

Percepties over economische ongelijkheid en voorkeur voor linkse partijen en de rol van percepties over immigranten.

Bachelorwerkstuk - herziene eindversie

Inleverdatum: 01-08-2025

Jurre Koers - S4918290

j.m.koers.1@student.rug.nl

Begeleider: Marinus Spreen

m.spreen@rug.nl

Tweede corrector: Rita Smaniotto

r.c.smaniotto@rug.nl

Abstract

In dit onderzoek wordt de relatie onderzocht tussen de houding die kiezers hebben ten opzichte van economische ongelijkheid en het stemmen op linkse politieke partijen in Nederland. Daarnaast wordt gekeken of deze relatie wordt beïnvloed door de houding van kiezers ten opzichte van immigranten. Op basis van gegevens uit het LISS-panel is een logistische regressieanalyse onder 1788 respondenten uitgevoerd, waarbij het wel of niet stemmen op een linkse partij de afhankelijke variabele is. De resultaten tonen aan dat kiezers met een negatievere kijk op ongelijkheid eerder op linkse partijen zullen stemmen. Tegelijkertijd blijkt dat een negatieve houding ten opzichte van immigranten deze kans significant verlaagt. Voor de verwachte invloed van de houding ten opzichte van immigranten op de relatie tussen ongelijkheid en stemgedrag bleek geen statistisch bewijs, maar kansberekeningen en visualisaties tonen een inhoudelijk relevant patroon: bij kiezers met een negatieve kijk op immigratie verdwijnt het positieve verband tussen negatieve kijk op inkomensongelijkheid en steun voor linkse partijen grotendeels. Vanwege de afwezigheid van statistisch bewijs kan dit patroon niet worden aangenomen als conclusie, maar het blijft een opvallende bevinding. Politieke voorkeuren lijken dus niet enkel economisch, maar ook deels cultureel en ideologisch gedreven.

Inhoudsopgave

Inleiding	3
Theoretisch kader	6
Economisch eigenbelang	6
Rechtvaardiging	7
Morele betrokkenheid	8
Populisme	9
Etnische dreiging	11
Alternatieve verklaringen	12
Methoden	14
Operationalisaties	15
Wat is een linkse partij?	17
Analyseplan	19
Data	22
Resultaten	26
Modevaluatie	26
Hypothesetoetsing	29
Conclusie	34
Discussie	35
Literatuur	36
Bijlagen	40
Bijlage 1 & 2	41
Bijlage 3	105
Bijlage 4 (procesverslag)	108
Toelichtende teksten	108
Conclusie (4-6-2025)	111
Verslag reparatie (01-08-2025)	112

Inleiding

In veel westerse democratieën staat de groeiende economische ongelijkheid hoog op de politieke agenda. Onder invloed van globalisering, flexibilisering van de arbeidsmarkt en neoliberale hervormingen is de kloof tussen arm en rijk toegenomen (Makhlouf, 2022), wat volgens diverse politieke en sociologische theorieën zou moeten leiden tot een toename in steun voor linkse partijen die herverdeling en economische gelijkheid vooropstellen (Petrova & Ranaldi, 2024; Kenworthy & McCall, 2007). Toch blijkt deze relatie in praktijk minder eenvoudig: niet alle mensen die economische ongelijkheid als problematisch ervaren stemmen op linkse partijen.

In het wetenschappelijke debat groeit de aandacht voor culturele factoren die dit verschil kunnen verklaren. Een van die verklaringen richt zich op rechtvaardiging in de vorm van het geloof in meritocratie, dat ongelijkheid als verdiend kan legitimeren (Mijs & Savage, 2020). Een tweede, veelbesproken verklaring is de rol van attitudes ten opzichte van immigratie. Steeds meer onderzoek laat zien dat negatieve opvattingen ten opzichte van immigranten solidariteit en steun voor herverdelingsbeleid kunnen verminderen, met name onder lager opgeleide of economisch kwetsbare groepen (Scheepers, Gijsberts & Coenders, 2002; Fox, 2004; Macdonald, 2020). Dit heeft belangrijke politieke gevolgen: wanneer mensen de herverdeling van middelen als gunstig voor ‘de ander’, met name migranten, zien, kan dit hun steun voor linkse partijen verminderen, zelfs als zij economische ongelijkheid afwijzen.

Hoewel het verband tussen houding ten opzichte van economische ongelijkheid en links stemgedrag vaak als vanzelfsprekend wordt verondersteld, toont de literatuur dat culturele percepties zoals de houding ten opzichte van immigranten de vertaling van economische onvrede naar politiek gedrag kunnen beïnvloeden (McVeigh et al., 2014; Roex, Huijts & Sieben, 2018). Dit onderzoek sluit aan bij deze discussie en onderzoekt in hoeverre de houding ten opzichte van immigranten van invloed is op

de relatie tussen hoe kiezers economische ongelijkheid ervaren en hun keuze om al dan niet op linkse partijen te stemmen.

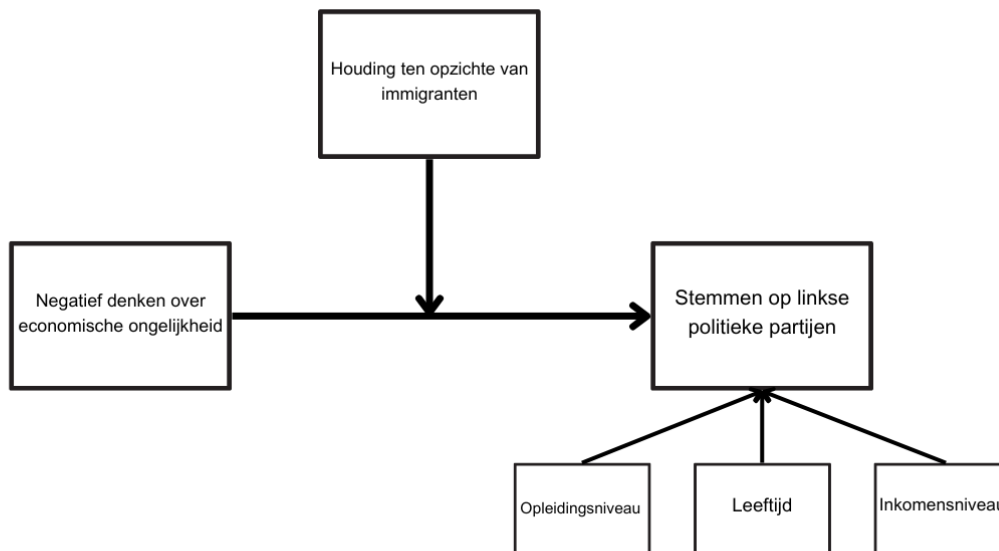
Deze keuze wordt verder bemoeilijkt doordat herverdelingsbeleid vaak wordt tegengewerkt door het idee dat andere groepen met een slechte economische positie concurreren om hulpbronnen. Vooral in de lagere economische klassen zijn gevoelens van culturele vervreemding en concurrentie over hulpbronnen aanwezig (Gidron & Hall, 2017). Hierdoor ontstaat een vreemde situatie: juist de groepen die het meeste voordeel zouden halen uit herverdelingsbeleid, keren zich soms af tegen partijen die deze willen bewerkstelligen.

In het publieke debat wordt deze situatie regelmatig zichtbaar in discussies over sociale voorzieningen, krapte op de huizenmarkt en de arbeidsmarkt, waarbij migranten vaak worden voorgesteld als gelukszoekers of belasting op het systeem. Tegelijk stellen linkse partijen zich vaak op als verdedigers van zowel herverdeling als culturele en etnische inclusie, wat tot twijfels kan leiden bij een achterban die economisch links, maar cultureel conservatief kan zijn (Stewart et al., 2020).

Door deze keuze te onderzoeken, levert dit onderzoek een bijdrage aan het begrijpen van de discrepantie tussen economische onvrede en links stemgedrag, en aan het debat over de sociale basis van politieke partijen in tijden van groeiende ongelijkheid en culturele polarisatie. De centrale vraag die daaruit volgt is:

In hoeverre beïnvloedt de houding ten opzichte van immigranten het verband tussen de perceptie van economische ongelijkheid en het stemmen op een linkse partij?

Figuur 1: het onderzoeksmodel van dit onderzoek.



Theoretisch kader

Zoals gezegd wordt hier onderzocht of het negatiever denken over economische ongelijkheid door mensen leidt tot het stemmen op linkse politieke partijen, en of een negatieve houding ten opzichte van immigranten deze mogelijke samenhang vermindert. Met het negatief denken over economische ongelijkheid wordt bedoeld dat mensen de verschillen in sociale en economische status erkennen, en deze ook als probleem binnen de samenleving beschouwen. Onder een negatieve houding ten opzichte van immigranten wordt bedoeld het idee dat immigranten over het algemeen meer negatieve invloeden op de samenleving hebben dan positieve. Het eerste verband, de invloed van een negatieve kijk op economische ongelijkheid op het stemmen op linkse partijen, wordt als eerste uitgewerkt.

Economisch eigenbelang

Een klassiek uitgangspunt in de sociologie is dat mensen handelen op basis van rationeel eigenbelang (Smith, 1776). Volgens deze benadering zullen individuen geneigd zijn om die beleidsmaatregelen en politieke partijen te steunen die hun economische of sociale positie verbeteren of beschermen. In de context van economische ongelijkheid betekent dit meestal dat mensen met lagere inkomens meer baat hebben bij herverdelingsbeleid, en dus eerder geneigd zijn om te stemmen op linkse partijen die dergelijke maatregelen ondersteunen (Kenworthy & McCall, 2007).

Empirisch onderzoek ondersteunt dit patroon: mensen met een lagere economische positie hebben gemiddeld meer steun voor herverdelingsbeleid (Roex, Huijts & Sieben, 2018). Daarnaast identificeren ze zich vaker met de arbeidersklasse, wat ook samengaat met een voorkeur voor progressief herverdelingsbeleid (Franko & Witko, 2023). Anderzijds zijn mensen met een hogere economische positie minder vaak voorstander van herverdeling. Enerzijds komt dit voort uit het directe materiële belang: herverdeling betekent voor hen mogelijk hogere belastingen, wat gepaard gaat met een verlies van hun economische positie, of andere vormen van herstructurering.

Het eigenbelang-mechanisme veronderstelt daarmee dat mensen met een gunstige economische positie een lagere (politieke) motivatie hebben om ongelijkheid te verminderen, omdat zij van de bestaande verhouding profiteren. Zelfs als deze groepen zich bewust zijn van de mate van ongelijkheid, leidt dat besef niet noodzakelijk tot steun voor linkse politiek: het kan zelfs leiden tot een acceptatie of actieve verdediging van de status quo (Lepianka et al., 2010; Mijs, 2015). Dit kan onder andere gebeuren via een mate van geloof in meritocratische principes.

Rechtvaardiging

De manier waarop mensen economische ongelijkheid moreel opvatten, speelt een erg belangrijke rol in de mate waarin ze bereid zijn politieke actie te ondernemen of partijen te steunen die voor herverdeling staan (over het algemeen linkse partijen). Niet alleen de waarneming van ongelijkheid is daarbij relevant, maar vooral de waardering ervan; in andere woorden of de bestaande ongelijkheid als legitiem of oneerlijk wordt beschouwd.

Onder de opvattingen van rechtvaardiging vallen overtuigingen over de oorzaken, terechtheid en onvermijdelijkheid van ongelijkheid. Een belangrijke vorm is het geloof in meritocratie: het idee dat sociale en economische uitkomsten het gevolg zijn van individuele inzet, talent en verdienste (Young, 1958). Wanneer mensen ongelijkheid opvatten als het gevolg van eerlijke competitie, zal de motivatie om hier politiek tegen in opstand te komen klein of niet-bestaand zijn (Mijs, 2015; Mijs & Savage, 2020). Ongelijkheid wordt dan niet als een fout in het systeem gezien, maar als een rechtvaardige beloning voor prestaties.

Hieronder vallen ook armoedeverklaringen: schrijven mensen armoede toe aan individuele schuld (bijvoorbeeld luiheid of antisociaal gedrag), of aan structurele factoren zoals werkloosheid, discriminatie of ongelijke onderwijskansen? Onderzoek laat zien dat mensen met een hogere economische positie vaker individuele schuldverklaringen hanteren (Lepianka, Gelissen & Van Oorschot, 2010; Oorschot & Halman, 2000), wat gepaard gaat met minder steun voor

herverdelingsbeleid. Dergelijke verklaringen impliceren dat ongelijkheid terecht is en dus geen politiek ingrijpen vereist.

Mensen gebruiken deze overtuigingen vaak om bestaande sociale verhoudingen te legitimeren. Dit sluit aan bij opvattingen uit de zogeheten “system justification theory”, die stelt dat mensen (en vooral degenen met een gunstige economische positie) geneigd zijn ongelijkheid niet alleen te accepteren, maar ook actief goed te praten om cognitieve dissonantie te vermijden (Jost et al., 2004). De opvatting dat het systeem werkt zoals het zou moeten, maakt de geobserveerde ongelijkheid mentaal acceptabel, ook voor mensen die erkennen dat ongelijkheid bestaat.

Morele betrokkenheid

Naast economische belangen en rechtvaardigingen speelt ook empathie, het vermogen om zich in te leven in anderen, een grote rol in het vormen van politieke opvattingen ten opzichte van ongelijkheid en herverdeling. Dit mechanisme gaat ervan uit dat mensen niet alleen handelen op basis van eigenbelang of politieke overtuigingen, maar ook op basis van hun morele verbondenheid met anderen. Wanneer mensen zich emotioneel betrokken voelen bij groepen die door ongelijkheid worden getroffen, vergroot dat de kans op steun voor beleid dat ongelijkheid vermindert, zoals herverdeling of stemmen op linkse partijen.

Empathie kan zich uiten in een algemeen gevoel van solidariteit, maar ook in morele verontwaardiging over ongelijkheid. Belangrijk is dat deze gevoelens niet alleen gericht zijn op de eigen groep, maar ook op anderen die (economisch) achtergesteld zijn. Tropp en Barlow (2018) laten zien dat empathie vaak ontstaat uit intergroep contact: persoonlijk contact met mensen uit achtergestelde groepen bevordert de motivatie om ongelijkheid te bestrijden. Dit geldt met name voor leden van bevoordeelde groepen (bijvoorbeeld rijke mensen), die via positieve interacties met minderheidsgroepen bewuster worden van ongelijkheid en gemotiveerd raken om daar iets aan te veranderen.

Ook Delhey en Dragolov (2013) tonen aan dat economische ongelijkheid gevoelens van wantrouwen, statusangst en conflict oproept, iets wat de sociale cohesie aantast. In omgevingen waar mensen meer vertrouwen hebben in elkaar, en zich meer verbonden voelen met de samenleving als geheel, is de steun voor herverdelingsbeleid aanzienlijk groter. Empathie en morele betrokkenheid werken op die manier als een soort brug tussen individuele waarnemingen van ongelijkheid en collectieve politieke actie.

Wat hieruit voortvloeit is de eerste centrale hypothese van dit onderzoek, namelijk dat een negatieve kijk op economische ongelijkheid vaak samengaat met het stemmen op een linkse partij, omdat het ook vaak samengaat met steun voor idealen als economische herverdeling; een ideaal dat linkse partijen vaak delen. De eerste hypothese uit dit onderzoek is dus:

Hypothese 1: Een negatieve houding tegenover inkomensongelijkheid leidt tot grotere kans op stemmen op een linkse partij.

Hoewel wordt verondersteld dat het bovenstaande verband bestaat, zijn er een aantal factoren die hier invloed op kunnen hebben. In dit onderzoek wordt specifiek verder ingegaan op de rol van de houding ten opzichte van immigranten bij dit verband.

Populisme

Hoewel economische ongelijkheid in klassieke modellen (bijvoorbeeld die van Marx) wordt gezien als een bron van klassenbewustzijn en steun voor herverdeling, blijkt in de praktijk dat onvrede over ongelijkheid niet altijd stemmen op linkse partijen als gevolg heeft. Een belangrijke verklaring hiervoor is het mechanisme van populistische mobilisatie: politieke bewegingen, vaak rechts-populistisch, framen economische onvrede in termen van culturele of etnische dreiging (McVeigh et al., 2014; Scheepers et al., 2002). Hierdoor worden niet de economische elite, maar de ‘anderen’ zoals migranten of stedelijke, ‘progressieve’ elites de zondebok.

Populistische ideologie benoemt een moreel onderscheid tussen het ‘echte volk’ en ‘corrupten’ of ‘indringers’, waarbij economische problemen worden toegeschreven aan externe groepen (zoals immigranten) of een elite die “het volk” in de steek laat. Hierin wordt economische ongelijkheid niet ontkend, maar anders geïnterpreteerd: niet als gevolg van structurele uitbuiting, maar als een symptoom van een “bedreigd nationaal belang” (McVeigh et al., 2014). Hierdoor kunnen ook mensen die ontevreden zijn over hun economische positie zich toch verzetten tegen herverdelingsbeleid, zeker als ze het idee hebben dat de voordelen daarvan vooral naar migranten of ‘anderen’ gaan (Lepianka et al., 2010; Oorschot & Halman, 2000).

Macdonald (2020) toont empirisch aan dat in de Verenigde Staten negatieve opvattingen ten opzichte van immigranten de opvattingen ten opzichte van inkomensongelijkheid onder witte Amerikanen verminderen. Met andere woorden: zelfs als mensen economische ongelijkheid waarnemen en afwijzen, kan een negatieve houding tegenover migranten ervoor zorgen dat zij niet kiezen voor linkse partijen die voor herverdeling staan, maar voor rechts-populistische alternatieven die economische en culturele onvrede combineren.

Daarnaast laten de werken van Tajfel (1982) en Tropp & Barlow (2018) zien dat het denken in in-groups en out-groups een belangrijke rol speelt bij het solidariteitsprincipe. Wanneer populistische retoriek migranten of andere minderheden frameert als buitenstaanders die oneerlijk profiteren van het systeem, kan dat de morele betrokkenheid bij herverdeling verminderen, vooral bij degenen die zich identificeren met de nationale of etnische meerderheidsgroep.

Een mindere mate van geloof in meritocratische principes kan echter ook ander stemgedrag tot gevolg hebben: (rechts)populistische partijen. Onder mensen met een zwakkere economische positie in de samenleving bestaat een groep met een gevoel van relatieve deprivatie; mensen met het gevoel dat de groep waarvan ze denken dat ze toebehoren op economisch gebied niet krijgt waar ze recht op hebben, terwijl anderen meer krijgen dan ze recht op hebben (Elchardus & Spruyt, 2012).

Populistische partijen spreken mensen waarbij dit gevoel heerst aan door dit gevoel te rechtvaardigen

en schuldigen aan te wijzen. De schuldigen zijn vaak ‘de elite’, maar bij deze partijen worden immigranten vaak ook als schuldigen aangewezen. Of deze immigranten wel of niet als schuldige worden aangewezen bepaalt vaak of de partij links-populistisch (bijvoorbeeld de Socialistische Partij in Nederland) of rechts-populistisch (bijvoorbeeld de Partij Voor de Vrijheid) is (Vasilopoulos & Jost, 2020). Het is dan ook logisch om te stellen dat wanneer de keuze tussen een links-populistische of rechts-populistische partij gemaakt moet worden, een belangrijke factor de kijk op immigranten is. Wanneer iemand dus vatbaar is voor de idealen van populistische partijen kan dit het stemmen op linkse partijen betekenen bij een neutrale of positieve kijk op immigratie, en het stemmen op rechtse partijen bij een negatieve kijk op immigranten.

Etnische dreiging

Het verband tussen percepties van economische ongelijkheid en steun voor linkse partijen wordt dus beïnvloed door hoe mensen de aanwezigheid van etnische minderheden, en in het bijzonder immigranten, ervaren. Volgens de ethnic threat theory ervaren mensen migranten als een bedreiging voor hun eigen economische, culturele of politieke positie (Scheepers et al., 2002). Deze dreiging wordt geuit onder meer in de vorm van competitie om banen, uitkeringen of huisvesting (Semyonov et al., 2006).

In deze context kan herverdelingsbeleid, wat vaak wordt voorgesteld door linkse partijen, worden gezien als een bevoordeling van migranten ten koste van de ‘eigen’ groep (Oorschot & Halman, 2000; Van Oorschot, 2006). In deze situatie wordt herverdeling niet afgewezen vanwege bezwaren tegen solidariteit, maar vanwege de richting waarin deze solidariteit wordt gestuurd. Herverdeling wordt slechts legitiem gevonden als het binnen de ‘in-group’ plaatsvindt (Tajfel, 1982).

Onderzoek toont aan dat negatieve houdingen ten opzichte van immigranten samenhangen met lagere steun voor universele welvaartsprogramma's (Fox, 2004). Met name in economisch kwetsbare groepen (ondanks dat ze in principe het meest zouden profiteren van herverdeling) kan de aanwezigheid van een grote groep immigranten leiden tot verminderde steun voor links beleid, omdat

herverdeling wordt gezien als iets dat de ‘oneerlijk bevoorrechte’ immigranten ten goede komt (Roex et al., 2018).

