



Universiteit
Leiden
The Netherlands

Bets beat polls: Averaged predictions of election outcomes

Hofstee, W.K.B.; Schaapman, H.

Citation

Hofstee, W. K. B., & Schaapman, H. (1990). Bets beat polls: Averaged predictions of election outcomes. *Acta Politica*, 25: 1990(3), 257-270. Retrieved from <https://hdl.handle.net/1887/3449799>

Version: Publisher's Version

License: [Leiden University Non-exclusive license](#)

Downloaded from: <https://hdl.handle.net/1887/3449799>

Note: To cite this publication please use the final published version (if applicable).

S2g3
SSW

Acta Politica tijdschrift voor politicologie

Kernredactie

Prof. dr. R.B. Andeweg, dr. C. van der Eijk, drs. W. Hout, drs. K. Koch (secretaris), dr. Percy B. Lehnig, dr. R.H. Lieshout.

Redactieraad

Jhr. dr. G. van Benthem van den Berg, drs. H.W. Blom, prof. dr. H.R. van Gunsteren, prof. dr. A. Hoogerwerf, dr. C.P. Middendorp, prof. dr. J. van Putten, prof. dr. U. Rosenthal, drs. O. Schmidt, prof. mr. dr. I.Th.M. Snellen, prof. dr. J.J.A. Thomassen.

Bijdragen

en andere mededelingen voor de redactie zende men aan de redactiesecretaris, drs. K. Koch, p.a. Vakgroep Politieke Wetenschappen, Rijksuniversiteit Leiden, Postbus 9555, 2300 RB Leiden; de kopij dient persklaar te worden ingezonden, in zesvoud en in machineschrift.

Boeken ter recensie

en boekbesprekingen zende men aan de boekenredacteur, dr. R.H. Lieshout, Universiteit Twente, Postbus 217, 7500 AE Enschede; terugzending van ongevraagde recensie-exemplaren kan niet plaatsvinden.

Administratie

Voor abonnementen en advertenties richt men zich tot de uitgever, Boompers boeken- en tijdschriftenuitgeverij bv, Postbus 1058, 7940 KB Meppel (05220-54306); postrekening 2756509.

Acta Politica verschijnt driemaandelijks; de abonnementsprijs bedraagt f 96,75, voor instellingen f 148,-, voor studenten f 78,75 (maximaal gedurende vier achtereenvolgende jaren). Abonnementen kunnen wel tussentijds ingaan maar niet tussentijds beëindigd worden.

Opzeggingen ten minste één maand voor het einde van de jaargang. Voor het buitenland gelden aangepaste tarieven die op aanvraag te verkrijgen zijn.

Advertentietarieven: 1/1 pagina f 440,-; 1/2 pagina f 260,-.

Indien men het abonnementsgeld gireert voor een ander gelieve men op het girostrookje naam en adres van de abonnee te vermelden, ter vermindering van dubbele incasso.

© Boom, Meppel en Amsterdam



AP 1990/3

Bets Beat Polls: Averaged Predictions of Election Outcomes*

W.K.B. Hofstee and H. Schaapman

'...when I was in high school I was in the R.O.T.C. and we had to read the *Manual of Arms* and in this fat book there was a little bit about artillery. Now, remember this was in 1936, long before radar and all the homing-in devices. In fact, the book was probably written for World War I, although it might have been compiled some time later, I'm not sure. Anyway, the way they figured how to lob an artillery shell was to take a consensus. The Captain would ask,

- O.K., Larry, how far away do you think the enemy is?

- 625 yards, sir.

- Mike?

- 400 yards, sir.

- Barney?

- 100 yards, sir.

- Slim?

- 800 yards, sir.

- Bill?

- 300 yards.

Then the Captain would add up the yards and divide by the number of men asked. In this case, the answer would be 445 yards. They'd log the shell and generally blow up a large proportion of the enemy.' (Charles Bukowski, *Hollywood*).

Can people predict election outcomes? Do polls add to what was already known? This study contains a comparison between predictions of interested lay persons and professional polls. Its context is methodological: the general issue is the comparison between intuitive judgment and systematic investigation. In this context, the prediction of election outcomes has a special place, because the criterion for the predictions is immediately available; other areas of human judgment, such as clinical diagnosis, personnel selection, peer review, and program evaluation, do not provide such hard criteria.

