



Universiteit
Leiden
The Netherlands

Verfransing onder de loep: Nederlands-Frans taalcontact (1500-1900) vanuit historisch-sociolinguïstisch perspectief

Assendelft, B.M.E.

Citation

Assendelft, B. M. E. (2023, May 24). *Verfransing onder de loep: Nederlands-Frans taalcontact (1500-1900) vanuit historisch-sociolinguïstisch perspectief*. LOT dissertation series. LOT, Amsterdam. Retrieved from <https://hdl.handle.net/1887/3618471>

Version: Publisher's Version

License: [Licence agreement concerning inclusion of doctoral thesis in the Institutional Repository of the University of Leiden](#)

Downloaded from: <https://hdl.handle.net/1887/3618471>

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

Hoofdstuk 7 Franse leenwoorden: periode en sociaal domein

7.1 Inleiding²⁵

In het vorige hoofdstuk kwam naar voren dat ongeveer een tiende van alle woorden met een van het Frans afgeleid suffix in het *LOL*-corpus productief gevormd is in het Nederlands. De overige negentig procent bestaat uit woorden die in zijn geheel uit het Frans (of een andere taal) zijn geleend. Waar de mogelijkheid tot het productief inzetten van leensuffixen voor woordvorming gezien kan worden als een verandering in de morfologie van het Nederlands, brengen woorden die worden overgenomen uit andere talen een verandering in het lexicon teweeg. De invloed die het Franse lexicon op het Nederlandse lexicon heeft gehad door middel van volledige woorden en hun betekenissen die het Nederlands aan het Frans heeft ontleend, staat centraal in dit hoofdstuk.

Ontlening van lexicale elementen is het meest voorkomende resultaat van taalcontact: verreweg het meeste ontleende materiaal is lexicaal en niet bijvoorbeeld fonologisch, morfologisch of syntactisch (Poplack 2018: 1; zie ook hoofdstuk 2, paragraaf 2.3 in dit proefschrift). In dit en de volgende twee hoofdstukken staan leenwoorden die het Nederlands uit het Frans heeft overgenomen centraal. Een ontlening heet een leenwoord als zowel de klank als de betekenis uit de brontaal wordt geleend. Er kan ook alleen een betekenis uit een andere taal worden overgenomen (betekenisontlening) of een vreemd woord kan worden vertaald en de betekenis van het vreemde woord krijgen (leenvertaling) (Van der Sijs 2005: 35). Een voorbeeld van een betekenisontlening is het woord *stem*: oorspronkelijk had dit alleen betrekking op ‘het geluid dat wordt geproduceerd bij spreken en zingen’, maar later is daar de betekenis ‘stem bij verkiezingen’ bij gekomen, onder invloed van het Franse *voix* ‘stem’, dat deze betekenis al had (Van der Sijs 2005: 549). Het woord *schoonmoeder* is een voorbeeld van een leenvertaling van het Franse *belle-mère*. In dit hoofdstuk wordt echter alleen gekeken naar leenwoorden en niet naar betekenisontleningen of leenvertalingen.

²⁵ Onderdelen van dit hoofdstuk zijn ook beschreven in Assendelft, Rutten & Van der Wal (te verschijnen-b).

Er bestaan verschillende theorieën over hoe leenwoorden precies van de ene taal in de andere worden overgenomen. Eén idee is dat een woord één op één wordt overgenomen uit de brontaal en vervolgens steeds vaker en door een steeds groter deel van de taalgemeenschap wordt gebruikt, waardoor het woord steeds meer wordt aangepast aan de ontvangende taal op het gebied van fonologie, morfologie en syntaxis en daar dan uiteindelijk ook onderdeel van wordt (Matras 2009: 172–173; Poplack & Dion 2012: 279). Poplack & Dion (2012: 308) hebben echter aangetoond dat morfologische en syntactische integratie van leenwoorden ook direct kan plaatsvinden, en niet pas als ze al wijdverspreid zijn. In elk geval resulteert lexicale ontlening niet alleen in één op één overgenomen woorden (zoals *rendez-vous* uit het Frans), maar ook in aangepaste vormen (bijvoorbeeld Japans *pokemon* van Engels *pocket monster*) zodat ze voldoen aan de structurele en semantische regels van de ontlenende taal (Winford 2019: 62).

Zoals besproken in hoofdstuk 2, paragraaf 2.3, is ontlening van lexicale elementen volgens ontleningshiërarchieën (Thomason & Kaufman 1992; Matras & Sakel 2007) makkelijker dan ontlening van morfologische elementen (waarbij lexicale morfologie makkelijker te ontlenen is dan flexionele morfologie). Ook binnen de categorie van leenwoorden is er verschil in de mate waarin ontleend wordt, afhankelijk van de woordklasse. Volgens Haugen (1950) worden substantieven het vaakst geleend, gevolgd door werkwoorden of adjectieven/bijwoorden, en voorzetsels en interjuncties zijn de woorden die het minst vaak worden ontleend (Matras 2009: 167–168; Field 2002: 35). Het *Loanword Typology Project* (Haspelmath & Tadmor 2009) vond een hogere frequentie van ontleende adjectieven dan werkwoorden in de talen van de wereld. De specifieke volgorde van alle categorieën behalve substantieven is echter afhankelijk van de specifieke taalcontactsituatie (Field 2002: 35).

Dat lexicale ontlening van het Frans naar het Nederlands veelvuldig plaats heeft gevonden in de geschiedenis van het Nederlands, is te zien aan de grote hoeveelheid Franse leenwoorden die in de loop van de tijd in de Nederlandse woordenschat is opgenomen. Volgens Van der Sijs (2009: 196) had het Nederlands tot het begin van de 21^e eeuw tussen de 4.000 en 5.000 woorden uit het Frans overgenomen. Daarmee is het Frans de belangrijkste brontaal van leenwoorden voor het Nederlands; het Latijn staat op een tweede plek met circa 2.000 leenwoorden. Dit betekent echter niet dat al deze leenwoorden op hetzelfde moment deel uitmaken van de Nederlandse woordenschat. Leenwoorden komen niet op hetzelfde moment de taal binnen en

sommige woorden verouderen na verloop van tijd ook weer. Het werkwoord *defroyeren* ‘onkosten vergoeden’ bijvoorbeeld werd in de vijftiende eeuw uit het Frans overgenomen en werd volgens het *WNT* tot het begin van de negentiende eeuw gebruikt in het Nederlands, maar nu is dat woord niet meer in de *Van Dale* te vinden. Aan de andere kant kwam een woord als *lokaliseren* volgens het *WNT* pas in het Nederlands voor vanaf ongeveer 1850. Zowel *defroyeren* als *lokaliseren* is een Frans leenwoord, maar ze zijn waarschijnlijk niet op hetzelfde moment in gebruik geweest in het Nederlands.

De Franse leenwoorden hebben zich in meer of mindere mate aangepast aan de grammatica van het Nederlands. Zo worden werkwoorden vervoegd volgens het Nederlandse werkwoordsysteem (bijv. *defroyeren* – *defroyeert* – *defroyeerde* – *gedefroyeerd*) en krijgen zelfstandig naamwoorden een lidwoord (en daarmee een grammaticaal genus): het is bijvoorbeeld *het reçu* en *de directeur*. Ook kunnen naamwoorden Nederlandse morfologie krijgen, zoals het diminutiefsuffix *-(t)je* (bijvoorbeeld in *reçuutje*). Sommige leenwoorden hebben een dermate grote verandering doorgemaakt, dat ze niet meer lijken op de oorspronkelijke Franse vorm en daardoor ook moeilijk herkenbaar zijn als Frans leenwoord. Het woord *toren* zou bijvoorbeeld ontleend zijn aan Oudfrans *tur* ‘hoog bouwwerk’ (Modern Frans *tour*), een woord dat volgens het *Oudnederlands Woordenboek* (ONW) al in de negende eeuw in het Nederlands voorkwam. Heel anders is dat bij het woord *depêche*, een zestiende-eeuws leenwoord. Hierbij is de huidige spelling vrijwel hetzelfde als de Franse, inclusief het accent circonflexe, waardoor het onmiskenbaar een Frans leenwoord is.

In dit en de volgende twee hoofdstukken worden de resultaten besproken van het leenwoordenonderzoek met het *LOL*-corpus. Het doel van dit onderzoek is om tot een zo compleet mogelijk beeld te komen van de invloed die het Frans op de Nederlandse woordenschat heeft gehad op basis van een representatief corpus. In het huidige hoofdstuk wordt behandeld hoe het gebruik van Franse leenwoorden zich ontwikkelt van de zestiende tot en met de negentiende eeuw en in de verschillende sociale domeinen, zowel voor de tokens als voor de types leenwoorden. Daarbij vergelijk ik de resultaten met die van het leensuffixonderzoek om te kijken in hoeverre die met elkaar overeenkomen. In de volgende twee hoofdstukken bespreek ik de resultaten van enkele andere analyses die ik met de leenwoorden heb uitgevoerd: de verhouding tussen verschillende woordsoorten, het jaar van ontlening, de frequentie van individuele leenwoorden, de verhouding

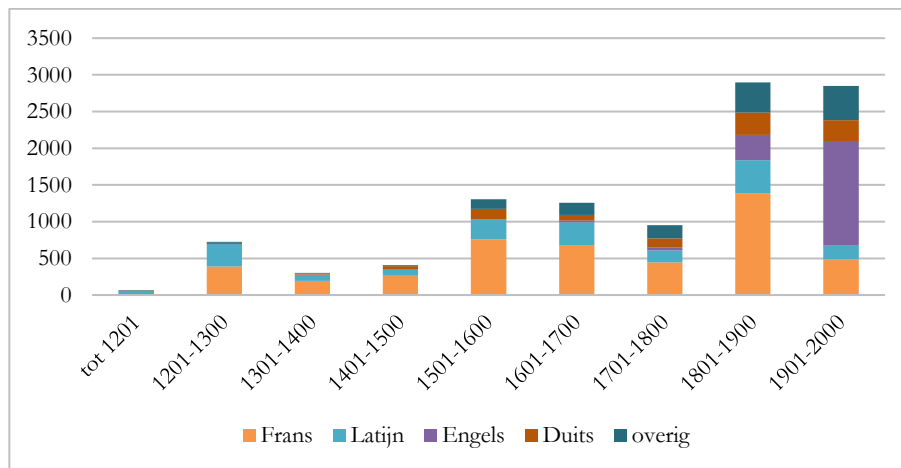
tussen terminologische en algemene woordenschat en de mate van integratie van de leenwoorden in de Nederlandse woordenschat. Deze analyses samen worden gebruikt om antwoord te geven op de vraag in hoeverre het Frans invloed heeft gehad op het Nederlandse taalgebruik op het niveau van het lexicon.