Deze vorm van dreiging is dus geen afwijzing van herverdeling op zichzelf, maar van inclusieve herverdeling. Daardoor wordt het verwachte verband tussen waargenomen ongelijkheid en steun voor linkse partijen, net zoals bij het mechanisme van populisme, gemodereerd door etnische percepties. Hieruit volgt de tweede centrale hypothese van dit onderzoek:

Hypothese 2: De relatie tussen houding over inkomensongelijkheid en stemgedrag is zwakker bij negatieve houding tegenover immigranten.

Alternatieve verklaringen

Om beter aan te kunnen tonen dat een mogelijk verband daadwerkelijk door de hiervoor genoemde verklaringen komt, moet er rekening gehouden worden met een aantal mogelijke alternatieve verklaringen. De eerste verklaring is opleidingsniveau: zo ervaren mensen met een hoger opleidingsniveau over het algemeen inkomensongelijkheid eerder als legitiem. Dit verschil bestaat vooral in landen met een hogere mate van economische ongelijkheid (Barr & Miller, 2020).

Daarnaast kan leeftijd een rol spelen; oudere mensen ervaren economische ongelijkheid vaker als ‘zeer onrechtmatig’ dan jongere mensen (Reyes & Gasparini, 2022), wat kan betekenen dat hun kijk op economische ongelijkheid over het algemeen negatiever is. Andersom zijn jongeren weer veel minder snel geneigd om te stemmen op rechtse partijen en neigen ze meer naar liberale en groene partijen (Rekker, 2024), wat invloed zou kunnen hebben op het verband tussen deze twee factoren.

Ten slotte wordt er rekening gehouden met inkomensniveau: zo zou het kunnen dat mensen met een lager inkomen negatiever denken over economische ongelijkheid dan mensen met een hoger inkomen, en dat de werkelijke verklaring het inkomensniveau is. Op basis van bestaand onderzoek lijkt dit niet het geval, zo hebben Duitsers een grote mate van steun voor economische herverdeling ongeacht hun

eigen economische positie (Engelhardt & Wagener, 2017), maar dit kan per populatie of land verschillen.

Methoden

Voor het uitvoeren van dit onderzoek wordt gebruikgemaakt van het LISS (Longitudinal Internet studies for the Social Sciences) data archive. Het LISS data archive is een databestand dat voortvloeit uit een willekeurige steekproef van leden van het LISS panel. Het LISS panel is een onderzoeksgroep bestaande uit 7.500 mensen van 16 jaar en ouder verdeeld over 5.000 huishoudens, waarbij deelnemers maandelijks online vragenlijsten invullen in ruil voor een financiële vergoeding. Dit panel is uitsluitend toegankelijk op basis van uitnodiging, en uitgenodigde huishoudens zonder toegang tot het internet worden voorzien van een simpele computer met toegang tot het internet, waardoor ook zij in het onderzoek meegenomen kunnen worden. Selectie van de deelnemers is gebaseerd op een zo representatief mogelijke verdeling van huishoudtype, leeftijd en etniciteit, en de benadering gebeurt op traditionele wijze (per telefoon, fysieke post of huiselijk bezoek). Van de geregistreerde huishoudens doet ongeveer 80% van de mogelijke mensen mee, waarbinnen de maandelijkse respons van deze participanten varieert tussen de 50- en 80%.

Binnen de groep van de gehele dataset is een selectie gemaakt op nonrespons, waarbij participanten die een voor het onderzoek relevante vraag niet hebben beantwoord ook niet meegenomen worden in het onderzoek. Daarnaast is het voor het onderzoek van fundamenteel belang dat we kunnen bepalen of een participant wel of niet op een linkse partij heeft gestemd, waardoor respondenten die hebben aangegeven dat ze op een andere partij hebben gestemd dan de gegeven lijst ook niet worden meegenomen. Dit is omdat het bij deze groep mensen niet valt te bepalen of de partij waarop ze zeggen te hebben gestemd wel of niet links is: partijen die niet op de lijst staan kunnen in feite alle mogelijke vormen aannemen. Participanten die blanco hebben gestemd worden wel gewoon meegenomen en horen bij de groep participanten die niet op een linkse partij hebben gestemd, omdat een blanco stem ook een stem is en niet een stem op een linkse partij. Het uiteindelijke, totale aantal respondenten is 1788.

Operationalisaties

Voor het onderzoek zijn een aantal variabelen samengevoegd om een zo goed mogelijk beeld te geven van het onderzochte concept. Ten eerste zijn voor de onafhankelijke variabele, de kijk op ongelijkheid, de antwoorden van twee stellingen samengevoegd tot één variabele. De stellingen, in de dataset genoemd '*ineq_1*' en '*ineq_5*', luiden als volgt:

- De verschillen in inkomen zijn in Nederland te groot.
- Ik maak me kwaad over verschillen in rijkdom tussen de rijken en de armen.

Hierbij konden respondenten kiezen uit de antwoordmogelijkheden *helemaal oneens*, *oneens*, *een beetje oneens*, *niet oneens en niet eens*, *een beetje eens*, *eens*, en *helemaal eens* (1-7). De Cronbach's Alpha voor deze twee variabelen is 0,724.

De reden dat er gekozen is voor deze twee stellingen en niet de overige stellingen over ongelijkheid uit de vragenlijst, is dat deze twee vragen het beste vragen naar specifiek de mening over of houding ten opzichte van ongelijkheid van de respondenten. De andere stellingen¹ gaan namelijk meer over in welke mate respondenten vinden dat er maatregelen genomen moeten worden om ongelijkheid tegen te gaan, terwijl maatregelen nodig vinden niet een essentieel onderdeel is van een negatieve houding ten opzichte van ongelijkheid. Daarnaast zijn maatregelen met betrekking tot ongelijkheid vaak een belangrijk kenmerk van de partijprogramma's van linkse partijen, waardoor anders gedeeltelijk het verband tussen het hebben van een links standpunt en het stemmen op een linkse partij wordt onderzocht, wat lang niet zo interessant en relevant is als het verband dit stemgedrag en een bepaalde houding. Voor de eindvariabele is het gemiddelde van de twee gebruikte variabele genomen, en is deze gecentreerd.

Voor de modererende variabele, de kijk op immigranten, zijn in totaal zes variabelen samengevoegd tot één variabele. Hierbij is gekeken naar welke stellingen het meest direct vertaald kunnen worden naar de houding ten opzichte van immigratie en immigranten. De stellingen - in de dataset genoemd ‘*etnocentr1*’, ‘*etnocentr2*’, ‘*etnocentr4*’, ‘*marxmigr1*’, ‘*marxmigr4*’ en ‘*marxmigr5*²’ - zijn:

- Het culturele leven in Nederland wordt over het algemeen verrijkt (beter gemaakt) door mensen uit andere landen die hier zijn komen wonen.
- Nederland is een betere plek geworden door mensen uit andere landen die hier zijn.
- Nederland had eigenlijk nooit gastarbeiders binnen moeten laten.
- Mensen met een migratieachtergrond nemen de banen in van mensen zonder een migratieachtergrond.
- De belangen van werknemers zonder een migratieachtergrond zijn belangrijker dan die van mensen met een migratieachtergrond.
- De meeste mensen met een migratieachtergrond zijn onderdeel van een lagere sociale klasse dan ik².

¹ De overige, niet geïnccludeerde stellingen over ongelijkheid zijn:

- De overheid moet werklozen een fatsoenlijke levensstandaard garanderen.
- De overheid moet maatregelen nemen om verschillen tussen arm en rijk kleiner te maken.
- Bedrijven moeten maatregelen nemen om verschillen in beloning tussen hun werknemers met hoge lonen en met lage lonen te verkleinen.

² De overige, niet geïnccludeerde stellingen over immigratie zijn:

- Mensen met een migratieachtergrond die in Nederland wonen moeten Nederlandse gewoonten en gebruiken overnemen.
- Het werk dat mensen zonder een migratieachtergrond niet willen doen wordt door mensen met een migratieachtergrond gedaan.
- Werknemers met een migratieachtergrond en werknemers zonder een migratieachtergrond hebben dezelfde belangen.
-

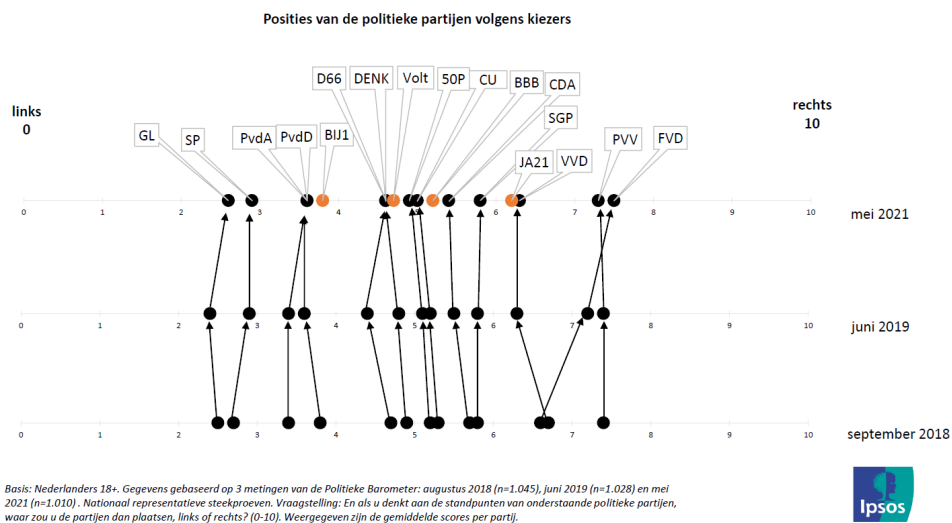
De antwoordmogelijkheden zijn hier hetzelfde als bij de vragen over ongelijkheid (*helemaal oneens, oneens, een beetje oneens, niet oneens en niet eens, een beetje eens, eens, en helemaal eens*). Omdat bij de eerste twee stellingen het meer eens zijn met de stelling een positievere kijk op immigranten betekent en bij de overige stellingen een negatievere kijk, zijn de antwoorden van de respondenten bij deze vraag bij het samenvoegen omgedraaid. Hieruit volgt een Cronbach's Alpha van 0,730.

Deze stellingen geven het beste beeld over de houding ten opzichte van immigranten, omdat ze het meest direct vragen naar een mening over immigranten, terwijl de andere vragen daar minder effectief in zijn. Zo vraagt *'etnocentr3'* naar de mening van respondenten over dat immigranten de gewoonten en gebruiken van Nederlanders moeten overnemen, waarbij het hier sterk mee eens zijn vaak samengaat met het negatief kijken ten opzichte van immigranten, maar dit niet per se zo hoeft te zijn. Positief kijken ten opzichte van immigranten en tegelijkertijd ook vinden dat immigranten gebruiken en gewoontes moeten overnemen is mogelijk, en maakt een positieve kijk ten opzichte van immigranten in principe niet minder valide. De overige, niet gekozen stellingen zijn vergelijkbaar in dat het sterk met de stellingen eens zijn vaak samengaat met een bepaalde kijk op immigranten, maar dat niet per se zo hoeft te zijn. In het geval van deze vragen is dat omdat ze ook opgevat kunnen worden als neutrale observaties, en dus niet per se relevant hoeven te zijn voor de mening van respondenten. Verder is hier voor de eindvariabele ook hier het gemiddelde van de gebruikte variabelen genomen, en is ook deze variabele is gecentreerd.

Wat is een linkse partij?

Voor het bepalen wat wel en wat niet valt onder het stemmen op een linkse partij, is het van belang om een classificering te maken. In de meest recente peiling over het plaatsen van partijen op een links-rechts schaal van Ipsos (Van Heck, 2021), is een aantal clusters van partijen op de schaal te zien.

Figuur 2: peiling van Ipsos over het plaatsen van partijen op een links-rechts schaal.



Waar de PVV en het FvD er duidelijk bovenuit steken als de partijen die als het meest rechts gezien worden, is er ook een duidelijk groepje linkse partijen te herkennen. Uit de kloof tussen BIJ1 en D66 blijkt dat mensen hier ongeveer een duidelijk onderscheid maken tussen linkse en centrumpartijen. Op basis hiervan concluderen we dan ook dat hier uitsluitend GroenLinks, SP, PvdA, PvdD en BIJ1 onder linkse partijen vallen. Sinds de peiling van Ipsos zijn GroenLinks en de PvdA gefuseerd tot één verkiesbare partij, en was BIJ1 geen antwoordmogelijkheid in de vragenlijst die gebruikt is voor dit onderzoek. Hierdoor wordt stemmen op de volgende partijen uiteindelijk gezien als stemmen op linkse partijen:

- GroenLinks-PvdA
- SP
- PvdD

De selectie hiervan is gebaseerd op de economische standpunten van de partijen (in welke mate en de manier waarop ze ongelijkheid willen bestrijden) en sociale standpunten (in welke mate ze pleiten voor een samenleving met moderne normen en waarden op gebieden als sociale gelijkheid). De antwoordmogelijkheden PVV, VVD, NSC, D66, BBB, CDA, Forum voor Democratie, ChristenUnie,

SGP, JA21 en Blanco worden geclassificeerd als het niet stemmen op een linkse partij. Zoals eerder benoemd worden antwoorden van ‘Andere partij’ behandeld als nonrespons, omdat hierbij niet te zeggen is of het gaat om het wel of niet stemmen op een linkse partij. De uiteindelijke variabele is gecodeerd tot een dummy variabele waarbij 0 het niet stemmen op een linkse partij is, en 1 het wel stemmen op een linkse partij is.

Ten slotte zijn nog drie controlevariabelen gebruikt in het onderzoek: inkomen, opleidingsniveau en leeftijd. Inkomen is gemeten door de variabele *nettohh_f*, het netto maandinkomen van het huishouden in euro's. Hier is voor gekozen omdat dit een beter beeld geeft van wat een persoon te besteden heeft dan persoonlijk inkomen, omdat vaste lasten verdeeld kunnen worden over meerdere leden van een huishouden. Deze variabele is voor het onderzoek ook verkleind (gedeeld door 1,000), om de interpretatie van de coëfficiënten inzichtelijker te maken. Opleidingsniveau is gemeten aan de hand van een categorische variabele met antwoordmogelijkheden in de 6 CBS-categorieën van onderwijs: basisonderwijs, vmbo, havo/vwo, mbo, hbo en wo, verdeeld van 1 tot 6 in deze volgorde. Leeftijd is gemeten in de leeftijd van het lid van het huishouden, omdat dit over de respondent zelf gaat, en niet over bijvoorbeeld het huishoud-hoofd. Al deze controlevariabelen zijn ook gecentreerd.

Analyseplan

Het onderzoek wordt uitgevoerd in de vorm van een logistische regressie, waarbij het effect van de onafhankelijke variabele ‘kijk op ongelijkheid’ wordt getoetst op de afhankelijke dummy voor het wel of niet stemmen op een linkse partij, en wordt gekeken of dit effect gemodereerd wordt door het toevoegen van de variabele ‘kijk op immigranten’.

Het toetsen van dit mogelijk modererende effect wordt gedaan aan de hand van regressieanalyses in meerdere modellen, in totaal vier. Model 1 bevat alleen de effecten van de controlevariabelen op het wel of niet stemmen op een linkse partij, model 2 bevat de effecten van zowel de controlevariabelen als de kijk op economische ongelijkheid, model 3 bevat de controlevariabelen, de kijk op

economische ongelijkheid en de kijk op immigranten, en tot slot bevat model 4 de controlevariabelen, de kijk op zowel economische ongelijkheid, de kijk op immigranten en een interactievariabele tussen de kijk op economische ongelijkheid en de kijk op immigranten.

Op basis van de uitgevoerde regressieanalyse wordt ook voor elk model een modevaluatie gedaan, om na te gaan of het model de data goed verklaart. Hierbij worden de -2 Log Likelihood (-2LL) en de model χ^2 -waardes gebruikt om de verklarende kracht van de modellen te evalueren, en om te kijken of het toevoegen van nieuwe voorspellers het verklarend vermogen van de modellen vergroot. Verder wordt ook de Hosmer-Lemeshow toets toegepast om de mate van overeenstemming tussen geobserveerde en voorspelde waarden te controleren, om te kunnen aantonen dat het model een goede fit heeft en de resultaten betrouwbaar zijn. Ook wordt er gecontroleerd voor multicollineariteit tussen de onafhankelijke variabelen, wat gedaan wordt door de Variance Inflation Factor (VIF) scores te berekenen in een aparte, lineaire regressieanalyse. Deze regressieanalyse is uitsluitend voor het verkrijgen van de VIF-scores uitgevoerd, en gebruikt een willekeurige afhankelijke variabele (in dit geval geboortjaar). Voor het bepalen of er sprake is van multicollineariteit wordt de vuistregel aangehouden dat een VIF-score van onder de 4 acceptabel is.

Verder wordt er ook rekening gehouden met uitbijters, dit omdat individuele respondenten met een bijzonder profiel een onevenredige invloed op het model kunnen hebben. Hiervoor is de Cook's Distance berekend, waarbij wordt aangehouden dat punten met een waarde van hoger dan $4/N$ ($4 / 1788 = 0,002$) worden beschouwd als potentieel invloedrijk en nader bekeken worden.

Ten slotte worden voor het interpreteren van de coëfficiënten van de regressietabel een aantal relevante log-odds vertaald naar procentuele kansen. Dit omdat de regressiecoëfficiënten alleen iets zeggen over de veranderingen in log-odds, wat inhoudelijk lastig te interpreteren is. Om dit interpreteren duidelijker te maken worden de kansen berekend voor het stemmen op een linkse partij bij bijvoorbeeld een gemiddelde score op alles, en een gemiddelde score op alles behalve een

maximale score op de kijk op ongelijkheid, om zo de verschillen in deze scores inhoudelijk vatbaar te maken.

Om inzicht te krijgen in de verdeling en samenhang van de variabelen die in dit onderzoek gebruikt zijn, wordt allereerst gekeken naar de beschrijvende statistieken (Tabel 1). De variabelen in dit onderzoek zijn onder andere de houding ten opzichte van inkomensongelijkheid, houding ten opzichte van immigranten, het maandinkomen, opleidingsniveau, leeftijd en een binaire variabele die aangeeft of iemand op een linkse partij heeft gestemd (0 = nee, 1 = ja). Voor de overzichtelijkheid en interpreteerbaarheid van de waarden zijn de waarden weergegeven in de tabel niet gecentreerd, en is de variabele van inkomen hier niet verkleind.

Data

De gebruikte dataset bestaat in totaal uit 2719 respondenten. Niet al deze reacties zijn echter volledig, en voor dit onderzoek zijn alleen de volledige reacties gebruikt. Om dat te doen is een simpele lineaire regressie met dezelfde afhankelijke en onafhankelijke variabelen als in het onderzoek uitgevoerd, waarvan de residuen zijn bewaard. Bij niet-volledige gevallen is de waarde voor deze residuen “missend”, op basis waarvan vervolgens een dummy-variabele is aangemaakt voor of een variabele volledig is of niet. Vervolgens zijn alleen de volledige gevallen geselecteerd, waarna 1793 respondenten overblijven.

Tabel 1: univariate statistieken van de gebruikte variabelen.

<i>Variabele</i>	<i>Gemiddelde (standaarddeviatie)</i>	<i>Minimum</i>	<i>Maximum</i>	<i>N (totaal)</i>
Houding ten opzichte van inkomens- ongelijkheid (Schaal 2 items)	4,405 (1,391)	1	7	1793
Houding ten opzichte van immigranten (Schaal 6 items)	3,523 (0,980)	1	7	1793
Inkomen	4498,63 (2411,740)	0	30509	1793
Opleidingsniveau	4,32 (1,313)	1	6	1793
Leeftijd	46,51 (14,016)	18	67	1793
Wel of niet stemmen op een linkse partij (0 = nee, 1 = ja)	67,8% niet, 32,2% wel	0	1	1793

Een opvallend verschil is zichtbaar in de gemiddelde scores op de twee houding-schalen. De gemiddelde score op de schaal over inkomensongelijkheid ligt met 4,405 relatief hoog, terwijl de gemiddelde score op de schaal over immigranten een lagere 3,523 bedraagt. Beide schalen lopen van 1 (positief) tot 7 (negatief), wat betekent dat respondenten gemiddeld genomen een negatievere houding hebben ten opzichte van inkomensongelijkheid dan van immigranten. Dit verschil is opvallend, vooral gezien het maatschappelijke debat waarin immigratie vaak gepaard gaat met sterke of controversiële meningen. Deze bevinding suggereert dat de zorgen over inkomensongelijkheid meer universeel gedeeld worden of sociaal wenselijker zijn om te uiten dan negatieve gevoelens ten opzichte van immigranten.

