any time-invariant discrimination. The country-by-round fixed effects δ_{cr} absorb every national shock, including the round-specific wording and translation of the item, the national political climate and fieldwork conditions. Identification therefore comes from a group's change in access to power relative to the other groups interviewed in the same country in the same round. We therefore refer to β throughout as a *relative representation effect*, and we reserve causal language for that quantity rather than for an absolute effect on the represented group.

Presidential transitions occur at the country level, and shocks that accompany them are likely to be correlated across the ethnic groups of a country precisely when a transition happens. Clustering on group-by-country does not accommodate that dependence, so we treat country-level inference as primary. Clustering on country gives a standard error of 0.0118 on the headline estimate of -0.0380 ($p = 0.0033$, 95 percent interval $[-0.0610, -0.0149]$), against 0.0093 when clustering on group-by-country. Because only eleven of twenty-eight countries carry treatment variation, we also report a wild cluster bootstrap at the country level ($p = 0.015$) and a leave-one-country-out range. All three procedures reject at conventional levels. Throughout the paper, country-clustered standard errors are reported in parentheses as the primary inference and group-by-country errors in brackets for comparability with the literature. The bootstrap is a wild cluster bootstrap- t with Rademacher weights drawn at the country level, 9,999 replications, and the null imposed through restricted residuals. Estimation drops the 41 observations in group-by-country cells containing a single respondent, so reported observation counts are estimation samples.

For the co-ethnicity treatment, 30 of 916 country-group cells switch status across rounds, covering 44,623 respondents in eleven countries (Benin, Ghana, Liberia, Malawi, Mali, Namibia, Niger, Nigeria, Senegal, South Africa, Zambia). This is the identifying variation, and it is worth being explicit about how narrow it is. Two things reassure us. First, dropping any single country moves the headline coefficient only within $[-0.043, -0.031]$ against a full-sample -0.038 (Table 19). Second, the cabinet-share treatment, which varies continuously for all 243 country-groups and does not depend on presidential turnover at all, gives the same sign and a proportionate magnitude (Table 13).

The country-by-round effects make β a relative quantity: the comparison group is the other groups interviewed in the same country and round, and at a transition those groups have just lost access. A reader may reasonably ask whether representation reallocates grievance rather than reducing it. Dropping the country-by-round effects, so that identification rests additionally on within-country variation over time, gives -0.0531 (s.e. 0.0175), larger than the baseline rather than smaller. That specification is not

an absolute causal effect either, since removing the country-by-round effects readmits national shocks, transition-related conflict, macroeconomic movements and country-specific trends. It establishes only that the pooled response is not mechanically zero-sum. Section 8 shows that in Cameroon, where no transition occurs, the within-country distribution moves while the total does not.

The parallel-trends assumption is that, absent the transition, a group's grievance would have evolved like that of other groups in the same country and round. Extending the panel to Round 3 means 34 of the 43 switch events now carry two or more pre-periods. The pooled event study in Figure A1 is imprecise about that assumption: the coefficient two rounds before the transition is -0.049 (s.e. 0.045) for gaining groups, which is indistinguishable from zero but also nearly as large as the estimated response itself, so it cannot by itself rule out a pre-existing group-specific trend.

We therefore report a stacked design. The 43 transitions in which a cell changes its access to the presidency, 22 gains and 21 losses across 30 cells in 11 countries, each form a separate stack. Within a stack the treated unit is the transitioning group and the controls are the cells in the same country that never change status over Rounds 3 to 9, and the fixed effects are event-by-group and event-by-round, so no comparison is ever made across stacks and no already-treated group serves as a control for a later event. Table 2 reports the result separately for gaining and losing groups.

The pre-period is now informative. For gaining groups the coefficient two rounds before the transition is -0.019 (s.e. 0.025), a third the size of the response at the transition and not distinguishable from zero ($p = 0.46$); for losing groups it is $+0.003$ (s.e. 0.017). The response for gaining groups is -0.061 (s.e. 0.028) at the transition and -0.060 (s.e. 0.028) one round later, both significant at 5 percent. For losing groups the movement is positive, $+0.022$ (s.e. 0.021), and not significant. The pooled column is attenuated because gains and losses move in opposite directions, which is the same relative structure the baseline estimand carries.

This does not make the transitions exogenous. It removes the specific concern that the baseline compares groups at different points in their own treatment histories, and it shows that the pre-period movement in the pooled design was an artefact of pooling rather than evidence of a group-specific trend.

The mode of the transition matters, and it matters a great deal. The 43 cell-events arise from 24 distinct country-by-round transitions, which we classify from the public record as twelve competitive electoral turnovers, five deaths in office or constitutional successions, five within-party successions, one coup and one event that spans more than one change. Table 3 estimates the stacked design separately by mode for gaining groups, and the result is not what a simple representation story predicts.

In the stacked design, the negative response is precisely estimated only for competitive electoral turnovers, where it is -0.1027 (s.e. 0.0217) at the transition, roughly 1.7 times the pooled estimate. The succession estimates are not detectably negative, at $+0.0103$ (s.e. 0.0240) for deaths in office and constitutional successions and $+0.0690$ (s.e. 0.1120) for within-party successions, although the small number of events, especially the five within-party successions, limits the precision of that comparison. Deaths and within-party successions change who occupies the presidency and the ethnic group they belong

Table 2: Stacked event study around presidential transitions

Event time	(1) All events	(2) Gaining groups	(3) Losing groups
$t - 2$	−0.0040 (0.0194)	−0.0193 (0.0253)	0.0033 (0.0172)
$t - 1$		<i>reference</i>	
$t = 0$	−0.0126 (0.0159)	−0.0609** (0.0275)	0.0221 (0.0208)
$t + 1$	−0.0163 (0.0190)	−0.0600** (0.0275)	0.0157 (0.0155)
Events	43	22	21
Countries	11	11	11
Observations	162,744	80,683	82,061
Event $\times$ group effects	Yes	Yes	Yes
Event $\times$ round effects	Yes	Yes	Yes

Notes. Outcome is an indicator for reporting that one's ethnic group is treated unfairly by the government often or always. Each of the 43 transitions in which a group-by-country cell changes its access to the presidency forms a separate stack; within a stack the treated unit is the transitioning group and the controls are the cells in the same country that never change status over Rounds 3 to 9. Event time is measured in survey rounds, with $t - 1$ the omitted reference. Fixed effects are event-by-group and event-by-round, so no comparison is made across stacks. Standard errors clustered on country. Treated respondents supporting each coefficient: gaining groups, 4,442 at $t - 2$ (19 events), 3,860 at $t - 1$ (19), 4,454 at $t = 0$ (22), 4,040 at $t + 1$ (17); losing groups, 4,731, 5,485, 5,696 and 4,691 across the same event times. A respondent in a never-switching cell can serve as a control in more than one stack (1.41 stacks on average), so countries with several transitions contribute control observations more than once; reweighting each observation by the inverse of the number of stacks it appears in gives -0.056 (s.e. 0.029) at $t = 0$ for gaining groups against -0.061 unweighted. * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.

to, but they do not involve a partisan victory, a campaign, or a mobilised electorate.

We read this in two directions and think both belong in the paper. It is informative about mechanism: if the fairness belief simply tracked the ethnic identity of the office-holder, the successions should move it, and they do not. It is also the clearest evidence available in these data that the treatment bundles descriptive ethnic representation with partisan victory, which is the confound we flag in Section 7 and which we cannot break without pre-treatment party affiliation. A reader who believes the belief responds to partisan victory rather than to ethnic representation will find the pattern in Table 3 consistent with that reading, and we do not have the evidence to rule it out.

Table 3: Stacked estimates by the mode of the presidential transition, gaining groups

Mode	Events	Countries	$t - 2$	$t = 0$	$t + 1$
Competitive electoral turnover	10	7	−0.0415 (0.0260)	−0.1027*** (0.0217)	−0.0925*** (0.0230)
Death in office or constitutional succession	4	3	0.0277 (0.0210)	0.0103 (0.0240)	−0.0066 (0.0120)
Within-party succession	5	4	0.1721 (0.1090)	0.0690 (0.1120)	0.0635 (0.1100)
All gaining groups	22	11	−0.0193 (0.0253)	−0.0609** (0.0275)	−0.0600** (0.0275)

Notes. Stacked design as in Table 2, split by the mode of the underlying presidential transition. The 43 cell-events arise from 24 distinct country-by-round transitions, classified from the public record as twelve competitive electoral turnovers, five deaths in office or constitutional successions, five within-party successions, one coup (Mali, 2020–21) and one event spanning more than one change (Nigeria, Obasanjo to Jonathan across an absent round). The coup and the spanning event are too few to estimate separately and are retained in the pooled row only. Twelve competitive transitions produce the ten gaining-group stacks shown, because two of them generate only losing-group events; likewise five successions produce four gaining-group stacks. Standard errors clustered on country. In a two-period version of the competitive design ($t = 0$ against $t - 1$), a country-level wild cluster bootstrap gives $p = 0.033$ and the jackknife (CRV3) standard error is 0.0249 ($p = 0.014$). * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.

How much differential trend would overturn the result? Following [Rambachan and Roth \(2023\)](#), we allow the post-treatment violation of parallel trends to be at most $\bar{M}$ times the violation observed in the pre-period and report the breakdown value, the largest $\bar{M}$ at which zero stays outside the confidence set. Table 4 reports it. For all gaining groups pooled the breakdown value is 0.10, which is to say the pooled estimate does not survive even a small allowance for differential trends. For competitive turnovers, where the pre-trend is larger in absolute terms but the effect is larger still, the breakdown value is 0.70. With a single pre-period in the stacked window this analysis has limited power, and we report it as a statement about what the design can certify rather than as evidence that differential trends exist.

Table 4: Sensitivity of the effect at $t = 0$ to violations of parallel trends

$\bar{M}$	All gaining groups	Competitive turnovers
0.00	$[-0.115, -0.007]$	$[-0.145, -0.060]$
0.25	$[-0.120, 0.008]$	$[-0.149, -0.041]$
0.50	$[-0.126, 0.024]$	$[-0.156, -0.020]$
1.00	$[-0.139, 0.056]$	$[-0.185, 0.024]$
Breakdown value $\bar{M}^*$	0.10	0.70

Notes. Following the relative-magnitudes restriction of [Rambachan and Roth \(2023\)](#), we allow the post-treatment violation of parallel trends to be at most $\bar{M}$ times the violation observed in the pre-period, and report the 95 percent confidence set for the effect at $t = 0$. Because the stacked window contains a single pre-period, the worst-case bias is a linear function of the estimated pre-trend coefficient, and we compute the interval from the joint country-clustered covariance matrix of the pre-trend and treatment coefficients, taking the wider of the two sign cases. This is a conservative approximation to the conditional procedure of Rambachan and Roth rather than that procedure itself. The breakdown value is the largest $\bar{M}$ at which zero remains outside the confidence set. The pooled estimate does not survive even a small allowance for differential trends; the competitive-turnover estimate survives an allowance of seven tenths of the observed pre-trend.