In dit hoofdstuk leg ik, na een bespreking van het bestaande onderzoek naar Franse leenwoorden in paragraaf 7.2, in paragraaf 7.3 uit op welke manier ik de Franse leenwoorden uit het *LOL*-corpus heb gehaald. Bij de resultaten die in paragraaf 7.4 worden besproken, ga ik in op de invloed van periode en domein op het gebruik van Franse leenwoorden: hier wordt de vraag beantwoord hoe het gebruik van Franse leenwoorden zich ontwikkelt gedurende de eeuwen die het *LOL*-corpus beslaat en hoe de verschillende domeinen zich tot elkaar verhouden met betrekking tot het gebruik van leenwoorden. Deze resultaten worden ook vergeleken met die van het leensuffixonderzoek dat is besproken in hoofdstuk 5. Na een kort besluit in paragraaf 7.5 volgen in hoofdstuk 8 en 9 de andere leenwoordanalyses die ik heb uitgevoerd.

7.2 Eerder onderzoek

Er bestaan al diverse overzichten van Franse leenwoorden die in de loop van de tijd in het Nederlands zijn opgenomen. Salverda de Grave (1906) geeft lijsten van Franse woorden in het Nederlands, die hij samenstelde op basis van zijn eigen observaties en op basis van diverse Nederlandse woordenboeken en woordenlijsten. Hij maakte een overzicht van alle Franse leenwoorden gegroepeerd naar verschillende semantische domeinen als ‘kunst en wetenschap’ en ‘mens in zijn openbaar leven’. Het *Groot Leenwoordenboek* van Van der Sijs (2005) geeft een overzicht van leenwoorden uit alle talen waar het Nederlands aan heeft ontleend, waaronder dus ook Franse leenwoorden. Hierin is echter alleen gekeken naar leenwoorden die anno 2005 nog tot de Nederlandse woordenschat behoorden (gebaseerd op onder meer de leenwoorden die in de *Grote Van Dale* staan (Van der Sijs 2005: 99)). In het *Chronologisch Woordenboek* van Van der Sijs (2001) tot slot wordt de diachrone ontwikkeling van de Nederlandse woordenschat in kaart gebracht, waarbij ook aandacht is voor de herkomst van de woorden. Van der Sijs (2001: 214) geeft hierbij een overzicht van het aantal leenwoorden dat per eeuw uit de vier belangrijkste brontalen – Latijn, Frans, Duits en Engels – is over-

genomen. In figuur 7.1 is duidelijk te zien dat het Frans van de dertiende tot en met de negentiende eeuw de voornaamste herkomsttaal van leenwoorden was. Tussen de vijftiende en zestiende eeuw is een toename te zien van het aantal ontleende Franse woorden, waarna het in de twee volgende eeuwen weer wat afneemt, om in de negentiende eeuw weer sterk toe te nemen. Op basis van het aantal ontlenen per eeuw zou de negentiende eeuw dus de belangrijkste eeuw qua invloed van het Frans zijn geweest.

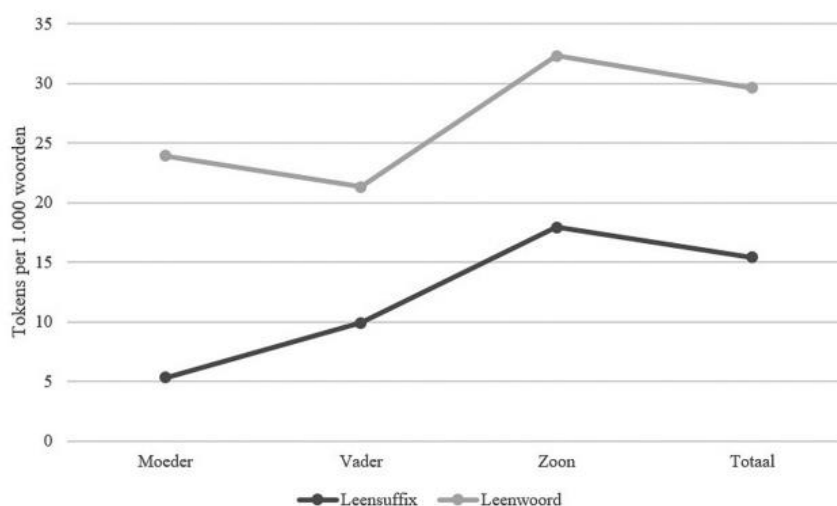


Figuur 7.1. Leenwoorden in het Nederlands naar eeuw van ontlening op basis van Van der Sijs 2001: 214.

Net als de overzichten van Salverda de Grave (1906) en Van der Sijs (2005) is ook dit onderzoek echter gebaseerd op woorden die in woordenboeken te vinden zijn, waarbij vaktaalwoorden en vreemde woorden veelal zijn weggelaten (Van der Sijs 2001: 32). Geen van de besproken overzichten van leenwoorden baseert zich op empirisch leenwoordenonderzoek, waarbij naar het daadwerkelijke gebruik van deze leenwoorden in teksten wordt gekeken.

Stevens (2019) is het enige empirische corpusonderzoek naar Franse leenwoorden in het Nederlands vanuit een historisch-sociolinguïstische invalshoek. Naast het leensuffixonderzoek, dat in hoofdstuk 5 is besproken, heeft Stevens (2019) ook onderzocht hoe veel Franse leenwoorden voorkomen in het briefencorpus van de familie Bijleveld. Hierbij is met behulp van de etymologische informatie die het *WNT* en de *Etymologiebank* bevatten, nagegaan welke woorden in het

Bijleveldcorpus Franse leenwoorden zijn. Deze woorden zijn met de hand uit het corpus gehaald. Vervolgens is de relatieve frequentie van de leenwoorden per woordsoort berekend (Stevens 2019: 146). Van de 179 typen leenwoorden die in het corpus voorkwamen, behoorden verreweg de meeste tot de zelfstandig naamwoorden, 96 in totaal; daarnaast waren er 44 werkwoorden, 27 bijvoeglijk naamwoorden, 3 tussenwerpsels en 2 voorzetsels (Stevens 2019: 150). In figuur 7.2 worden het totaal aantal tokens leensuffixen per 1.000 woorden en het aantal tokens leenwoorden per 1.000 woorden vergeleken.



Figuur 7.2. Vergelijking leensuffixen en leenwoorden per 1.000 woorden in het Bijleveldcorpus. Uit Stevens 2019: 151.

Figuur 7.2 laat zien dat er in het corpus meer Franse leenwoorden dan Franse leensuffixen te vinden zijn: in totaal zitten er 15,4 leensuffixen en 29,6 leenwoorden per 1.000 woorden in het corpus (Stevens 2019: 151). Dit is voor alle woordsoorten en ook voor elk van de drie briefschrijvers het geval. Stevens (2019: 151) concludeert dan ook dat ‘een handmatige analyse van leenwoorden in een klein corpus een waardevolle aanvulling vormt op een automatisch uitgevoerde analyse van leensuffixen in een groot corpus’.

Het corpusonderzoek van Stevens (2019) kan gezien worden als een exploratieve studie en beperkte zich dan ook tot slechts één sociaal domein (Privé), één periode (begin negentiende eeuw) en tot enkele

leden van één familie. In het huidige onderzoek ga ik met het *LOL*-corpus na hoe het aantal Franse leenwoorden in verschillende sociale domeinen zich tot elkaar verhoudt en hoe het gebruik van deze leenwoorden zich ontwikkelt over een periode van vier eeuwen, om op die manier te achterhalen wat de daadwerkelijke invloed van het Frans op het Nederlandse lexicon is geweest. Net als Stevens (2019) zal ik de resultaten van het leenwoordenonderzoek vergelijken met de resultaten van het leensuffixonderzoek die in hoofdstuk 5 zijn gepresenteerd.

7.3 Methode

7.3.1 Werkwijze datavergaring

Een grote uitdaging bij het huidige leenwoordenonderzoek was de datavergaring: de vraag was hoe ik zo veel mogelijk Franse leenwoorden uit het *LOL*-corpus kon filteren op een manier die haalbaar zou zijn. Aangezien het corpus gedigitaliseerd is, ligt een automatische zoekmethode voor de hand, waarbij net als bij het leensuffixonderzoek met een zoekprogramma als *AntConC* wordt gezocht op woorden of woorddelen. Bij het zoeken naar Franse leenwoorden is het echter de vraag welke automatische zoekmethode kan worden gebruikt, aangezien er duizenden Franse leenwoorden zijn die weinig gemeenschappelijke kenmerken hebben, wat automatisch zoeken lastig maakt. Een mogelijkheid is om een leenwoordenboek als Van der Sijs (2001) of (2005) erbij te pakken en alle woorden die daar als Franse leenwoorden zijn opgetekend, in *AntConC* in te voeren. Het kost echter veel tijd om honderden woorden op deze manier met de hand in te voeren en daarnaast is er een grote kans dat allerlei voorkomens gemist worden vanwege afwijkende spellingen. Ook hoeft het niet zo te zijn dat alle Franse leenwoorden in de woordenboeken staan. Het woord *comparant* staat bijvoorbeeld niet in Van der Sijs (2005), terwijl dat wel een (zeer) frequent leenwoord in het *LOL*-corpus is.