Ook andere variabelen bieden interessante inzichten. Zo ligt het gemiddelde maandinkomen op ongeveer €4500, met een standaarddeviatie van ruim €2400, wat wijst op aanzienlijke inkomensverschillen binnen de steekproef. Het gemiddelde opleidingsniveau is 4,32 op een schaal van 1 tot 6, wat betekent dat een groot deel van de respondenten minstens een hbo-opleiding heeft gevolgd. De gemiddelde leeftijd ligt rond de 46,5 jaar, met een spreiding van 18 tot 67 jaar. Wat betreft stemgedrag geeft 32,2% aan op een linkse partij te hebben gestemd, tegenover 67,8% die dat niet deed.

Tabel 2: correlatiematrix tussen alle gebruikte variabelen.

	stemmen op een linkse partij	inkomen	opleidings- niveau	leeftijd	inkomens- ongelijkheid	immigranten
stemmen op een linkse partij	-					
inkomen	-0,071**	-				
opleidings- niveau	0,177**	0,203**	-			
leeftijd	-0,058**	-0,033	-0,106**	-		
inkomens- ongelijkheid	0,321**	-0,217**	-0,082**	0,009	-	
immigranten	-0,378**	-0,037	-0,243**	0,009	-0,169**	-

* significant bij $p < 0,05$, ** significant bij $p < 0,01$

In Tabel 2 is een correlatiematrix weergegeven om de bivariate verbanden tussen de belangrijkste variabelen te tonen. Hieruit blijkt dat bijna alle correlaties statistisch significant zijn, wat waarschijnlijk te wijten is aan de grote populatie (1788). Een paar correlaties zijn bijzonder opvallend, zo is er een negatieve correlatie tussen het stemmen op een linkse partij en inkomen ($r = -0,071$, $p <$

.01). Dit suggereert dat mensen met een hoger inkomen minder geneigd zijn om op een linkse partij te stemmen, al is deze correlatie relatief zwak en vallen er geen conclusies uit te trekken.

Daarnaast correleert opleidingsniveau positief met stemmen op een linkse partij ($r = 0,177, p < .01$) en met inkomen ($r = 0,203, p < .01$), maar negatief met houding ten opzichte van immigranten ($r = -0,243, p < .01$), wat suggereert dat hoger opgeleiden gemiddeld positiever zijn over immigratie.

Verder zijn de houding ten opzichte van inkomensongelijkheid en die ten opzichte van immigranten negatief gecorreleerd ($r = -0,169, p < .01$), wat erop wijst dat mensen die negatief denken over het ene onderwerp niet per se negatief zijn over het andere onderwerp.

Belangrijk in het licht van de onderzoeksvraag zijn vooral de correlaties tussen stemgedrag en de houdingen: een positieve correlatie met houding over inkomensongelijkheid ($r = 0,321, p < .01$) en een negatieve correlatie met houding over immigranten ($r = -0,378, p < .01$). Dit ondersteunt de verwachting dat de wijze waarop mensen denken over deze thema's hun stemgedrag beïnvloedt.

Resultaten

Modevaluatie

Om de effecten van de onafhankelijke variabelen op het stemgedrag te analyseren, zijn vier logistische regressiemodellen opgesteld. De kwaliteit van deze modellen is beoordeeld aan de hand van verschillende maatstaven: de -2 Log Likelihood (-2LL) in combinatie met de χ^2 -toetsen (Chi-kwadraat uit omnibus-tests), en de Hosmer-Lemeshow toetsen.

Het eerste model (Block 1) bevat alleen de controlevariabelen, en wordt vergeleken met een leeg nulmodel (Block 0). De daling van de -2LL naar 2133,634 is aan de hand van de bijbehorende Chi-kwadraat toets ($\Delta\chi^2 = 85,131, p < 0,01$) significant, wat aantoont dat deze controlevariabelen bijdragen aan het voorspellend vermogen van het model. Desondanks geeft de Hosmer-Lemeshow toets een significante p-waarde ($p < 0,01$), wat erop wijst dat de controlevariabelen op zichzelf geen goed voorspellend model geven.

In model 2 (Block 2) wordt de variabele voor de kijk op inkomensongelijkheid toegevoegd. De -2LL daalt verder, naar 1931,478, wat in combinatie met de Chi-kwadraat toets ($\Delta\chi^2 = 202,157, p < 0,01$) een significante daling is. Op basis hiervan valt te zeggen dat de variabele voor inkomensongelijkheid een sterke voorspellende kracht in dit model heeft, aangezien het model opnieuw significant verbetert. In tegenstelling tot model 1 geeft de Hosmer-Lemeshow toets bij dit model geen significante p-waarde ($p = 0,711$), wat betekent dat er in ieder geval niet op basis van deze toets gesteld kan worden dat dit model een zwakke voorspellende kracht heeft.

Model 3 (Block 3) is opnieuw een verbetering van het model, met een verdere verlaging van de -2LL naar 1743,828 die wederom significant is ($\Delta\chi^2 = 187,650, p < 0,01$), en ook hier geeft de Hosmer-Lemeshow toets geen significante p-waarde ($p = 0,719$), wat wederom wijst op een goed voorspellend vermogen van het model.

Deze trend zet echter niet door in model 4 (Block 4), waarbij de interactieterm wordt toegevoegd. De -2LL daalt, maar wel erg licht: naar 1741,676. De Chi-kwadraat toets van deze daling toont aan dat deze daling dan ook niet significant is ($\Delta\chi^2 = 0,343$, $p = 0,558$). Dit betekent dat het toevoegen van de interactieterm geen significante verbetering van het model geeft. Dit betekent dan ook dat de interactie tussen de houdingen ten opzichte van economische ongelijkheid en immigratie niet merkbaar bijdraagt aan het voorspellend vermogen van het model.

De data voor dit onderzoek is afkomstig uit het LISS-panel en werd verkregen via het bijbehorende LISS Data Archive. Het LISS-panel is een representatief Nederlands panel dat longitudinale surveydata verzamelt op huishoudniveau. Dit betekent dat meerdere personen uit één huishouden kunnen deelnemen aan het onderzoek.

Hierdoor bestaat de mogelijkheid dat de assumptie van onafhankelijke observaties, die ten grondslag ligt aan regressieanalyses zoals logistische regressie, wordt geschonden. Observaties van personen uit hetzelfde huishouden kunnen namelijk onderling gecorreleerd zijn door gedeelde (sociale) omstandigheden (bijv. economische status of achtergrond). Deze huishoudelijke afhankelijkheid zou kunnen leiden tot een onderschatting van de standaardfouten en daarmee tot een overschatting van statistische significantie. Om hier rekening mee te houden zijn alle belangrijke analyses nogmaals uitgevoerd na het selecteren van één persoon per huishouden en vergeleken met de originele resultaten. Bijlage 3 bevat een verdere uitleg en de resultaten van deze nieuwe analyses, maar de belangrijkste bevinding is dat alle gemaakte conclusies op basis van de originele resultaten nog steeds opgaan, en er vallen geen nieuwe belangrijke conclusies te trekken.

Ook is er gecontroleerd op multicollineariteit door het berekenen van de VIF-scores. Deze scores liggen allemaal ruim onder de gebruikelijke grenswaarden ($\max = 1,308$), wat erop wijst dat de onafhankelijke variabelen geen problematische overlap vertonen. De variabelen kunnen dus betrouwbaar naast elkaar worden opgenomen in de modellen.

Wat uitbijters betreft zijn er een aantal punten met een relatief hoge Cook's Distance, maar deze worden op basis van nadere inspectie niet als problematisch beschouwd en dus zijn er geen outliers verwijderd (zie bijlage 3 voor meer info).

Hypothesetoetsing

Tabel 4: regressietabel voor de uitgevoerde logistische regressie met de dummy voor het niet of wel stemmen op een linkse partij als afhankelijke variabele.

	Model 1		Model 2		Model 3		Model 4		VIF
	<i>b</i>	<i>SE</i>	<i>b</i>	<i>SE</i>	<i>b</i>	<i>SE</i>	<i>b</i>	<i>SE</i>	
Constante	-0,850**	0,053	-0,960**	0,059	-1,022**	-0,063	-1,018**	0,063	1,322
Inkomen	-0,113**	0,024	-0,051*	0,026	-0,068*	0,028	-0,068*	0,028	1,094
Opleidingsniveau	0,354**	0,045	0,395**	0,046	-0,246**	0,048	0,244**	0,049	1,156
Leeftijd	-0,006	0,004	-0,008*	0,004	-0,009*	0,004	-0,009*	-0,004	1,017
Houding ten opzichte van economische ongelijkheid			0,599**	0,046	0,540**	0,048	0,532**	0,050	1,204
Houding ten opzichte van immigranten					-0,929**	0,075	-0,915**	0,078	1,228
Economische ongelijkheid * immigranten							-0,033	0,055	1,039
ΔChi-kwadraat	85,131**		202,157**		187,650**		0,343		
-2LL	2133,634		1931,478		1743,828		1743,484		
<i>N</i>	1793		1793		1793		1793		

* significant bij $p < 0,05$, ** significant bij $p < 0,01$

De hypothesen worden getoetst aan de hand van de uitgevoerde logistische regressie, weergegeven in tabel 4. Model 1 bevat alleen de controlevariabelen. De coëfficiënten voor zowel inkomen als opleidingsniveau zijn significant, wat wijst op een negatief verband tussen de hoogte van het inkomen en de kans op het stemmen op een linkse partij, en een positief verband tussen het opleidingsniveau en de kans op het stemmen op een linkse partij. De coëfficiënt van leeftijd wordt pas in de modellen waarbij de houding ten opzichte van economische ongelijkheid wordt toegevoegd significant. Dit wijst op een suppressie-effect: er bestaat waarschijnlijk een zekere mate van overlap in variantie tussen inkomen en de kijk op economische ongelijkheid die vanaf model 2 wordt gecorrigeerd.

Hypothese 1: Een negatieve houding over inkomensongelijkheid leidt tot grotere kans op stemmen op een linkse partij.

In model 2 wordt de variabele houding ten opzichte van inkomensongelijkheid toegevoegd. De coëfficiënt voor deze variabele is positief en significant ($b = 0,600, p < .01$), wat betekent dat een negatievere houding over inkomensongelijkheid gepaard gaat met een grotere kans op stemmen op een linkse partij.

De logistische regressie geeft ons coëfficiënten in log-odds, die lastig te interpreteren zijn zonder verdere berekeningen. Daarom zijn in Tabel 5 kansen berekend op basis van extreme waarden. Bij gemiddelde waarden voor alle voorspellers (gestandaardiseerde waarde = 0), is de kans op stemmen op een linkse partij 26,48%. Bij een maximale negatieve houding over inkomensongelijkheid (waarde = 2,6), stijgt deze kans tot 81,74%. Dit is een substantiële stijging en lijkt dus de hypothese te ondersteunen. Deze stijgingen zijn ook te zien in model 3 (24,79% naar 79,49%) en in model 4 (27,01% naar 72,65%). Een kritische kijk op deze stijgingen is echter belangrijk, omdat dit kansen zijn op basis van extreme situaties en niet per se systematische werkingen. Hierom kan alleen de significante helling worden gezien als bewijs van de hypothese.

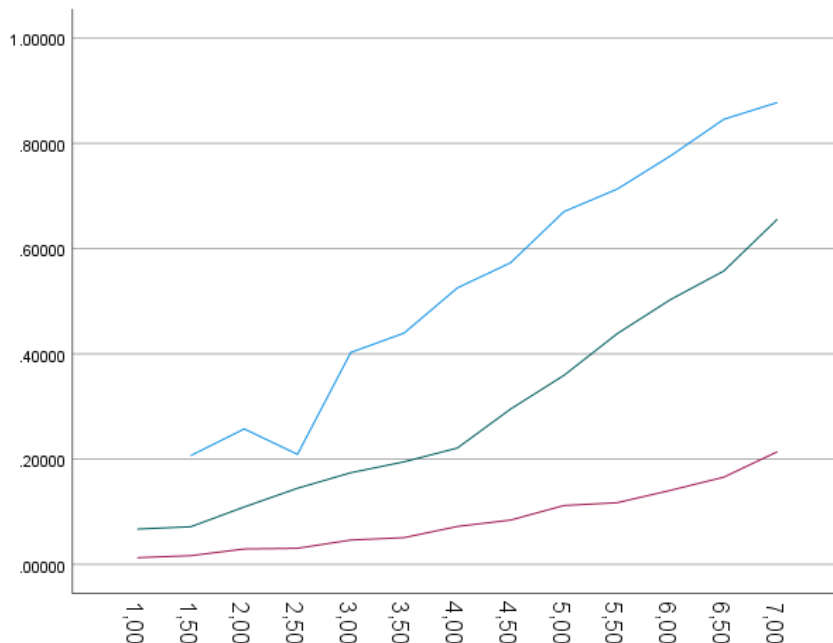
Hypothese 2: De relatie tussen houding over inkomensongelijkheid en stemgedrag is zwakker bij negatieve houding tegenover immigranten.

Model 4 test het modererende effect door een interactieterm toe te voegen tussen de houdingen over inkomensongelijkheid en immigranten. De coëfficiënt van deze interactie is -0,065 en niet significant, wat betekent dat er statistisch gezien geen bewijs is voor moderatie. Deze hypothese wordt dus niet ondersteund door de data en kan niet worden aangenomen.

Toch zijn er aanwijzingen in de data dat er wel sprake zou kunnen zijn van interactie. Bij kansberekeningen zien we dat de kans op stemmen op een linkse partij bij een maximale negatieve houding over inkomensongelijkheid 59,03% is (model 4), maar als dit gecombineerd wordt met een maximale negatieve houding over immigranten daalt de kans tot 5,30%. Dit duidt op een mogelijk modererend effect dat statistisch niet overtuigend blijkt uit de regressieanalyse. Hoewel er dus geen statistisch bewijs is voor een interactie-effect, lijkt de kijk op immigranten wel degelijk invloed te hebben op het effect van de kijk op ongelijkheid op het wel of niet stemmen op een linkse partij.

Om dit visueel inzichtelijk te maken is Figuur 3 opgesteld. Respondenten zijn ingedeeld in drie groepen op basis van hun houding tegenover immigranten (positief, gemiddeld, negatief). Voor elke groep is het effect van houding over inkomensongelijkheid op het wel of niet stemmen op een linkse partij weergegeven. De grafiek laat zien dat het effect van inkomensongelijkheid op stemgedrag sterk is bij respondenten met een positieve of gemiddelde houding tegenover immigranten, maar grotendeels afwezig bij respondenten met een negatieve houding. Dit lijkt, ondanks het ontbreken van statistisch bewijs, toch op een vorm van een moderatie-effect.

Figuur 3: het effect van de houding ten opzichte van inkomensongelijkheid (x-as) op de voorspelde waarde van het niet of wel stemmen op een linkse partij (y-as). Blauwe lijn = positief, groene lijn = gemiddeld, rode lijn = negatief.



Tabel 5 geeft verdere detaillering van hoe de kans op stemmen op een linkse partij verandert bij extreme waarden van verschillende voorspellers. Zo laat de tabel bijvoorbeeld zien dat een maximaal opleidingsniveau de kans verhoogt tot 35,25% (model 4), terwijl het hoogste inkomen uit de dataset deze verlaagt tot 5,80%. Relevant voor dit onderzoek is dat een maximale score op ongelijkheid, dus een maximaal negatieve houding ten opzichte van economische ongelijkheid, de kans op het stemmen op een linkse partij vergroot tot 59,03%. Opvallend zijn de kansen voor een maximale score op alleen kijk op immigranten en die voor een maximale score op zowel kijk op ongelijkheid als immigranten. Hoewel de kans enigszins toeneemt (van 1,38% naar 5,30%) is dit nog steeds een vrij lage kans, zeker vergeleken met de kans bij alleen een maximale score op ongelijkheid in dit model (59,03%). Hoewel het interactie-effect niet significant was, wijzen dus zowel de kansberekeningen als de visuele analyse op relevante verschillen tussen groepen.

Tabel 5: Kansberekening op het stemmen op een linkse partij op basis van bepaalde maximale waarden aan de hand van de log-odds.

Ingevoerde scores	Model 2		Model 3		Model 4	
	Log-odds	Kans (%)	Log-odds	Kans (%)	Log-odds	Kans (%)
Alles = 0	-0,960	26,46%	-1,022	24,79%	-1,018	26,54%
Alles = 0, inkomen = max	-2,287	9,22%	-2,791	5,78%	-2,787	5,80%
Alles = 0, opleidingsniveau = max	-0,296	42,64%	-1,436	19,23%	-0,608	35,25%
Alles = 0, leeftijd = max	-0,796	31,08%	-1,206	23,04%	-1,100	24,98%
Alles = 0, ongelijkheid = max	1,506	81,85%	1,336	79,18%	0,365	59,03%
Alles = 0, immigranten = max			-4,320	1,31%	-4,266	1,38%
Alles = 0, ongelijkheid & immigranten = max			-2,916	5,14%	-2,883	5,30%

Samenvattend bieden de resultaten overtuigend bewijs voor de eerste hypothese en lijkt de theorie ondersteund te worden: houdingen ten opzichte van inkomensongelijkheid en immigranten zijn beide significante voorspellers van het stemmen op een linkse partij. De tweede hypothese, die een modererend effect veronderstelt, wordt niet bevestigd met behulp van een interactieterm in de regressieanalyse en kan dus niet worden aangenomen. Desondanks lijkt een soortgelijk interactie-effect wel te bestaan op basis van de berekende kansen bij extreme situaties, maar dit kan niet als bewijs voor een effect worden gezien.

Conclusie

Dit onderzoek had als doel om te analyseren hoe de houding ten opzichte van economische ongelijkheid samenhangt met stemmen op linkse partijen, en in hoeverre deze relatie beïnvloed wordt door de houding ten opzichte van immigranten. Op basis van data uit het LISS-panel (N = 1793) is aangetoond dat beide houdingen afzonderlijk een duidelijke rol spelen. Mensen die economische ongelijkheid als problematisch ervaren, stemmen vaker links. Tegelijkertijd stemmen mensen met een negatievere kijk op immigratie juist minder vaak op een linkse partij. Dit bevestigt eerdere theorieën waarin politieke keuzes voortkomen uit zowel economische als sociaal-culturele overwegingen (Mijs & Savage, 2020; Macdonald, 2020; Scheepers et al., 2002).

Andere benaderingen, zoals kansberekeningen en groepsindelingen, lijken wel een mogelijk inhoudelijk relevant patroon te tonen. Bij respondenten met een negatieve houding ten opzichte van immigratie lijkt het verband tussen economische onvrede en links stemgedrag zwakker of zelfs afwezig. Bij mensen met een neutrale of positieve kijk op immigranten is dit verband juist sterker. Hoewel dit patroon niet statistisch significant is, lijkt het wel het theoretische idee te steunen dat culturele overtuigingen (zoals houding ten opzichte van immigratie) invloed kunnen hebben op het verband tussen de houding over ongelijkheid en steun voor linkse partijen. (Roex et al., 2018; Ares et al., 2017).

Een verklaring voor het uitblijven van een significant interactie-effect kan liggen in de complexiteit van politiek gedrag. Mogelijk is de invloed van de houding ten opzichte van immigranten vooral zichtbaar bij respondenten met zeer uitgesproken standpunten, die in een statistisch model niet goed worden weergegeven. Toekomstig onderzoek kan dit verder verkennen met methoden die dergelijke complexe patronen beter kunnen weergeven.

Discussie

Dit onderzoek kent wel enkele beperkingen. Politieke voorkeur werd gemeten via stemgedrag bij de vorige verkiezingen, waarbij de indeling in links en niet-links gebaseerd is op een indeling die altijd subjectief zal zijn. Het is mogelijk dat kiezers partijen anders indelen, bijvoorbeeld op basis van specifieke standpunten of politici. Ook blijft de vragenlijst gevoelig voor sociaal wenselijke antwoorden, zoals bij de vragen over ongelijkheid en vooral de vragen over immigratie.

Verder is er in dit onderzoek gecontroleerd voor opleidingsniveau, leeftijd en inkomen, maar andere variabelen als politieke betrokkenheid, invloed van de media of het vertrouwen in de democratie zijn niet meegenomen. Deze kunnen relevant zijn voor een vollediger beeld van stemgedrag. Ten slotte blijft de keuze voor partijen een complexe keuze die niet altijd een directe weerspiegeling is van houdingen, maar ook beïnvloed kan worden door bijvoorbeeld strategisch stemmen.

Samenvattend laat dit onderzoek zien dat een negatieve kijk op economische ongelijkheid op zichzelf niet in staat is om steun voor linkse partijen in zijn geheel te verklaren. Culturele opvattingen, in het bijzonder over immigratie, spelen een belangrijke rol in hoe economische onvrede politiek wordt geuit. Hoewel dit onderzoek geen statistisch bewijs levert voor een interactie-effect, wijzen de inhoudelijke patronen wel op een potentiële invloed van immigratieperceptie op deze relatie. De enige conclusie die getrokken kan worden is dus weliswaar dat er geen bewijs is gevonden voor dit moderatie-effect, maar inhoudelijk interessante patronen zijn er wel.