5 Results

5.1 Perceived discrimination and access to the presidency

Table 6 reports estimates of equation (2) with the co-ethnicity treatment. The estimate is stable across specifications and outcome codings. Individual controls change the coefficient by less than 0.001, which is expected: the treatment varies at the group level, and the group-by-country fixed effects already absorb the group composition of the sample.

The estimand also depends on how respondents are weighted, and the paper’s baseline is unweighted at the respondent level. Table 5 reports four alternatives. Applying the Afrobarometer sampling weight, rescaled within country-round, leaves the estimate at -0.0382 (s.e. 0.0122). Weighting countries equally gives -0.0335 (s.e. 0.0115). Weighting every group-by-country cell equally gives -0.0180 (s.e. 0.0152), which is not significant. That last estimand answers a different question, the unweighted average across cells rather than across people, and it down-weights precisely the large groups from which presidents are drawn; we report it because it shows the limit of what these data support, not because we think it is the quantity of interest. Readers who want the average across cells rather than across respondents should treat the result as suggestive.

Geography is a further margin. Adding region-by-round fixed effects to the full estimation sample, 1,938 cells, gives -0.0254 (s.e. 0.0123): comparisons made entirely within the same region and the same round retain two thirds of the estimate, which is consistent with the finer cluster-level decomposition in Section 6 and says that a third of the baseline estimate travels through cross-region composition within countries.

Table 20 reports the same specification country by country for each of the 11 countries that contain any switching cell, using ethnic-group and round effects within the country. The estimate is negative in 8 of the eleven and significantly so in South Africa, Ghana, Benin, Namibia and Zambia; Mali, Niger

Table 5: The headline estimate under four weighting schemes

Estimand	Coefficient	Observations
Unweighted, respondent-level (baseline)	−0.0380*** (0.0118)	196,728
Afrobarometer weights, normalised within country-round	−0.0382*** (0.0122)	196,728
Equal weight per country	−0.0335*** (0.0115)	196,728
Equal weight per group-by-country cell	−0.0180 (0.0152)	196,728

Notes. All rows estimate equation (2) with group-by-country and country-by-round fixed effects and standard errors clustered on country. Row 1 is the specification reported throughout the paper. Row 2 applies the Afrobarometer within-country sampling weight, rescaled to mean one within each country-round. Row 3 additionally rescales so that every country contributes equally. Row 4 rescales so that every group-by-country cell contributes equally, which down-weights precisely the large groups from which presidents are drawn. The respondent-level and survey-weighted estimands are the ones the paper interprets; the equal-cell estimand answers a different question, namely the unweighted average across cells, and it is not estimated precisely enough to be informative. Row 2 rescales the Afrobarometer weight to mean one within each country-round; the robustness table’s “Survey weights (unnormalised)” row applies the raw weight with group-clustered errors, which is why it reads −0.0385 there. Observation counts are the estimation samples after 41 singleton cells are dropped. * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.

and Nigeria are positive and insignificant. The dispersion is wide, and the leave-one-country-out range in Table 19 is the relevant summary: no single country’s removal moves the pooled estimate outside $[-0.043, -0.031]$.

5.2 The response ladder

Table 7 contains the paper’s central comparison. Each row reports a separate regression of a different outcome on the same treatment, with the same fixed effects, on essentially the same respondents; Figure 1 plots the coefficients in standard deviations of each outcome.

The pattern is striking. Beliefs about the government move by roughly four-tenths of a standard deviation. The belief about the fairness of the government’s treatment of one’s own group moves by about a tenth of a standard deviation. Measured local public goods and measured household deprivation do not move, and their confidence intervals exclude anything close to the belief response: the public-goods interval is $[-0.10, +0.03]$ SD and the lived-poverty interval $[-0.05, +0.03]$ SD. These contrasts are not merely visual: stacked specifications that estimate two responses jointly in standard deviations, with outcome-specific fixed effects and clustering on the group, reject equality of the grievance and public-goods responses ($p = 0.002$) and of the grievance and approval responses ($p < 0.001$); the grievance and lived-poverty responses differ at marginal significance ($p = 0.058$).

The two material rows do not carry equal weight, and saying so matters more than the stars. The country-clustered confidence intervals carry the claim: they rule out a favourable lived-poverty response larger than 0.050 standard deviations, about half of the 0.098 belief response, and a favourable local-

Table 6: Perceived ethnic discrimination and holding the presidency

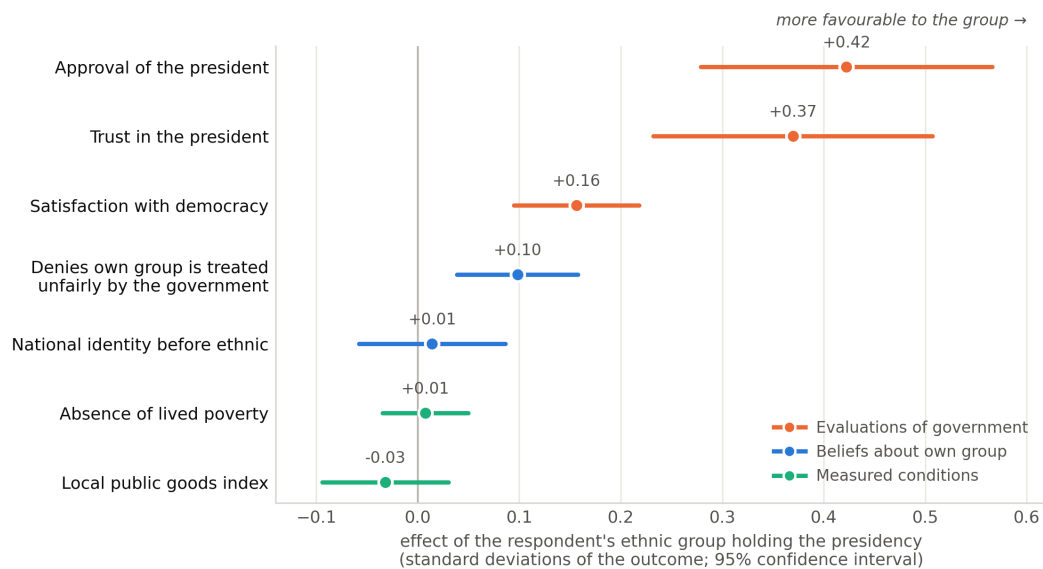
	Unfairly treated, often/always			Ever	0–3 scale
	(1)	(2)	(3)	(4)	(5)
Co-ethnic with head of state	−0.0380*** (0.0118) [0.0093]	−0.0385*** (0.0118) [0.0093]	−0.0311** (0.0149) [0.0093]	−0.0606** (0.0239) [0.0162]	−0.1175*** (0.0398) [0.0290]
Individual controls	No	Yes	No	No	No
Region fixed effects	No	No	Yes	No	No
Group × country FE	Yes	Yes	Yes	Yes	Yes
Country × round FE	Yes	Yes	Yes	Yes	Yes
Mean of dependent variable	0.183	0.183	0.183	0.426	0.687
Observations	196,728	193,748	196,727	196,728	196,728

Afrobarometer Rounds 3–9. The treatment is an indicator for the respondent’s ethnic group being the group of the head of state in office on the interview date. Individual controls are age, age squared, sex, education and an urban indicator. Standard errors clustered on country in parentheses (the primary inference, since transitions occur at the country level); standard errors clustered on ethnic-group-by-country in brackets. Stars follow the country-clustered errors. Group-by-country cells containing a single respondent (41 observations) are dropped in estimation. * $p < 0.1$, ** $p < 0.05$, *** $p < 0.01$.

Table 7: The response ladder: what moves when a group takes the presidency

Outcome	Coef.	(s.e.)	[s.e.]	In SD	95% CI in SD	N
Approval of the president	+0.4236***	(0.0735)	[0.0578]	+0.422	[+0.279, +0.566]	187,441
Trust in the president	+0.4263***	(0.0810)	[0.0569]	+0.370	[+0.232, +0.507]	161,046
Satisfaction with democracy	+0.1558***	(0.0314)	[0.0244]	+0.156	[+0.095, +0.218]	183,338
Group treated unfairly, often/always	−0.0380***	(0.0118)	[0.0093]	−0.098	[−0.158, −0.038]	196,728
Group treated unfairly, ever	−0.0606**	(0.0239)	[0.0162]	−0.123	[−0.217, −0.028]	196,728
Ethnic identity before national	−0.0049	(0.0128)	[0.0083]	−0.014	[−0.086, +0.058]	195,189
Local public goods index	−0.8375	(0.8314)	[0.8823]	−0.032	[−0.094, +0.030]	196,716
Lived poverty index	−0.0069	(0.0199)	[0.0197]	−0.008	[−0.050, +0.035]	126,207

Each row is a separate estimate of equation (2) with ethnic-group-by-country and country-by-round fixed effects. Standard errors clustered on country in parentheses (the primary inference, since transitions occur at the country level); standard errors clustered on ethnic-group-by-country in brackets. Stars follow the country-clustered errors. “In SD” divides the coefficient by the standard deviation of the outcome in the estimation sample; the interval uses the country-clustered errors. * $p < 0.1$, ** $p < 0.05$, *** $p < 0.01$.



Afrobarometer Rounds 3-9, 28 countries, 196,728 respondents in the estimation sample. Each row is a separate regression of the outcome on an indicator for the respondent's ethnic group being the head of state's group, with ethnic-group-by-country and country-by-round fixed effects; 95% intervals from standard errors clustered on country. Signs are oriented so that positive means more favourable to the respondent's group.

Figure 1: The response ladder. Effect of the respondent's ethnic group holding the presidency on seven outcomes, each in standard deviations of that outcome, with 95% confidence intervals. Signs are oriented so that positive means more favourable to the respondent's group. Colour marks the three families of outcome: evaluations of the government, beliefs about the respondent's own group, and conditions the enumerator or the household reports rather than evaluates. The rows are ordered and grouped by family, and each point carries its own value, so the grouping does not rest on colour alone.

public-goods response larger than 0.030 standard deviations, about a third of it. The minimum detectable effects at conventional power, 0.060 standard deviations for lived poverty and 0.094 for public goods, are ex ante power calculations and are reported as supplementary information rather than as bounds on the true effects. We state the result in these terms rather than as an absence.