To the present study, pollsters may object in advance that polls are not meant to be predictive, and that therefore the comparison is irrelevant. The standard question of pollsters to the respondent is: 'If elections were held today, what party if any would you vote for?' To discourage predictive interpretation, many pollsters refuse to translate vote percentages into seats, even though the translation is unequivocal. We realize that polls have functions that are not manifestly predictive; most notably, poll

results may serve as a dependent variable in gauging the political effect of particular events and actions. Firstly, however, even in that context the poll results derive their significance partly from their eventual electoral connotations, at least in parliamentary as opposed to direct-democratic systems (see also Van der Eijk 1988). Secondly, polls that are published in the media in the weeks preceding an election, are almost universally interpreted in a predictive manner. Such polls are at issue here.

Another possible objection is that lay persons would have no other basis for their predictions than published polls, so that again the comparison would make no sense. We deal extensively with this objection and its corollary, namely, that deviations of lay predictions from polls would rest on wishful thinking or strategic manoeuvring, so that their averaged predictions would be biased because it is unlikely that they would constitute a representative sample from the electorate. From an experimental point of view, the ideal situation would be one in which polls were administered but kept secret to everyone, notably, to politicians and journalists, to prevent direct or indirect contamination. Such bans are sometimes advocated and, to some extent, enforced in countries other than ours. The argument is not experimental contamination but band-wagon effects, that is, contamination of the actual election outcome. Empirically, that argument is tenuous (e.g., Irwin and Van Holsteyn 1988). Even if it were not, we have no inclination to plead for the abolishment of a democratic pastime that is very representative of the interplay between the social and behavioral sciences and their public. Therefore, we concentrate on indirect means to face the contamination problem.

Finally, a real problem is that elections are infrequent natural experiments; for a more or less definitive answer to our question, one would therefore have to wait a long time. We attempt to compensate for the lack of data with methodological argumentation. We hope that the data and the arguments will be found interesting enough to inspire replications and extensions of our study.

Study I

Method – Five weeks before the election of the provincial councils and indirectly, the First Chamber of the Dutch parliament on March 18., 1987, the first author put an announcement in the Groningen University Weekly. It presented the actual composition of the First Chamber, and invited the readers to bet upon its composition after the elections. A prize of DFL. 250 was put up for the best prediction. The procedure was repeated one

week before the election, with a separate prize of DFL. 250. The number of participants in the first round was 165, in the second, 76.

Results and Discussion – Table 1 shows the averaged predictions and the old and new compositions of the First Chamber. Three small confessional parties have been grouped together. In this case the translation of poll percentages into seats was difficult because of the indirect nature of the election and, especially, the unpredictable role played by specific provincial parties; no such translations were available and a direct comparison between polls and averaged bets is therefore not possible. An indirect comparison can be made as follows: under ideal conditions, that is, direct elections and random sampling on the day of the election itself, the standard error of a vote proportion p is an estimated $\{p(1-p)/N\}^{1/2}$; for the usual poll sample size in the order of $N=1000$ and a number of seats of 75, this expected error exceeds 1 seat per party for the larger parties. Roughly, a discrepancy (sum of absolute differences between numbers of seats per party) between poll and election outcome in the order of 4 to 5 seats would be expected under these unrealistic ideal circumstances. The averaged predictions, with observed discrepancies of 6 and 5 seats, respectively for the two rounds, would thus in all likelihood have beaten the polls.

Table 1: Predictions of Seat Composition of the First Chamber, 1987

	CDA ¹	Pvda ²	VVD ³	D66 ⁴	SC ⁵	GrL ⁶	Discrepancy
Old Composition	26	17	16	6	4	6	18
First Prediction	26	23	13	6	4	3	6
Second Prediction	26	24	13	6	4	3	5
New Composition	26	26	12	5	3	3	0

¹ Christian-democratic

² Social-democratic

³ Liberal-democratic (right-wing)

⁴ Liberal-democratic (left-wing)