Literatuur

- Barr, A., & Miller, L. (2020). The effect of education, income inequality and merit on inequality acceptance. *Journal of Economic Psychology*, 80, 102276.
<https://doi.org/10.1016/j.joep.2020.102276>
- Delhey, J., & Dragolov, G. (2013). Why inequality makes Europeans less happy: the role of distrust, status anxiety, and perceived conflict. *European Sociological Review*, 30(2), 151–165.
<https://doi.org/10.1093/esr/jct033>
- Elchardus, M., & Spruyt, B. (2012). Populisme en de zorg over de samenleving. *Sociologie*, 8(1). <https://doi.org/10.5117/soc2012.1.elch>
- Engelhardt, C., & Wagener, A. (2017). What do Germans think and know about income inequality? A survey experiment. *Socio-Economic Review*, 16(4), 743–767.
<https://doi.org/10.1093/ser/mwx036>
- Fox, C. (2004). The Changing Color of Welfare? How Whites' Attitudes toward Latinos Influence Support for Welfare. *American Journal of Sociology*, 110(3), 580–625.
<https://doi.org/10.1086/422587>
- Franko, W. W., & Witko, C. (2022). Unions, class identification, and policy attitudes. *The Journal of Politics*, 85(2), 553–567. <https://doi.org/10.1086/722347>
- Gidron, N., & Hall, P. A. (2017). The politics of social status: economic and cultural roots of the populist right. *British Journal of Sociology*, 68(S1). <https://doi.org/10.1111/1468-4446.12319>

- Jost, J. T., Banaji, M. R., & Nosek, B. A. (2004). A decade of system justification Theory: accumulated evidence of conscious and unconscious bolstering of the status quo. *Political Psychology*, 25(6), 881–919. <https://doi.org/10.1111/j.1467-9221.2004.00402.x>
- Lepianka, D., Gelissen, J., & Van Oorschot, W. (2010). Popular explanations of poverty in Europe. *Acta Sociologica*, 53(1), 53–72. <https://doi.org/10.1177/0001699309357842>
- Kenworthy, L., & McCall, L. (2007). Inequality, public opinion and redistribution. *Socio-Economic Review*, 6(1), 35–68. <https://doi.org/10.1093/ser/mwm006>
- Macdonald, D. (2020). Immigration attitudes and white Americans' responsiveness to rising income inequality. *American Politics Research*, 49(2), 132–142. <https://doi.org/10.1177/1532673x20972104>
- Makhlouf, Y. (2022). Trends in Income Inequality: Evidence from Developing and Developed Countries. *Social Indicators Research*, 165(1), 213–243. <https://doi.org/10.1007/s11205-022-03010-8>
- McVeigh, R., Beyerlein, K., Vann, B., & Trivedi, P. (2014). Educational Segregation, Tea Party Organizations, and Battles over Distributive Justice. *American Sociological Review*, 79(4), 630–652. <https://doi.org/10.1177/0003122414534065>
- Mijs, J. J. B. (2015). The Unfulfillable Promise of Meritocracy: Three lessons and their implications for Justice in education. *Social Justice Research*, 29(1), 14–34. <https://doi.org/10.1007/s11211-014-0228-0>
- Mijs, J. J., & Savage, M. (2020). Meritocracy, elitism and inequality. *The Political Quarterly*, 91(2), 397–404. <https://doi.org/10.1111/1467-923x.12828>

- Oorschot, W. V., & Halman, L. (2000). BLAME OR FATE, INDIVIDUAL OR SOCIAL? *European Societies*, 2(1), 1–28. <https://doi.org/10.1080/146166900360701>
- Petrova, B., & Ranaldi, M. (2024). Determinants of income composition inequality. *Comparative Politics*, 56(4), 449–471. <https://doi.org/10.5129/001041524x17045519842791>
- Reyes, G., & Gasparini, L. (2022). Are fairness perceptions shaped by income inequality? evidence from Latin America. *The Journal of Economic Inequality*, 20(4), 893–913. <https://doi.org/10.1007/s10888-022-09526-w>
- Rekker, R. (2024). Electoral change through generational replacement: An age-period-cohort analysis of vote choice across 21 countries between 1948 and 2021. *Frontiers in Political Science*, 6. <https://doi.org/10.3389/fpos.2024.1279888>
- Roex, K. L., Huijts, T., & Sieben, I. (2018b). Attitudes towards income inequality: ‘Winners’ versus ‘losers’ of the perceived meritocracy. *Acta Sociologica*, 62(1), 47–63. <https://doi.org/10.1177/0001699317748340>
- Scheepers, P., Gijsberts, M., & Coenders, M. (2002). Ethnic Exclusionism in European Countries. Public Opposition to Civil Rights for Legal Migrants as a Response to Perceived Ethnic Threat. *European Sociological Review*, 18(1), 17–34. <https://doi.org/10.1093/esr/18.1.17>
- Semyonov, M., Raijman, R., & Gorodzeisky, A. (2006b). The rise of anti-foreigner sentiment in European societies, 1988-2000. *American Sociological Review*, 71(3), 426–449. <https://doi.org/10.1177/000312240607100304>

- Smith, A. (1776). *An Inquiry into the Nature and Causes of the Wealth of Nations*. In *Oxford University Press eBooks*. <https://doi.org/10.1093/oseo/instance.00043218>
- Stewart, A. J., McCarty, N., & Bryson, J. J. (2020). Polarization under rising inequality and economic decline. *Science Advances*, 6(50). <https://doi.org/10.1126/sciadv.abd4201>
- Tajfel, H. (1982). Social Psychology of intergroup Relations. *Annual Review of Psychology*, 33(1), 1–39. <https://doi.org/10.1146/annurev.ps.33.020182.000245>
- Tropp, L. R., & Barlow, F. K. (2018). Making advantaged racial groups care about inequality: Intergroup contact as a route to psychological investment. *Current Directions in Psychological Science*, 27(3), 194–199. <https://doi.org/10.1177/0963721417743282>
- Vasilopoulos, P., & Jost, J. T. (2020). Psychological similarities and dissimilarities between left-wing and right-wing populists: Evidence from a nationally representative survey in France. *Journal of Research in Personality*, 88, 104004. <https://doi.org/10.1016/j.jrp.2020.104004>
- Young, M. D. (1958). *The rise of the meritocracy*. <https://ci.nii.ac.jp/ncid/BA22002648>

Bijlagen

Omdat het operationaliseren van de variabelen en het uitvoeren van de analyses chronologisch een beetje door elkaar heen loopt, zijn bijlagen 1 en 2 samengevoegd tot een grote bijlage. Het exporteren van een SPSS-output naar een Word-bestand en deze openen in Google Docs werkte niet helemaal, dus deze uiteindelijke bijlage is redelijk aangepast om leesbaar te kunnen zijn. Belangrijk is vooral dat de tabel met ‘notes’ steeds is verwijderd om ruimte te besparen, verder is alle relevante informatie uit de SPSS-output behouden.

Bijlage 1 & 2

Beschrijving originele variabelen

Dit is de verdeling van (bijna) alle originele, onbewerkte variabelen die gebruikt zijn in dit onderzoek.

Voor het inzicht in de variabelen zijn ook histogrammen gemaakt, maar deze zijn van verwaarloosbaar belang voor deze bijlage en zijn dus niet bijgevoegd.

FREQUENCIES VARIABLES= ineq_1 ineq_5 etnocentr1 etnocentr2 etnocentr4 marxmigr1

marxmigr4

marxmigr5 nettohh_f oplcat leeftijd cv24p307

/FORMAT=NOTABLE

/NTILES=4

/STATISTICS=STDDEV MINIMUM MAXIMUM MEAN

/HISTOGRAM

/ORDER=ANALYSIS.

Statistics													
De verschillen in inkomens zijn in Nederland te groot.		Ik maak me kwaad over verschillen in rijkdom tussen de rijken en de armen.	Het culturele leven in Nederland wordt over het algemeen verrijkt (beter gemaakt)	Nederland is een betere plek geworden door mensen uit andere landen die hier zijn	Nederland had eigenlijk nooit gastarbeiders binnen moeten laten.	Mensen met een migratieachtergrond nemen de banen in van mensen zonder een migra	De belangen van werknemers zonder een migratieachtergrond zijn belangrijker dan	De meeste mensen met een migratieachtergrond zijn onderdeel van een lagere socia	Netto inkomen van het huishouden, geschat	Opleiding in CBS-categorieën	Leeftijd van lid van huishouden	For which party did you vote in the parliamentary elections of 22 November 2023?	
N	Valid	2694	2694	2692	2692	2692	2688	2688	2688	2418	2719	2719	2119
	Missing	25	25	27	27	27	31	31	31	301	0	0	600
Mean		4.90	4.00	4.29	3.93	3.11	3.14	2.86	4.35	4326.79	4.18	45.08	9.52
Std. Deviation		1.403	1.687	1.579	1.568	1.620	1.457	1.481	1.367	2484.519	1.393	14.140	8.738
Minimum		1	1	1	1	1	1	1	1	0	1	18	-9
Maximum		7	7	7	7	7	7	7	7	30509	6	67	21
Percentiles	25	4.00	3.00	3.00	3.00	2.00	2.00	2.00	4.00	2650.00	3.00	33.00	2.00
	50	5.00	4.00	4.00	4.00	3.00	3.00	2.50	4.00	4100.00	4.00	46.00	6.00
	75	6.00	5.00	5.00	5.00	4.00	4.00	4.00	5.00	5600.00	5.00	58.00	20.00

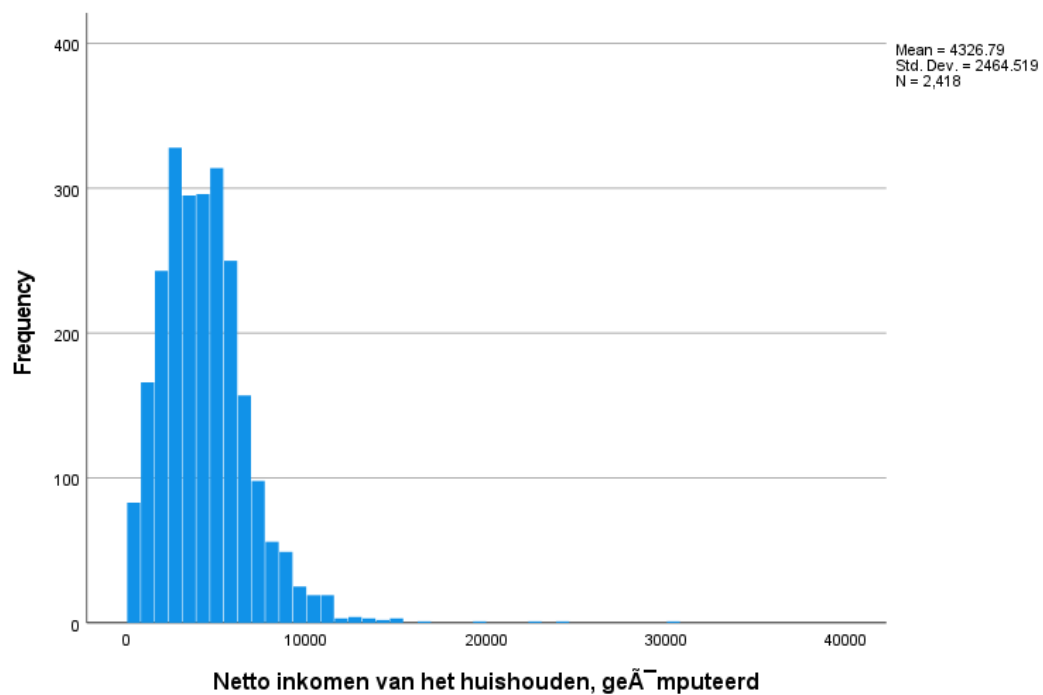
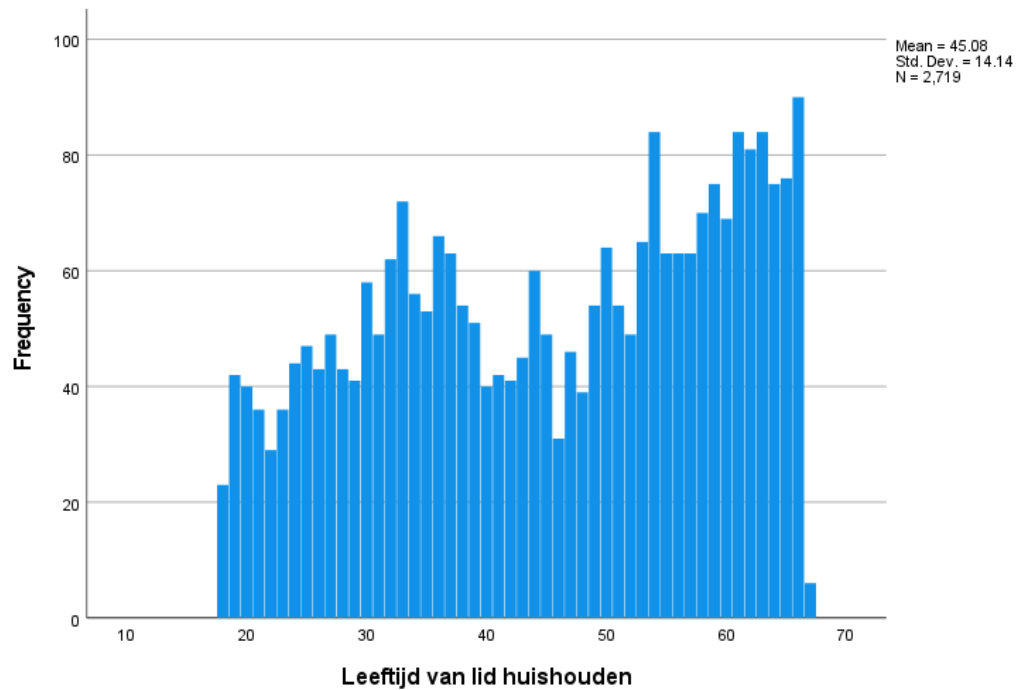
Twee variabelen staan niet in deze tabel: leeftijd en inkomen. Dit omdat deze variabelen te veel antwoordmogelijkheden hebben om zinvol in een tabel weergegeven te kunnen worden. Daarom is voor deze variabelen apart een histogram gemaakt.

GRAPH

/HISTOGRAM=leeftijd.

GRAPH

/HISTOGRAM=nettohh_f.



Aanpassingen aan de variabelen

Bij de variabelen 'etnocentr1' en 'etnocentr2' betekende, in tegenstelling tot de rest van de vragen over immigratie, een hogere score een positievere kijk op immigranten. Om alle variabelen samen te kunnen gebruiken moeten deze eerst worden omgedraaid.

```
COMPUTE entocentr1_flipped=8 - etnocentr1.
```

```
VARIABLE LABELS entocentr1_flipped 'entocentr 1 maar dan omgedraaid'.
```

```
EXECUTE.
```

```
COMPUTE entocentr2_flipped=8 - etnocentr2.
```

```
VARIABLE LABELS entocentr2_flipped 'entocentr 2 maar dan omgedraaid'.
```

```
EXECUTE.
```

Voor het samenvoegen van de gebruikte variabelen is de Cronbach's Alpha van de groepen variabelen berekend.

```
RELIABILITY
```

```
  /VARIABLES=ineq_1 ineq_5
```

```
  /SCALE('ALL VARIABLES') ALL
```

```
  /MODEL=ALPHA.
```

```
RELIABILITY
```

```
  /VARIABLES=etnocentr4 marx migr1 marx migr4 marx migr5 entocentr1_flipped  
entocentr2_flipped
```

```
  /SCALE('ALL VARIABLES') ALL
```

```
  /MODEL=ALPHA.
```

Case Processing Summary

		N	%
Cases	Valid	2694	99.1
	Excluded ^a	25	.9
	Total	2719	100.0

a. Listwise deletion based on all variables in the procedure.

Reliability Statistics

Cronbach's Alpha	N of Items
.709	2

Case Processing Summary

		N	%
Cases	Valid	2688	98.9
	Excluded ^a	31	1.1
	Total	2719	100.0

a. Listwise deletion based on all variables in the procedure.

Reliability Statistics

Cronbach's Alpha	N of Items
.730	6

Vervolgens zijn de relevante variabelen samengevoegd tot de uiteindelijk gebruikte variabelen.

```
COMPUTE inkomensongelijkheid=(ineq_1 + ineq_5) / 2.
```

```
VARIABLE LABELS inkomensongelijkheid 'houding tenopzichte van inkomensongelijkheid'.
```

```
EXECUTE.
```

```
COMPUTE immigranten=(etnocentr4 + marx migr1 + marx migr4 + marx migr5 +  
entocentr1_flipped +
```

entocentr2_flipped) / 6.

VARIABLE LABELS immigranten 'houding ten opzichte van immigranten'.

EXECUTE.

RECODE cv24p307 (6=1) (9=1) (12=0) (16=0) (20=1) (1=0) (2=0) (3=0) (4=0) (8=0) (11=0)
(13=0) (14=0)

(17=0) (18=0) (21=0) (15=SYSMIS) INTO partij_dummy.

VARIABLE LABELS partij_dummy 'dummy voor het stemmen op een niet of wel linkse
partij'.

EXECUTE.

De interactieterm is ook hier gemaakt.

COMPUTE interactie=inkomensongelijkheid * immigranten.

VARIABLE LABELS interactie 'Interactievariabele tussen kijk op ongelijkheid en kijk op '+
'immigranten'.

EXECUTE.

De variabele voor inkomen is voor de interpreteerbaarheid verkleind door hem te delen door
1000.

COMPUTE huishoudinkomen_verkleind=nettohh_f / 1000.

VARIABLE LABELS huishoudinkomen_verkleind 'Interactievariabele tussen kijk op
ongelijkheid en '+
'kijk op immigranten'.

EXECUTE.

Tot slot zijn alle variabelen gecentreerd.

```
COMPUTE inkomen_verkleind_gecentreerd=huishoudinkomen_verkleind - 4.50358.  
VARIABLE LABELS inkomen_gecentreerd 'Gecentreerde variabele voor het verkleind  
inkomen van het huishouden'.  
EXECUTE.
```

```
COMPUTE opleidingsniveau_gecentreerd=oplcat - 4.32.  
VARIABLE LABELS opleidingsniveau_gecentreerd 'Gecentreerde variabele voor  
opleidingsniveau'.  
EXECUTE.
```

```
COMPUTE leeftijd_gecentreerd=leeftijd - 46.53.  
VARIABLE LABELS leeftijd_gecentreerd 'Gecentreerde variabele voor leeftijd'.  
EXECUTE.
```

```
COMPUTE inkomensongelijkheid_gecentreerd=inkomensongelijkheid - 4.404.  
VARIABLE LABELS inkomensongelijkheid_gecentreerd 'Gecentreerde variabele voor de  
houding ten '+  
    'opzichte van inkomensongelijkheid'.  
EXECUTE.
```

```
COMPUTE immigranten_gecentreerd=immigranten - 3.454.  
VARIABLE LABELS immigranten_gecentreerd 'Gecentreerde variabele voor de houding ten  
opzichte '+  
    'van immigranten'.  
EXECUTE.
```

```
COMPUTE interactie=inkomensongelijkheid_gecentreerd * immigranten_gecentreerd.
```

VARIABLE LABELS interactie 'Interactievariabele tussen de gecentreerde kijk op ongelijkheid en '+'

'kijk op immigranten'.

EXECUTE.

Onderzoek

Voor het onderzoek zijn alleen de 'complete' cases geselecteerd en gebruikt.

LOGISTIC REGRESSION VARIABLES partij_dummy

/METHOD=ENTER leeftijd oplcat nettohh_f inkomensongelijkheid immigranten

/SAVE=RESID

/CRITERIA=PIN(.05) POUT(.10) ITERATE(20) CUT(.5).

RECODE RES_1 (SYSMIS=0) (ELSE=1) INTO dummy_observaties.

VARIABLE LABELS dummy_observaties 'dummy voor of een observatie niet of wel compleet '+'

'geobserveerd is'.

EXECUTE.

USE ALL.

COMPUTE filter_\$=(dummy_observaties = 1).

VARIABLE LABELS filter_\$ 'dummy_observaties = 1 (FILTER)'.

VALUE LABELS filter_\$ 0 'Not Selected' 1 'Selected'.

FORMATS filter_\$ (f1.0).

FILTER BY filter_\$.

EXECUTE.

Vervolgens zijn de simpele correlaties en de descriptieve statistieken berekend.

CORRELATIONS

/VARIABLES=opleidingsniveau_gecentreerd leeftijd_gecentreerd

inkomensongelijkheid_gecentreerd

immigranten_gecentreerd inkomen_verkleind_gecentreerd partij_dummy

/PRINT=TWOTAIL NOSIG FULL

/MISSING=PAIRWISE.