Two further qualifications belong here rather than in a footnote. The amenity index describes the respondent’s enumeration area, not the territory the ethnic group occupies, and the median group has only 44 percent of its respondents in its own largest region, so the measure is mis-targeted for a group-level treatment. And it is contemporaneous: the favouritism literature we cite in Section 2 estimates allocation over horizons of five to twenty years, and [Franck and Rainer \(2012\)](#) across birth cohorts. A null within a survey round bounds how fast the material response arrives, not whether it arrives. Our claim is about the wedge at this horizon, and we do not extend it.

Table 8 shows where in the response distribution the movement occurs. The share answering “never” rises by 6.1 points and each of the three positive categories falls by about two, so the effect is a shift across the whole scale rather than a compression at one end, which is what a change in scale use at the top would produce.

Table 8: Where in the response distribution the effect appears

	Effect	(s.e.)	[s.e.]	Sample share
Never	+0.0606**	(0.0239)	[0.0162]	57.4%
Sometimes	−0.0226	(0.0144)	[0.0095]	24.3%
Often	−0.0189**	(0.0072)	[0.0055]	10.4%
Always	−0.0190***	(0.0054)	[0.0057]	7.8%

Each row is a linear probability model for one response category of the unfair-treatment item, with the baseline fixed effects. Standard errors clustered on country in parentheses (the primary inference, since transitions occur at the country level); standard errors clustered on ethnic-group-by-country in brackets. Stars follow the country-clustered errors. * $p < 0.1$, ** $p < 0.05$, *** $p < 0.01$.

Table 9 re-estimates the ladder on the intersection of respondents for whom every outcome is observed, which we report in full. The ordering is preserved, but the common sample is 85,091 respondents, less than half the estimation sample, and its minimum detectable effects exceed the belief response for every row. It can rank the outcomes; it cannot separate them. We therefore keep the maximal-sample estimates as the headline and report the common sample as the more conservative and much less informative check.

Because the paper examines several outcomes, Table 10 adjusts across the family of eight outcomes estimable on the common design. The headline result survives Holm’s step-down procedure at 5 percent (adjusted $p = 0.017$) and the Benjamini–Hochberg procedure at 1 percent (adjusted $p = 0.007$). The two material outcomes remain far from significance under either adjustment.

It is tempting to summarise this as a ratio, the standardised belief response over the standardised allocation response, and to call that ratio a grievance multiplier. We do not, for the following reason. The denominator is a small number whose confidence interval contains zero, so the ratio is not identified in

Table 9: The response ladder on the common sample, with minimum detectable effects

	Effect in SD	95% CI in SD	MDE	N
Approval of the president	+0.485***	[+0.189, +0.781]	0.422	85,080
Trust in the president	+0.365***	[+0.149, +0.581]	0.308	85,080
Satisfaction with democracy	+0.201**	[+0.023, +0.380]	0.255	85,080
Group treated unfairly, often/always	−0.086	[−0.261, +0.089]	0.250	85,080
Group treated unfairly, ever	−0.127	[−0.324, +0.070]	0.281	85,080
Ethnic identity before national	−0.052	[−0.127, +0.023]	0.107	85,080
Local public goods index	−0.031	[−0.129, +0.067]	0.140	85,080
Lived poverty index	+0.013	[−0.082, +0.108]	0.136	85,080

Estimated on the intersection of respondents for whom every outcome is observed. MDE is the minimum effect detectable at 80% power and 5% size, $2.80 \times$ the standard error, in standard deviations of the outcome. The common sample is less than half the size of the estimation sample and its minimum detectable effects exceed the belief response for every outcome, so it cannot separate the rows; we report it because a common sample is the right check and this is what it shows. * $p < 0.1$, ** $p < 0.05$, *** $p < 0.01$.

Table 10: Multiplicity adjustment across the outcome family

Outcome	Coefficient	s.e.	p	Holm	Benjamini–Hochberg
Presidential approval	0.4174	0.0743	< 0.001	< 0.001	< 0.001
Trust in the president	0.4218	0.0814	< 0.001	< 0.001	< 0.001
Satisfaction with democracy	0.1606	0.0309	< 0.001	< 0.001	< 0.001
Frequent unfair treatment (state)	−0.0380	0.0118	0.003	0.017	0.007
Unfair treatment, 0–3 scale	−0.1175	0.0398	0.006	0.026	0.010
Any unfair treatment (state)	−0.0606	0.0239	0.018	0.053	0.023
Local public goods index	−0.8314	0.7678	0.289	0.577	0.330
Lived poverty	−0.0081	0.0191	0.673	0.673	0.673

Notes. Every row is equation (2) with the same fixed effects and country-clustered standard errors, estimated on its maximal sample. Adjustments treat the eight outcomes as one family. The headline result survives Holm’s step-down procedure at the 5 percent level and the Benjamini–Hochberg procedure at the 1 percent level. The unfair-treatment-by-other-citizens item is excluded because it is asked in a single round, so country-by-round effects absorb the treatment; it is reported separately in Table 15 on a matched sample.

any useful sense: its point estimate is about three against public goods and about twelve against lived poverty, and neither is stable to reasonable changes in the sample. Every claim we make is therefore a claim about bounds on the two responses separately, as stated above and in the decomposition of Table 11, and not about their quotient.

5.3 A descriptive comparison with presidential approval

The obvious alternative reading is that the fairness item is an alternative measure of government approval, and that co-ethnics of a sitting president are simply more pro-government about everything. Table 11 describes how far the data go on this. Conditioning on the public-goods index absorbs essentially none of the coefficient; conditioning on presidential approval reduces it by about 55 percent, and adding satisfaction with democracy adds little. Because approval is itself affected by presidential co-ethnicity, this exercise does not identify a component of the response independent of generic favourability; it establishes only that the two responses are not perfectly collinear.

Table 11: Descriptive conditioning on government evaluations

	(1)	(2)	(3)	(4)	(5)
Co-ethnic with head of state	−0.0385*** (0.0124) [0.0096]	−0.0388*** (0.0125) [0.0096]	−0.0173 (0.0113) [0.0091]	−0.0164 (0.0109) [0.0089]	−0.0165 (0.0110) [0.0089]
Local public goods index	No	Yes	No	No	Yes
Approval of the president	No	No	Yes	Yes	Yes
Satisfaction with democracy	No	No	No	Yes	Yes
Share of column (1) explained	—	−0.6%	55.2%	57.5%	57.1%
Observations	175,530	175,530	175,530	175,530	175,530

Dependent variable: group treated unfairly often or always. All columns include ethnic-group-by-country and country-by-round fixed effects and are estimated on the common sample for which every control is observed. Approval and satisfaction are themselves affected by co-ethnicity, so columns (3) to (5) condition on post-treatment variables: they describe how much of the coefficient is accounted for, not a causal decomposition. Standard errors clustered on country in parentheses (the primary inference, since transitions occur at the country level); standard errors clustered on ethnic-group-by-country in brackets. Stars follow the country-clustered errors. * $p < 0.1$, ** $p < 0.05$, *** $p < 0.01$.

5.4 Cabinet representation

Table 13 repeats the exercise with the continuous cabinet-share treatment on Rounds 3–8. Because this treatment varies for every group in every month, it does not rely on the small set of presidential transitions. It is not, however, cleaner variation. Cabinet composition is chosen by the executive and may itself respond to ethnic grievance; it is correlated with presidential ethnicity, measured on a coarser ethnic classification, unavailable in Round 9, and observed only for the respondents whose groups can be mapped to ACPED. We therefore read the cabinet estimates as a complementary association rather than as independent causal

validation, and Table 12 subjects that association to the controls it would need to survive to be more than that.

It does not survive them. With country-clustered errors the baseline association is -0.0080 (s.e. 0.0049). Controlling for presidential co-ethnicity leaves the cabinet coefficient at -0.0072 (s.e. 0.0069) while co-ethnicity itself is unchanged at -0.0382 (s.e. 0.0139). Excluding the president's own group, so that the identifying variation is cabinet representation short of the presidency, gives -0.0035 (s.e. 0.0035), a precise near-zero that rules out effects larger than about -0.010 at conventional levels. Cabinet share relative to the group's population share is mechanically identical to the baseline here, because the population share is absorbed by the group fixed effect. The association in Table 13 is therefore carried by the president's own group, and cabinet representation short of the presidency has no detectable effect on the fairness belief.

Table 12: The cabinet association under the controls it should survive

Specification	Cabinet share (10 pp)	Presidential co-ethnicity	<i>N</i>
1. Baseline	-0.0080 (0.0049)		166,889
2. Controlling for presidential co-ethnicity	-0.0072 (0.0069)	-0.0382^{***} (0.0139)	142,736
3. Excluding the president's group	-0.0035 (0.0035)		135,725
4. Current share, holding the 12-month lag fixed	-0.0055 (0.0049)		166,889

Notes. Outcome is the frequent-unfair-treatment indicator; all rows use ACPED-group and country-by-round fixed effects with standard errors clustered on country, which is why the baseline row is weaker than the group-clustered version in Table 13. Row 2 is estimated on the subsample where presidential co-ethnicity is defined. Cabinet share relative to the group's population share gives coefficients identical to rows 1 and 3, because the population share is time-invariant within a group and is absorbed by the group fixed effect; it is therefore not a separate check and is not reported as one. In row 4 the coefficient on the 12-month lag, holding the current share fixed, is -0.0047 (s.e. 0.0040). * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.

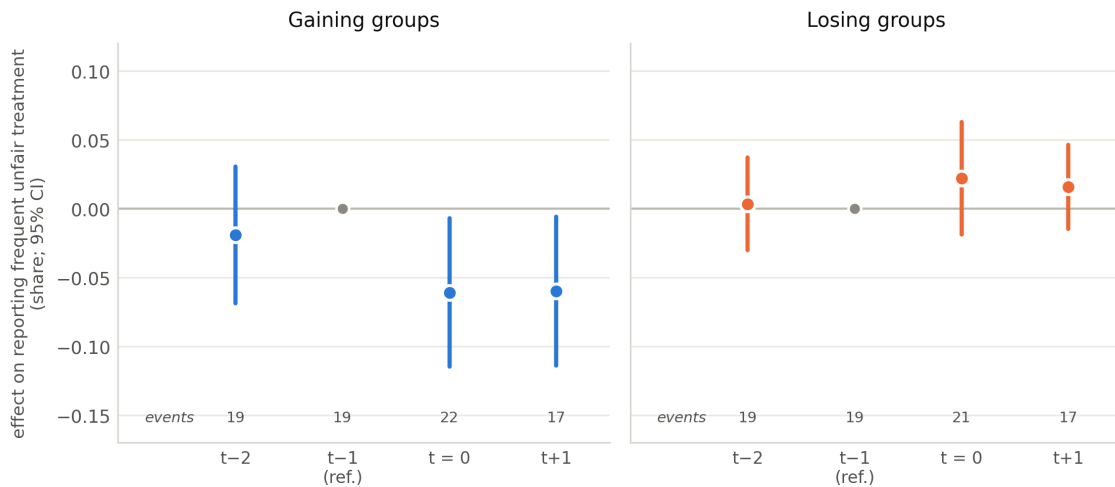
The reading we take from Tables 7, 13 and 12 together is narrower than the one an earlier draft took. The presidency is a visible, symbolic prize: winning it moves approval, trust, satisfaction with democracy and the fairness belief together. Cabinet seats short of the presidency move nothing we can detect. Combined with the transition-mode result of Table 3, in which successions that deliver the presidency without an electoral victory also move nothing, the evidence points consistently at one treatment: a group's candidate winning the presidency in a competitive election. Representation in the executive short of that, whether a cabinet seat or the office acquired by succession, does not detectably move the belief. The stacked test's failure to reject equality of the standardised grievance and approval responses under the cabinet treatment ($p = 0.23$) is consistent with this, and the framework's prediction that cabinet composition should move the fairness belief more than approval is not supported.