Een andere mogelijkheid om op geautomatiseerde wijze zo veel mogelijk Franse leenwoorden uit het *LOL*-corpus te filteren, is door te zoeken op veel voorkomende letters of lettercombinaties in Franse woorden. De *c* en de *q* zijn bijvoorbeeld letters die oorspronkelijk niet voorkomen in Nederlandse (en meer algemeen Germaanse) woorden, maar wel veel in Franse (Van der Sijs 2005: 63). Zoeken op de *c* en de *q* in *AntConC* zou dus veel Franse leenwoorden op moeten leveren en een minimale ruis. Het probleem is hierbij echter dat er alsnog veel woorden

die in het Frans met een *c* of een *q* worden gespeld, worden gemist, omdat de spelling is aangepast aan het Nederlands. Zo wordt de *c* vaak een *k* (vergelijk Frans *copie*, *colossal* met de Nederlandse spellingsvarianten *kopie*, *kolossaal*) en de *q* een *k* of *kw* (Frans *quitance*, *publique* en *request* versus Nederlands *kwitantie*, *publiek* en *rekest/rekwest*). Weliswaar heeft een deel van dit soort woorden die ik in het *LOL*-corpus heb aangetroffen de Franse *c* en *q* behouden, maar daarnaast komen de Nederlandse spellingen ook voor. Bovendien zijn er genoeg Franse leenwoorden die geen van deze klanken bevatten, wat het niet aannemelijk maakt dat een zoekopdracht met *c* en *q* een representatief beeld oplevert van de Franse leenwoorden die in het *LOL*-corpus zitten.

Een derde manier om zo veel mogelijk leenwoorden uit het corpus te halen, is op basis van de leensuffixanalyse die in het vorige hoofdstuk is besproken. De meeste Nederlandse woorden die een Frans suffix bevatten, zijn immers oorspronkelijk in zijn geheel leenwoorden uit het Frans. Het zou kunnen zijn dat de inventaris van de leensuffixen uit het vorige hoofdstuk al het grootste deel van de leenwoorden omvat. Stevens (2019) toonde echter al aan dat in haar corpus slechts ongeveer de helft van de Franse leenwoorden een leensuffix bevat, en concludeerde dan ook dat het voor leenwoordenonderzoek waardevol is om een corpus handmatig door te nemen in aanvulling op automatisch uitgevoerd suffixonderzoek (zie paragraaf 7.2). Het corpus van Stevens (2019) bevatte echter slechts 15.000 woorden, waardoor het handmatig doornemen ervan te overzien was. Anders is dit bij een groot corpus als het *LOL*-corpus, dat ruim 250.000 woorden bevat: dit hele corpus met de hand doornemen is erg arbeidsintensief en zou veel tijd in beslag nemen. Om die reden is ervoor gekozen om eerst een pilotstudy te doen, waarbij ik van elk van de zeven domeinen twee tijdvakken uitkoos (1600–1649 en 1800–1849, dus een vroege en een late periode; vanwege het ontbreken van materiaal voor de periode 1600–1649 bij Publieke Opinie zijn daar de tijdvakken 1650–1699 en 1850–1899 genomen). Van deze tijdvakken nam ik alle teksten met de hand door en filterde ik alle leenwoorden eruit. Afhankelijk van de duur van de analyse en de uitkomsten zou ik besluiten hoe ik de andere tijdvakken zou aanpakken. De pilotstudy zou bijvoorbeeld tot de conclusie kunnen leiden dat de suffixanalyse toch al een groot deel van de Franse leenwoorden had opgeleverd: Stevens (2019) had immers leensuffixen als *-ment* en *-(er)ij/-(er)ie* niet meegenomen, terwijl die in het suffixonderzoek met het *LOL*-corpus wel bij de analyse zijn betrokken. Mijn aanpak van de suffixanalyse zou dus al meer Franse leenwoorden moeten hebben

opgeleverd. De Franse leenwoorden zonder leensuffix die in de pilotstudy werden gevonden, zouden dan vervolgens met behulp van *AntConC* in de overige tijdvakken, die geen deel uitmaakten van de pilot, kunnen worden opgezocht. Op deze manier kon de datavergaring bij de rest van het corpus dus gebaseerd worden op de resultaten van het leensuffixonderzoek plus zoekacties met *AntConC* op leenwoorden zonder leensuffix die uit de pilot waren gekomen.

De conclusie na de pilot was dat het met de hand doornemen van het corpus weliswaar veel tijd kost, maar dat dit toch de enige manier is om een goed beeld te krijgen van het leenwoordgebruik gedurende alle eeuwen die het corpus beslaat. Het bleek dat er in alle perioden en domeinen meer leenwoorden zonder Franse suffixen dan met Franse suffixen zaten, en ook waren er grote verschillen te zien in de specifieke leenwoorden die werden gebruikt in de onderzochte teksten, zowel tussen domeinen als binnen één domein tussen twee tijdvakken. Bij het domein Publieke Opinie bijvoorbeeld kwamen in de periode 1650–1699 leenwoorden voor als *plakkaat*, *tartaan*, *patent* en *fluweel*, die in geen enkele andere periode voorkwamen; bij 1850–1899 waren er dan weer woorden als *flanel*, *schaats*, *pers* en *Sire*, die alleen in deze periode in krantenartikelen werden gebruikt. Het ligt dus voor de hand dat in andere perioden ook veel unieke leenwoorden worden gebruikt, die dan niet in de analyse worden meegenomen wanneer je je zou baseren op een leenwoordenlijst van één zeventiende-eeuwse en één negentiende-eeuwse periode.

Bij bovenstaande voorbeeldwoorden had het in principe wel gekund dat die in de andere perioden voor zouden komen: de unieke leenwoorden uit 1650–1699 bestaan in het huidige Nederlands nog steeds, en de leenwoorden uit de periode 1850–1899 kwamen ook al voor in Nederlandse teksten uit de eerste helft van de zeventiende eeuw. Het is echter ook mogelijk dat er in een bepaalde periode leenwoorden in de teksten zitten die alleen in die periode of in een beperkt aantal perioden voor kunnen komen, omdat ze in een volgende periode weer uit de taal zijn verdwenen of omdat ze in eerdere perioden nog niet in het Nederlands bestonden. In het domein Publieke Opinie periode 1850–1899 komen bijvoorbeeld de woorden *desinfectie* en *coöperatief* voor, die beide pas na 1850 zijn geleend. Het is dus onmogelijk om deze leenwoorden te vinden met behulp van leenwoordenlijsten gebaseerd op eerdere perioden.

Naast deze bezwaren tegen de methode om met leenwoordenlijsten de leenwoorden uit de rest van het corpus te halen, was het ook nog de vraag of deze methode echt zo tijdbesparend zou zijn als die lijkt.

Het handmatig invoeren in *AntConC* van alle leenwoorden die in de pilotstudy waren gevonden en het filteren van alle ruis is erg bewerkelijk en zou mogelijk even veel werk zijn als het handmatig doornemen van alle teksten. Om deze redenen heb ik er dan ook voor gekozen om de methode van de pilotstudy voort te zetten voor de rest van het corpus en alle leenwoorden handmatig uit het corpus te halen. Dit was zoals verwacht erg arbeidsintensief: afhankelijk van het domein en tijdvak (sommige domeinen en tijdvakken bevatten nu eenmaal meer leenwoorden dan andere) kon dit wel twee volle dagen per domein en per tijdvak kosten. Het voordeel is echter dat er vanuit kan worden gegaan dat vrijwel alle Franse leenwoorden uit het corpus konden worden gehaald. Er kan niet uitgesloten worden dat er tijdens het doornemen van de teksten een woord over het hoofd is gezien, maar dat zal vermoedelijk tot een enkel geval beperkt zijn gebleven.

7.3.2 Welke woorden zijn Franse leenwoorden?

Bij het extraheren van de Franse leenwoorden uit het corpus is het uiteraard cruciaal om vast te stellen of een woord daadwerkelijk een Frans leenwoord is of niet. Om dit te bepalen, heb ik gebruikgemaakt van verschillende etymologische woordenboeken: het *WNT* en de woordenboeken die zijn opgenomen in de *Etymologiebank* (met name Philippa e.a. 2003-2009; Van der Sijs 2001; Van der Sijs 2005; Van Veen & Van der Sijs 1997). Voor elk potentieel Frans leenwoord zocht ik dit woord in beide online bronnen op middels de zoekfunctie op de respectievelijke sites. Vervolgens keek ik per woordenboek welke etymologische brontaal daarbij genoemd werd. Bij het woord *plezier* bijvoorbeeld staat zowel in het *WNT* als in alle woordenboeken in de *Etymologiebank* dat het ontleend is aan het Franse woord *plaisir*. Dit is dus duidelijk een Frans leenwoord.

Er zijn echter ook zeer veel woorden waarbij niet onomstotelijk vaststaat dat het woord uit het Frans afkomstig is. Over het woord *feest* bijvoorbeeld zeggen vijf etymologische woordenboeken die in de *Etymologiebank* staan dat het uit het Frans afkomstig is; twee geven aan dat het uit het Frans of uit het Latijn komt, en één woordenboek geeft alleen het Latijn als brontaal. Het *WNT* geeft dan weer alleen het Franse *fête* (Oudfrans *feste*) als oorsprong. Het woord *feest* staat daarin zeker niet op zichzelf: ook over woorden als *familie*, *meester*, *plaats*, *staat*, *uur*, *advies*, *persoon* en vele meer bestaat in de etymologische woordenboeken geen

consensus of het Frans de directe brontaal is voor het Nederlands (waarbij het Frans dit woord eerder al aan het Latijn heeft ontleend) of dat het woord direct uit het Latijn in het Nederlands is overgenomen. Ook is er een klein groepje woorden dat niet het Latijn als tweede mogelijke brontaal naast het Frans heeft, maar bijvoorbeeld het Engels (zoals bij *conquest*) of het Italiaans (zoals bij *compliment* en *post*). Bij verreweg de meeste woorden is echter alleen het Latijn een alternatieve bron voor het leenwoord, naast het Frans.