DESCRIPTIVES VARIABLES=leeftijd nettohh_f oplcat inkomensongelijkheid immigranten

/STATISTICS=MEAN STDDEV MIN MAX.

Correlations							
		Gecentreerde variabele voor opleidingsnive au	Gecentreerde variabele voor leeftijd	Gecentreerde variabele voor de houding ten opzichte van inkomensonge lijkheid	Gecentreerde variabele voor de houding ten opzichte van immigranten	Gecentreerde variabele voor het verkleind inkomen van het huishouden	dummy voor het stemmen op een niet of wel linkse partij
Gecentreerde variabele voor opleidingsniveau	Pearson Correlation	1	-.106**	-.082**	-.243**	.203**	.177**
	Sig. (2-tailed)		<.001	<.001	<.001	<.001	<.001
	N	1793	1793	1793	1793	1793	1793
Gecentreerde variabele voor leeftijd	Pearson Correlation	-.106**	1	.031	.009	-.033	-.058*
	Sig. (2-tailed)	<.001		.195	.713	.168	.014
	N	1793	1793	1793	1793	1793	1793
Gecentreerde variabele voor de houding ten opzichte van inkomensongelijkheid	Pearson Correlation	-.082**	.031	1	-.169**	-.217**	.321**
	Sig. (2-tailed)	<.001	.195		<.001	<.001	<.001
	N	1793	1793	1793	1793	1793	1793
Gecentreerde variabele voor de houding ten opzichte van immigranten	Pearson Correlation	-.243**	.009	-.169**	1	-.037	-.378**
	Sig. (2-tailed)	<.001	.713	<.001		.118	<.001
	N	1793	1793	1793	1793	1793	1793
Gecentreerde variabele voor het verkleind inkomen van het huishouden	Pearson Correlation	.203**	-.033	-.217**	-.037	1	-.071**
	Sig. (2-tailed)	<.001	.168	<.001	.118		.003
	N	1793	1793	1793	1793	1793	1793
dummy voor het stemmen op een niet of wel linkse partij	Pearson Correlation	.177**	-.058*	.321**	-.378**	-.071**	1
	Sig. (2-tailed)	<.001	.014	<.001	<.001	.003	
	N	1793	1793	1793	1793	1793	1793

** . Correlation is significant at the 0.01 level (2-tailed).

* . Correlation is significant at the 0.05 level (2-tailed).

Descriptive Statistics

	N	Minimum	Maximum	Mean	Std. Deviation
Leeftijd van lid huishouden	1793	18	67	46.51	14.016
Netto inkomen van het huishouden, geïmputeerd	1793	0	30509	4498.63	2411.740
Opleiding in CBS-categorieën	1793	1	6	4.32	1.313
houding ten opzichte van inkomensongelijkheid	1793	1.00	7.00	4.4046	1.39138
houding ten opzichte van immigranten	1793	1.00	7.00	3.5227	.98025
Valid N (listwise)	1793				

Statistics

		Gecentreerde variabele voor opleidingsniveau	Gecentreerde variabele voor leeftijd	Gecentreerde variabele voor de houding ten opzichte van inkomensongelijkheid	Gecentreerde variabele voor de houding ten opzichte van immigranten	Gecentreerde variabele voor het verkleind inkomen van het huishouden
N	Valid	1793	1793	1793	1793	1793
	Missing	0	0	0	0	0
Mean		-.0043	-.0247	.0006	.0687	-.0050
Std. Deviation		1.31257	14.01574	1.39138	.98025	2.41174
Minimum		-3.32	-28.53	-3.40	-2.45	-4.50
Maximum		1.68	20.47	2.60	3.55	26.01
Percentiles	25	-.3200	-12.5300	-.9040	-.6207	-1.6936
	50	-.3200	2.4700	.0960	.0460	-.2036
	75	.6800	12.4700	1.0960	.7127	1.2464

Voor het verkrijgen van de VIF-scores is een lineaire regressie uitgevoerd.

REGRESSION

/MISSING LISTWISE

/STATISTICS COLLIN TOL

/CRITERIA=PIN(.05) POUT(.10)

/NOORIGIN

/DEPENDENT gebjaar

/METHOD=ENTER opleidingsniveau_gecentreerd leeftijd_gecentreerd

inkomensongelijkheid_gecentreerd

immigranten_gecentreerd inkomen_verkleind_gecentreerd partij_dummy.

Coefficients^a

Model		Collinearity Statistics	
		Tolerance	VIF
1	Gecentreerde variabele voor opleidingsniveau	.865	1.156
	Gecentreerde variabele voor leeftijd	.983	1.017
	Gecentreerde variabele voor de houding ten opzichte van inkomensongelijkheid	.831	1.204
	Gecentreerde variabele voor de houding ten opzichte van immigranten	.815	1.228
	Gecentreerde variabele voor het verkleind inkomen van het huishouden	.914	1.094
	Interactievariabele tussen de gecentreerde kijk op ongelijkheid en kijk op immigranten	.963	1.039
	dummy voor het stemmen op een niet of wel linkse partij	.757	1.322

a. Dependent Variable: Geboortejaar

Vervolgens is hier de logistische regressie uitgevoerd.

LOGISTIC REGRESSION VARIABLES partij_dummy

/METHOD=ENTER inkomen_verkleind_gecentreerd leeftijd_gecentreerd
opleidingsniveau_gecentreerd

/METHOD=ENTER opleidingsniveau_gecentreerd leeftijd_gecentreerd
inkomen_verkleind_gecentreerd inkomensongelijkheid_gecentreerd

/METHOD=ENTER inkomen_verkleind_gecentreerd immigranten_gecentreerd
inkomensongelijkheid_gecentreerd opleidingsniveau_gecentreerd leeftijd_gecentreerd

/METHOD=ENTER inkomensongelijkheid_gecentreerd immigranten_gecentreerd interactie
inkomen_verkleind_gecentreerd opleidingsniveau_gecentreerd
leeftijd_gecentreerd

/SAVE=PRED COOK LEVER DFBETA

/PRINT=GOODFIT

/CRITERIA=PIN(0.05) POUT(0.10) ITERATE(20) CUT(0.5).

Logistic Regression

Notes		
Output Created		30-JUL-2025 16:37:29
Comments		
Input	Data	C:\Users\vxia\Desktop\liss_baw_class.sav
	Active Dataset	DataSet5
	Filter	dummy_observaties = 1 (FILTER)
	Weight	<none>
	Split File	<none>
	N of Rows in Working Data File	1793
Missing Value Handling	Definition of Missing	User-defined missing values are treated as missing
Syntax		LOGISTIC REGRESSION VARIABLES partij_dummy /METHOD=ENTER inkomen_verkleind_g ecentreerd leeftijd_gecentreerd opleidingsniveau_gec entreerd /METHOD=ENTER opleidingsniveau_gec

		entreeerd leeftijd_gecentreerd inkomen_verkleind_g ecentreerd inkomensongelijkheid _gecentreerd /METHOD=ENTER inkomen_verkleind_g ecentreerd immigranten_gecentr eerd inkomensongelijkheid _gecentreerd opleidingsniveau_gec entreerd leeftijd_gecentreerd /METHOD=ENTER inkomensongelijkheid _gecentreerd immigranten_gecentr eerd interactie inkomen_verkleind_g ecentreerd opleidingsniveau_gec entreerd leeftijd_gecentreerd /SAVE=PRED COOK LEVER DFBETA /PRINT=GOODFIT /CRITERIA=PIN(0.05) POUT(0.10) ITERATE(20) CUT(0.5).
Resources	Processor Time	00:00:00,02
	Elapsed Time	00:00:00,11

Variables Created or Modified	PRE_1	Predicted probability
	COO_1	Analog of Cook's influence statistics
	LEV_1	Leverage value
	DFB0_1	DFBETA for constant
	DFB1_1	DFBETA for Gecentreerde variabele voor het verkleind inkomen van het huishouden
	DFB2_1	DFBETA for Gecentreerde variabele voor leeftijd
	DFB3_1	DFBETA for Gecentreerde variabele voor opleidingsniveau
	DFB4_1	DFBETA for Gecentreerde variabele voor de houding ten opzichte van inkomensongelijkheid
	DFB5_1	DFBETA for Gecentreerde variabele voor de houding ten opzichte van immigranten
	DFB6_1	DFBETA for Interactievariabele tussen de gecentreerde kijk op ongelijkheid en kijk op immigranten

Case Processing Summary			
Unweighted Cases ^a		N	Percent
Selected Cases	Included in Analysis	1793	100.0
	Missing Cases	0	.0
	Total	1793	100.0
Unselected Cases		0	.0
Total		1793	100.0

a. If weight is in effect, see classification table for the total number of cases.

Dependent Variable Encoding	
Original Value	Internal Value
.00	0
1.00	1

Block 0: Beginning Block

Classification Table ^{a,b}					
	Observed		Predicted		
			dummy voor het stemmen op een niet of wel linkse partij		Percentage Correct
			.00	1.00	
Step 0	dummy voor het stemmen op een niet of wel linkse partij	.00	1238	0	100.0

		1.00	555	0	.0
	Overall Percentage				69.0

a. Constant is included in the model.

b. The cut value is .500

Variables in the Equation							
		B	S.E.	Wald	df	Sig.	Exp(B)
Step 0	Constant	-.802	.051	246.655	1	<,001	.448

Variables not in the Equation					
			Score	df	Sig.
Step 0	Variables	Gecentreerde variabele voor het verkleind inkomen van het huishouden	9.100	1	.003
		Gecentreerde variabele voor leeftijd	6.046	1	.014
		Gecentreerde variabele voor opleidingsniveau	56.336	1	<,001
	Overall Statistics		80.779	3	<,001

Block 1: Method = Enter

Omnibus Tests of Model Coefficients
--

		Chi-square	df	Sig.
Step 1	Step	85.131	3	<,001
	Block	85.131	3	<,001
	Model	85.131	3	<,001

Model Summary			
Step	-2 Log likelihood	Cox & Snell R Square	Nagelkerke R Square
1	2133.634 ^a	.046	.065

a. Estimation terminated at iteration number 4 because parameter estimates changed by less than .001.

Hosmer and Lemeshow Test			
Step	Chi-square	df	Sig.
1	44.347	8	<,001

Contingency Table for Hosmer and Lemeshow Test						
		dummy voor het stemmen op een niet of wel linkse partij = .00		dummy voor het stemmen op een niet of wel linkse partij = 1.00		Total
		Observed	Expected	Observed	Expected	
Step 1	1	135	153.056	44	25.944	179
	2	141	144.472	38	34.528	179
	3	135	136.607	44	42.393	179
	4	146	130.393	33	48.607	179

5	146	125.385	33	53.615	179
6	118	120.792	61	58.208	179
7	130	116.740	49	62.260	179
8	110	111.684	69	67.316	179
9	99	105.058	80	73.942	179
10	78	93.812	104	88.188	182

Classification Table ^a					
	Observed		Predicted		
			dummy voor het stemmen op een niet of wel linkse partij		Percentage Correct
			.00	1.00	
Step 1	dummy voor het stemmen op een niet of wel linkse partij	.00	1219	19	98.5
		1.00	512	43	7.7
	Overall Percentage				70.4

a. The cut value is .500

Variables in the Equation						
		B	S.E.	Wald	df	Sig.
Step 1 ^a	Gecentreerde variabele voor het verkleind inkomen van het huishouden	-.113	.024	21.707	1	<,001
	Gecentreerde variabele voor leeftijd	-.006	.004	2.755	1	.097
	Gecentreerde variabele voor opleidingsniveau	.354	.045	62.941	1	<,001
	Constant	-.850	.053	253.465	1	<,001

Variables in the Equation		
		Exp(B)
Step 1 ^a	Gecentreerde variabele voor het verkleind inkomen van het huishouden	.893
	Gecentreerde variabele voor leeftijd	.994
	Gecentreerde variabele voor opleidingsniveau	1.425
	Constant	.428

a. Variable(s) entered on step 1: Gecentreerde variabele voor het verkleind inkomen van het huishouden, Gecentreerde variabele voor leeftijd, Gecentreerde variabele voor opleidingsniveau.

Block 2: Method = Enter

Omnibus Tests of Model Coefficients				
		Chi-square	df	Sig.
Step 1	Step	202.157	1	<,001
	Block	202.157	1	<,001
	Model	287.287	4	<,001

Model Summary			
Step	-2 Log likelihood	Cox & Snell R Square	Nagelkerke R Square
1	1931.478 ^a	.148	.209

a. Estimation terminated at iteration number 4

because parameter estimates changed by less than .001.

Hosmer and Lemeshow Test

Step	Chi-square	df	Sig.
1	5.429	8	.711

Contingency Table for Hosmer and Lemeshow Test

		dummy voor het stemmen op een niet of wel linkse partij = .00		dummy voor het stemmen op een niet of wel linkse partij = 1.00		Total
		Observed	Expected	Observed	Expected	
Step 1	1	165	166.394	14	12.606	179
	2	156	157.848	23	21.152	179
	3	146	149.923	33	29.077	179
	4	146	142.341	33	36.659	179
	5	134	134.545	45	44.455	179
	6	131	125.419	48	53.581	179
	7	122	114.869	57	64.131	179
	8	97	100.984	82	78.016	179
	9	77	84.700	102	94.300	179
	10	64	60.978	118	121.022	182

Classification Table^a

Observed		Predicted		
		dummy voor het stemmen op een niet of wel linkse partij		Percentage Correct
		.00	1.00	

Step 1	dummy voor het stemmen op een niet of wel linkse partij	.00	1121	117	90.5
		1.00	355	200	36.0
	Overall Percentage				73.7

a. The cut value is .500

Variables in the Equation						
		B	S.E.	Wald	df	Sig.
Step 1 ^a	Gecentreerde variabele voor het verkleind inkomen van het huishouden	-.051	.026	3.992	1	.046
	Gecentreerde variabele voor leeftijd	-.008	.004	4.012	1	.045
	Gecentreerde variabele voor opleidingsniveau	.395	.046	72.076	1	<,001
	Gecentreerde variabele voor de houding ten opzichte van inkomensongelijkheid	.599	.046	170.056	1	<,001
	Constant	-.960	.059	267.803	1	<,001

Variables in the Equation		
		Exp(B)
Step 1 ^a	Gecentreerde variabele voor het verkleind inkomen van het huishouden	.950
	Gecentreerde variabele voor leeftijd	.992
	Gecentreerde variabele voor opleidingsniveau	1.484
	Gecentreerde variabele voor de houding ten opzichte van inkomensongelijkheid	1.821
	Constant	.383

a. Variable(s) entered on step 1: Gecentreerde variabele voor het verkleind inkomen van het huishouden, Gecentreerde variabele voor leeftijd, Gecentreerde variabele voor opleidingsniveau, Gecentreerde variabele voor de houding ten opzichte van inkomensongelijkheid.

Block 3: Method = Enter

Omnibus Tests of Model Coefficients				
		Chi-square	df	Sig.
Step 1	Step	187.650	1	<,001
	Block	187.650	1	<,001
	Model	474.937	5	<,001

Model Summary			
Step	-2 Log likelihood	Cox & Snell R Square	Nagelkerke R Square
1	1743.828 ^a	.233	.328

a. Estimation terminated at iteration number 5 because parameter estimates changed by less than .001.

Hosmer and Lemeshow Test			
Step	Chi-square	df	Sig.

	e		
1	5.354	8	.719

Contingency Table for Hosmer and Lemeshow Test						
		dummy voor het stemmen op een niet of wel linkse partij = .00		dummy voor het stemmen op een niet of wel linkse partij = 1.00		Total
		Observed	Expected	Observed	Expected	
Step 1	1	173	172.916	6	6.084	179
	2	164	164.772	15	14.228	179
	3	157	157.487	22	21.513	179
	4	159	150.128	20	28.872	179
	5	135	140.363	44	38.637	179
	6	129	129.383	50	49.617	179
	7	117	115.053	62	63.947	179
	8	92	95.529	87	83.471	179
	9	68	71.825	111	107.175	179
	10	44	40.546	138	141.454	182

Classification Table ^a					
	Observed		Predicted		
			dummy voor het stemmen op een niet of wel linkse partij		Percentage Correct
			.00	1.00	
Step 1	dummy voor het stemmen op een niet of wel linkse partij	.00	1105	133	89.3
		1.00	287	268	48.3
	Overall Percentage				76.6

a. The cut value is .500

Variables in the Equation						
		B	S.E.	Wald	df	Sig.
Step 1 ^a	Gecentreerde variabele voor het verkleind inkomen van het huishouden	-.068	.028	6.003	1	.014
	Gecentreerde variabele voor leeftijd	-.009	.004	4.565	1	.033
	Gecentreerde variabele voor opleidingsniveau	.246	.048	25.685	1	<,001
	Gecentreerde variabele voor de houding ten opzichte van inkomensongelijkheid	.540	.048	124.116	1	<,001
	Gecentreerde variabele voor de houding ten opzichte van immigranten	-.929	.075	155.012	1	<,001
	Constant	-1.022	.063	261.924	1	<,001

Variables in the Equation		
		Exp(B)
Step 1 ^a	Gecentreerde variabele voor het verkleind inkomen van het huishouden	.935
	Gecentreerde variabele voor leeftijd	.991
	Gecentreerde variabele voor opleidingsniveau	1.278
	Gecentreerde variabele voor de houding ten opzichte van inkomensongelijkheid	1.716
	Gecentreerde variabele voor de houding ten opzichte van immigranten	.395
	Constant	.360

a. Variable(s) entered on step 1: Gecentreerde variabele voor het verkleind inkomen van het huishouden, Gecentreerde variabele voor leeftijd, Gecentreerde variabele voor opleidingsniveau, Gecentreerde variabele voor de houding ten opzichte van inkomensongelijkheid, Gecentreerde variabele voor de houding ten opzichte van immigranten.

Block 4: Method = Enter

Omnibus Tests of Model Coefficients				
		Chi-square	df	Sig.
Step 1	Step	.343	1	.558
	Block	.343	1	.558
	Model	475.281	6	<.001

Model Summary			
Step	-2 Log likelihood	Cox & Snell R Square	Nagelkerke R Square
1	1743.484 ^a	.233	.328

a. Estimation terminated at iteration number 5 because parameter estimates changed by less than .001.

Hosmer and Lemeshow Test

Step	Chi-square	df	Sig.
1	4.324	8	.827

Contingency Table for Hosmer and Lemeshow Test						
		dummy voor het stemmen op een niet of wel linkse partij = .00		dummy voor het stemmen op een niet of wel linkse partij = 1.00		Total
		Observed	Expected	Observed	Expected	
Step 1	1	174	172.504	5	6.496	179
	2	163	164.424	16	14.576	179
	3	158	157.353	21	21.647	179
	4	157	150.060	22	28.940	179
	5	135	140.733	44	38.267	179
	6	129	129.989	50	49.011	179
	7	118	115.614	61	63.386	179
	8	93	95.966	86	83.034	179
	9	69	71.718	110	107.282	179
	10	42	39.639	140	142.361	182

Classification Table ^a					
Observed			Predicted		
			dummy voor het stemmen op een niet of wel linkse partij		Percentage Correct
			.00	1.00	
Step 1	dummy voor het stemmen op een niet of wel linkse partij	.00	1105	133	89.3
		1.00	288	267	48.1
	Overall Percentage				76.5

a. The cut value is .500

Variables in the Equation						
		B	S.E.	Wald	df	Sig.
Step 1 ^a	Gecentreerde variabele voor het verkleind inkomen van het huishouden	-.068	.028	6.077	1	.014
	Gecentreerde variabele voor leeftijd	-.009	.004	4.493	1	.034
	Gecentreerde variabele voor opleidingsniveau	.244	.049	25.321	1	<,001
	Gecentreerde variabele voor de houding ten opzichte van inkomensongelijkheid	.532	.050	112.152	1	<,001
	Gecentreerde variabele voor de houding ten opzichte van immigranten	-.915	.078	138.240	1	<,001
	Interactievariabele tussen de gecentreerde kijk op ongelijkheid en kijk op immigranten	-.033	.056	.343	1	.558
	Constant	-1.018	.063	259.265	1	<,001

Variables in the Equation		
		Exp(B)
Step 1 ^a	Gecentreerde variabele voor het verkleind inkomen van het huishouden	.934
	Gecentreerde variabele voor leeftijd	.991
	Gecentreerde variabele voor opleidingsniveau	1.277

	Gecentreerde variabele voor de houding ten opzichte van inkomensongelijkheid	1.702
	Gecentreerde variabele voor de houding ten opzichte van immigranten	.400
	Interactievariabele tussen de gecentreerde kijk op ongelijkheid en kijk op immigranten	.968
	Constant	.361

a. Variable(s) entered on step 1: Gecentreerde variabele voor het verkleind inkomen van het huishouden, Gecentreerde variabele voor leeftijd, Gecentreerde variabele voor opleidingsniveau, Gecentreerde variabele voor de houding ten opzichte van inkomensongelijkheid, Gecentreerde variabele voor de houding ten opzichte van immigranten, Interactievariabele tussen de gecentreerde kijk op ongelijkheid en kijk op immigranten.