Table 13: Cabinet representation and perceived discrimination, Rounds 3–8

	Unfair, often (1)	Unfair, ever (2)	Unfair, 0–3 (3)	Approve president (4)	Public goods (5)	Lived poverty (6)
Cabinet share (10 pp)	−0.0080* (0.0041)	−0.0135** (0.0057)	−0.0281** (0.0118)	0.0493* (0.0293)	−0.2316 (0.4241)	0.0047 (0.0085)
Effect in SD	−0.021	−0.028	−0.030	+0.049	−0.009	+0.005
Mean of dependent variable	0.179	0.412	0.668	2.684	54.711	1.320
Observations	166,889	166,889	166,889	156,756	166,857	96,917

Treatment is the respondent's ethnic group's share of cabinet posts in the month of the interview, in ten-percentage-point units, from ACPED v2. Group-by-country and country-by-round fixed effects; standard errors clustered on group-by-country. * $p < 0.1$, ** $p < 0.05$, *** $p < 0.01$.

5.5 Dynamics and asymmetry



Stacked design; each of the 43 transitions in which a group-by-country cell changes access to the presidency forms its own stack; controls are the cells in the same country that never change status over Rounds 3–9, with event-by-group and event-by-round fixed effects, so no comparison crosses stacks. Event time in survey rounds, $t-1$ omitted. 95% intervals from standard errors clustered on country. Counts beneath each point are the transitions contributing at that event time.

Figure 2: Stacked event study around presidential transitions. Each transition forms its own stack; the controls are the cells in the same country that never change status, with event-by-group and event-by-round fixed effects. Event time is measured in survey rounds, with the round before the transition omitted; intervals use standard errors clustered on country. Counts beneath the points are the transitions contributing at each event time.

Figure 2 plots the stacked estimates of Table 2. Gaining the presidency produces a decline that is present in the first survey round in which the group holds the office (−6.1 pp, s.e. 2.8) and essentially unchanged one round later (−6.0 pp, s.e. 2.8), with a pre-period coefficient a third that size and indistinguishable from zero. Losing the presidency produces a small positive movement in the transition round (+2.2 pp, s.e. 2.1) that is not significant, with a flat pre-period. Because event time is measured in survey rounds whose spacing and distance from the actual transition vary, these are statements about the first

observed round, not about calendar speed. The pooled event study of the earlier draft is retained as a diagnostic in Appendix Figure A1; its pooled lead coefficient (-4.9 pp, s.e. 4.5) reflects the pooling of groups at different points in their treatment histories rather than a pre-trend, which is what the stacked design corrects. Any asymmetry between gaining and losing is a comparison between coefficients that are individually imprecise, and we have not tested their differences; the dynamics beyond the first observed round are the least secure result in the paper.

Table 14 reports the cabinet-share version of the asymmetry, splitting the twelve-month change in a group's cabinet share into losses and gains. On the full panel neither side is distinguishable from zero (loss $+0.00$ pp, gain -0.30 pp per ten points); restricted to Rounds 5–8 the loss side is $+1.58$ pp against a gain side of -0.01 . The event-study asymmetry in Table 2 is the version that survives the longer panel; the cabinet version does not, and we no longer lean on it.

5.6 Electoral timing

Using the exact interview date and a register of 150 national elections held between 2003 and 2023, we compare respondents interviewed within six months of an election to those interviewed further away, within group-by-country, country and round fixed effects. On Rounds 5–9 alone, respondents interviewed in the six months *after* an election are 2.8 pp more likely to say their group is treated unfairly often or always ($p = 0.008$), with nothing before. Adding Rounds 3–4 attenuates the post-election coefficient to 1.2 pp ($p = 0.45$) and the asymmetry between before and after disappears. We report the full-panel estimates in Table 14 and flag the post-electoral pattern as specific to the 2011–2023 rounds: possibly real, since electoral disputes have sharpened over the period, but not a robust feature of the two decades pooled.

Table 14: Electoral timing and the asymmetry of the response

	Coef.	(s.e.)	<i>N</i>
Interviewed within 6 months <i>before</i> an election	+0.0143	(0.0105)	253,815
Interviewed within 6 months <i>after</i> an election	+0.0122	(0.0160)	253,815
Ten-point <i>loss</i> of cabinet share over 12 months	+0.0002	(0.0120)	166,889
Ten-point <i>gain</i> of cabinet share over 12 months	−0.0030	(0.0055)	166,889

Dependent variable: group treated unfairly often or always. Upper panel: group-by-country, country and round fixed effects; the omitted category is being interviewed more than six months from any national election. Lower panel: group-by-country and country-by-round fixed effects, with the twelve-month change in cabinet share split into its negative and positive parts. Standard errors clustered on group-by-country. * $p < 0.1$, ** $p < 0.05$, *** $p < 0.01$.

Where the pattern does appear it is the mirror image of the canonical result: Eifert et al. (2010) show that ethnic *identification* rises as a competitive election approaches, whereas grievance, when it moves, moves after the outcome is known. The contrast between the two items, one mobilised in anticipation and one a verdict after the fact, is itself the empirical case for taking the grievance item seriously as an object in its own right, but on the pooled two decades the post-electoral movement is not reliably there.

6 Identification concerns

Table 15 collects the tests that could overturn the result.

Table 15: Threats, placebos and corrections

	Coef.	(s.e.)	<i>N</i>
Probability of “don’t know” / refusal	+0.0049	(0.0046)	205,361
Group’s share of the country-round sample (%)	+0.5946*	(0.3259)	2,282
Placebo: <i>religious</i> group treated unfairly	−0.0267	(0.0209)	26,393
Unfair treatment by other citizens, not the state	−0.0340**	(0.0171)	55,594
Effect of a <i>co-ethnic interviewer</i> on the item	+0.0214***	(0.0080)	98,961
Co-ethnicity, interviewer-ethnicity subsample	−0.0569***	(0.0145)	98,961
controlling for interviewer co-ethnicity	−0.0571***	(0.0143)	98,961
with interviewer fixed effects	−0.0179**	(0.0087)	98,928

Rows 1–4 use the co-ethnicity treatment with ethnic-group-by-country and country-by-round fixed effects; row 2 is estimated at the group-by-country-by-round level. The interviewer rows use Rounds 6, 7 and 9, the rounds in which Afrobarometer records the enumerator’s ethnic group. Standard errors clustered on ethnic-group-by-country. * $p < 0.1$, ** $p < 0.05$, *** $p < 0.01$.

Endogenous group membership. Müller-Crepon (2025) shows that ethnic self-identification in Africa responds to administrative and political boundaries, and Green (2021) shows the specific movement that threatens us: under autocracy, citizens switch their stated identity *toward* the new president’s group. If people re-identify with a winning group, our estimate would be contaminated by composition rather than belief change. Regressing a group’s share of its own country-round sample on its co-ethnic status, with the same fixed effects, gives +0.63 pp ($p = 0.07$): the sign Green predicts, and marginally significant on the seven-round panel. What that implies depends on the denominator, and the relevant denominator is not the national mean group share but the share in the cells that actually switch, which is 15.6 percent. An increase of +0.63 points then makes entrants 3.9 percent of the post-transition group, so even under the extreme assumption that every entrant reports no unfair treatment at all, composition accounts for 0.71 points of the 3.8-point estimate. That is a real contribution rather than a negligible one, and it is a bound rather than an estimate. Green also finds switching only in autocracies and in the president’s home region; our identifying transitions are concentrated in electoral regimes (Nigeria, Ghana, Senegal, South Africa, Malawi, Zambia).

Differential non-response. If co-ethnics of a sitting president became more reluctant to answer a sensitive item, the estimated effect would be a selection artefact. The probability of a “don’t know” or refusal moves by +0.49 pp (s.e. 0.46, $p = 0.28$). We read this as no detectable selection rather than as a clean zero: the interval admits a shift of up to 1.4 points, which is a third of the headline estimate.

Generic souring on government. Two parallel items test whether the effect is a general shift in how co-ethnics answer sensitive questions rather than a specific belief about the state. Both are asked in only one or two rounds, where the group-by-country and country-by-round effects are collinear, so both are estimated with the coarser ACPED-group effects; the comparison is only meaningful against the grievance item estimated the same way on the same respondents, which Table 16 reports alongside.

Table 16: Placebo items, against the grievance item on the same sample and specification

	Placebo item		Grievance item		N
	Effect	s.e.	Effect	s.e.	
Unfair treatment by other citizens	−0.0340**	(0.0171)	−0.0565**	(0.0283)	55,594
Unfair treatment of the religious group	−0.0267	(0.0209)	−0.0819**	(0.0340)	26,393

Both placebo items are asked in only one or two rounds, where group-by-country and country-by-round effects are collinear, so both are estimated with ACPED-group and country-by-round effects. The grievance columns apply the same specification to the same respondents, which is the only meaningful comparison. Standard errors clustered on ethnic-group-by-country. * $p < 0.1$, ** $p < 0.05$, *** $p < 0.01$.

The religious-group placebo gives −0.0267 against a matched grievance estimate of −0.0819, and its confidence interval excludes the matched estimate: the effect is specific to the ethnic referent. The item on unfair treatment by *other citizens* rather than by the state gives −0.0340 against a matched grievance estimate of −0.0565. That is three-fifths of the response to the state item, and we cannot reject equality. This is the paper’s least comfortable result and we do not present it as support. Two readings survive it. A group whose standing rises may genuinely be treated better by neighbours as well as by officials, which is the recognition interpretation of Section 7. Or co-ethnics of a sitting president may answer all questions about their group’s treatment more favourably, which is a response-style shift and would inflate our estimate. Nothing in these data separates them, and a design that could, most likely an anchoring-vignette module, does not yet exist for this item.