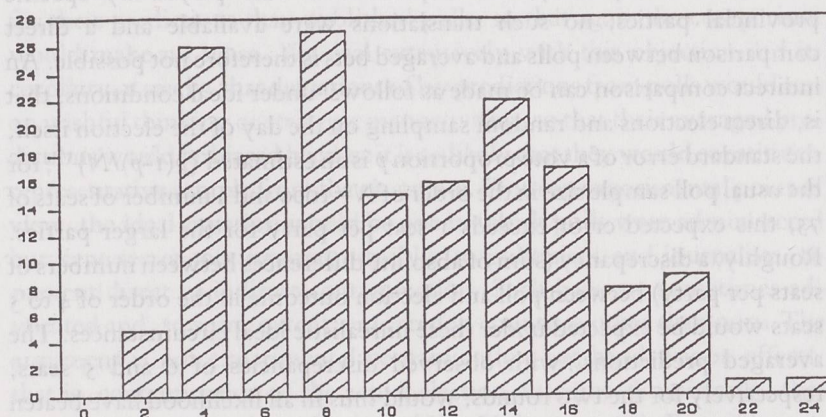
⁵ Small christian parties grouped together

⁶ Green left.

Two further analyses were performed that are relevant for the explanation of the result. The first is outlier analysis: for each participant, the discrepancy between his or her prediction and the group mean prediction was determined. Deletion of participants with high discrepancy scores

did not at all change the rounded (into seats) averaged prediction. The outlier analysis thus testifies to the extreme robustness of the averaged prediction even with these modest numbers of participants. The second analysis focussed on the discrepancies between individual predictions and the election outcome. Figure 1 contains the frequency distribution of

Figure 1: Frequency Distribution of Individual Discrepancies with Election Outcome, First Chamber, 1987, First Round



these individual discrepancies for the first round. Clearly, the average individual does worse than the group average. Only 31 participants have a discrepancy with the election outcome of less than 6, whereas 116 have a discrepancy of more than 6. The average discrepancy is in the order of 10. This superiority of the average over the individual is known in psychometrics as the Spearman-Brown effect (see, e.g., Gulliksen, 1950, p. 104): the more items a test has, or the more raters are employed, the more valid is the prediction, other things being equal. The appropriate formula is: $R_{12} = r_{12} N^{1/2} \{1 + (N-1)r_{11}\}^{-1/2}$, with R_{12} the validity of the averaged prediction, r_{12} the average individual validity, r_{11} the average correlation between two participants, and N the number of participants. Together, these analyses suggest a strictly quantitative interpretation of the success of the averaged bets method.

Study 2

Method – Direct elections for the Second Chamber of the Netherlands Parliament were held on September 5., 1989. We enlisted the cooperation of an advertisement paper, the 'Groninger Gezinsbode', with a circulation

of 135,000 in the Groningen region, which carried our election bet on its editorial front page. The reason for approaching this paper was that its recipients are not subject to political self-selection because every address gets it free. It should be realized, however, that the regional electorate is not at all representative for the Netherlands, that the predictions were collected during the summer vacation, and that many recipients do not read such papers at all. The prizes for the predictions were kept at a modest DFL 100, 50, and 25 for fear of an overwhelming response. The readers were informed about the agenda of the investigators: our presumption that the averaged prediction would be at least as good as the polls' was quoted on the front page. The closing date for submitting the bets was set at August 13.

The number of participants was 393, of which 361 were valid. (Multiple bets and predictions that did not sum up to the appropriate total of 150 seats were discarded). A number of 100 participants were subsequently approached by mail for a second and third round. They received two forms, each containing the actual composition of the Second Chamber, the averaged prediction from the first round, and two empty columns. In the first, they were requested to copy the outcome of a particular poll; in the second, to give their own prediction. The closing dates for these rounds were the days after the polls in question, namely, August 27. and September 3. The response for the second round was 85; the poll to which the third round referred did not take place, but 45 participants responded nonetheless, with reference to another poll.

Results and Discussion – Table 2 gives the old composition of the Second Chamber, the three averaged predictions, the poll outcomes, and the election outcome, in chronological order. The last two columns contain the discrepancy scores with respect to the old composition and the election outcome, respectively. As we expected, the first-round averaged prediction is superior to the early, and even later, polls; only at a late stage the polls equal or surpass the averaged prediction of three weeks before. The before-last column of Table 1 shows that the total shift between old and new composition is systematically overpredicted by the polls, and underpredicted by our participants.