De volgende stappen zijn uitgevoerd om de grafiek voor verschillen in groepen op basis van de kijk op immigranten te maken (alleen de tweede van de twee gemaakte grafieken is uiteindelijk gebruikt).

RECODE inkomensongelijkheid_gecentreerd (Lowest thru -1.39083=1) (1.39103 thru Highest=3) (ELSE=2)

INTO inkomensongelijkheid_groep_gecentreerd.

VARIABLE LABELS inkomensongelijkheid_groep_gecentreerd 'Gecentreerde groepsvariabele voor de '+'

'kijk op inkomensongelijkheid'.

EXECUTE.

RECODE immigranten_gecentreerd (Lowest thru -1.0080=1) (1.00776 thru Highest=3) (ELSE=2)

INTO immigranten_groep_gecentreerd.

VARIABLE LABELS immigranten_groep_gecentreerd 'Gecentreerde groepsvariabele voor de '+'

'kijk op immigranten'.

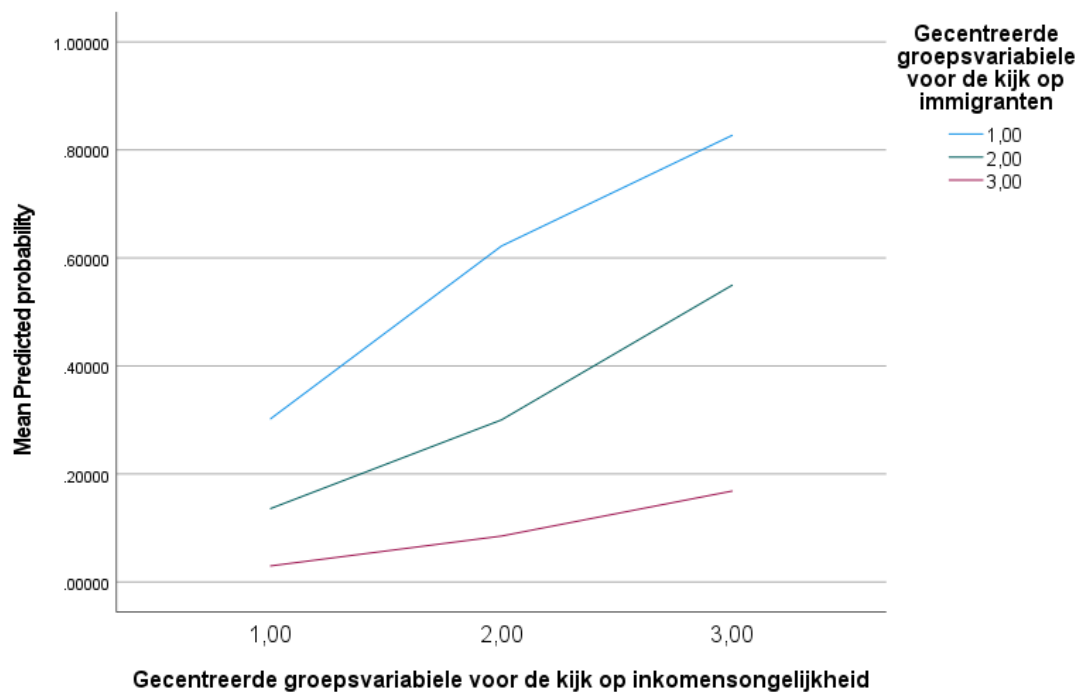
EXECUTE.

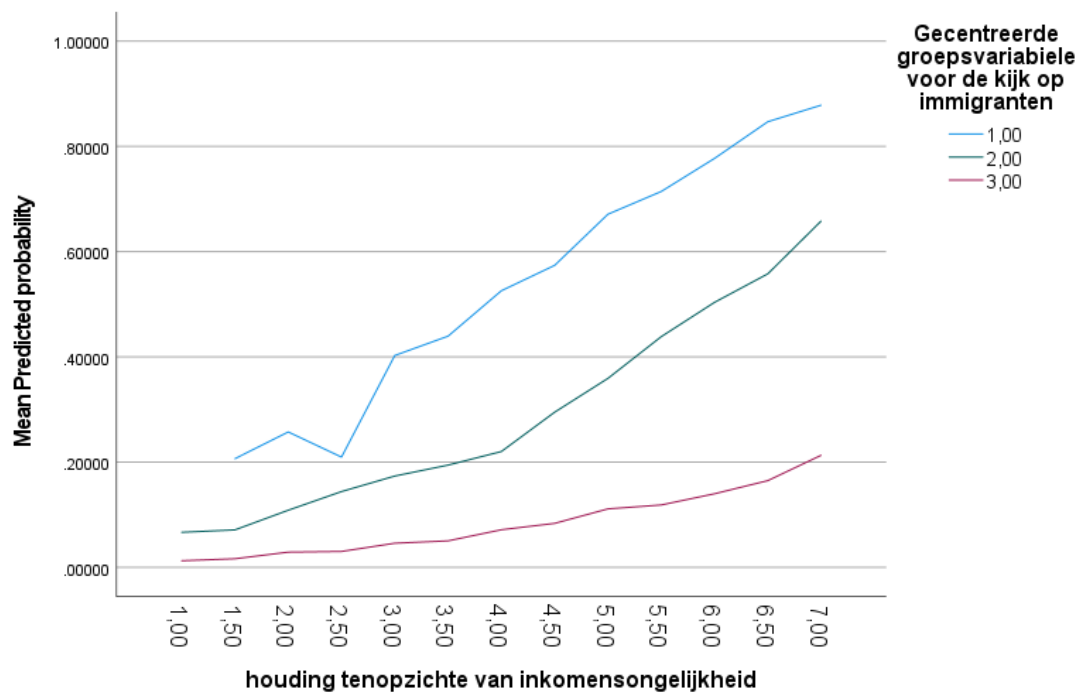
GRAPH

/LINE(MULTIPLE)=MEAN(PRE_1) BY inkomensongelijkheid_groep_gecentreerd BY
immigranten_groep_gecentreerd.

GRAPH

/LINE(MULTIPLE)=MEAN(PRE_1) BY inkomensongelijkheid BY
immigranten_groep_gecentreerd.





Uitbijters

Om na te gaan of de uitbijters een (te) grote invloed hebben gehad op de uiteindelijke resultaten, zijn hier alle uitbijters verwijderd en is de regressie nog een keer uitgevoerd.

USE ALL.

COMPUTE filter_\$=(COO_1 <= 0.0022).

VARIABLE LABELS filter_\$ 'COO_1 <= 0.0022 (FILTER)'.

VALUE LABELS filter_\$ 0 'Not Selected' 1 'Selected'.

FORMATS filter_\$ (f1.0).

FILTER BY filter_\$.

EXECUTE.

LOGISTIC REGRESSION VARIABLES partij_dummy

/METHOD=ENTER inkomen_verkleind_gecentreerd leeftijd_gecentreerd
opleidingsniveau_gecentreerd

/METHOD=ENTER opleidingsniveau_gecentreerd leeftijd_gecentreerd
inkomen_verkleind_gecentreerd inkomensongelijkheid_gecentreerd

/METHOD=ENTER inkomen_verkleind_gecentreerd immigranten_gecentreerd
inkomensongelijkheid_gecentreerd opleidingsniveau_gecentreerd leeftijd_gecentreerd

```

/METHOD=ENTER inkomensongelijkheid_gecentreerd immigranten_gecentreerd interactie
inkomen_verkleind_gecentreerd opleidingsniveau_gecentreerd
leeftijd_gecentreerd
/SAVE=PRED COOK LEVER DFBETA
/PRINT=GOODFIT
/CRITERIA=PIN(0.05) POUT(0.10) ITERATE(20) CUT(0.5).

```

Logistic Regression

Notes		
Output Created		30-JUL-2025 16:41:17
Comments		
Input	Data	C:\Users\vexia\Desktop\liss_baw_class.sav
	Active Dataset	DataSet5
	Filter	COO_1 <= 0.0022 (FILTER)
	Weight	<none>
	Split File	<none>
	N of Rows in Working Data File	1146
Missing Value Handling	Definition of Missing	User-defined missing values are treated as missing
Syntax		LOGISTIC REGRESSION VARIABLES partij_dummy /METHOD=ENTER inkomen_verkleind_g ecentreerd leeftijd_gecentreerd

	opleidingsniveau_gec entreerd /METHOD=ENTER opleidingsniveau_gec entreerd leeftijd_gecentreerd inkomen_verkleind_g ecentreerd inkomensongelijkheid _gecentreerd /METHOD=ENTER inkomen_verkleind_g ecentreerd immigranten_gecentr eerd inkomensongelijkheid _gecentreerd opleidingsniveau_gec entreerd leeftijd_gecentreerd /METHOD=ENTER inkomensongelijkheid _gecentreerd immigranten_gecentr eerd interactie inkomen_verkleind_g ecentreerd opleidingsniveau_gec entreerd leeftijd_gecentreerd /SAVE=PRED COOK LEVER DFBETA /PRINT=GOODFIT /CRITERIA=PIN(0.05) POUT(0.10)
--	--

		ITERATE(20) CUT(0.5).
Resources	Processor Time	00:00:00,00
	Elapsed Time	00:00:00,12
Variables Created or Modified	PRE_2	Predicted probability
	COO_2	Analog of Cook's influence statistics
	LEV_2	Leverage value
	DFB0_2	DFBETA for constant
	DFB1_2	DFBETA for Gecentreerde variabele voor het verkleind inkomen van het huishouden
	DFB2_2	DFBETA for Gecentreerde variabele voor leeftijd
	DFB3_2	DFBETA for Gecentreerde variabele voor opleidingsniveau
	DFB4_2	DFBETA for Gecentreerde variabele voor de houding ten opzichte van inkomensongelijkheid
	DFB5_2	DFBETA for Gecentreerde variabele voor de houding ten opzichte van immigranten
	DFB6_2	DFBETA for Interactievevariabele tussen de gecentreerde kijk op

		ongelijkheid en kijk op immigranten
--	--	--

Case Processing Summary			
Unweighted Cases ^a		N	Percent
Selected Cases	Included in Analysis	1146	100.0
	Missing Cases	0	.0
	Total	1146	100.0
Unselected Cases		0	.0
Total		1146	100.0

a. If weight is in effect, see classification table for the total number of cases.

Dependent Variable Encoding	
Original Value	Internal Value
.00	0
1.00	1

Block 0: Beginning Block

Classification Table ^{a,b}			
	Observed	Predicted	
		dummy voor het stemmen op een niet of wel linkse	Percentage Correct

			partij		
			.00	1.00	
Step 0	dummy voor het stemmen op een niet of wel linkse partij	.00	978	0	100.0
		1.00	168	0	.0
	Overall Percentage				85.3

a. Constant is included in the model.

b. The cut value is .500

Variables in the Equation							
		B	S.E.	Wald	df	Sig.	Exp(B)
Step 0	Constant	-1.762	.084	444.889	1	<,001	.172

Variables not in the Equation					
			Score	df	Sig.
Step 0	Variables	Gecentreerde variabele voor het verkleind inkomen van het huishouden	10.299	1	.001
		Gecentreerde variabele voor leeftijd	26.807	1	<,001
		Gecentreerde variabele voor opleidingsniveau	133.159	1	<,001
	Overall Statistics		177.585	3	<,001

Block 1: Method = Enter

Omnibus Tests of Model Coefficients				
		Chi-square	df	Sig.
Step 1	Step	229.538	3	<,001
	Block	229.538	3	<,001
	Model	229.538	3	<,001

Model Summary			
Step	-2 Log likelihood	Cox & Snell R Square	Nagelkerke R Square
1	725.677 ^a	.182	.321

a. Estimation terminated at iteration number 6 because parameter estimates changed by less than .001.

Hosmer and Lemeshow Test			
Step	Chi-square	df	Sig.
1	1.860	8	.985

Contingency Table for Hosmer and Lemeshow Test						
		dummy voor het stemmen op een niet of wel linkse partij = .00		dummy voor het stemmen op een niet of wel linkse partij = 1.00		Total
		Observed	Expected	Observed	Expected	
Step	1	115	114.671	0	.329	115

1						
	2	114	113.944	1	1.056	115
	3	111	111.830	4	3.170	115
	4	110	109.452	5	5.548	115
	5	109	106.821	6	8.179	115
	6	103	103.003	12	11.997	115
	7	95	97.803	20	17.197	115
	8	91	91.097	24	23.903	115
	9	79	79.705	36	35.295	115
	10	51	49.674	60	61.326	111

Classification Table ^a					
	Observed		Predicted		
			dummy voor het stemmen op een niet of wel linkse partij		Percentage Correct
			.00	1.00	
Step 1	dummy voor het stemmen op een niet of wel linkse partij	.00	953	25	97.4
		1.00	130	38	22.6
	Overall Percentage				86.5

a. The cut value is .500

Variables in the Equation						
		B	S.E.	Wald	df	Sig.
Step 1 ^a	Gecentreerde variabele voor het verkleind inkomen van het huishouden	-.311	.051	36.888	1	<,001
	Gecentreerde variabele voor leeftijd	-.023	.008	9.458	1	.002

	Gecentreerde variabele voor opleidingsniveau	1.378	.127	117.022	1	<,001
	Constant	-2.463	.139	312.593	1	<,001

Variables in the Equation		
		Exp(B)
Step 1 ^a	Gecentreerde variabele voor het verkleind inkomen van het huishouden	.732
	Gecentreerde variabele voor leeftijd	.977
	Gecentreerde variabele voor opleidingsniveau	3.969
	Constant	.085

a. Variable(s) entered on step 1: Gecentreerde variabele voor het verkleind inkomen van het huishouden, Gecentreerde variabele voor leeftijd, Gecentreerde variabele voor opleidingsniveau.

Block 2: Method = Enter

Omnibus Tests of Model Coefficients				
		Chi-square	df	Sig.
Step 1	Step	444.153	1	<,001
	Block	444.153	1	<,001
	Model	673.691	4	<,001

Model Summary

Step	-2 Log likelihood	Cox & Snell R Square	Nagelkerke R Square
1	281.523 ^a	.444	.786

a. Estimation terminated at iteration number 8 because parameter estimates changed by less than .001.

Hosmer and Lemeshow Test			
Step	Chi-square	df	Sig.
1	.866	8	.999

Contingency Table for Hosmer and Lemeshow Test						
		dummy voor het stemmen op een niet of wel linkse partij = .00		dummy voor het stemmen op een niet of wel linkse partij = 1.00		Total
		Observed	Expected	Observed	Expected	
Step 1	1	115	115.000	0	.000	115
	2	115	114.998	0	.002	115
	3	115	114.990	0	.010	115
	4	115	114.956	0	.044	115
	5	115	114.862	0	.138	115
	6	115	114.501	0	.499	115
	7	113	113.104	2	1.896	115
	8	106	105.257	9	9.743	115
	9	61	62.614	54	52.386	115
	10	8	7.718	103	103.282	111

Classification Table ^a		
	Observed	Predicted

			dummy voor het stemmen op een niet of wel linkse partij		Percentage Correct
			.00	1.00	
Step 1	dummy voor het stemmen op een niet of wel linkse partij	.00	953	25	97.4
		1.00	34	134	79.8
	Overall Percentage				94.9

a. The cut value is .500

Variables in the Equation						
		B	S.E.	Wald	df	Sig.
Step 1 ^a	Gecentreerde variabele voor het verkleind inkomen van het huishouden	-.209	.097	4.614	1	.032
	Gecentreerde variabele voor leeftijd	-.032	.013	5.816	1	.016
	Gecentreerde variabele voor opleidingsniveau	2.333	.248	88.750	1	<,001
	Gecentreerde variabele voor de houding ten opzichte van inkomensongelijkheid	2.994	.275	118.517	1	<,001
	Constant	-4.835	.401	145.295	1	<,001

Variables in the Equation		
		Exp(B)
Step 1 ^a	Gecentreerde variabele voor het verkleind inkomen van het huishouden	.812
	Gecentreerde variabele voor leeftijd	.969

	Gecentreerde variabele voor opleidingsniveau	10.309
	Gecentreerde variabele voor de houding ten opzichte van inkomensongelijkheid	19.962
	Constant	.008

a. Variable(s) entered on step 1: Gecentreerde variabele voor het verkleind inkomen van het huishouden, Gecentreerde variabele voor leeftijd, Gecentreerde variabele voor opleidingsniveau, Gecentreerde variabele voor de houding ten opzichte van inkomensongelijkheid.

Block 3: Method = Enter

Omnibus Tests of Model Coefficients				
		Chi-square	df	Sig.
Step 1	Step	237.882	1	<,001
	Block	237.882	1	<,001
	Model	911.573	5	<,001

Model Summary			
Step	-2 Log likelihood	Cox & Snell R Square	Nagelkerke R Square
1	43.642 ^a	.549	.970

a. Estimation terminated at iteration number 12 because parameter estimates changed by less than .001.

Hosmer and Lemeshow Test			
Step	Chi-square	df	Sig.
1	.119	8	1.000

Contingency Table for Hosmer and Lemeshow Test						
		dummy voor het stemmen op een niet of wel linkse partij = .00		dummy voor het stemmen op een niet of wel linkse partij = 1.00		Total
		Observed	Expected	Observed	Expected	
Step 1	1	115	115.000	0	.000	115
	2	115	115.000	0	.000	115
	3	115	115.000	0	.000	115
	4	115	115.000	0	.000	115
	5	115	115.000	0	.000	115
	6	115	115.000	0	.000	115
	7	115	114.997	0	.003	115
	8	115	114.893	0	.107	115
	9	58	58.102	57	56.898	115
	10	0	.008	111	110.992	111

Classification Table ^a					
	Observed		Predicted		
			dummy voor het stemmen op een niet of wel linkse partij		Percentage Correct
			.00	1.00	
Step 1	dummy voor het stemmen op een niet of wel linkse partij	.00	974	4	99.6
		1.00	4	164	97.6

	Overall Percentage			99.3
--	--------------------	--	--	------

a. The cut value is .500

Variables in the Equation						
		B	S.E.	Wald	df	Sig.
Step 1 ^a	Gecentreerde variabele voor het verkleind inkomen van het huishouden	-.307	.397	.598	1	.440
	Gecentreerde variabele voor leeftijd	-.071	.050	2.044	1	.153
	Gecentreerde variabele voor opleidingsniveau	1.898	.743	6.531	1	.011
	Gecentreerde variabele voor de houding ten opzichte van inkomensongelijkheid	4.063	.975	17.349	1	<.001
	Gecentreerde variabele voor de houding ten opzichte van immigranten	-9.104	1.923	22.410	1	<.001
	Constant	-9.047	1.659	29.727	1	<.001

Variables in the Equation		
		Exp(B)
Step 1 ^a	Gecentreerde variabele voor het verkleind inkomen van het huishouden	.736
	Gecentreerde variabele voor leeftijd	.932
	Gecentreerde variabele voor opleidingsniveau	6.671
	Gecentreerde variabele voor de houding ten opzichte van inkomensongelijkheid	58.127
	Gecentreerde variabele voor de houding ten	.000

	opzichte van immigranten	
	Constant	.000

a. Variable(s) entered on step 1: Gecentreerde variabele voor het verkleind inkomen van het huishouden, Gecentreerde variabele voor leeftijd, Gecentreerde variabele voor opleidingsniveau, Gecentreerde variabele voor de houding ten opzichte van inkomensongelijkheid, Gecentreerde variabele voor de houding ten opzichte van immigranten.

Block 4: Method = Enter

Omnibus Tests of Model Coefficients				
		Chi-square	df	Sig.
Step 1	Step	.038	1	.845
	Block	.038	1	.845
	Model	911.611	6	<.001

Model Summary			
Step	-2 Log likelihood	Cox & Snell R Square	Nagelkerke R Square
1	43.603 ^a	.549	.970

a. Estimation terminated at iteration number 13 because parameter estimates changed by less than .001.