Enumerator effects. Adida et al. (2016) show that African survey respondents answer ethnically sensitive and government-evaluation items differently depending on the enumerator’s ethnicity. This item is precisely their failure case, and no published study using it controls for enumerator identity. Afrobarometer records the interviewer’s ethnic group in Rounds 6, 7 and 9; in those rounds 32.0% of respondents were interviewed by a co-ethnic. A co-ethnic interviewer raises the probability of reporting frequent unfair treatment by 2.1 pp ($p = 0.007$), the first estimate of this effect on this item.

Controlling for interviewer co-ethnicity does not change our coefficient. Adding full interviewer fixed effects does: on that subsample it falls from −0.0569 to −0.0179. Because Afrobarometer assigns each interviewer to a small number of enumeration areas, this collapse is consistent with two very different stories. It may be the enumerator, in which case our estimate is contaminated by who asks the question. Or it may be the locality, in which case absorbing the interviewer also absorbs the neighbourhood, and the comparison that remains is the sharper test of our own mechanism: representation is a national object, so co-ethnics and their neighbours in the same locality should differ by the full amount.

The two can be partially distinguished under a proxy construction. Afrobarometer does not release a primary-sampling-unit identifier, but the sample design places eight respondents in each enumeration area, and the enumerator’s amenity observations are recorded once per area, so respondents sharing a country, round, locality name and complete amenity vector are grouped into a reconstructed local cluster. The reconstruction recovers 12,646 clusters with a median of exactly eight respondents, which matches the design.

Table 17 reports the decomposition on a common sample. Region effects, the geography the paper would otherwise control for, change the estimate by two percent. Reconstructed-cluster effects move it to -0.0365 , statistically indistinguishable from our full-sample headline of -0.0380 : comparing a co-ethnic of the president to the neighbour interviewed in the same place reproduces the paper’s estimate. Interviewer effects alone move it to -0.0179 , and absorbing both leaves -0.0246 . The movement across these specifications is consistent with interviewer effects standing in largely for locality at a scale finer than any administrative unit, though the proxy construction does not support a formal variance decomposition and we do not claim one. The defensible range for this subsample is -0.0246 to -0.0365 , and we prefer the reconstructed-cluster comparison because it is the test our own theory implies.

Table 17: Separating the enumerator from the reconstructed local cluster

	Co-ethnic with head of state	(s.e.)	[s.e.]	<i>N</i>
Baseline: group and country-round effects	-0.0569^{***}	(0.0202)	[0.0145]	98,961
adding region effects	-0.0556^{**}	(0.0218)	[0.0142]	98,961
adding local-cluster effects	-0.0365^{**}	(0.0153)	[0.0110]	98,907
adding interviewer effects	-0.0179	(0.0110)	[0.0087]	98,928
adding both	-0.0246^{**}	(0.0118)	[0.0094]	98,874

Estimated on the respondents in Rounds 6, 7 and 9 for whom the interviewer’s ethnic group and a reconstructed local cluster are both observed. Afrobarometer does not release a primary-sampling-unit identifier; respondents are assigned to the same cluster when they share a country, round, locality name and the enumerator’s complete amenity vector, which is recorded once per enumeration area. The reconstruction yields a median of eight respondents per cluster, matching the sample design; Table 18 reports its diagnostics. Standard errors clustered on country in parentheses (the primary inference, since transitions occur at the country level); standard errors clustered on ethnic-group-by-country in brackets. Stars follow the country-clustered errors. * $p < 0.1$, ** $p < 0.05$, *** $p < 0.01$.

Afrobarometer does not publish a primary-sampling-unit identifier in the merged files, so we cannot validate the reconstruction against the true units and we call them local clusters rather than enumeration areas. What we can do is check them against the known sample design, and Table 18 does so. The size distribution matches: 56.2 percent of reconstructed clusters contain exactly the eight respondents the design places in an area, 79.6 percent contain between six and ten, and only 1.4 percent exceed sixteen. Just over half of locality-rounds split into more than one cluster, with a median of two, so the amenity vector is doing work beyond the locality name rather than merely reproducing it. The estimate is stable across the natural sensitivities: -0.0357 (s.e. 0.0157) when clusters larger than sixteen are dropped, and -0.0320 (s.e. 0.0149) under cluster-by-interviewer effects. Panel B also shows why cluster and interviewer effects are separately identified here: only 3.4 percent of clusters are worked by a single

interviewer and 90 percent are worked by more than two.

Table 18: Diagnostics for the reconstructed local clusters

<i>Panel A. Size distribution</i>	
Reconstructed clusters	12,646
Median size	8
Share with exactly eight respondents	56.2%
Share with six to ten respondents	79.6%
Share larger than sixteen (twice the design size)	1.4%
Respondents in clusters larger than sixteen	5.7%
<i>Panel B. Interviewers and localities</i>	
Clusters covered by a single interviewer	3.4%
Clusters with more than two interviewers	90.4%
Locality-rounds split into more than one cluster	54.4%
Median clusters per locality-round	2
Distinct amenity profiles per country-round (median)	107
<i>Panel C. The within-cluster estimate under sensitivities</i>	
Baseline on this subsample	−0.0569 (0.0202)
Adding cluster effects	−0.0365 (0.0153)
Adding cluster effects, dropping clusters above sixteen	−0.0357 (0.0157)
Adding cluster-by-interviewer effects	−0.0320 (0.0149)
Adding region effects instead	−0.0556 (0.0218)

Notes. Clusters are reconstructed by grouping respondents who share a country, round, locality name and the complete vector of enumerator-recorded amenities. Panel C is estimated on the 98,969 respondents for whom a cluster and the outcome are both observed, with country-clustered standard errors. Afrobarometer does not publish a primary-sampling-unit identifier in the merged files, so these are consistency checks against the known sample design rather than a validation against the true units, and we call the resulting units local clusters rather than enumeration areas. Panel B shows that the reconstruction does not simply recover interviewer assignments: 90 percent of clusters are worked by more than two interviewers, which is what allows cluster and interviewer effects to be separated in Panel C.

6.1 Robustness

Because cross-country designs invite the objection that institutions and survey operations differ across countries, Table 20 re-estimates the specification one country at a time, in each of the eleven countries that contain a switching cell, using group and round fixed effects within the country. Section 5.1 discusses the country-level pattern: the estimate is negative in 8 of the eleven, the dispersion is wide, and no single country’s removal moves the pooled estimate outside the leave-one-country-out range. The pan-African estimate is not the artefact of any one country’s survey machinery, and the country-level heterogeneity is shown rather than pooled away.

A related check classifies regimes mechanically: a country counts as electoral if at least one head-of-state change between 2005 and 2023 was not a coup and followed a national election within twelve months. This splits the sample into 18 electoral and 10 non-electoral countries, and all 30 switching cells

Table 19: Robustness of the headline co-ethnicity coefficient

Specification	Coef.	(s.e.)	<i>N</i>
Baseline	−0.0380***	(0.0093)	196,728
Individual controls	−0.0385***	(0.0093)	193,748
Region fixed effects	−0.0311***	(0.0093)	196,727
Religion fixed effects	−0.0385***	(0.0095)	196,726
Survey weights (unnormalised)	−0.0385***	(0.0092)	196,728
Clustered on country	−0.0380***	(0.0118)	196,728
Urban only	−0.0329***	(0.0113)	79,145
Rural only	−0.0402***	(0.0117)	115,776
Rounds 5–8 only	−0.0390***	(0.0107)	163,288
Cells with ≥ 200 respondents	−0.0382***	(0.0097)	165,647
Leave-one-country-out, range	[−0.0430, −0.0311] across 28 drops		

Dependent variable: group treated unfairly often or always. Every row includes ethnic-group-by-country and country-by-round fixed effects. Standard errors clustered on ethnic-group-by-country unless stated. * $p < 0.1$, ** $p < 0.05$, *** $p < 0.01$.

Table 20: The design within single countries

	Co-ethnic with head of state	s.e.	<i>N</i>
Benin	−0.0692***	(0.0235)	7,946
Ghana	−0.1147***	(0.0429)	12,787
Liberia	−0.0346***	(0.0020)	6,676
Malawi	−0.0189	(0.0154)	10,225
Mali	+0.0205	(0.0181)	4,680
Namibia	−0.0371***	(0.0118)	7,920
Niger	+0.0815*	(0.0425)	5,876
Nigeria	+0.0244	(0.0448)	12,426
Senegal	−0.0206	(0.0124)	8,024
South Africa	−0.0905***	(0.0077)	11,130
Zambia	−0.0459**	(0.0216)	7,820

Each row estimates the baseline specification within one country: the grievance indicator on co-ethnicity with the head of state, with ethnic-group and round fixed effects, standard errors clustered on ethnic group. The rows cover all eleven countries that contain a switching cell. * $p < 0.1$, ** $p < 0.05$, *** $p < 0.01$.

lie in the electoral set: the estimate is identified entirely from electoral regimes, with the non-electoral countries (Cameroon, Uganda, Togo, Gabon and Zimbabwe among them) contributing comparison groups only. This matters for the composition threat above, because [Green \(2021\)](#) finds identity switching in autocracies and none in electoral regimes; the variation that identifies our estimate comes precisely from the regimes where his mechanism is absent.

We find no heterogeneity worth reporting, and the way we reach that conclusion is worth stating because the obvious specification gives the opposite answer. The effect does not differ by group size (interaction -0.006 , s.e. 0.016 , for groups above 20% of the national population). Interacting treatment with urban residence while holding the fixed effects common across town and countryside suggests the response is 1.1 points larger in urban areas, and the same specification suggests a 1.4-point education gradient. But that specification constrains the group and country-round effects to be identical in the two subgroups, which is a strong restriction when urban and rural samples differ in composition. Interacting the fixed effects with the splitting variable, which reproduces the subsample estimates exactly, removes both differentials: $+0.0074$ (s.e. 0.0134) for urban residence and $+0.0103$ (s.e. 0.0214) for post-secondary education. Table 21 reports both specifications side by side. The belief response is the same for every subgroup we can distinguish, which is evidence against readings that require the response to be concentrated among the politically attentive.

Table 21: Heterogeneity under two specifications of the fixed effects

	Common fixed effects		Fixed effects interacted	
	Interaction	s.e.	Interaction	s.e.
Urban residence	-0.0112^{**}	(0.0045)	$+0.0074$	(0.0134)
Post-secondary education	-0.0141^{**}	(0.0055)	$+0.0103$	(0.0214)

Each row reports the coefficient on the interaction of co-ethnicity with the splitting variable. The left pair holds the group-by-country and country-by-round effects common across the two subgroups; the right pair interacts both sets with the splitting variable, which reproduces separate subsample estimates. Standard errors clustered on ethnic-group-by-country. $*p < 0.1$, $**p < 0.05$, $***p < 0.01$.