The frequency distribution of the individual discrepancy scores of the first round vis-a-vis the election outcome is given in Figure 2. Again, the Spearman-Brown effect is confirmed: only 22 participants do better than the group average, and 288 do worse, notwithstanding the fact that a number of ludicrous predictions, represented in the right hand tail of the frequency distribution, have gone into the average. At least as striking,

Table 2: Predictions of Seat Composition of the Second Chamber, 1989

	CDA	PvdA	VVD	D66	SC	GrL	CD ¹	Dis.1	Dis.2
Old Compos.	54	52	27	9	5	3	0	0	16
First Pred.	55	51	22	10	5	7	0	12	8
Poll, 13-8 ²	58	49	20	8	7	8	0	22	14
Poll, 27-8 ²	55	51	20	9	6	9	0	16	12
Second Pred.	55	51	22	9	5	8	0	12	10
Poll, 1-9 ³	53	46	22	12	6	11	0	24	10
Poll, 3-9 ²	53	49	22	12	6	8	0	18	4
Third Pred.	54	50	22	10	6	8	0	14	6
Poll, 5-9 ⁴	54	48	22	10	6	10	0	18	8
New Compos.	54	49	22	12	6	6	1	16	0

¹ Ultra-right; for other parties, see Table 1.

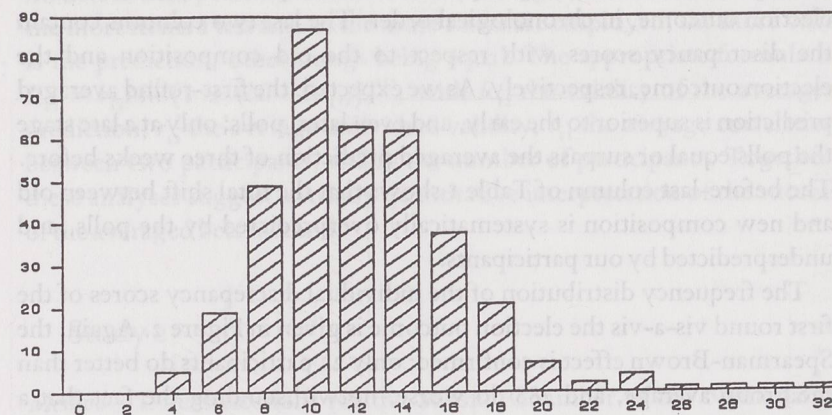
² Agency: Inter View

³ Agency: NSS

⁴ Agency: NIPO

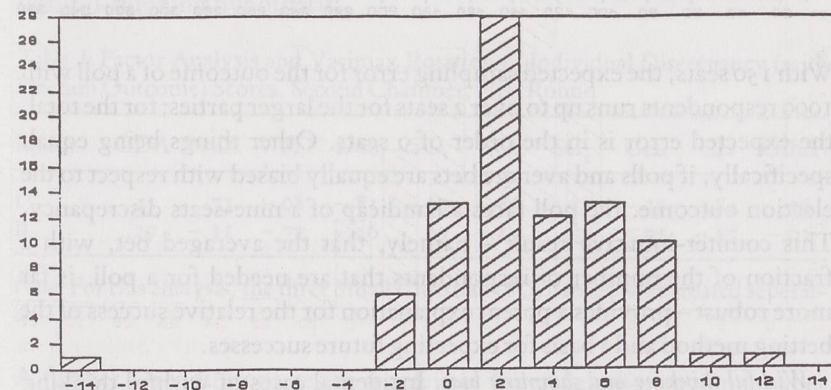
however, are two unexpected findings. One is that not even by chance one participant predicts perfectly. The other is that the modal, median, and mean individual participant, with discrepancy scores of 10, 12, and 12.2, respectively, all do better than the synchronous poll. If this result would be generalizable, the rational thing to do would be to randomly ask an interested lay person rather than spend a large sum on a poll.

Figure 2: Frequency Distribution of Individual Discrepancies with Election Outcome, Second Chamber, 1989, First Round