Hosmer and Lemeshow Test			
Step	Chi-square	df	Sig.
1	.115	8	1.000

Contingency Table for Hosmer and Lemeshow Test						
		dummy voor het stemmen op een niet of wel linkse partij = .00		dummy voor het stemmen op een niet of wel linkse partij = 1.00		Total
		Observed	Expected	Observed	Expected	
Step 1	1	115	115.000	0	.000	115
	2	115	115.000	0	.000	115
	3	115	115.000	0	.000	115
	4	115	115.000	0	.000	115
	5	115	115.000	0	.000	115
	6	115	115.000	0	.000	115
	7	115	114.997	0	.003	115
	8	115	114.898	0	.102	115
	9	58	58.096	57	56.904	115
	10	0	.010	111	110.990	111

Classification Table ^a					
	Observed		Predicted		
			dummy voor het stemmen op een niet of wel linkse partij		Percentage Correct
			.00	1.00	
Step 1	dummy voor het stemmen op een niet of wel linkse partij	.00	974	4	99.6
		1.00	4	164	97.6

	Overall Percentage			99.3
--	--------------------	--	--	------

a. The cut value is .500

Variables in the Equation						
		B	S.E.	Wald	df	Sig.
Step 1 ^a	Gecentreerde variabele voor het verkleind inkomen van het huishouden	-.310	.399	.604	1	.437
	Gecentreerde variabele voor leeftijd	-.073	.051	2.031	1	.154
	Gecentreerde variabele voor opleidingsniveau	1.933	.770	6.306	1	.012
	Gecentreerde variabele voor de houding ten opzichte van inkomensongelijkheid	4.410	2.043	4.660	1	.031
	Gecentreerde variabele voor de houding ten opzichte van immigranten	-9.446	2.627	12.929	1	<.001
	Interactievariabele tussen de gecentreerde kijk op ongelijkheid en kijk op immigranten	.536	2.704	.039	1	.843
	Constant	-9.325	2.232	17.457	1	<.001

Variables in the Equation		
		Exp(B)
Step 1 ^a	Gecentreerde variabele voor het verkleind inkomen van het huishouden	.733

	Gecentreerde variabele voor leeftijd	.930
	Gecentreerde variabele voor opleidingsniveau	6.912
	Gecentreerde variabele voor de houding ten opzichte van inkomensongelijkheid	82.248
	Gecentreerde variabele voor de houding ten opzichte van immigranten	.000
	Interactievariabele tussen de gecentreerde kijk op ongelijkheid en kijk op immigranten	1.710
	Constant	.000

a. Variable(s) entered on step 1: Gecentreerde variabele voor het verkleind inkomen van het huishouden, Gecentreerde variabele voor leeftijd, Gecentreerde variabele voor opleidingsniveau, Gecentreerde variabele voor de houding ten opzichte van inkomensongelijkheid, Gecentreerde variabele voor de houding ten opzichte van immigranten, Interactievariabele tussen de gecentreerde kijk op ongelijkheid en kijk op immigranten.

Deze resultaten verschillen van de originele resultaten, vooral in de zin dat de constanten en coëfficiënten een stuk extremer zijn, maar de belangrijke conclusies die hieruit volgen verschillen niet. Hierom is ervoor gekozen om verder geen aandacht te besteden aan de uitbijters, en dus worden alle (complete) cases weer geselecteerd.

FILTER OFF.

USE ALL.

EXECUTE.

USE ALL.

COMPUTE filter_\$=(dummy_observaties = 1).

VARIABLE LABELS filter_\$ 'dummy_observaties = 1 (FILTER)'.

VALUE LABELS filter_\$ 0 'Not Selected' 1 'Selected'.

FORMATS filter_\$ (f1.0).

FILTER BY filter_\$.

EXECUTE.

Onafhankelijkheid

Vervolgens wordt er voor de assumptie van onafhankelijke waarden gekeken naar de resultaten als slechts één persoon per huishouden wordt geselecteerd.

```
SORT CASES BY nohouse_encr(A).
```

```
COMPUTE willekeurig=RV.UNIFORM(0,1).
```

```
EXECUTE.
```

```
SORT CASES BY nohouse_encr(A) willekeurig(A).
```

```
IF ($CASENUM = 1 OR nohouse_encr ~= LAG(nohouse_encr)) uniek_huishouden=1.
```

```
EXECUTE.
```

```
IF (NOT ($CASENUM = 1 OR nohouse_encr ~= LAG(nohouse_encr)))
```

```
uniek_huishouden=0.
```

```
EXECUTE.
```

```
COMPUTE casefilter=herkomstgroep = 0 AND dummy_observaties = 1 AND
```

```
uniek_huishouden = 1.
```

```
EXECUTE.
```

```
USE ALL.
```

```
COMPUTE filter_$=(casefilter = 1).
```

```
VARIABLE LABELS filter_$ 'casefilter = 1 (FILTER)'.  
VALUE LABELS filter_$ 0 'Not Selected' 1 'Selected'.  
FORMATS filter_$ (f1.0).  
FILTER BY filter_$.  
EXECUTE.
```

```
LOGISTIC REGRESSION VARIABLES partij_dummy
```

```
  /METHOD=ENTER inkomen_verkleind_gecentreerd leeftijd_gecentreerd
```

```
  opleidingsniveau_gecentreerd
```

```

/METHOD=ENTER opleidingsniveau_gecentreerd leeftijd_gecentreerd
inkomen_verkleind_gecentreerd inkomensongelijkheid_gecentreerd
/METHOD=ENTER inkomen_verkleind_gecentreerd immigranten_gecentreerd
inkomensongelijkheid_gecentreerd opleidingsniveau_gecentreerd leeftijd_gecentreerd
/METHOD=ENTER inkomensongelijkheid_gecentreerd immigranten_gecentreerd interactie
inkomen_verkleind_gecentreerd opleidingsniveau_gecentreerd
leeftijd_gecentreerd
/SAVE=PRED COOK LEVER DFBETA
/PRINT=GOODFIT
/CRITERIA=PIN(0.05) POUT(0.10) ITERATE(20) CUT(0.5).

```

Logistic Regression

Notes		
Output Created		30-JUL-2025 16:45:40
Comments		
Input	Data	C:\Users\vexia\Desktop\liss_baw_class.sav
	Active Dataset	DataSet5
	Filter	casefilter = 1 (FILTER)
	Weight	<none>
	Split File	<none>
	N of Rows in Working Data File	1202
Missing Value Handling	Definition of Missing	User-defined missing values are treated as missing
Syntax		LOGISTIC REGRESSION VARIABLES partij_dummy

	/METHOD=ENTER inkomen_verkleind_g ecentreerd leeftijd_gecentreerd opleidingsniveau_gec entreerd /METHOD=ENTER opleidingsniveau_gec entreerd leeftijd_gecentreerd inkomen_verkleind_g ecentreerd inkomensongelijkheid _gecentreerd /METHOD=ENTER inkomen_verkleind_g ecentreerd immigranten_gecentr eerd inkomensongelijkheid _gecentreerd opleidingsniveau_gec entreerd leeftijd_gecentreerd /METHOD=ENTER inkomensongelijkheid _gecentreerd immigranten_gecentr eerd interactie inkomen_verkleind_g ecentreerd opleidingsniveau_gec entreerd leeftijd_gecentreerd /SAVE=PRED COOK LEVER DFBETA
--	---

		/PRINT=GOODFIT /CRITERIA=PIN(0.05) POUT(0.10) ITERATE(20) CUT(0.5).
Resources	Processor Time	00:00:00,06
	Elapsed Time	00:00:00,11
Variables Created or Modified	PRE_4	Predicted probability
	COO_4	Analog of Cook's influence statistics
	LEV_4	Leverage value
	DFB0_4	DFBETA for constant
	DFB1_4	DFBETA for Gecentreerde variabele voor het verkleind inkomen van het huishouden
	DFB2_4	DFBETA for Gecentreerde variabele voor leeftijd
	DFB3_4	DFBETA for Gecentreerde variabele voor opleidingsniveau
	DFB4_4	DFBETA for Gecentreerde variabele voor de houding ten opzichte van inkomensongelijkheid
	DFB5_4	DFBETA for Gecentreerde variabele voor de houding ten opzichte van immigranten

	DFB6_4	DFBETA for Interactievariabele tussen de gecentreerde kijk op ongelijkheid en kijk op immigranten
--	--------	--

Case Processing Summary			
Unweighted Cases ^a		N	Percent
Selected Cases	Included in Analysis	1202	100.0
	Missing Cases	0	.0
	Total	1202	100.0
Unselected Cases		0	.0
Total		1202	100.0

a. If weight is in effect, see classification table for the total number of cases.

Dependent Variable Encoding	
Original Value	Internal Value
.00	0
1.00	1

Block 0: Beginning Block

Classification Table ^{a,b}					
	Observed		Predicted		
			dummy voor het stemmen op een niet of wel linkse partij		Percentage Correct
			.00	1.00	
Step 0	dummy voor het stemmen op een niet of wel linkse partij	.00	832	0	100.0
		1.00	370	0	.0
	Overall Percentage				69.2

a. Constant is included in the model.

b. The cut value is .500

Variables in the Equation							
		B	S.E.	Wald	df	Sig.	Exp(B)
Step 0	Constant	-.810	.062	168.168	1	<,001	.445

Variables not in the Equation					
			Score	df	Sig.
Step 0	Variables	Gecentreerde variabele voor het verkleind inkomen van het huishouden	6.950	1	.008
		Gecentreerde variabele voor leeftijd	3.817	1	.051
		Gecentreerde variabele voor opleidingsniveau	41.032	1	<,001
	Overall Statistics		58.821	3	<,001

Block 1: Method = Enter

Omnibus Tests of Model Coefficients				
		Chi-square	df	Sig.
Step 1	Step	62.796	3	<,001
	Block	62.796	3	<,001
	Model	62.796	3	<,001

Model Summary			
Step	-2 Log likelihood	Cox & Snell R Square	Nagelkerke R Square
1	1421.303 ^a	.051	.072

a. Estimation terminated at iteration number 4 because parameter estimates changed by less than .001.

Hosmer and Lemeshow Test			
Step	Chi-square	df	Sig.
1	43.806	8	<,001

Contingency Table for Hosmer and Lemeshow Test
--

		dummy voor het stemmen op een niet of wel linkse partij = .00		dummy voor het stemmen op een niet of wel linkse partij = 1.00		Total
		Observed	Expected	Observed	Expected	
Step 1	1	90	103.969	30	16.031	120
	2	92	97.864	28	22.136	120
	3	94	91.737	26	28.263	120
	4	102	87.648	18	32.352	120
	5	98	84.012	22	35.988	120
	6	83	81.198	37	38.802	120
	7	84	78.424	36	41.576	120
	8	75	74.594	45	45.406	120
	9	69	70.266	51	49.734	120
	10	45	62.288	77	59.712	122

Classification Table ^a					
Observed			Predicted		
			dummy voor het stemmen op een niet of wel linkse partij		Percentage Correct
			.00	1.00	
Step 1	dummy voor het stemmen op een niet of wel linkse partij	.00	821	11	98.7
		1.00	334	36	9.7
	Overall Percentage				71.3

a. The cut value is .500

Variables in the Equation						
		B	S.E.	Wald	df	Sig.
Step	Gecentreerde	-.126	.030	17.285	1	<,001

1 ^a	variabele voor het verkleind inkomen van het huishouden					
	Gecentreerde variabele voor leeftijd	-.002	.005	.260	1	.610
	Gecentreerde variabele voor opleidingsniveau	.400	.059	46.349	1	<,001
	Constant	-.884	.067	175.145	1	<,001

Variables in the Equation		
		Exp(B)
Step 1 ^a	Gecentreerde variabele voor het verkleind inkomen van het huishouden	.882
	Gecentreerde variabele voor leeftijd	.998
	Gecentreerde variabele voor opleidingsniveau	1.491
	Constant	.413

a. Variable(s) entered on step 1: Gecentreerde variabele voor het verkleind inkomen van het huishouden, Gecentreerde variabele voor leeftijd, Gecentreerde variabele voor opleidingsniveau.

Block 2: Method = Enter

Omnibus Tests of Model Coefficients				
		Chi-square	df	Sig.
Step 1	Step	147.321	1	<,001
	Block	147.321	1	<,001
	Mode	210.117	4	<,001

	I			
--	---	--	--	--

Model Summary			
Step	-2 Log likelihood	Cox & Snell R Square	Nagelkerke R Square
1	1273.982 ^a	.160	.226

a. Estimation terminated at iteration number 4 because parameter estimates changed by less than .001.

Hosmer and Lemeshow Test			
Step	Chi-square	df	Sig.
1	4.808	8	.778

Contingency Table for Hosmer and Lemeshow Test						
		dummy voor het stemmen op een niet of wel linkse partij = .00		dummy voor het stemmen op een niet of wel linkse partij = 1.00		Total
		Observed	Expected	Observed	Expected	
Step 1	1	110	112.724	10	7.276	120
	2	108	107.128	12	12.872	120
	3	99	101.727	21	18.273	120
	4	98	96.315	22	23.685	120
	5	93	90.927	27	29.073	120
	6	85	84.135	35	35.865	120
	7	84	76.507	36	43.493	120
	8	65	66.623	55	53.377	120
	9	52	56.117	68	63.883	120
	10	38	39.798	84	82.202	122

Classification Table ^a					
	Observed		Predicted		
			dummy voor het stemmen op een niet of wel linkse partij		Percentage Correct
			.00	1.00	
Step 1	dummy voor het stemmen op een niet of wel linkse partij	.00	752	80	90.4
		1.00	229	141	38.1
	Overall Percentage				74.3

a. The cut value is .500

Variables in the Equation						
		B	S.E.	Wald	df	Sig.
Step 1 ^a	Gecentreerde variabele voor het verkleind inkomen van het huishouden	-.076	.032	5.547	1	.019
	Gecentreerde variabele voor leeftijd	-.004	.005	.605	1	.437
	Gecentreerde variabele voor opleidingsniveau	.451	.062	53.737	1	<,001
	Gecentreerde variabele voor de houding ten opzichte van inkomensongelijkheid	.637	.058	121.969	1	<,001
	Constant	-.997	.074	181.853	1	<,001

Variables in the Equation

		Exp(B)
Step 1 ^a	Gecentreerde variabele voor het verkleind inkomen van het huishouden	.927
	Gecentreerde variabele voor leeftijd	.996
	Gecentreerde variabele voor opleidingsniveau	1.570
	Gecentreerde variabele voor de houding ten opzichte van inkomensongelijkheid	1.890
	Constant	.369

a. Variable(s) entered on step 1: Gecentreerde variabele voor het verkleind inkomen van het huishouden, Gecentreerde variabele voor leeftijd, Gecentreerde variabele voor opleidingsniveau, Gecentreerde variabele voor de houding ten opzichte van inkomensongelijkheid.

Block 3: Method = Enter

Omnibus Tests of Model Coefficients				
		Chi-square	df	Sig.
Step 1	Step	162.893	1	<,001
	Block	162.893	1	<,001
	Model	373.010	5	<,001

Model Summary			
Step	-2 Log likelihood	Cox & Snell R Square	Nagelkerke R Square
1	1111.088 ^a	.267	.376

a. Estimation terminated at iteration number 5 because parameter estimates changed by less than .001.

Hosmer and Lemeshow Test			
Step	Chi-square	df	Sig.
1	8.348	8	.400

Contingency Table for Hosmer and Lemeshow Test						
		dummy voor het stemmen op een niet of wel linkse partij = .00		dummy voor het stemmen op een niet of wel linkse partij = 1.00		Total
		Observed	Expected	Observed	Expected	
Step 1	1	116	117.223	4	2.777	120
	2	110	112.448	10	7.552	120
	3	108	107.814	12	12.186	120
	4	112	102.727	8	17.273	120
	5	94	95.749	26	24.251	120
	6	88	87.957	32	32.043	120
	7	73	77.440	47	42.560	120
	8	64	62.637	56	57.363	120
	9	43	44.993	77	75.007	120
	10	24	23.011	98	98.989	122

Classification Table ^a			
	Observed	Predicted	
		dummy voor het stemmen op een niet of wel linkse partij	Percentage Correct

			.00	1.00	
Step 1	dummy voor het stemmen op een niet of wel linkse partij	.00	739	93	88.8
		1.00	177	193	52.2
	Overall Percentage				77.5

a. The cut value is .500

Variables in the Equation						
		B	S.E.	Wald	df	Sig.
Step 1 ^a	Gecentreerde variabele voor het verkleind inkomen van het huishouden	-.098	.035	7.648	1	.006
	Gecentreerde variabele voor leeftijd	-.007	.006	1.614	1	.204
	Gecentreerde variabele voor opleidingsniveau	.267	.065	17.086	1	<,001
	Gecentreerde variabele voor de houding ten opzichte van inkomensongelijkheid	.601	.062	92.960	1	<,001
	Gecentreerde variabele voor de houding ten opzichte van immigranten	-1.085	.096	126.463	1	<,001
	Constant	-1.052	.081	169.512	1	<,001

Variables in the Equation		
		Exp(B)
Step 1 ^a	Gecentreerde variabele voor het verkleind inkomen van het huishouden	.907

	Gecentreerde variabele voor leeftijd	.993
	Gecentreerde variabele voor opleidingsniveau	1.307
	Gecentreerde variabele voor de houding ten opzichte van inkomensongelijkheid	1.824
	Gecentreerde variabele voor de houding ten opzichte van immigranten	.338
	Constant	.349

a. Variable(s) entered on step 1: Gecentreerde variabele voor het verkleind inkomen van het huishouden, Gecentreerde variabele voor leeftijd, Gecentreerde variabele voor opleidingsniveau, Gecentreerde variabele voor de houding ten opzichte van inkomensongelijkheid, Gecentreerde variabele voor de houding ten opzichte van immigranten.

Block 4: Method = Enter

Omnibus Tests of Model Coefficients				
		Chi-square	df	Sig.
Step 1	Step	.105	1	.746
	Block	.105	1	.746
	Model	373.115	6	<.001

Model Summary			
Step	-2 Log likelihood	Cox & Snell R Square	Nagelkerke R Square
1	1110.983 ^a	.267	.376

a. Estimation terminated at iteration number 5 because parameter estimates changed by less than .001.

Hosmer and Lemeshow Test			
Step	Chi-square	df	Sig.
1	8.207	8	.413

Contingency Table for Hosmer and Lemeshow Test						
		dummy voor het stemmen op een niet of wel linkse partij = .00		dummy voor het stemmen op een niet of wel linkse partij = 1.00		Total
		Observed	Expected	Observed	Expected	
Step 1	1	116	117.372	4	2.628	120
	2	110	112.607	10	7.393	120
	3	108	107.910	12	12.090	120
	4	112	102.765	8	17.235	120
	5	93	95.605	27	24.395	120
	6	88	87.704	32	32.296	120
	7	75	77.197	45	42.803	120
	8	63	62.416	57	57.584	120
	9	43	45.082	77	74.918	120
	10	24	23.341	98	98.659	122

Classification Table ^a			
	Observed	Predicted	
		dummy voor het stemmen op een niet of wel linkse partij	Percentage Correct

			.00	1.00	
Step 1	dummy voor het stemmen op een niet of wel linkse partij	.00	739	93	88.8
		1.00	177	193	52.2
	Overall Percentage				77.5

a. The cut value is .500

Variables in the Equation						
		B	S.E.	Wald	df	Sig.
Step 1 ^a	Gecentreerde variabele voor het verkleind inkomen van het huishouden	-.098	.035	7.600	1	.006
	Gecentreerde variabele voor leeftijd	-.007	.006	1.649	1	.199
	Gecentreerde variabele voor opleidingsniveau	.268	.065	17.172	1	<,001
	Gecentreerde variabele voor de houding ten opzichte van inkomensongelijkheid	.606	.065	87.901	1	<,001
	Gecentreerde variabele voor de houding ten opzichte van immigranten	-1.098	.105	109.853	1	<,001
	Interactievariabele tussen de gecentreerde kijk op ongelijkheid en kijk op immigranten	.024	.074	.105	1	.746
	Constant	-1.057	.082	164.381	1	<,001

Variables in the Equation		
		Exp(B)
Step 1 ^a	Gecentreerde variabele voor het verkleind inkomen van het huishouden	.907
	Gecentreerde variabele voor leeftijd	.993
	Gecentreerde variabele voor opleidingsniveau	1.307
	Gecentreerde variabele voor de houding ten opzichte van inkomensongelijkheid	1.834
	Gecentreerde variabele voor de houding ten opzichte van immigranten	.334
	Interactievariabele tussen de gecentreerde kijk op ongelijkheid en kijk op immigranten	1.024
	Constant	.347

a. Variable(s) entered on step 1: Gecentreerde variabele voor het verkleind inkomen van het huishouden, Gecentreerde variabele voor leeftijd, Gecentreerde variabele voor opleidingsniveau, Gecentreerde variabele voor de houding ten opzichte van inkomensongelijkheid, Gecentreerde variabele voor de houding ten opzichte van immigranten, Interactievariabele tussen de gecentreerde kijk op ongelijkheid en kijk op immigranten.

Bijlage 3

Bij dit onderzoek is gebruikgemaakt van logistische regressie, waarbij de onafhankelijkheid van de variabelen de enige relevante assumptie is. Deze assumptie houdt in dat de waarnemingen onafhankelijk van elkaar zijn. De gegevens uit het LISS-panel zijn in principe aselekt verzameld, maar het kan dat meerdere respondenten uit hetzelfde huishouden deelnemen. Hierdoor is de kans aanwezig dat deze assumptie wordt geschonden. Om te onderzoeken of dit invloed heeft op de resultaten van de regressieanalyse, is de analyse herhaald na een bewerking waarin per huishouden slechts één persoon is geselecteerd. Hieronder volgt de syntax voor deze bewerking:

```
SORT CASES BY nohouse_encr(A).
```

```
COMPUTE willekeurig=RV.UNIFORM(0,1).
```

```
EXECUTE.
```

```
SORT CASES BY nohouse_encr(A) willekeurig(A).
```

```
IF ($CASENUM = 1 OR nohouse_encr ~= LAG(nohouse_encr)) uniek_huishouden=1.
```

```
EXECUTE.
```

```
IF (NOT ($CASENUM = 1 OR nohouse_encr ~= LAG(nohouse_encr)))
```

```
uniek_huishouden=0.
```

```
EXECUTE.
```

```
COMPUTE casefilter=herkomstgroep = 0 AND dummy_observaties = 1 AND
```

```
uniek_huishouden = 1.
```

```
EXECUTE.
```

USE ALL.