7 Interpretation

Three interpretations are consistent with the evidence, and they carry different implications. We present each, together with the findings that bear on it.

Recognition. Holding executive office is a public statement about a group’s standing. On this reading the fairness item measures status rather than allocation, and its movement is not an error: citizens are correctly reporting a change in their group’s position in the national hierarchy, which is simply not the same thing as a change in their clinic. The finding that unfair treatment by *other citizens* also falls supports this, since the state cannot plausibly have changed how neighbours behave within one survey round.

Information and attribution. Citizens have poor information about who gets what, and use the identity of the executive as a heuristic for it. Under this reading the multiplier is an information failure and would shrink if allocation were more legible. The heterogeneity in Section 6 bears on this reading without settling it. Once the fixed effects are interacted with the splitting variable, the response is no larger among urban or post-secondary respondents, the people for whom the executive's composition is most visible. An information mechanism running through political attentiveness would predict such a gradient, and we do not find one. What we cannot yet test is the cross-country version (a larger multiplier where media are weaker and the state's distributive footprint smaller), which requires country-level media and fiscal-capacity data we have not merged.

Anticipation. Beliefs may be forward-looking: a group that has just taken the presidency expects better treatment, and the survey catches the expectation before the allocation. Our event study is the natural test, but it is too imprecise to distinguish anticipation from persistence. The point estimate is larger one round (roughly two years) after the transition than at the transition itself, which is not the pattern anticipation predicts, but no individual event-time coefficient is significant at conventional levels and their differences are untested.

A fourth possibility is not a reading of the belief at all but a statement about what the treatment contains. Becoming co-ethnic with the president is, in ten of the twelve transitions that drive the result, also becoming aligned with a winning party. The mode split in Table 3 shows that the response appears only in those transitions and not in the successions that change the president's ethnicity without a partisan victory. Distinguishing the two requires pre-treatment party affiliation or prior vote choice, which Afrobarometer records in some rounds and which we have not yet merged. Until that is done, the estimate should be read as the effect of a co-ethnic group winning power, not of ethnic representation holding partisanship fixed.

These readings cannot be fully separated with the data in hand. What the data establish is that the reported belief moves substantially more than the two contemporaneous material measures we observe, that it is not a simple restatement of presidential approval, and that it persists for at least one further round.

One further implication is worth stating as a hypothesis rather than a finding. To the extent that symbolic inclusion improves perceived group standing independently of short-run material allocation, representative appointments may affect legitimacy through a channel distinct from redistribution. We observe perceptions over a single survey interval and observe neither conflict nor allocation at the horizon over which the favouritism literature works, so we cannot evaluate that hypothesis here. What the results do support is narrower and more durable: measured horizontal inequality and perceived horizontal inequality are distinct objects, and the large literature which uses the first as a proxy for the second is making a substantive assumption rather than a measurement convenience.

8 Case study: Cameroon

Cameroon is instructive precisely because it does not fit the pan-African design: Paul Biya has held the presidency since 1982, so there is no transition to exploit. What Cameroon offers instead is the reverse experiment. In the pan-African design, representation changes and we ask whether beliefs follow; in Cameroon, representation is frozen and the political salience of the cleavages changes, twice, within our sample window. If the framework of Section 2 is right, the same signature should appear: beliefs about the state moving along the newly salient cleavage, with measured conditions and identity salience moving little, though imprecisely. This section examines that prediction; the pattern appears, and the pre-period sensitivity documented below qualifies it.

The two episodes are distinct. The first is the Anglophone crisis: protests by lawyers and teachers in the North-West and South-West regions in late 2016 escalated, after a violent state response, into an armed separatist conflict that continues today. The second is the October 2018 presidential election and its aftermath: Maurice Kamto, whose support base is read as Bamiléké, took an official 14.2 percent, claimed victory, and was arrested in January 2019; the “sardinard”/“tontinard” slurs entered mass circulation, and clashes followed at Obala in April 2019 and Sangmelima in October 2019 ([International Crisis Group, 2020](#)). Ethnicised commentary on both episodes framed them as evidence that the state belongs to some groups and not to others.

Cameroon has the second-highest perceived ethnic discrimination of the 36 countries asked the item in Round 9: 71.3% say their group is treated unfairly at least sometimes, behind only Nigeria (77.9%) and against a 36-country mean of 42.0%. The within-country pattern tracks the October 2018 presidential election and its aftermath with unusual precision (Table 22). Round 7 was fielded 7 May to 4 June 2018, four months before the election, and in that round Bamiléké and Beti respondents were statistically indistinguishable. By Round 8 (February–March 2021), after Maurice Kamto’s 14.2% and arrest, the “sardinard”/“tontinard” framing, the Obala clashes of April 2019 and the Sangmelima riots of October 2019 ([International Crisis Group, 2020](#)), the gap was 22 points.

Aggregate Cameroonian grievance on the “often/always” measure did not rise relative to the 23-country balanced panel over the same interval (−0.7 pp, s.e. 1.3). Aggregate grievance does not change detectably, while its distribution across groups does, moving against the opposition’s groups and in favour of the president’s; the aggregate estimate is imprecise, so this is a statement about the distribution, not a demonstration that the total was fixed.

Before the ladder, the choice of pre-period must be put on the table rather than in a note, because it changes the answer. Taking Round 7 alone as the pre-period gives +0.204; Rounds 6 and 7 give +0.163; adding Round 5 gives +0.05 and it is not significant, because the 2013 gap was larger than the post-2018 gap. The series is therefore not a clean opening but a re-opening, and any reader should treat the Cameroon estimate as conditional on where the pre-period starts. We report the Rounds 6–7 version below and state the alternatives here so the choice is visible. The coefficient in Table 22 differs in the fourth decimal from the one in Table 23 because the first conditions on a group indicator and the second on round effects; nothing turns on it.

Table 22: Cameroon: share saying their group is treated unfairly often or always

Round (fieldwork)	Bamiléké	Beti	Others	Gap	(s.e.)	n_B	n_{Bt}
5 (Mar 2013)	0.472	0.109	0.203	+0.363	(0.045)	176	156
6 (Jan 2015)	0.302	0.210	0.172	+0.093	(0.050)	162	143
7 (May 2018)	0.232	0.212	0.242	+0.020	(0.055)	112	118
8 (Feb 2021)	0.365	0.145	0.212	+0.220	(0.044)	200	152
9 (Mar 2022)	0.406	0.178	0.240	+0.228	(0.045)	212	169
<i>Difference-in-differences, Bamiléké vs Beti, pre-period vs Rounds 8–9</i>							
Round 7 pre-period				+0.2037***	(0.0634)	963	
Rounds 6–7 pre-period				+0.1609***	(0.0487)	1,268	

“Gap” is the Bamiléké minus Beti difference in the share answering often or always, with the standard error of the difference in proportions. The difference-in-differences rows use heteroskedasticity-robust standard errors. Round 5 codes the group as “Bameleké” and Rounds 6–9 as “Bamiléké”. * $p < 0.1$, ** $p < 0.05$, *** $p < 0.01$.

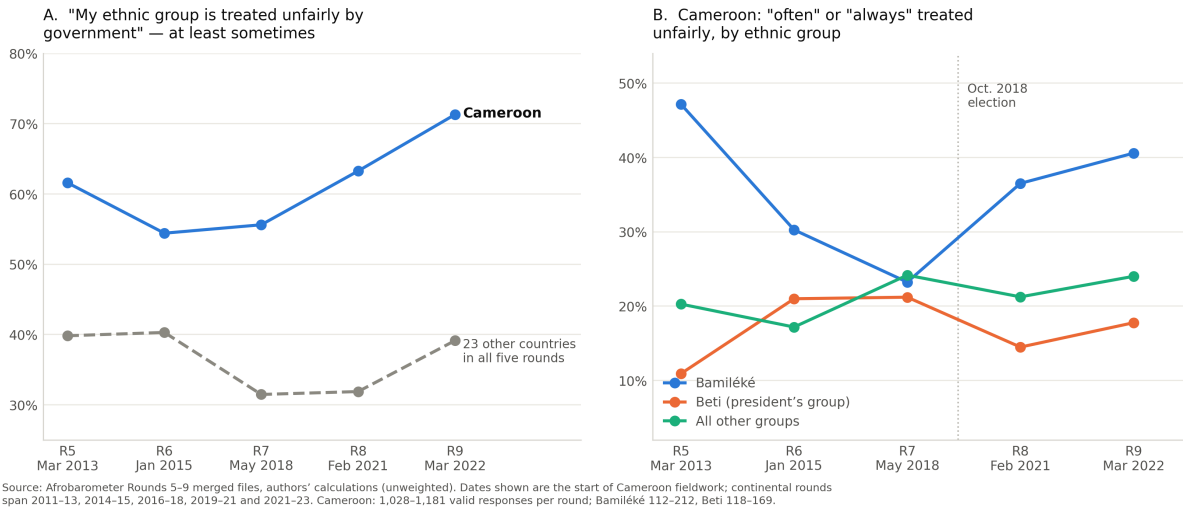


Figure 3: Cameroon. Panel A: perceived ethnic discrimination in Cameroon against the 23 countries surveyed in all of Rounds 5 to 9 (Cameroon is not in Rounds 3 and 4). Panel B: the Bamiléké–Beti gap reopens after October 2018, having also been open in 2013.

Table 23 runs the full outcome ladder through the same difference-in-differences, comparing Bamiléké to Beti respondents before and after the 2018 election, and it reproduces the pan-African pattern within a single country. Beliefs about the state move, and move together: relative to the Beti, Bamiléké reports of frequent unfair treatment rise by 16.3 percentage points (0.37 standard deviations), approval of the president falls by 0.60 standard deviations, trust in the president by 0.36 and satisfaction with democracy by 0.32. What does not move is exactly what does not move in the pan-African design: the enumerator-observed public-goods index (+0.13 SD, $p = 0.24$), lived poverty (+0.21 SD, $p = 0.16$) and, notably, ethnic-versus-national identification (-0.11 SD, $p = 0.31$). The salience of the cleavage transformed the fairness belief without transforming either measured conditions or the identity itself, which is the distinction between grievance and salience that Section 1 argued the literature has conflated.

One caution belongs with the table. Approval falls by more than grievance rises, 0.60 standard deviations against 0.37, so the Cameroon case does not by itself establish that the belief is separable from generic favourability; the pan-African decomposition in Table 11 is what carries that claim, and the case study corroborates the belief response rather than its independence from approval.