Ik heb besloten om voor de leenwoordanalyse alle leenwoorden mee te nemen die een link met het Frans (kunnen) hebben, dus zowel de leenwoorden waarbij alleen het Frans als brontaal wordt genoemd door de etymologische woordenboeken als de leenwoorden waarbij naast het Frans ook nog een of meer andere talen als brontaal kunnen hebben gefungeerd. Voor de analyse maak ik echter wel onderscheid tussen de groep leenwoorden waarvan zeker is dat die uit het Frans komen en de groep woorden die mogelijk een andere brontaal hebben. Ik heb de leenwoorden dus ingedeeld in de groep ‘zeker uit het Frans’ of ‘mogelijk uit het Frans’, afhankelijk van de etymologische informatie die beschikbaar was. Als alle woordenboeken voor een bepaald woord aangaven dat het woord uit het Frans kwam, zonder een andere brontaal te noemen, zoals bij *plezier*, dan categoriseerde ik dit leenwoord als ‘zeker Frans’. Op het moment dat ten minste één van de woordenboeken aangaf dat het woord ofwel uit het Frans kon komen ofwel uit een andere taal, of wanneer het Frans door minstens één woordenboek helemaal niet als brontaal werd genoemd en hetzelfde woord tegelijkertijd door ten minste één ander woordenboek wel als Frans leenwoord werd benoemd, dan deelde ik dit woord in bij de categorie ‘mogelijk uit het Frans’.

Er is echter wel een verschil tussen de leenwoorden in de hoeveelheid etymologische informatie die beschikbaar is. Zo zijn er woorden die in alle woordenboeken voorkomen, zoals het (zeer frequente) woord *plaats*, maar er zijn ook minder frequente woorden die niet door alle woordenboeken worden vermeld. Er zijn zelfs woorden die in slechts één woordenboek voorkomen. Een voorbeeld hiervan is *contentement*, dat niet in de *Etymologiebank* voorkomt maar alleen in het *WNT*. Aangezien het *WNT* dit woord aanmerkt als Frans leenwoord en dit de enige bron is waarop ik mij kan baseren, is dit woord in de categorie ‘zeker Frans’ geplaatst. Het woord *feest* dat hierboven is besproken, komt in negen verschillende woordenboeken voor, waarvan er zes aangeven dat het Frans is, twee dat het Frans of Latijn is en één dat het alleen Latijn is. Er gaan dus duidelijk meer stemmen op voor het

Frans als brontaal, maar aangezien door sommige bronnen het Latijn ook als optie wordt gegeven, is dit woord bij de leenwoorden ‘mogelijk uit het Frans’ ingedeeld.

Naast het bepalen van de brontaal zijn de etymologische woordenboeken ook gebruikt om het jaar van ontlening te bepalen. Ook hier is niet altijd overeenstemming tussen de verschillende woordenboeken, maar hier heb ik ervoor gekozen om het oudste jaartal dat in de woordenboeken wordt genoemd als uitgangspunt te nemen. Het is niet verwonderlijk dat het jaar van ontlening af kan wijken: de etymologische woordenboeken zelf zijn samengesteld gedurende de gehele twintigste en eenentwintigste eeuw en zijn dan ook gebaseerd op verschillende bronnen. Naarmate de tijd vordert kunnen er natuurlijk nieuwe teksten opduiken met vroegere attestaties van een bepaald woord, waardoor de etymologische informatie kan worden aangescherpt. Ook in het *LOL*-corpus komen enkele leenwoorden voor in teksten die uit een eerdere periode komen dan de vroegste attestatie van dit leenwoord in de etymologische woordenboeken; in deze gevallen is het jaartal van de tekst uit het *LOL*-corpus als jaar van ontlening beschouwd. Van enkele leenwoorden die alleen in het *WNT* voorkomen, heb ik zelf op basis van het citatenmateriaal van het *WNT* de oudste attestatie van het betreffende woord moeten bepalen, aangezien het *WNT* vaak geen jaar van ontlening bij de lemma's geeft. Deze werkwijze komt overeen met die van veel etymologische woordenboeken in de *Etymologiebank*: deze geven als vindplaats voor de oudste attestatie van een woord vaak het *WNT*.

In principe zijn alle woorden die als Franse leenwoorden beschouwd kunnen worden meegenomen in het onderzoek, inclusief afleidingen en samenstellingen met Franse leenwoorden. *Plaats* is bijvoorbeeld een Frans leenwoord met Nederlandse afleidingen als *plaatsen* en *(ver)plaatsing* en samenstellingen als *bewaarplaats*. Al dit soort voorkomens zijn meegeteld bij de resultaten van het leenwoord *plaats*. Ook afgekorte leenwoorden zijn meegenomen. Bij het leensuffix-onderzoek in de vorige twee hoofdstukken zijn afkortingen niet meegeteld, omdat bij afkortingen vaak het suffix is weggelaten, terwijl in dat onderzoek het suffix juist centraal stond. Afkortingen van Franse leenwoorden die in niet-afgekorte vorm een leensuffix bevatten, zoals *maj.* voor *majesteit* en *comp.* voor *comparant*, zijn in het huidige leenwoordenonderzoek echter wel meegenomen, omdat het woord als geheel een Frans leenwoord is.

Enkele tientallen woorden die uit het Frans of uit het Latijn afkomstig kunnen zijn, vertonen Latijnse morfologie in het corpus bij de

domeinen Academie en Religie. Voorbeelden zijn *collegii* en *synodus*. De woorden *college* en *synode* zijn uit het Frans of uit het Latijn geleend en zouden daarom in de categorie ‘mogelijk uit het Frans’ kunnen worden geplaatst, maar vanwege de morfologische uitgangen is het waarschijnlijker dat deze woorden in deze vorm uit het Latijn afkomstig zijn en niet uit het Frans. Deze voorkomens zijn daarom niet meegenomen in de analyse.

7.4 Resultaten: periode en sociaal domein

In deze paragraaf bespreek ik de resultaten van het leenwoordenonderzoek per periode (7.4.1), per domein (7.4.2) en per periode per domein (7.4.3). Hierbij vergelijk ik de resultaten ook met de resultaten van de leensuffixen uit hoofdstuk 5 om te zien of het patroon dat uit de twee onderzoeken komt hetzelfde is. Hoofdstuk 8 en 9 gaan in op enkele andere onderzoeksvragen die de leenwoordanalyse oproept, waaronder de verdeling van de leenwoorden over de verschillende woordsoorten, het jaar van ontlening en de mate van integratie van de leenwoorden.

7.4.1 Leenwoorden per periode

Zoals toegelicht in paragraaf 7.3.2 zijn de resultaten van de leenwoorden opgedeeld in de categorieën ‘zeker uit het Frans’ en ‘mogelijk uit het Frans’. Voor alle perioden en alle domeinen bij elkaar zijn voor de categorie ‘zeker uit het Frans’ 8.767 tokens gevonden en voor de categorie ‘mogelijk uit het Frans’ 6.419 tokens. De data in beide categorieën zijn deels gelijk aan die van de leensuffixanalyse uit het vorige hoofdstuk: de meeste woorden die een uit het Frans afkomstig suffix bevatten, zijn Franse leenwoorden en tellen dus ook mee bij de resultaten van de leenwoordanalyse die in dit hoofdstuk worden besproken. Het is echter niet zo dat alle voorkomens uit het leensuffixonderzoek in de leenwoordanalyse bij de categorie ‘zeker uit het Frans’ zijn geteld. Ten eerste staat van sommige woorden met een Frans suffix niet vast dat ze uit het Frans afkomstig zijn. *Predikant*, *principaal*, *universiteit*, *regeren* en *fundament* bijvoorbeeld komen volgens de etymologische woordenboeken uit het Frans of uit het Latijn. Deze woorden zijn voor de leenwoordanalyse daarom in de categorie ‘mogelijk uit het Frans’ geplaatst, terwijl andere woorden met een Frans suffix als *officier*, *passage*, *majesteit*, *informer* en *logement*, waarvan wel zeker is dat die

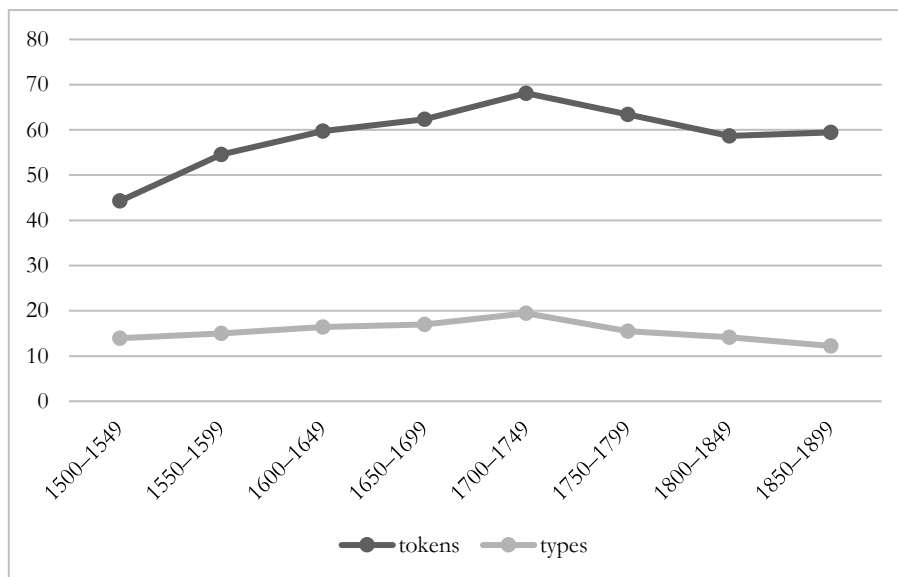
woorden als geheel uit het Frans zijn geleend, in de categorie ‘zeker uit het Frans’ zijn opgenomen. Ten tweede zijn er woorden als *student* en *communiceren* die wel een Frans suffix bevatten, maar die volgens de etymologische woordenboeken in het geheel niet uit het Frans maar uit het Latijn afkomstig zijn. Deze woorden zijn in de leensuffixanalyse wel meegeteld, maar worden voor het leenwoordenonderzoek buiten beschouwing gelaten. Hetzelfde geldt voor productief gevormde woorden met een Frans suffix, waarbij het suffix aan een Germaanse stam wordt gekoppeld, zoals het geval is bij *lekkage*, *koetsier* en *mennoniet*. Aangezien het ook hier niet om Franse leenwoorden gaat, zijn ze dus ook niet in de leenwoordanalyse betrokken. Wanneer echter een woord productief is gevormd in het Nederlands op basis van een Franse stam en een Frans leensuffix en wanneer deze Franse stam op zichzelf ook een leenwoord is dat het Nederlands uit het Frans heeft overgenomen, dan is dit woord wel meegeteld bij de leenwoorden, aangezien samenstellingen en afleidingen met leenwoorden ook zijn meegenomen in de leenwoordanalyse (zie paragraaf 7.3.2 hierboven). Zo is het woord *telegraferen* een nieuwvorming in het Nederlands, maar op basis van het woord *telegraaf*, dat een leenwoord is uit het Frans; het woord *telegraferen* is dus opgenomen in de leenwoordanalyse als afleiding van het woord *telegraaf*.