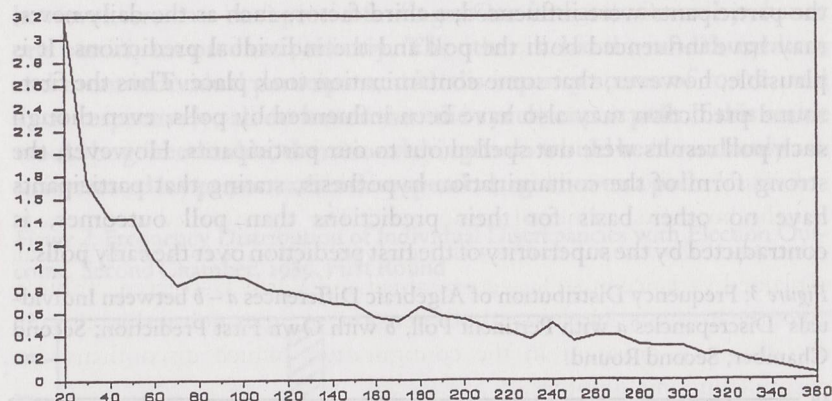
Contamination of predictions by polls. The second and third round were designed to gain more insight into the extent to which participants are influenced by polls. The results are mixed. On the one hand, not one of the 85 respondents to the second round gave the relevant poll results as his or her own prediction. Also, the second-round average has a lower discrepancy (2 seats) with the first-round average than with the poll in question (4 seats); the shift of two seats between the first and second round, however, is in the direction of the pertinent poll. (Ironically, both of these one-seat shifts worsen the second-round prediction). At the individual level, we checked whether a participant's second prediction was closer to the pertinent poll than to his or her own first-round prediction, that is: we took the algebraic difference between the participant's discrepancy between second prediction and pertinent poll, and that participant's discrepancy between his or her second and first prediction. The frequency distribution of these algebraic differences is given in Figure 3. For the large majority of the participants, the difference is positive, meaning that the second prediction looked more like the poll than like that participant's first prediction. That result does not prove that the participants were influenced; a third factor, such as the daily news, may have influenced both the poll and the individual predictions. It is plausible, however, that some contamination took place. Thus the first-round prediction may also have been influenced by polls, even though such poll results were not spelled out to our participants. However, the strong form of the contamination hypothesis, stating that participants have no other basis for their predictions than poll outcomes, is contradicted by the superiority of the first prediction over the early polls.

Figure 3: Frequency Distribution of Algebraic Differences $a - b$ between Individuals' Discrepancies a with Pertinent Poll, b with Own First Prediction; Second Chamber, Second Round.



Robustness of averaged predictions. The outlier analysis described in Study 1 was repeated with identical results: deleting – or, more significantly, adding – of outliers did not change the rounded average at all. To establish the minimum number of participants needed to secure a robust average, random subsamples were drawn from the pool of 361 participants, and discrepancies were calculated between the averaged predictions of the subsamples and the total pool. Figure 4 presents the discrepancies, averaged over four subsamples per subsample size, in function of subsample size. The empirical constant appears to be about 60, that is: adding participants over and above that number does not shift the averaged prediction by one seat. (It should be realized that with rounded predictions the unit of measurement is two seats, given a fixed total number of seats). The number of 60 probably depends on parameters like the total number of seats and the number of parties, and can therefore not be generalized to other situations.

Figure 4: Discrepancy between Subsample Prediction and Total Sample Prediction, as a Function of Subsample Size; Second Chamber, First Round.



With 150 seats, the expected sampling error for the outcome of a poll with 1000 respondents runs up to over 2 seats for the larger parties; for the total, the expected error is in the order of 9 seats. Other things being equal, specifically, if polls and average bets are equally biased with respect to the election outcome, the poll faces a handicap of a nine-seats discrepancy. This counter-intuitive result – namely, that the averaged bet, with a fraction of the number of respondents that are needed for a poll, is far more robust – provides a potent explanation for the relative success of the betting method and a basis for expecting future successes.

Wishful thinking and sampling bias. Incidental cases of wishful thinking

are clearly present in the data. A handful of participants predicted up to ten seats for far-out parties that could not reasonably be expected to enter parliament. As noted above, however, these outlier predictions had no influence at all on the averaged prediction. Clearly, the average prediction itself is slightly biased, as its discrepancy with the election outcome cannot be explained by error. As in the case of polls, the bias can partly be explained by real shifts in electoral mood between the two points in time. Wishful thinking, however, or counter-wishful thinking (in the manner of a side-bet, so that the participants secure a higher chance of winning a prize if their worst fears come true), or strategic manoeuvring, are competing explanations.

A reasonable deduction from the hypothesis is the following: if wishful thinking were a massive phenomenon, some clustering of the individual predictions along a left-right or other politically relevant dimension should be expected. This deduction was tested through principal component analysis.