Table 23: Cameroon: the Bamiléké–Beti difference-in-differences across outcomes

	Bamiléké $\times$ post-2018	s.e.	In SD	<i>N</i>
Group treated unfairly, often/always	+0.1629***	(0.0486)	+0.367	1,268
Unfair treatment, 0–3 scale	+0.3150***	(0.1016)	+0.342	1,268
Approval of the president	−0.5388***	(0.0974)	−0.603	1,248
Trust in the president	−0.4003***	(0.1146)	−0.364	1,301
Satisfaction with democracy	−0.2749***	(0.0953)	−0.324	1,281
Ethnic identity before national	−0.0323	(0.0320)	−0.109	1,312
Local public goods index	+2.1573	(1.8326)	+0.127	1,327
Lived poverty index	+0.1793	(0.1262)	+0.211	981

Each row is a separate difference-in-differences regression on Bamiléké and Beti respondents in Rounds 6–9: the outcome on a Bamiléké indicator, its interaction with a post-2018 indicator (Rounds 8–9), and round fixed effects, with heteroskedasticity-robust standard errors. Round 5 is excluded because its gap predates the 2018 election; including it attenuates the grievance coefficient to +0.05 ($p = 0.27$), which is the re-opening caution discussed in the text. * $p < 0.1$, ** $p < 0.05$, *** $p < 0.01$.

The Anglophone crisis leaves its mark on a different margin, and the difference is itself informative about what the survey question measures. Because the item asks about the respondent's *ethnic* group while Anglophone identity is linguistic and regional, the crisis appears not in the often/always share but in the extensive margin of the item: among respondents in the North-West and South-West, the share reporting *any* unfair treatment of their group jumps from 57.2% in Round 7 (fielded mid-2018) to 80.1% in Round 8 (2021), against 55.3% and 60.1% among francophone respondents over the same interval, and stays at 81.6% in Round 9. A twenty-point relative rise on the extensive margin, concentrated exactly where the conflict is, with the intensive margin flat, is what one expects when a regional grievance is filtered through an ethnic questionnaire item; we flag it rather than lean on it, since the two Anglophone

regions contribute only about 195 respondents per round and fieldwork there during an armed conflict raises selection concerns we cannot rule out.

Two cautions. The Round 5 gap (2013) was in fact larger than the post-2018 gap, so the pre-trend is not flat across the whole series and the case should be presented as a re-opening rather than as an opening. And with roughly 150 respondents per group per round, Cameroon alone cannot carry an identification claim; it is the illustration, not the evidence.

Cameroon also offers two institutions the pan-African design cannot exploit, which we flag for future work. The *équilibre régional* quota system (operatively, Décision 0015/MINFOPRA/CAB of 20 August 1992, building on decrees of 1975 and 1982) fixes explicit regional shares in every public competitive examination and has not been revised since 1992, so the wedge between quota shares and population shares has widened for three decades, a slow-moving, codified allocation rule against which perceived fairness could be tested directly. And the interior ministry publishes complete registries of gazetted chieftaincies (79 first-degree, 867 second-degree, roughly 13,000 third-degree), which map to Afrobarometer regions and would allow chiefly density to be tested as a mediator.

9 Conclusion

Taken together, our findings show that access to executive power substantially changes the relative distribution of reported ethnic grievance across groups within a country: gaining the presidency lowers the share of a group's members reporting frequent unfair treatment by about a fifth relative to other groups interviewed in the same country and round. Over the same interval, the two material measures we observe, the schools, clinics, water and roads of the respondent's own locality and their household deprivation, move substantially less, within the bounds these data can reject. Perceived group treatment is not an interchangeable proxy for measured ethnic allocation.

Three limitations bound what this draft can claim. First, descriptive ethnic representation cannot be separated from partisan electoral victory: the response is precisely estimated only in competitive turnovers, and pre-treatment party affiliation or prior vote choice is not yet merged. Second, the material outcomes are contemporaneous and imperfectly targeted to ethnic groups: they describe the respondent's locality over a single survey interval, while the favouritism literature estimates allocation over five to twenty years, and the group-level version (night-time luminosity over GeoEPR homelands) requires the restricted Afrobarometer geocodes, which cover Rounds 1 to 6 only and require an application. Third, changes in response-scale use cannot be ruled out: the group fixed effects difference out level differences in scale use but not changes, anchoring vignettes (King et al., 2004) are the known remedy, and no vignette-corrected measure of perceived ethnic discrimination yet exists in the African survey literature. ACPED's June 2021 end date also leaves the cabinet channel unobserved in Round 9, though this matters less now that the cabinet association does not survive the natural controls.

Three implications follow. For measurement, perceived and measured horizontal inequality should not be used interchangeably, and studies that proxy one with the other are making a substantive assumption.

For measurement practice, the enumerator's ethnicity is a first-order determinant of what this item records and should be standard to control for, though controlling for it leaves the representation coefficient essentially unchanged. And for policy, if representation improves perceived group standing without a commensurate short-run change in allocation, then the link between perceived unfairness and actual unfairness is weaker than the literature proxying one with the other assumes. We observe neither conflict nor long-run allocation, so we offer that as a hypothesis for further work rather than as a policy conclusion.

References

- Adida, C.L., Ferree, K.E., Posner, D.N., Robinson, A.L., 2016. Who's asking? Interviewer coethnicity effects in African survey data. *Comparative Political Studies* 49, 1630–1660.
- Ali, M., Fjeldstad, O.H., Jiang, B., Shifa, A.B., 2019. Colonial legacy, state-building and the salience of ethnicity in sub-Saharan Africa. *The Economic Journal* 129, 1048–1081.
- Baldwin, K., Huber, J.D., 2010. Economic versus cultural differences: Forms of ethnic diversity and public goods provision. *American Political Science Review* 104, 644–662.
- Bandyopadhyay, S., Green, E., 2025. Explaining ethno-regional favouritism in sub-Saharan Africa. *World Development* 188, 106886.
- Bobo, L., Hutchings, V.L., 1996. Perceptions of racial group competition: Extending Blumer's theory of group position to a multiracial social context. *American Sociological Review* 61, 951–972.
- Bomprezzi, P., Dreher, A., Fuchs, A., Hailer, T., Kammerlander, A., Kaplan, L., Marchesi, S., Masi, T., Robert, C., Unfried, K., 2024. Political leaders' affiliation database (PLAD), version 7.0. Harvard Dataverse, doi:10.7910/DVN/YUS575.
- Burgess, R., Jedwab, R., Miguel, E., Morjaria, A., Padró i Miquel, G., 2015. The value of democracy: Evidence from road building in Kenya. *American Economic Review* 105, 1817–1851.
- Cederman, L.E., Gleditsch, K.S., Buhaug, H., 2013. *Inequality, Grievances, and Civil War*. Cambridge University Press, New York.
- Cederman, L.E., Weidmann, N.B., Gleditsch, K.S., 2011. Horizontal inequalities and ethnonationalist civil war: A global comparison. *American Political Science Review* 105, 478–495.
- de Chaisemartin, C., D'Haultfœuille, X., 2020. Two-way fixed effects estimators with heterogeneous treatment effects. *American Economic Review* 110, 2964–2996.
- Chlouba, V., 2026. The precolonial origins of African nationalism. *Comparative Political Studies* 59, 231–271.
- Cruces, G., Perez-Truglia, R., Tetaz, M., 2013. Biased perceptions of income distribution and preferences for redistribution: Evidence from a survey experiment. *Journal of Public Economics* 98, 100–112.
- De Luca, G., Hodler, R., Raschky, P.A., Valsecchi, M., 2018. Ethnic favoritism: An axiom of politics? *Journal of Development Economics* 132, 115–129.
- Depetris-Chauvin, E., Durante, R., Campante, F., 2020. Building nations through shared experiences: Evidence from African football. *American Economic Review* 110, 1572–1602.

- Dickens, A., 2018. Ethnolinguistic favoritism in African politics. *American Economic Journal: Applied Economics* 10, 370–402.
- Eifert, B., Miguel, E., Posner, D.N., 2010. Political competition and ethnic identification in Africa. *American Journal of Political Science* 54, 494–510.
- Franck, R., Rainer, I., 2012. Does the leader's ethnicity matter? Ethnic favoritism, education, and health in sub-Saharan Africa. *American Political Science Review* 106, 294–325.
- Gadjanova, E., 2022. Competitive elections, status anxieties, and the relative strength of ethnic versus national identification in Africa. *Political Behavior* 44, 1731–1757.
- Gay, C., 2002. Spirals of trust? the effect of descriptive representation on the relationship between citizens and their government. *American Journal of Political Science* 46, 717–732.
- Gimpelson, V., Treisman, D., 2018. Misperceiving inequality. *Economics and Politics* 30, 27–54.
- Green, E., 2021. The politics of ethnic identity in sub-Saharan Africa. *Comparative Political Studies* 54, 1197–1226.
- Hodler, R., Raschky, P.A., 2014. Regional favoritism. *Quarterly Journal of Economics* 129, 995–1033.
- International Crisis Group, 2020. Easing Cameroon's Ethno-political Tensions, On and Offline. Technical Report. Africa Report N°295.
- Isaksson, A.S., 2015. Corruption along ethnic lines: A study of individual corruption experiences in 17 African countries. *Journal of Development Studies* 51, 80–92.
- King, G., Murray, C.J.L., Salomon, J.A., Tandon, A., 2004. Enhancing the validity and cross-cultural comparability of measurement in survey research. *American Political Science Review* 98, 191–207.
- Kramon, E., Posner, D.N., 2013. Who benefits from distributive politics? How the outcome one studies affects the answer one gets. *Perspectives on Politics* 11, 461–474.
- Letsa, N.W., Wilfahrt, M., 2020. The mechanisms of direct and indirect rule: Colonialism and economic development in Africa. *Quarterly Journal of Political Science* 15, 539–577.
- Major, B., Quinton, W.J., McCoy, S.K., 2002. Antecedents and consequences of attributions to discrimination: Theoretical and empirical advances, in: *Advances in Experimental Social Psychology*. Academic Press. volume 34, pp. 251–330.
- Mamdani, M., 1996. *Citizen and Subject: Contemporary Africa and the Legacy of Late Colonialism*. Princeton University Press.
- McNamee, L., 2019. Indirect colonial rule and the salience of ethnicity. *World Development* 122, 142–156.