De categorie ‘zeker uit het Frans’ bevat dus de woorden met een Frans suffix die als geheel Franse leenwoorden zijn plus Franse leenwoorden die geen Frans leensuffix bevatten. In de categorie ‘mogelijk uit het Frans’ zijn alle woorden met een Frans leensuffix opgenomen die het Frans mogelijk (maar niet zeker) als brontaal hebben plus alle leenwoorden zonder Frans suffix die mogelijk uit het Frans, maar mogelijk ook uit een andere taal zijn overgenomen.

Figuur 7.3 op de volgende pagina toont allereerst de resultaten van de categorie ‘zeker uit het Frans’ en ‘mogelijk uit het Frans’ samen – dat wil zeggen: alle leenwoorden die het Nederlands waarschijnlijk uit het Frans heeft geleend of die door het Frans zijn beïnvloed – per periode voor alle domeinen samen per 1.000 woorden. Net als bij de leensuffixanalyse worden zowel de resultaten voor de tokens als voor de types gegeven.

Wanneer alle leenwoorden samen worden genomen, is te zien dat er een toename is van het aantal leenwoorden tot en met de periode 1700–1749 bij zowel de tokens als de types. Het aantal tokens neemt toe van circa 44 in de periode 1500–1549 tot ruim 68 tokens in de periode 1700–1749, en het aantal types stijgt van ongeveer 14 tot ruim 19 tussen

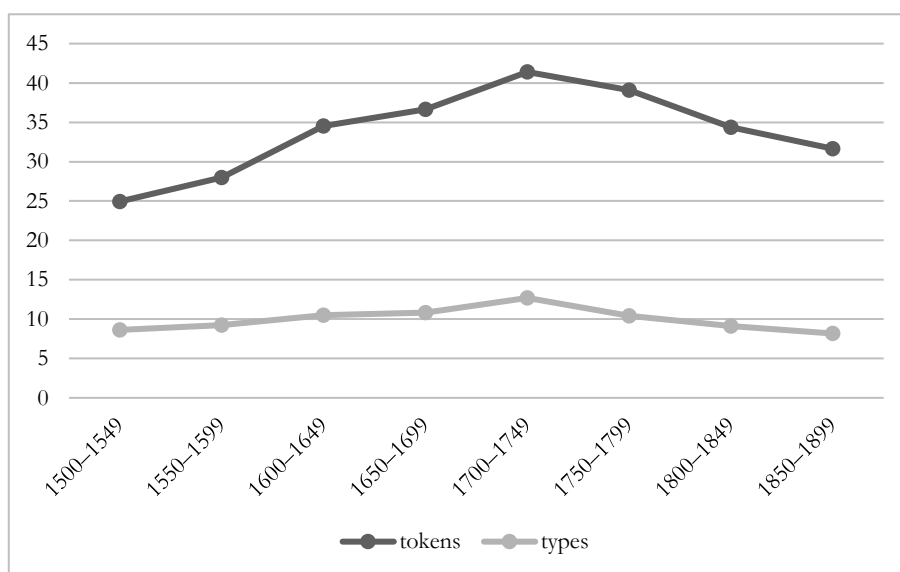
1500–1549 en 1700–1749. Vervolgens daalt de typefrequentie geleidelijk weer tot circa 12 per 1.000 woorden in de tweede helft van de negentiende eeuw, dus nog onder het niveau van 1500–1549. Bij de tokens is een ander patroon te zien: deze laten na een geleidelijke stijging tussen 1500–1549 en 1700–1749 en een daling in de perioden 1750–1799 en 1800–1849 een stabilisatie van het aantal tokens Franse leenwoorden zien in de periode 1850–1899: het gebruik neemt licht toe in deze periode, maar blijft redelijk op hetzelfde niveau als in de eerste helft van de negentiende eeuw. In elk geval zit de tokenfrequentie gedurende de hele negentiende eeuw op een lager niveau ten opzichte van de achttiende en de tweede helft van de zeventiende eeuw; de tokenfrequentie is dan vergelijkbaar met die in de periode 1600–1649 (er zijn ongeveer 59,5 tokens in 1850–1899 tegenover circa 59,8 tokens in 1600–1649).



Figuur 7.3. Aantal leenwoorden dat (zeker of mogelijk) uit het Frans komt (tokens en types) per 1.000 woorden per periode in het LOL-corpus.

Bovenstaande figuur heeft betrekking op alle leenwoorden die (mogelijk) uit het Frans afkomstig zijn. Hierbij zijn echter ook veel woorden meegeteld waarvan het de vraag is of deze woorden uit het Frans komen. Naast de eerder genoemde woorden met een Frans leensuffix als *predikant*, *regeren* en *universiteit* zijn er zeer frequente woorden als *familie*,

meester en *persoon*, die mogelijk uit het Latijn geleend zijn. Om een beeld te krijgen van de token- en typefrequenties van de leenwoorden waarvan vaststaat dat die uit het Frans komen, zijn in figuur 7.4 het aantal tokens en types per periode opgenomen waarbij alleen de leenwoorden zijn opgenomen die zeker uit het Frans komen.



Figuur 7.4. Aantal leenwoorden dat zeker uit het Frans komt (tokens en types) per 1.000 woorden per periode in het *LOL*-corpus.

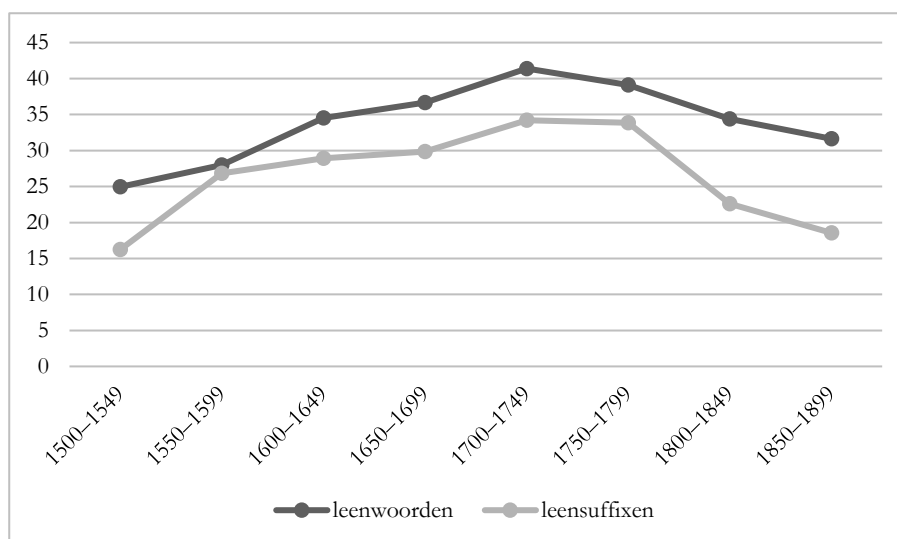
In grote lijnen lijkt het verloop van de token- en typefrequentie in figuur 7.4 op die in figuur 7.3. Bij de tokenfrequentie in figuur 7.4 is er gedurende de zestiende tot en met de eerste helft van de achttiende eeuw een geleidelijke toename van het aantal Franse leenwoorden te zien. Het gebruik van Franse leenwoorden stijgt van bijna 25 leenwoorden per 1.000 woorden in de periode 1500–1549 tot ruim 41 leenwoorden per 1.000 woorden in de periode 1700–1749. In de tweede helft van de achttiende eeuw zet de daling in en komen er ongeveer 39 tokens per 1.000 woorden voor; de daling zet door in de negentiende eeuw, wanneer de tokenfrequentie uitkomt op circa 34 in de periode 1800–1849 en nog geen 32 tokens per 1.000 woorden in de periode 1850–1899. Op dit laatste punt wijkt figuur 7.4 af van figuur 7.3: wanneer alle leenwoorden samen worden genomen, stijgt de tokenfrequentie juist weer in de periode 1850–1899, zoals te zien is in figuur 7.3. In

tegenstelling tot de tokenfrequentie van de periode 1850–1899 in figuur 7.3 zit de tokenfrequentie van dezelfde periode in figuur 7.4 nog onder die van de periode 1600–1649.

De lichte stijging of stabilisatie van het aantal tokens in de periode 1850–1899 ten opzichte van de periode 1800–1849 die zichtbaar was in figuur 7.3, wordt dus volledig veroorzaakt door de leenwoorden die mogelijk uit het Frans komen en niet door de leenwoorden die zeker uit het Frans komen. Opvallend is dat bij het domein Academie een grote toename in het aantal tokens bij de categorie ‘mogelijk uit het Frans’ te zien is, wat vooral wordt veroorzaakt door het veelvuldig voorkomen van (de afkorting van) het woord *numero*, dat uit het Frans of uit het Italiaans komt (130 tokens in 1850–1899 tegenover 34 tokens in hetzelfde domein in de periode 1800–1849). Ook het woord *minister* komt in het domein Academie vaker voor in vergelijking met de voorgaande perioden; bij zowel Academie als Religie valt daarnaast een opvallende toename bij het woord *notulen* te zien in de periode 1850–1899. Dit soort frequentietoenames bij een aantal veel gebruikte mogelijk Franse leenwoorden zou de opleving van het aantal tokens in de periode 1850–1899 in figuur 7.3 kunnen verklaren.