COMPUTE filter_\$=(casefilter = 1).

VARIABLE LABELS filter_\$ 'casefilter = 1 (FILTER)'.

VALUE LABELS filter_\$ 0 'Not Selected' 1 'Selected'.

FORMATS filter_\$ (f1.0).

FILTER BY filter_\$.

EXECUTE.

LOGISTIC REGRESSION VARIABLES partij_dummy

/METHOD=ENTER inkomen_verkleind_gecentreerd leeftijd_gecentreerd
opleidingsniveau_gecentreerd

/METHOD=ENTER opleidingsniveau_gecentreerd leeftijd_gecentreerd
inkomen_verkleind_gecentreerd inkomensongelijkheid_gecentreerd

/METHOD=ENTER inkomen_verkleind_gecentreerd immigranten_gecentreerd
inkomensongelijkheid_gecentreerd opleidingsniveau_gecentreerd leeftijd_gecentreerd

/METHOD=ENTER inkomensongelijkheid_gecentreerd immigranten_gecentreerd interactie
inkomen_verkleind_gecentreerd opleidingsniveau_gecentreerd

leeftijd_gecentreerd

/SAVE=PRED COOK LEVER DFBETA

/PRINT=GOODFIT

/CRITERIA=PIN(0.05) POUT(0.10) ITERATE(20) CUT(0.5).

Uit deze regressieanalyse volgen dezelfde conclusies als in de originele analyse, waardoor het lijkt dat de mogelijke onderlinge afhankelijkheid van de variabelen geen (groot genoeg) invloed heeft gehad op de resultaten. Op basis hiervan wordt geconcludeerd dat deze mogelijke afhankelijkheid geen problemen oplevert.

Wat uitbijters betreft zijn er een aantal punten met een relatief hoge Cook's Distance. Onder een relatief hoge Cook's Distance verstaan we punten met een score hoger dan de algemene drempelwaarde van $4/N$. In dit model is dit namelijk een drempelwaarde van 0,0022, dus bij logistische regressie is het te verwachten dat er redelijk wat cases zijn die hier boven liggen. Bij nadere inspectie blijken dit vooral punten te zijn met een onverwachte samenstelling van houdingen ten opzichte van inkomensongelijkheid, immigranten en het wel of niet stemmen op een linkse partij. De case met de hoogste Cook's Distance (0,097) heeft namelijk een bovengemiddeld positieve kijk op inkomensongelijkheid, en heeft toch op een linkse partij gestemd. De case met de op één na grootste CD heeft dat ook, en daarnaast een bovengemiddeld negatieve kijk op immigranten, maar heeft toch op een linkse partij gestemd. Deze punten zijn dus wel degelijk invloedrijk op de fit van het model, maar hebben niet een dusdanig hoge afstand tot het model dat het een probleem zou vormen. Er zijn dan ook geen 'outliers' verwijderd uit het model, want hier lijkt geen noodzaak voor. Dit wordt bevestigd door de herhaalde regressieanalyse waarbij alle uitbijters zijn verwijderd en de conclusies intact blijven.

Bijlage 4 (procesverslag)

Toelichtende teksten

Deelopdracht 1 (19-2-2025)

Tot nu toe heb ik de bovenstaande artikelen alleen kort doorgenomen en gescand om te bekijken of de artikelen relevant voor mijn werkstuk zijn. Verder heb ik voor de statistische toetsen de voorbereiding voor de variabelen gedaan: ik heb de variabelen die ik ga toetsen operationeel gemaakt, en kort onderzocht hoe deze variabelen er uitzien.

Ik heb nog niet de nodige toetsen gedaan om te bekijken of de variabelen daadwerkelijk gecombineerd kunnen worden tot één variabele (Cronbach's Alpha) zoals ik dat van plan ben vanwege tijds limitaties, maar ik verwacht dit gedaan te hebben tegen de tijd ik daadwerkelijk bezig ga met de statistische testen, en ook verwacht ik geen problemen met het samenvoegen. Ik heb namelijk, zo denk ik, bij beide gevallen van het samenvoegen van variabelen gekozen voor de variabelen die het meest direct antwoord geven op de vragen in welke mate iemand negatief kijkt naar economische ongelijkheid of immigranten. Er waren vragen die ook inzicht gaven in deze meningen, maar omdat ze ook anders geïnterpreteerd konden worden of minder direct antwoord gaven heb ik ervoor gekozen om daar geen gebruik van te maken.

Daarnaast heb ik vooral uit eigen redenering bedacht hoe dit onderzoek te werk gaat en is het huidige onderzoeksmodel gebaseerd op mijn eigen intuïtie, en ik verwacht ook dat dit goed gaat kloppen en op basis van literatuur te onderbouwen valt.

Als hulpbronnen heb ik met name de richtlijnen van de opdracht zoals deze op Brightspace staat gebruikt, de literatuurlijst zoals die staat bij het document van de vraagstellingen en de libguides van sociologie van de rug. Voor de statistiek heb ik gebruik gemaakt van IBM SPSS 28, en om daar

genoeg gebruik van te kunnen maken heb ik ook de SPSS-syllabus van Mark Huisman en Frans Siero, degene die beschikbaar is op de cursuspagina van het vak Voortgezette Statistiek, geraadpleegd.

Deelopdracht 2 (3-3-2025)

Ik heb voor deze opdracht - op basis van de literatuur die ik op tijd heb kunnen bestuderen - een opzet gemaakt voor de theoretische onderbouwing van mijn bachelorwerkstuk. Ik denk dat ik de huidige onderbouwingen nog verder kan (en ga) onderbouwen en aanvullen met een uitgebreide bestudering van aanvullende literatuur, maar op dit moment is dit waar ik tijd voor had. Voor het vinden van relevante artikelen heb ik voornamelijk de literatuurlijst van het onderzoeksvoorstel gebruikt, en aanvullende literatuur heb ik gezocht via de SocIndex, en voor het laatste artikel heb ik de Scholar GPT module van ChatGPT gebruikt om op te zoeken. Hiernaast heb ik als hulpmiddel ChatGPT de gebruikte artikelen laten samenvatten *nadat* ik ze zelf heb gelezen en bestudeerd. Ik heb notities gemaakt van elk artikel dat ik heb gelezen, om vervolgens deze notities aan de hand van de samenvattingen compleet te maken. De verwijzingen komen dus vanuit de teksten zelf, *niet* vanuit de AI-gegenereerde samenvattingen. Buiten deze toepassingen heb ik geen gebruik gemaakt van AI. Met name de artikelen van Engelhardt & Wagener, Macdonald en Bartram waren belangrijke instrumenten voor het formuleren van het theoretisch kader, en hoewel ik er in de tekst niet naar verwijs is het artikel van La Roex et al. erg belangrijk geweest voor het verkrijgen van relevante achtergrondkennis. Met de feedback op het onderzoeksmodel ben ik vooral weer gaan nadenken over hoe alles tot stand is gekomen en hoe het gedefinieerd wordt, zo ben ik bij nader inzien tot de conclusie gekomen dat het onlogisch is om D66 te includeren bij de linkse partijen, omdat er eigenlijk heel weinig is wat bij hun programma aansluit op hoe linkse partijen in het onderzoek gedefinieerd en herkend worden. Voor de feedback: feedback op de schrijfstijl is niet per se nutteloos, maar ik ben zelf ook op de hoogte van het feit dat die nog niet zo wetenschappelijk is en lekker loopt als zou moeten. Voor nu heb ik de focus ook vooral gelegd op de inhoud en niet zo zeer de presentatie ervan; ik ben van plan om dat later beter uit te werken.

Deelopdracht 3 (2-4-2025)

Voor deelopdracht 3 heb ik een deel van de feedback op mijn deelopdracht 2 verwerkt, zo heb ik vooral de tekst opnieuw ingedeeld om een betere structuur te geven en nieuwe bronnen opgezocht om zo meer wetenschappelijk te kunnen onderbouwen. Hiervoor heb ik vooral gebruik gemaakt van de feedback van mijn begeleider, en daarnaast de literature guides van sociologie van de RUG, SmartCat en Google Scholar gebruikt om artikelen te vinden. Voor de methoden paragraaf heb ik uiteengezet welke variabelen gebruikt worden voor het onderzoek en waarom, en daarnaast welke operationalisaties hebben plaatsgevonden voor het onderzoek. Hiervoor heb ik gebruik gemaakt van IBM SPSS statistics, het LISS data archive en de SPSS syllabus van het vak Voortgezette Statistiek. Het kiezen van de variabelen is een belangrijke afweging geweest, maar ik denk dat deze in de tekst (met uitzondering van de last-minute wijziging van keuze in linkse partijen) goed genoeg heb onderbouwd om te kunnen rechtvaardigen.

Wat betreft AI heb ik een paar keer instrumenteel gebruik gemaakt van OpenAI's ChatGPT om te vragen of bepaalde elementen in gevonden artikelen voorkwamen, om vervolgens zelf in de artikelen dit te controleren en de artikelen zelf als bron van informatie te gebruiken. Verder heb ik geen gebruik gemaakt van AI-tools.

Deelopdracht 4 (2-5-2025)

Voor deze opdracht heb ik aan de hand van de feedback van de opdracht van voortgezette statistiek een eerste opzet van een resultatenparagraaf gemaakt. De inrichting van de paragraaf is nog niet final, maar de inhoudelijke interpretatie klopt nu als het goed is wel. Sinds de statistiek opdracht heb ik een beter inzicht gekregen in wat logistische regressie is en hoe je het interpreteert, en de 'grote fouten' uit de opdracht zijn als het goed is opgelost.

Deelopdracht 5 (19-5-2025)

Wat betreft de ontvangen feedback was ik het eens met de gegeven punten en heb ik ze op dit moment globaal geïntegreerd. Deze week was ik van plan de feedback nog verder structureel te verwerken, maar de belangrijkste punten zijn ondertussen verwerkt. Dat geldt ook voor de resultatenparagraaf, waarop ik het liefst de meest uitbundige feedback ontvang, omdat ik daar het meeste heb veranderd en niet zeker weet of die nu aan de vormvereisten voldoet. Een paar veranderingen zijn:

- Het veranderen van elke keer dat de term “sociaaleconomische ongelijkheid” o.i.d. is gebruikt naar enkel “economische ongelijkheid” om beter aan te sluiten bij de daadwerkelijke toetsing.
- Het verplaatsen van de bepaling wat wel of geen linkse partij is naar de methodenparagraaf.
- Het verminderen van subjectief taalgebruik.
- Algemene snelle verwerkingen van feedback, bijvoorbeeld taalfouten.

Zoals gezegd zou ik vooral graag feedback op de resultatenparagraaf krijgen: is die nu duidelijk genoeg? Ontbreekt er nog iets en zijn de bevindingen concreet genoeg uitgelegd? Is er nog iets anders dat veranderd moet worden? Ik zie het graag tegemoet.

Conclusie (4-6-2025)

Ik heb uiteindelijk nu een eindversie gemaakt van mijn bachelorwerkstuk. Ik heb het gevoel dat de feedback die ik heb gekregen bruikbaar was en dat ik die goed heb verwerkt, dus ik hoop dan ook dat deze aanpassingen uiteindelijk hebben gezorgd voor een goed bachelorwerkstuk. Ongeveer na de tweede deelopdracht heb ik een andere begeleider gekregen, van Jaap Nieuwenhuis naar Marinus Spreen, vanwege het vaderschapsverlof van Jaap. Deze verandering heeft niet erg veel last met zich meegebracht, zo was het namelijk ook wel prettig om feedback te krijgen van twee verschillende begeleiders op de eerdere onderdelen, maar uiteraard is het nog steeds een beetje lastig inschatten of beide begeleiders hetzelfde denken, of ze allebei gelijk hebben en of de nieuwe begeleider de eerdere onderdelen op dezelfde manier had aangepakt als de eerdere begeleider. Ik denk niet dat dit heel stroef is verlopen, maar het zou kunnen dat enige verschillen invloed hebben gehad op het uiteindelijke resultaat van mijn bachelorwerkstuk.

Terugblikkend ben ik tevreden met mijn onderwerp, dit is namelijk een onderwerp die mijn interesses goed liggen en ik had ook meteen het idee dat ik begreep met welk onderzoek ik bezig was, iets wat

niet altijd een gegeven is. Ik denk nu dat de feedback die ik heb gekregen goed en goed te begrijpen en verwerken was, maar dat moet natuurlijk nog wel blijken uit de uiteindelijke beoordeling. Als er namelijk onderdelen mankeren uit mijn onderzoek dat ik nergens in de feedback heb gehoord is dat namelijk wel enigszins onprettig, maar ik heb voor nu geen reden om daar vanuit te gaan, dus dat doe ik ook niet. Al met al vond ik dit een intensief maar erg leerzaam traject. Het eindresultaat is het veruit meeste werk dat ik ooit in een project heb gestopt, en ik hoop dat ik er later tevreden op kan terugkijken.

Verslag reparatie (01-08-2025)

Op basis van het beoordelingsformulier en de mondelinge feedback ben ik achter een grote hoeveelheid probleempunten van dit bachelorwerkstuk gekomen. Naar mijn idee was een van de grootste punten dat de conclusies niet klopten omdat de kansberekeningen verkeerd gedaan waren, maar die bleken alleen verkeerd gerapporteerd. Op basis van wat er in de tekst stond zouden ze inderdaad verkeerd gedaan zijn, alleen was ik de tekst simpelweg vergeten aan te passen. Uiteraard was dit niet het enige probleem, dus hier volgt per onderdeel wat ik heb veranderd/gedaan.

Abstract:

Ik heb deze samenvatting deels herschreven om te conclusies beter te nuanceren, want het leek inderdaad alsof er werd gezegd dat er op basis van ‘aanvullende analyses’ geconcludeerd kon worden dat er sprake was van moderatie. Dat is niet zo: er kan niet geconcludeerd worden dat er sprake is van moderatie, alleen lijkt op basis van visualisaties en kansberekeningen dat er wel verschillen bestaan tussen de groepen (zoals werd verondersteld). Hopelijk is dat nu helder, want anders wordt er inderdaad de verkeerde conclusie getrokken.

Inleiding:

Ik ben begonnen met het dieper onderzoek doen naar de artikelen die door Jaap zijn toegevoegd bij de vraagstellingen. Op basis hiervan ben ik tot (deels) nieuwe inzichten gekomen, waardoor de inleiding en het theoretisch kader opnieuw geschreven moesten worden. Ik heb dit gedaan op een manier

waarbij ik probeer zoveel mogelijk rekening te houden met de bestaande literatuur en hier voldoende naar te refereren. Verder heb ik (net als bij alle herschreven delen) geprobeerd de zinsformulering zo helder en begrijpelijk mogelijk te maken, zonder inhoud te verliezen. Tot slot heb ik de vraagstelling veranderd en duidelijk afgesplitst, zodat die nu wel kloppend is.

Theoretisch kader:

Ik heb het hele theoretisch kader opnieuw geschreven, omdat ik er toch achter kwam dat het grootste deel van de zinnen en de algemene structuur lastig te begrijpen en ongefocust was. Om dit anders aan te pakken ben ik begonnen met annotaties maken bij de literatuur, heb ik mechanismen geïdentificeerd en uitgewerkt en heb ik een soort schrijfplan voor de paragraaf gemaakt. Op basis hiervan kon ik een stuk helderder schrijven (denk ik) en heb ik er op gelet dat mechanismen en argumenten op basis van de literatuur zijn geformuleerd. Ook door de literatuur ben ik achter een compleet nieuw mechanisme voor de eerste hypothese gekomen, en zijn de mechanismen voor de tweede hypothese nu beter afgesplitst en kloppend. Ook heb ik het stuk van de controlevariabelen aangepast zodat nu de juiste variabelen worden geïntroduceerd (opleidingsniveau i.p.v. stedelijkheid).

Methoden:

Om deze paragraaf wat helderder te maken heb ik een paar onderdelen van de beschrijving van de variabelen aangepast om duidelijker en informatiever te zijn, en hopelijk ook meer volgens de richtlijnen van de Syllabus Schrijfvaardigheid te zijn. Eerlijk gezegd begreep ik niet helemaal hoe die volledig niet volgens de richtlijnen was dus het zou kunnen dat die wat dat betreft nog steeds niet helemaal klopt, maar inhoudelijk is die nu als het goed is in ieder geval wel juist.

Resultaten:

Ik heb aan de modevaluatie een stuk over de assumptie van onafhankelijke waarnemingen toegevoegd, die verder wordt uitgewerkt in bijlage 3. Op basis van een kritischere blik naar deze assumptie heb ik ook een manier gevonden om slechts één persoon per huishouden te selecteren en de analyse opnieuw uit te voeren, om na te gaan of dat andere resultaten opleverde. Dat was niet zo,

waardoor ik de keuze heb gemaakt om verder niet de centrale analyse aan te passen met deze variant (met een kleinere N), maar slechts te benoemen dat deze resultaten niet verschillen en ze in de bijlage te laten zien.

Verder heb ik ook hier geprobeerd om duidelijker te maken dat ik geen conclusie wil trekken op basis van de kansberekeningen of grafiek voor groepen, maar dat ze wel een interessant verschil met de regressieanalyse tonen. Verder heb ik de tekst aangepast om de correcte maxima te beschrijven (en dus niet te zeggen dat de ingevulde score 7 was), en de tabel voor de kansberekeningen overzichtelijker gemaakt. Ook heb ik de interpretatie van het interactie-effect aangepast om daadwerkelijk naar model 4 te kijken (net zoals alle belangrijke observaties). Ook heb ik aan het begin kort iets geschreven over de rol van de controlevariabelen, zodat dat ook besproken is. Verder heb ik de tabel aangepast om de correcte waarden te tonen en dat de bijlage ook klopt.

Notitie: de N is veranderd ten opzichte van de originele eindversie omdat ik achter een fout in mijn syntax kwam bij mijn afhankelijke variabele. Origineel rekende ik de partij Volt nog mee met linkse partijen, maar dat heb ik later veranderd. In de syntax had ik deze verandering niet doorgevoerd, waardoor de afhankelijke variabele niet klopte en de selectie van volledige waarden op basis daarvan dus ook niet. Hierdoor zijn alle coëfficiënten en kansen ook net anders dan de vorige versie, maar inhoudelijk verschillen de conclusies niet.

Conclusie en discussie

Ook deze heb ik herschreven omdat de redeneringen vaak inderdaad lastig te volgen waren, en de conclusie niet voldoende genuanceerd werd. Als het goed is is dit nu wel helder en zijn de zinnen wat makkelijk te volgen, ook als dit misschien iets ten koste is gegaan van de lengte van de conclusie. Ik denk namelijk wel dat de conclusie nu inhoudelijk sterker is, ondanks de kortere tekst.

Bijlage

Ik bleek inderdaad geen syntax aan mijn bijlage hebben toegevoegd, en dat was natuurlijk niet de bedoeling. Ik heb de hele bijlage opnieuw samengesteld door elk onderdeel kort toe te lichten, de

relevante syntax te tonen en vervolgens de relevante output van SPSS bij te voegen. In het geval van de logistische regressies heb ik voor de zekerheid gewoon de hele SPSS-output bijgevoegd, om er zeker van te zijn dat er geen mogelijk belangrijke onderdelen missen. De analyses zijn nu dus wel controleerbaar, en als het goed is komen ze overeen met de output in de bijlage.

Korte toevoeging over het AI-gebruik (uit het beoordelingsformulier): ik heb ChatGPT destijds de artikelen laten samenvatten nadat ik ze gelezen heb, puur om voor mezelf een makkelijke manier te hebben om een overzicht te maken van wat er in welk artikel gezegd werd. Sindsdien heb ik er voor gekozen om de artikelen zelf samen te vatten, ook omdat dit een goede manier is om de teksten daadwerkelijk te begrijpen en ze te kunnen toepassen, maar dat is wat de overweging was. Voor deze versie heb ik ook AI gebruikt, maar alleen voor de artikelen die de mogelijke verbanden van de controlevariabelen aantonen, omdat deze relatief moeilijk te vinden waren voor het minder grote belang van deze verbanden ten opzichte van de hoofdverbanden. Uiteraard heb ik wel gecontroleerd waar deze bronnen vandaan kwamen en of ze betrouwbaar waren, maar daar leek geen problemen mee te zijn. Verder heb ik geen gebruik gemaakt van AI-tools voor deze versie.

De literatuurlijst heb ik nu ook alfabetisch geordend.