- Müller-Crepon, C., 2025. Building tribes: How administrative units shaped ethnic groups in Africa. *American Journal of Political Science* 69, 406–422.
- Nunn, N., Wantchekon, L., 2011. The slave trade and the origins of mistrust in Africa. *American Economic Review* 101, 3221–3252.
- Raleigh, C., Wigmore-Shepherd, D., 2020. Elite coalitions and power balance across African regimes: Introducing the African cabinet and political elite data project (ACPED). *Ethnopolitics* 19, 43–73.
- Rambachan, A., Roth, J., 2023. A more credible approach to parallel trends. *Review of Economic Studies* 90, 2555–2591.
- Robinson, A.L., 2014. National versus ethnic identification in Africa: Modernization, colonial legacy, and the origins of territorial nationalism. *World Politics* 66, 709–746.
- Stewart, F. (Ed.), 2008. *Horizontal Inequalities and Conflict: Understanding Group Violence in Multiethnic Societies*. Palgrave Macmillan, Basingstoke.
- Tiku, E., Sylwester, K., 2024. Differences in ethnic favoritism across countries and over time: The case of Africa. *Journal of Economic Development* 49, 21–31.
- Tuki, D., 2025. You're not like us! Ethnic discrimination and national belonging in Nigeria. *Self and Identity* .
- Villamil, F., 2023. Civilian victimization and ethnic attitudes in Africa. *European Political Science Review* 15, 617–627.
- Vogt, M., Bormann, N.C., Rüegger, S., Cederman, L.E., Hunziker, P., Girardin, L., 2015. Integrating data on ethnicity, geography, and conflict: The ethnic power relations data set family. *Journal of Conflict Resolution* 59, 1327–1342.

A Measurement and coding

A.1 The grievance item across rounds

Table A1: The grievance item across rounds

Round	Variable	Fieldwork window (pooled)	Valid responses
3	q81	Mar 2005 – Mar 2006	21,136
4	Q82	Mar 2008 – Jun 2009	23,893
5	Q85a	Oct 2011 – Jun 2013	41,150
6	Q88a	Mar 2014 – Nov 2015	43,340
7	Q85a	Sep 2016 – Sep 2018	38,636
8	Q82a	Jul 2019 – Jul 2021	41,736
9	Q84b	Oct 2021 – Jul 2023	45,768

Wording is stable across rounds. The companion “experienced discrimination” item is not: it changes referent between Round 7 (discrimination based on ethnicity) and Round 8 (unfair treatment by other people based on ethnicity), and we never pool the two without a wording indicator.

A.2 From 42 countries to 28

Table A2: Country accounting

	Countries	Running total
Asked the grievance item in at least one round, 3–9	42	42
less: ethnicity not asked or not a political cleavage (Sudan, Tunisia, Algeria, Egypt, Morocco, Cabo Verde, São Tomé and Príncipe, Seychelles, Mauritius, Burundi)	–10	32
less: effectively mono-ethnic (Eswatini, Lesotho)	–2	30
less: co-ethnic share uninterpretable (Mauritania)	–1	29
less: no leader spell with a documented group in any sampled round	–1	28
Co-ethnicity estimation sample	28	
of which contribute switching cells	11	
of which classified electoral	18	
Cabinet sample (ACPED coverage, Rounds 3–8)	29	
Leader register, countries with any coded spell	38	

The leader register covers 38 countries because it was built before the drop rules were applied and retains countries that later leave the estimation sample. The estimation sample is 28 countries and 916 ethnic-group-by-country cells; 30 of those cells change co-ethnic status across rounds, and those cells are what identify β .

A.3 PLAD errors corrected

Table A3: Corrections to PLAD v7.0 in the estimation sample

Leader	PLAD v7.0	Corrected to
Blaise Compaoré	Birthplace Ouagadougou, Kadiogo	Ziniaré, Plateau-Central
George Weah	Ethnicity Krahn	Kru
John Magufuli	Ethnicity “Bassari”	not documented; spell dropped
Yahya Jammeh	Born 1942	1965
Michel Kafando	Born 1957	1942
Patrice Talon	Born 1985	1958
Ernest Bai Koroma	Start 17 Sep 2007	17 Nov 2007; ethnicity contested, spell dropped
Nana Akufo-Addo	Ethnicity “Akan”, precision 3	Akyem (Akan); retained as Akan
Rupiah Banda	Ethnicity Chewa; Ngoni second	Chewa (was entered Ngoni)
Kgalema Motlanthe	Ethnicity “Tswana”, precision 1	undocumented; spell dropped
Festus Mogae	Tswana; Kalanga second	contested; spell dropped
Benjamin Mkapa	Ethnicity “Makua”	contested; spell dropped
Amadou Toumani Touré	Fula; Mande second	contested; spell dropped

Errors identified by cross-checking every leader spell in the 38-country sample against the reference record. PLAD’s own `ethnicity_precision` flag runs 1 (explicit source) to 3 (inferred); several sample leaders are precision 3 or missing.

A.4 Coding decisions that materially affect the estimates

- **Angola.** “Mbundu” must exclude “Ovimbundu”; naive substring matching conflates two different groups and would code 32% of the sample as co-ethnic with dos Santos.
- **Mauritania.** Afrobarometer’s “Maure” pools Beydane and Haratine, giving a 90.5% co-ethnic share. The country is dropped.
- **Nigeria.** Hausa and Fulani are separate Afrobarometer labels while the political bloc is Hausa-Fulani. The baseline codes Buhari narrowly as Fulani; a broad Hausa-Fulani coding is a planned robustness check.
- **Zimbabwe.** 2,760 respondents answer the generic “Shona” alongside sub-group options, which ACPED does not carry. We treat these as unmapped in the ACPED specification rather than assigning them to the largest sub-group.
- **Gabon.** Ali Bongo’s Batéké affiliation spans three Afrobarometer labels (Bateke, Mbédè, Obamba); the baseline uses the narrow coding.
- **Zambia.** Michael Sata is coded Bemba, though both his parents were Bisa and Afrobarometer offers Bisa separately. The Bisa are a Bemba-speaking group inside the Bemba bloc and Zambian politics treated Sata as a Bemba president, so Bemba is the category a respondent would have answered with, which is the estimand. The genealogical coding is the alternative.

- **Contested leaders are dropped, not guessed.** Five of the eleven exclusions arise because PLAD, ACPED and [Green \(2021\)](#) disagree: Mogae (Tswana or Kalanga, a plurality or a 7% minority), Mkapa (Makua or Makonde), Touré (Songhai, Fula, or mixed Malinké–Peulh), Barrow and Koroma. Motlanthe is excluded on the different ground that no source attributes any group to him at all. Each exclusion costs a country-round of data and buys the guarantee that no result rests on our choosing between two sourced attributions.
- **Every one of these judgements is tested.** Re-estimating the headline specification under each alternative coding moves the coefficient only within $[-0.036, -0.049]$ around the baseline -0.038 : Sata coded Bisa by parentage rather than Bemba, -0.038 (0.010); Nigeria’s Fulani presidents coded as broad Hausa-Fulani, -0.049 (0.012); Gabon’s Teke broadened to Batéké–Mbédè–Obamba, -0.038 (0.009); Zimbabwe recoded at the Shona sub-group level (Mugabe Zezuru, Mnangagwa Karanga, generic “Shona” answers set to missing), -0.036 (0.009); and Rupiah Banda returned to the folk Ngoni attribution, -0.039 (0.009). No coding choice we made is doing the work.
- **Incommensurable pairs.** The Afrobarometer and ACPED classifications cannot be reconciled for Zimbabwe, Tanzania, Côte d’Ivoire and Morocco. These should be dropped from the cabinet specification rather than patched; we report them included and excluded.

A.5 The within-round event design is not available at scale

We had intended the paper’s cleanest identification to come from a regression discontinuity in interview date: Afrobarometer publishes the exact date of every interview, fieldwork in a country runs three to eight weeks, and political events happen on a single announced day. Respondents interviewed just before and just after such an event are comparable in everything but exposure.

Enumerating every national election, runoff and successful coup between 2003 and 2023 against every country-round fieldwork window, we find exactly one event inside a window: the Burkina Faso coup of 30 September 2022, which fell in the middle of Round 9 fieldwork (20 September to 12 October 2022), with 638 respondents interviewed before and 531 after. Reported grievance was 6.6% before and 8.5% after, a difference of +1.9 pp (s.e. 1.6) that survives a linear date trend but disappears once region fixed effects are added. With 1,169 respondents and twelve days on each side, the design is under-powered and we treat it as inconclusive.

We report this because it is a fact about the data that the field should know. Afrobarometer’s fieldwork calendars deliberately avoid election periods, so the within-round event discontinuity, which reads as an obvious opportunity on paper, is not available at scale even with seven rounds. Anyone planning to build a paper on it should check the event register first; ours is in the replication package.

A.6 The pooled event study

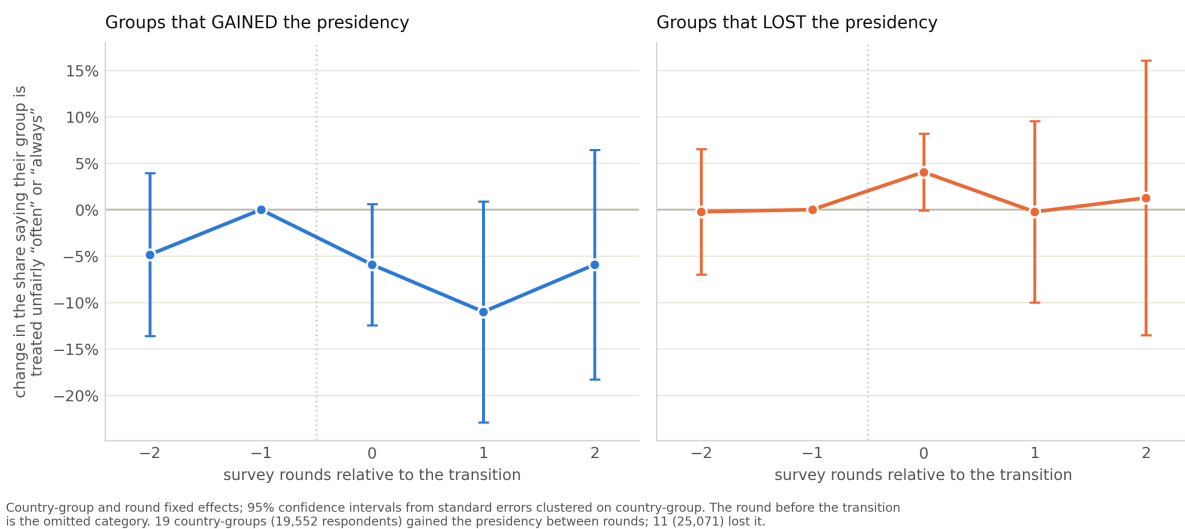


Figure A1: The pooled event study, retained as a diagnostic. Left: the nineteen country-group cells that gained the presidency; right: the eleven that lost it (a cell can transition more than once, so these counts differ from the 43 stacked events). Country-group and round fixed effects; the round before the transition is the omitted category. The stacked design of Figure 2 and Table 2 is the primary evidence: this pooled version compares groups at different points in their treatment histories, which is why its lead coefficient is unreliable.