De typefrequentie in figuur 7.4 blijft net als in figuur 7.3 gedurende de gehele periode van vier eeuwen relatief stabiel, met rond de 10 types per 1.000 woorden, maar ook hier is net als bij de tokenfrequentie een piek te zien in de periode 1700–1749. Daarnaast is net als in figuur 7.3 ook in figuur 7.4 de typefrequentie in de tweede helft van de negentiende eeuw het laagst van alle perioden (8,2 per 1.000 woorden, tegenover 8,6 per 1.000 woorden in de periode 1500–1549 en 9,1 per 1.000 woorden in de periode 1800–1849: de twee perioden die na 1850–1899 de laagste typefrequenties hebben). De tokenfrequentie daarentegen blijft in de tweede helft van de negentiende eeuw ruim boven het niveau van de zestiende eeuw.

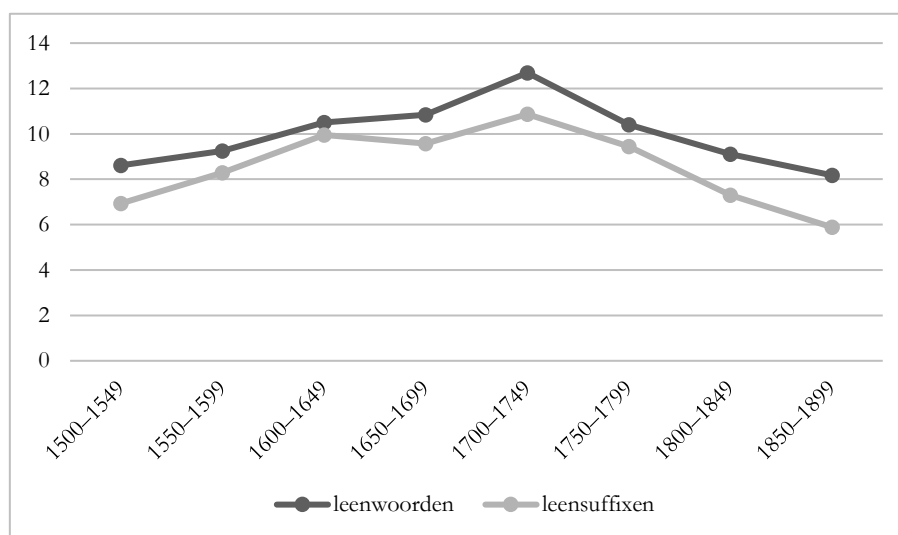
Het algemene beeld dat hierboven is geschetst van het gebruik van Franse leenwoorden in het Nederlands gedurende de zestiende tot en met de negentiende eeuw komt in hoge mate overeen met de resultaten van het onderzoek naar de leensuffixen die in hoofdstuk 5 zijn besproken. Om dit te illustreren, is in figuur 7.5 op de volgende pagina nogmaals de tokenfrequentie weergegeven van de leenwoorden die zeker uit het Frans komen (de donkergrijze lijn in figuur 7.4 hierboven en de donkergrijze lijn in figuur 7.5) en daarnaast is de tokenfrequentie van de leensuffixen opgenomen (de donkergrijze lijn in hoofdstuk 5, figuur 5.3, die correspondeert met de lichtgrijze lijn in figuur 7.5).



Figuur 7.5. Aantal leenwoorden dat zeker uit het Frans komt (tokens) en aantal Franse leensuffixen (tokens) per 1.000 woorden per periode in het *LOL*-corpus.

Figuur 7.5 laat zien dat het verloop van de tokenfrequentie redelijk gelijk opgaat bij de leensuffixen en de leenwoorden. De tokenfrequentie is bij de leenwoorden telkens hoger dan bij de leensuffixen, met het kleinste verschil in de periode 1550–1599 met een verschil van ongeveer 1,2 tokens per 1.000 woorden, en het grootste verschil in de periode 1850–1899 met een verschil van circa 13 tokens per 1.000 woorden. In het algemeen is de stijging van het aantal leensuffixen in de zestiende eeuw en de daling van het aantal leensuffixen in de negentiende eeuw scherper dan die van de leenwoorden in dezelfde perioden. Het verschil tussen leenwoorden en leensuffixen in figuur 7.5 is kleiner dan bij Stevens (2019: 151), die in haar corpus gemiddeld een verschil van ongeveer 14 tokens per 1.000 woorden vond tussen de leenwoorden en leensuffixen (zie figuur 7.2 in paragraaf 7.2 hierboven). Een verklaring voor het verschil zou kunnen zijn dat Stevens (2019) minder leensuffixen heeft meegenomen in haar onderzoek dan ik, waardoor de resultaten voor de leensuffixen bij haar automatisch lager uitvallen.

In figuur 7.6 is de typefrequentie te zien van de leenwoorden en leensuffixen in het *LOL*-corpus per 1.000 woorden, waar het overeenkomende patroon van het gebruik van de leenwoorden en leensuffixen over de tijd nogmaals wordt geïllustreerd.

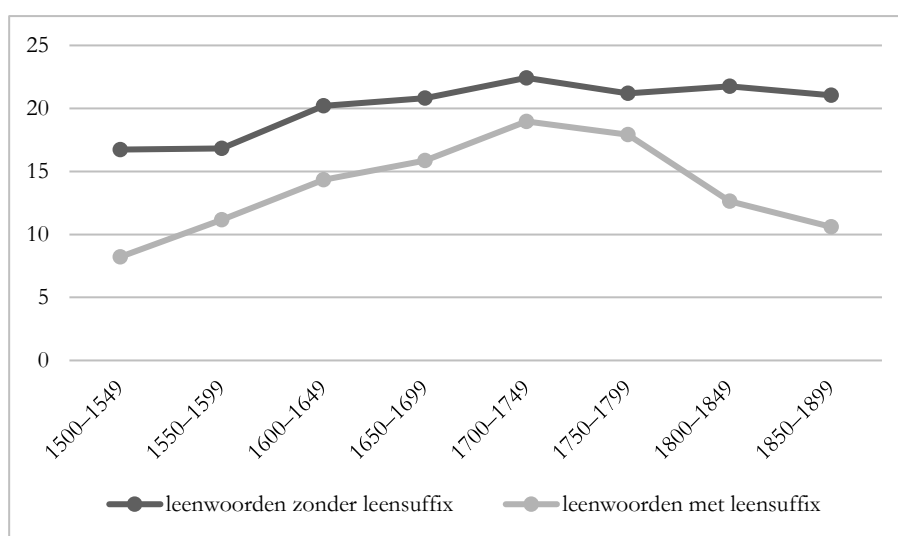


Figuur 7.6. Aantal leenwoorden dat zeker uit het Frans komt (types) en aantal Franse leensuffixen (types) per 1.000 woorden per periode in het *LOL*-corpus.

Op basis van het algemene beeld van de leenwoorden en leensuffixen is het dus duidelijk dat het gebruik van deze Franse leenwoorden en leensuffixen grote gelijkenissen vertoont: over het algemeen is het zo dat er meer leensuffixen worden gebruikt als er ook meer Franse leenwoorden worden gebruikt, en omgekeerd. Het is uiteraard wel zo dat de resultaten van de leenwoorden deels zijn opgebouwd uit Franse leenwoorden die ook een Frans leensuffix bevatten en die in bovenstaande twee grafieken dus zijn meegeteld bij zowel de resultaten voor de leensuffixen als bij de resultaten voor de leenwoorden. Om de verhouding te bepalen tussen het aantal Franse leenwoorden met en zonder leensuffix, zijn deze twee categorieën gesplitst en weergegeven in figuur 7.7 op de volgende pagina.

In figuur 7.7 is te zien dat, ook wanneer de leenwoorden waarvan vaststaat dat ze uit het Frans komen worden opgedeeld in leenwoorden met een Frans leensuffix en leenwoorden zonder Frans leensuffix, beide tokenfrequenties een stijging laten zien van het aantal leenwoorden tot de eerste helft van de achttiende eeuw, en daarna een daling. Een belangrijk verschil tussen de twee lijnen in figuur 7.7 is echter dat de frequenties van de Franse leenwoorden zonder leensuffix veel stabielier blijven gedurende de eeuwen dan de frequenties van de Franse leenwoorden met een leensuffix: de stijgingen en dalingen tussen de verschillende perioden zijn over het algemeen veel scherper bij de

leenwoorden met een leensuffix. Het aantal leenwoorden zonder leensuffix blijft ook min of meer stabiel in de zestiende eeuw en in de laatste drie perioden van het corpus, terwijl de leenwoorden met leensuffix wel een duidelijke stijging laten zien tussen de twee perioden van de zestiende eeuw en een scherpe daling tussen de perioden 1750–1799 en 1800–1849 en tussen de perioden 1800–1849 en 1850–1899.



Figuur 7.7. Aantal leenwoorden zonder leensuffix dat zeker uit het Frans komt (tokens) en aantal leenwoorden met leensuffix dat zeker uit het Frans komt (tokens) per 1.000 woorden per periode in het *LOL*-corpus.