Raw-score principal component analysis of the participants $\times$ parties matrix yielded a first principal component accounting for 99.7% of the variance. This finding reflects the extremely high correlation between participants over parties due to the different orders of magnitude of these parties. For a microscopic analysis of the remaining variance, the individual predictions were transformed into algebraic deviations from the election outcome (that is, the election outcome was subtracted from each individual prediction). This operation roughly corresponds to partialling out the huge first principal component. The transformed matrix was subjected to principal component analysis. Two factors with a total explained variance of 52% were retained and Varimax-rotated. Table 3 gives the loadings. No left-right dimension, or other politically relevant dimension, is apparent in the configuration, which disconfirms the wishful-thinking hypothesis.

Table 3: Factor Analysis and Varimax Rotation of Individual Discrepancy (with Election Outcome) Scores, Second Chamber, First Round

Factor	CDA	PvdA	VVD	D66	SC ₁	SC ₂	SC ₃	GrL	CD	Other
I	.43	.73	-.08	-.81	-.60	-.86	-.11	-.36	-.84	.36
II	.27	-.11	-.77	-.26	.27	-.11	.58	.71	-.17	-.03

Note: For this analysis, the three Small Christian (SC) parties were treated separately.

A further serendipitous result is relevant here. Another regional paper

covering the North-Eastern part of the province of Friesland copied our procedure, and produced an average prediction that was virtually identical to our Groningen result. The electoral composition of the two readerships differs systematically. The finding again pleads against massive influence of wishful thinking.

A more direct and powerful test of the hypothesis would consist of asking the participants about their own party preference and relating these preferences to the predictions. We considered and rejected this design, as we did not wish to raise any misunderstandings with the participants about the character of the betting approach. Generally, we doubt whether one can direct such a dual relationship with unknown respondents without running the risk of spoiling the predictions by second-guessing and other strategic moves on the part of the participants.

A final question under this heading is how the selfselection of the participants came about. The participants are a highly selected group in a quantitative sense: in both studies, fewer than a half percent of the potential readers participated. One might therefore surmise that the participants formed a highly knowledgeable elite. Inspection of the names and addresses of the participants does not support that hypothesis. We spotted several family members, staff and students of the Psychology department, and members of the senior author's bridge club. Of the large majority of names unknown to us, the addresses were concentrated in districts where few intellectuals live. The University of Groningen does not have a department of Political Science. In sum, the self-selection probably took place in terms of political interest, pleasure in making predictions, expectation of monetary gain, and a number of irrelevant factors, rather than political scientific erudition.

General Discussion

Election Bets – Can averaged lay predictions be more generally expected to outperform commercial polls? The empirical evidence is insufficient to justify that claim. In the context of election research, some results have emerged that are tangentially supportive of the present findings. Irwin and Van Holsteyn (1988), for example, have reported on a large-scale study in The Netherlands in which one of the questions was about the respondents' expectations with respect to the future of the seated coalition; in that case also, the averaged intuition was superior to the polls. At the present stage, however, the argument for or against averaged bets must be largely methodological. Three such arguments will be reviewed here.

A first argument, which is strongly in favor of averaged predictions, is the psychometric argument that hinges upon the classical Spearman-Brown result dating from the beginning of the century. In the version that is pertinent here, the Spearman-Brown formula writes the validity of an averaged prediction as a negatively accelerated but monotone increasing function of the number of participants; the asymptote depends on the average individual validity. In averaged bets as well as in polls, there are two sources of error. One is bias due to time lag and other factors, the other is error due to limited numbers of respondents. With only about 60 participants, the latter component is reduced to essentially zero in the betting approach (see Figure 4). Pollsters would need prohibitively large samples to achieve an equally robust prediction.

A second argument pertains to the tuning of the respondent. In the betting situation, the target is the election outcome. Strategic considerations may play a role; in our studies, a small minority of the respondents may have pursued a band-wagon effect by predicting an excessive number of seats for their favored splinter party, hoping to slant the averaged prediction to be published. In the future, such moves can be discouraged by truthfully announcing that the attempts will be in vain. It can also be shown that the scoring rule (taking the sum of the absolute deviations per party) is not strictly incentive-compatible (c.q., proper, using the psychometric term); psychologically, however, the announcement of a proper scoring rule can hardly be expected to make a difference in this situation. In spite of these subtleties, the procedure is largely unambiguous, and most participants are simply motivated to predict as well as they can. This property may not hold for polls. Firstly, there is no rational motivation to refrain the respondent from strategic preoccupations as there is no premium upon responding correctly or honestly; secondly, the poll outcome is indeed sensitive to slanting efforts, so that the respondent can realistically hope to give off a message to the electorate or to politicians; for example, warning his or her own party that losses may be impending if no action is taken. We are even inclined to think that this is the most intelligent way of using the opportunity provided by a pollster. Thirdly, most polls suffer from a self-imposed handicap by asking the counterfactual question of how the respondent would vote if the elections were held today, *quod non*, instead of having the respondents predict their own factual voting. We are aware that this counterfactual tuning of the respondents is consistent with the professed time-slice nature of polls. To ensure a fair competition between polls and averaged bets, however, the handicap had better be done without.