Een verklaring voor de stabiele lijn bij de leenwoorden zonder suffix in de negentiende eeuw is lastig te geven. Wel is het zo dat het domein Literatuur een stijging van het aantal leenwoorden laat zien in de periode 1800–1849 ten opzichte van 1750–1799 en ook in 1850–1899 ten opzichte van 1800–1849, waardoor het aantal leenwoorden in 1850–1899 meer dan het dubbele is van dat in 1750–1799. Er lijken echter geen specifieke, hoogfrequente woorden in dit domein te zijn die de stijging kunnen verklaren, behalve dat in de toneelteksten uit de periode 1850–1899 relatief veel legertermen worden gebezigd die uit het Frans afkomstig zijn (zie verdere discussie in 7.4.3). Ook laat het domein Academie een stijging zien van het aantal leenwoorden zonder leensuffix in de periode 1850–1899 ten opzichte van 1800–1849 (van 25,8 naar 32,8 per 1.000 woorden) en bij het domein Religie stijgt de tokenfrequentie

van 12,5 per 1.000 woorden in 1750–1799 tot 19,4 per 1.000 woorden in 1800–1849, en stijgt vervolgens verder in 1850–1899 tot 20,9 per 1.000 woorden. Dit terwijl het domein Academie een daling in het aantal leenwoorden met een leensuffix laat zien in 1850–1899 ten opzichte van 1800–1849, van 21,7 naar 16,7, en het aantal leenwoorden met een leensuffix in dezelfde periode ook daalt bij het domein Religie, van 12,9 naar 7,4. Bij de andere domeinen is dit niet het geval. Bij Religie valt een toename van de woorden *missive* en *plaats* op in de periode 1800–1849 ten opzichte van 1750–1799 (van 1 naar 10 tokens en van 4 naar 16 tokens, respectievelijk). Frequentie woorden in de periode 1850–1899 in zowel Religie als in Academie zijn *missive*, *adres* en *florijn* (42, 19 respectievelijk 18 van de in totaal 166 tokens bij Academie en 20, 11 respectievelijk 17 van de in totaal 109 tokens bij Religie).

Uit figuur 7.7 blijkt ook dat het aantal leenwoorden zonder leensuffix het aantal leenwoorden met leensuffix in alle perioden overstijgt. Dit was ook het geval in figuur 7.5, waar alle leenwoorden die uit het Frans komen (met en zonder leensuffix) werden vergeleken met de resultaten van de leensuffixen uit het vorige hoofdstuk. De vergelijkingen van leenwoorden en leensuffixen tonen dus aan dat leenwoordenonderzoek en leensuffixonderzoek enigszins andere resultaten opleveren. De patronen van de leenwoorden en de leensuffixen zijn weliswaar over het algemeen hetzelfde, maar de hoge frequentie van het aantal leenwoorden zonder leensuffix geeft aan dat er een hoop woorden gemist worden als er alleen naar leensuffixen wordt gekeken. Bovendien laat het feit dat de ontwikkeling van het gebruik van leenwoorden en leensuffixen niet altijd parallel loopt, zoals het geval is in de negentiende eeuw in figuur 7.7, zien dat leenwoorden met en zonder leensuffix zich niet altijd hetzelfde gedragen. Op basis van deze onderzoeksresultaten kan ik dus met Stevens (2019: 151) concluderen dat leenwoordenonderzoek een waardevolle aanvulling is op leensuffixonderzoek. Ook in de komende paragrafen zal blijken dat het gebruik van Franse leenwoorden en Franse leensuffixen lang niet altijd met elkaar overeenkomt.

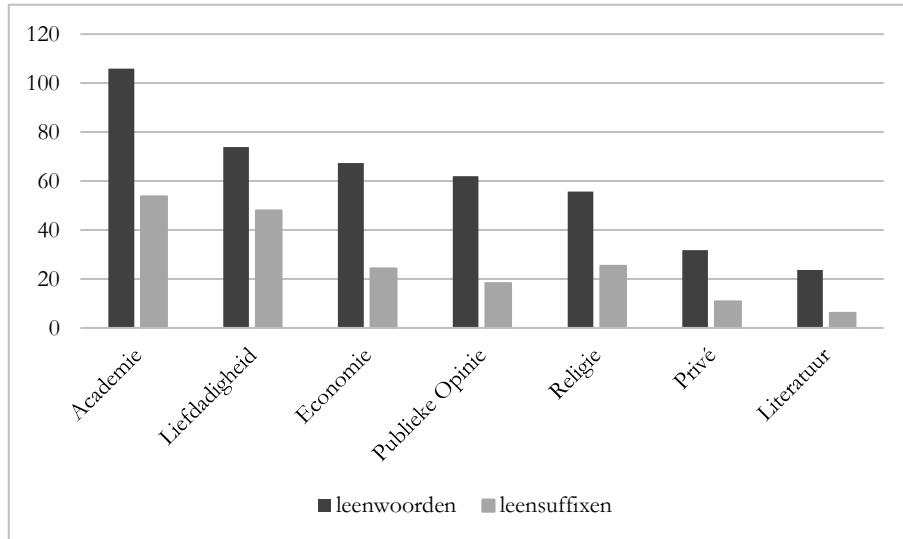
Aan de hand van de resultaten van het leenwoorden- en leensuffixonderzoek zoals dat in deze paragraaf is besproken, kan het volgende beeld worden geschetst van de invloed die het Frans op de Nederlandse taal heeft gehad: in de zestiende en zeventiende eeuw is er sprake van toenemende invloed van het Frans, die zijn hoogtepunt bereikt in de eerste helft van de achttiende eeuw, waarna de Franse invloed in de tweede helft van de achttiende eeuw en in de negentiende

eeuw weer afneemt, getuige de (over het algemeen) lagere frequentie van zowel leenwoorden als leensuffixen. Dit beeld komt overeen met de literatuur, die het hoogtepunt van de verfransing in de achttiende eeuw situeert (zie Frijhoff 1989: 594–598). De resultaten lijken echter in te gaan tegen het beeld dat Van der Sijs (2001: 214) schetst op basis van het aantal ontleningen aan verschillende talen per eeuw (zie figuur 7.1 in paragraaf 7.2 hierboven). Volgens Van der Sijs (2001) werden de meeste Franse leenwoorden in de negentiende eeuw ontleend, na een daling van het aantal ontleningen in de zeventiende en achttiende eeuw. Als naar het gebruik van Franse leenwoorden wordt gekeken, gebeurt echter precies het tegenovergestelde: daar neemt het aantal leenwoorden juist toe in de zeventiende en het begin van de achttiende eeuw en daalt daarna weer. Niet de negentiende eeuw, maar de zeventiende en begin achttiende eeuw lijken dus cruciale perioden in de invloed die het Frans op het Nederlands heeft gehad. Hier moet wel bij gezegd worden dat de resultaten in deze paragraaf zijn gebaseerd op de leenwoorden van alle domeinen samen. Zoals in paragraaf 7.4.3 zal blijken, waar de resultaten worden opgesplitst naar periode en domein, geldt het algemene beeld uit deze paragraaf voor lang niet alle domeinen. Voor ik overga op een bespreking van de resultaten per periode en domein, worden echter in de volgende paragraaf eerst de resultaten gegeven per domein.

7.4.2 Leenwoorden per sociaal domein

In figuur 7.8 op de volgende pagina is het aantal tokens leenwoorden per 1.000 woorden uitgesplitst naar domein. De domeinen zijn gerangschikt op frequentie, waarbij Academie het hoogste aantal Franse leenwoorden heeft per 1.000 woorden en Literatuur het laagste. Bij de resultaten van de leenwoorden zijn zowel de leenwoorden meegeteld die zeker uit het Frans komen als de leenwoorden die mogelijk uit het Frans komen. Om een goede vergelijking met het leensuffixonderzoek mogelijk te maken, is voor elk domein ook de tokenfrequentie van het aantal Franse suffixen per 1.000 woorden (figuur 5.5 uit hoofdstuk 5) in figuur 7.8 opgenomen.

Alle domeinen hebben een hogere tokenfrequentie bij de leenwoorden dan bij de leensuffixen en de frequentievolgorde van de domeinen is nagenoeg dezelfde als bij de leensuffixen. Het domein Religie vormt hierop een uitzondering: dit domein heeft een lagere tokenfrequentie dan Economie en Publieke Opinie, terwijl het bij de leensuffixen meer tokens per 1.000 woorden heeft dan die domeinen.



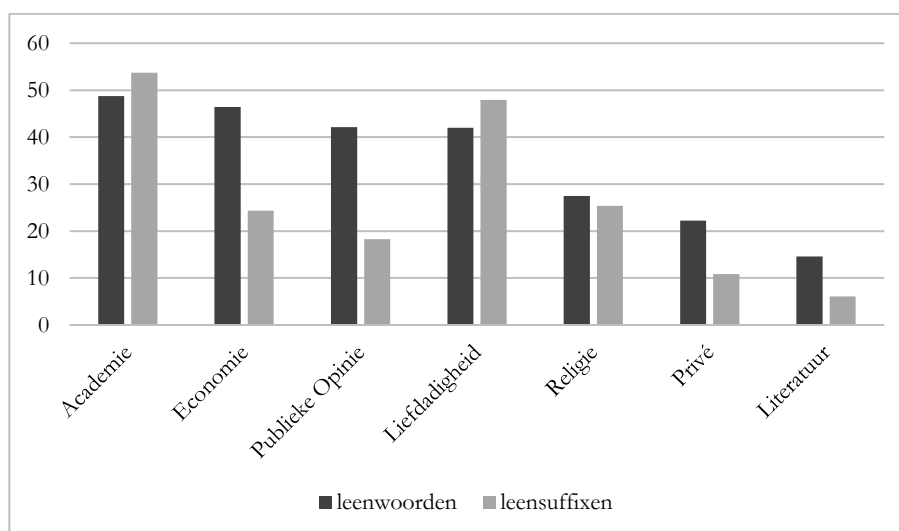
Figuur 7.8. Aantal leenwoorden dat (zeker of mogelijk) uit het Frans komt (tokens) en aantal leensuffixen (tokens) per 1.000 woorden per domein in het *LOL*-corpus.