A third reason to expect superiority of bets over polls is that the com-

mercial pollsters represented in this study have apparently not applied Bayes' theorem. When predictions are imperfectly reliable, as is the case with polls and individual intuitions, the rational Bayesian strategy is to make use of prior or collateral information such as the old composition of the parliament in question, which is itself highly predictive of the new composition. The pollsters represented in Table 2 could have done much better by taking an average between their outcome and the old composition, as can easily be verified from that table. Again, this apparent failure to apply Bayes' theorem is not an intrinsic shortcoming of polls, but just another self-imposed handicap.

To our participants, the old composition was spelled out before asking for their predictions. We do not assert that they were aware of Bayesian statistics. Rather, they may have been the victims of an involuntary anchoring effect. Their strategy thus represents an example of the functional nature of judgmental heuristics and biases (cf. Funder 1987). In discussing Table 2 we briefly noted that the size of the electoral shift was systematically underestimated by the averaged lay predictions. The suspicion might arise that the participants gave too much weight to the old composition, and that they would have been at a disadvantage if the electoral shift would have been radical indeed. However, that interpretation would be incorrect. The *individual* participants' first round prediction of the shift size varies from 6 to 34 seats, with a mean of 16.4, which is slightly above the actual shift. The shift size of only 12 seats implied in the *averaged* prediction is just another illustration of the Speatman-Brown effect: because the individual predictions correlate with the old composition, the averaged prediction has an even higher correspondence with that composition. To say that the averaged prediction is 'conservative' would be an illustration of the holistic fallacy (see, e.g., Elster 1986).

In sum, the emphasis here is not upon any intrinsic superiority of lay intuitions. Conversely, there seem to be no cogent reasons why these intuitions should suffer from anything but idiosyncratic fallibility, that is, the kind of error that is taken care of by the law of large numbers. The question remains whether that conclusion can be generalized beyond the domain of election predictions.

Methodological status of averaged predictions – Intuitive judgment has an ambivalent scientific status. On the one hand, it stands in opposition to systematic study, and functions as the rock of offence for many scientists. An important part of the history of psychology, for example, can be written in terms of the debunking of intuitive judgment. On the other hand, the peer review process by which the very plans and products of scientific

research are judged, is structurally very similar to the procedure of the present study, namely, the taking of a consensual intuitive judgment. There is a difference in that peer reviewers tend not to be aware of the predictive nature of their task; indeed, the criterion of future scientific fruitfulness is less tangible than an election outcome. Also, peer reviewers may prefer to interact instead of having their consensus (or relative lack of it) taken by an outsider. But these differences are superficial, and the fact remains that science is governed by the kind of intuitive consensual predictions that it would like to replace by systematic objective procedure. Other areas where intuitive judgment is ineradicable are personnel selection, clinical diagnosis, and, to some extent, academic grading; many more examples could be found.

The question what the place is of intuitive prediction in scientific judgment formation itself should be regarded as an empirical question. On the one hand, areas of research can be pointed at where outcomes tend to be overwhelmingly counter-intuitive. Clearly there is no substitute for rigorous research in those areas. Even there, however, it is worthwhile to take stock of intuitive predictions before the fact, if only to counteract the popular hindsight that is rampant with respect to, especially, the social and behavioral sciences. On the other hand, certain areas of investigation may indeed produce outcomes that were already expected. In such areas, pooled intuitions would seem to be a formidable competitor of laborious research. In sum, the gathering of advance predictions is generally recommendable.

Finally, the consensus procedure is reminiscent of social-methodological debates that were at their peak in the late sixties, stressing open or participatory research as an alternative to deception and authoritarianism in dealing with the human subject (Jourard 1968; Kelman 1968). The alternative essentially consisted of full disclosure of the experimental hypothesis and procedure to the subject, and asking the counterfactual question of how he or she would react if unaware. The participatory approach has not been taken into the mainstream of social and behavioral research, and the debate has subsided. The betting method may be viewed as a post-mature child of the sixties. However, it does not ask counterfactual questions; it does satisfy the most stringent operational criterion for participatory research, namely, that the investigator himself or herself could legitimately and credibly function as one of the subjects; and it is not meant to be a substitute for rigorous research until that substitution is shown to be feasible. In our assessment, the relationship between the social and behavioral sciences, their respondents, and their public is no less problematic than before. The betting approach may have the side-effect of improving the public relations of our enterprise.

References

- Eijck, C. van der (1988), Peilingen als voorspellingen van een verkiezingsuitslag (Polls as predictions of an election outcome). In: R.B. Andeweg (red.), *Tussen Steekproef en stembus* (Between Sample and Ballot-Box). DSWO, Leyden.
- Elster, J. (1986), *An introduction to Karl Marx*. Cambridge University Press, Cambridge.
- Funder, D.C. (1987), Errors and mistakes: Evaluating the accuracy of social judgment. *Psychological Bulletin* 101, p. 75-90.
- Gulliksen, H. (1950), *Theory of Mental Tests*. Wiley, New York.
- Irwin, G.A., and J.J.M. van Holsteyn (1988), Opiniepeiling en stemgedrag (Opinion polls and voting behavior). In: R.B. Andeweg (red.), *Tussen steekproef en stembus* (Between Sample and Ballot-Box). DSWO, Leyden.
- Jourard, S.M. (1986), *Disclosing Man to Himself*. Van Nostrand, Princeton, N.J.
- Kelman, H.C. (1968), *A Time to Speak: On Human Values and Social Research*. Jossey-Bass, San Francisco.

Author Note

* The content of this study was reported to the Social-Scientific Council (SWR) of the Netherlands Academy of Sciences on February 17., 1990, and at the Institute for Social Research at Oslo on March 12., 1990. The authors are indebted to Hans Daudt, Jon Elster, and the Editorial Board of *Acta Politica* for their comments.

'Het georganiseerde bedrijfsleven': Een ongerechtvaardigd monopolie op belangenbehartiging¹

A.H. Peterse

I. Inleiding

Inmiddels al weer een jaar of tien is een internationale groep onderzoekers, onder wie veel sociologen, doende de politieke organisatie van het 'het bedrijfsleven' in geïndustrialiseerde landen van West-Europa en Noord-Amerika te bestuderen. Zij kunnen worden aangeduid als 'neocorporatisten'. Het onderzoek wordt gecoördineerd door initiatiefnemers Wolfgang Streeck en Philippe C. Schmitter. Uit deze groep komen uitspraken over hun kenobject, waartoe ook een als politiek te kwalificeren proces van onderhandeling tussen de een netwerk behorende organisaties en uitvoering van gemaakte afspraken moet worden gerekend, die ook voor de politieke en beleidswetenschap relevant zijn. In dat proces worden verschillende fasen van het beleidsproces herkend. Daar blijft het echter niet bij. Sommige schrijvers kwalificeren hun kenobject ook als een te onderscheiden manier van beleidsvoering, met specifieke voor- en nadelen.

De bevindingen van de neocorporatisten zijn om twee redenen politiek relevant. Ten eerste omdat er beleidsproblemen op de politieke agenda staan waarvoor geldt dat ondernemingshandelen zowel onderdeel van het probleem als van de oplossing is. Te denken valt onder andere aan milieu-beleid en werkgelegenheidsbeleid. Beschrijving en evaluerende analyse van interorganisatorische beleidssystemen tussen overheid en bedrijven zijn dan ook een noodzakelijk soort kennis om beleidsmakers deugdelijk te informeren over de wijze waarop de genoemde problemen moeten worden aangepakt. Ten tweede is daar de minder tijdgebonden verontrusting over de democratisch-normatieve bezwaren die aan politieke besluitvorming binnen corporatistische structuren kleven.

Tot de nadelen rekenen ook de neocorporatistische schrijvers de democratisch-normatieve kritiek die vanouds al op het corporatisme wordt geleverd. Ze² erkennen dat er iets in dit opzicht niet in de haak is, als de behartigers van particularistische belangen een vinger in de pap van vor-