Wellicht wordt deze afwijking veroorzaakt doordat bij Economie en Publieke Opinie enkele frequente woorden voorkomen die wel uit het Frans komen, maar geen Frans leensuffix bevatten (zoals *grein*, *ras* en *baai* bij Economie en *markies* en *kapitein* bij Publieke Opinie; zie voor verdere discussie van deze domeinen verderop in deze paragraaf). Bij Religie echter waren enkele zeer frequente woorden die specifiek zijn voor dit domein leenwoorden met een Frans leensuffix, zoals *predikant* en *predicatie*, terwijl er geen opvallend frequente domein-specifieke woorden zijn bij Religie die geen leensuffix bevatten. Dit kan dus een reden zijn dat Religie een relatief hoge frequentie heeft bij de leensuffixen en een relatief lage frequentie bij de leenwoorden, terwijl voor Economie en Publieke Opinie precies het omgekeerde het geval is.

Opvallend is ook dat het domein Liefdadigheid vergeleken met de andere domeinen een relatief klein verschil heeft tussen het aantal leenwoorden en het aantal leensuffixen per 1.000 woorden: waar bij de andere domeinen het aantal leenwoorden ten minste het dubbele is van het aantal leensuffixen per 1.000 woorden (bij Academie is dat iets minder: 106 leenwoorden tegenover 54 leensuffixen), is het aantal leenwoorden bij Liefdadigheid maar iets meer dan de helft groter dan het aantal leensuffixen (74 leenwoorden tegenover 48 suffixen). Een verklaring hiervoor is dat een groot deel van de (mogelijk) Franse

leenwoorden die in het domein Liefdadigheid voorkomen ook een leensuffix heeft. Woorden als *testament*, *comparant(e)* en werkwoorden als *comparer* en *resideren* zijn zeer frequent in de testamenten en hebben ook allemaal een Frans leensuffix. Dit is anders bij bijvoorbeeld het domein Economie, waar frequent voorkomende leenwoorden juist geen leensuffix hebben. Hier kom ik verderop op terug.

In figuur 7.9 zijn de resultaten weergegeven voor de leenwoorden die zeker uit het Frans komen, samen met de leensuffixen die ook in figuur 7.8 zijn weergegeven in de lichtgrijze kolommen. Ook hier zijn de domeinen in volgorde geplaatst op basis van de leenwoordfrequentie.



Figuur 7.9. Aantal leenwoorden dat zeker uit het Frans komt (tokens) en aantal leensuffixen (tokens) per 1.000 woorden per domein in het *LOL*-corpus.

Het beeld van figuur 7.9 is iets anders dan dat van figuur 7.8: waar het domein Liefdadigheid in de vorige figuur nog de tweede plek innam qua leenwoordfrequentie, achter Academie, is nu Economie het domein met de hoogste tokenfrequentie van de leenwoorden na Academie. Daarnaast is bij de domeinen Economie, Publieke Opinie, Privé en Literatuur een duidelijk verschil te zien tussen de leenwoordfrequentie en de leensuffixfrequentie, waarbij de leenwoordfrequentie (veel) hoger is dan de leensuffixfrequentie, terwijl dit niet het geval is bij Academie, Liefdadigheid en Religie. Deze laatste drie domeinen hebben ofwel een vrijwel gelijke tokenfrequentie bij de Franse leenwoorden en de

leensuffixen, met iets meer tokens voor de leenwoorden dan voor de suffixen (Religie), of de tokenfrequentie van de suffixen is hoger dan van de Franse leenwoorden (Academie en Liefdadigheid). Bij de leensuffix-analyse hadden deze drie domeinen de hoogste tokenfrequenties, met het domein Academie op de eerste plaats, gevolgd door Liefdadigheid en daarna Religie, maar wanneer naar Franse leenwoorden wordt gekeken, is het beeld anders. Academie heeft nog steeds de hoogste tokenfrequentie van alle domeinen, maar wordt op de voet gevolgd door het domein Economie: met een frequentie van 48,7 tegenover 46,4 tokens per 1.000 woorden is er weinig verschil tussen deze domeinen. Nog duidelijker geldt dit voor Publieke Opinie en Liefdadigheid, die respectievelijk een frequentie van 42,1 en 42,0 hebben. Economie en Publieke Opinie bevatten dus relatief veel leenwoorden die zeker uit het Frans komen in vergelijking met Academie en Liefdadigheid.

De relatief hoge tokenfrequenties bij Economie en Publieke Opinie enerzijds en de vergeleken met de leensuffixen relatief lage tokenfrequenties bij Academie, Liefdadigheid en Religie anderzijds kan worden verklaard door enkele zeer frequente, domein-specifieke leenwoorden die in elk van deze vier domeinen voorkomen. Bij Economie allereerst worden vaak woorden gebruikt die met de textielnijverheid te maken hebben: namen voor specifieke stoffen als *ras*, *baai*, *saai*, *grein* en *fustein* en andere veel gebruikte (wevers)termen als *kalander* (een soort mangel waarmee stoffen glad en glanzend worden gemaakt). Deze woorden komen uit het Frans en zijn zeer frequent in de ordonnanties en rekwesten die het domein Economie in het *LOL*-corpus representeren. Dit soort woorden vormt dan ook een belangrijke reden dat dit domein een hoge tokenfrequentie heeft. Bij Publieke Opinie zijn woorden die met het militaire leven en de scheepvaart te maken hebben frequent, bijvoorbeeld *kolonel*, *soldaat*, *troep* en *kapitein*, evenals adellijke titels als *baron* en *markies* en het bijvoeglijke naamwoord *sint* in frequent voorkomende eilandennamen als *Sint Maarten*. Ook hier gaat het om Franse leenwoorden die geen Frans leensuffix bevatten en dus niet zijn meegeteld bij het leensuffixonderzoek, maar wel bij het leenwoordenonderzoek horen.

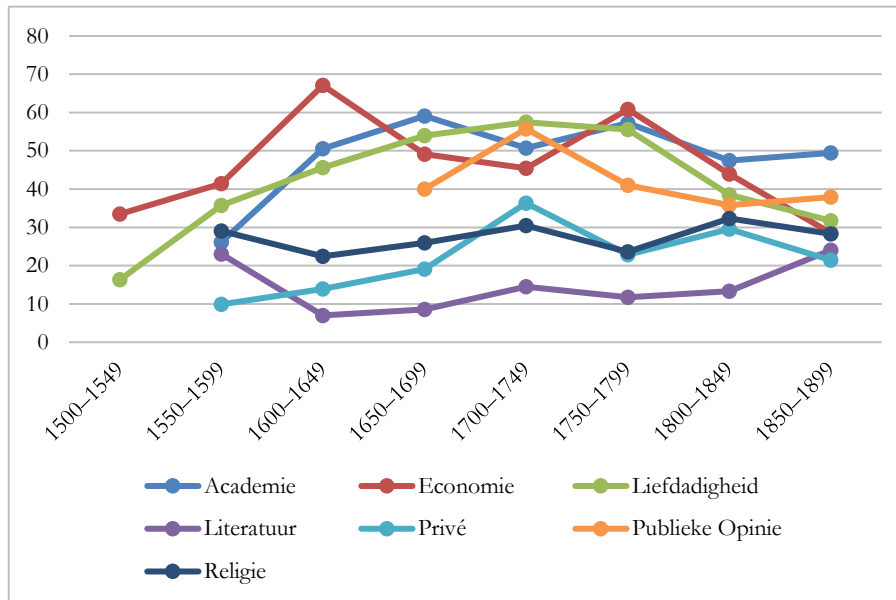
Academie, Liefdadigheid en Religie daarentegen haalden hoge frequenties bij de leensuffixanalyse door veelvoorkomende woorden als *universiteit* en *student* (bij Academie), *testament* en *admitteren* (bij Liefdadigheid) en *predikant* en *resolveren* (bij Religie). Van *universiteit*, *testament*, *predikant* en *resolveren* is echter niet zeker dat het Franse leenwoorden zijn – ze kunnen ook aan het Latijn ontleend zijn – en van

student en *admitteren* staat vast dat het leenwoorden uit het Latijn zijn, wat de reden is dat deze woorden niet in de leenwoordanalyse zijn meegenomen. Ook van veel andere woorden uit de domeinen Academie, Liefdadigheid en Religie, die wel leenwoorden zijn maar geen Frans leensuffix hebben, is niet zeker dat ze uit het Frans komen, zoals het zeer frequente woord *college* bij de domeinen Academie en Religie en *interest* en *rente* bij Liefdadigheid. Dit zorgt er al met al voor dat de frequenties van deze drie domeinen lager uitvallen.

Wanneer de resultaten worden uitgesplitst naar domein, zijn er dus verschillen te zien tussen de resultaten waarbij alleen wordt gekeken naar de leenwoorden die zeker uit het Frans komen en de resultaten waarbij ook de leenwoorden die mogelijk uit het Frans zijn geleend worden meegenomen. In elk geval maken de resultaten uit deze paragraaf duidelijk, net als die van het leensuffixonderzoek in hoofdstuk 5, dat het van het domein afhangt hoe veel Franse leenwoorden er worden gebruikt. De resultaten hebben ook laten zien dat alle domeinen een hogere tokenfrequentie hebben bij de leenwoorden dan bij de leensuffixen (althans als wordt gekeken naar de leenwoorden die zeker en mogelijk uit het Frans komen) en dat de frequentievolgorde van de domeinen bij de leenwoorden vrijwel dezelfde is als bij de leensuffixen. De mate van invloed van het Frans op het Nederlandse lexicon verschilt dus per sociaal domein en is bijvoorbeeld groter geweest in domeinen als Academie en Economie en minder groot in Privé en Literatuur. Net als bij het leensuffixonderzoek lijken ook bij het leenwoordenonderzoek domein-specifieke woorden een grote rol te spelen. Dit roept de vraag op wat de verhouding bij de leenwoorden is tussen domein-specifieke termen en algemene woordenschat. Hier wordt op ingegaan in hoofdstuk 9. In de volgende paragraaf worden de variabelen periode en domein gecombineerd en worden de leenwoorden uitgesplitst naar zowel periode als domein. Daarbij richt ik mij alleen op de leenwoorden die zeker uit het Frans komen.

7.4.3 Leenwoorden per periode en sociaal domein

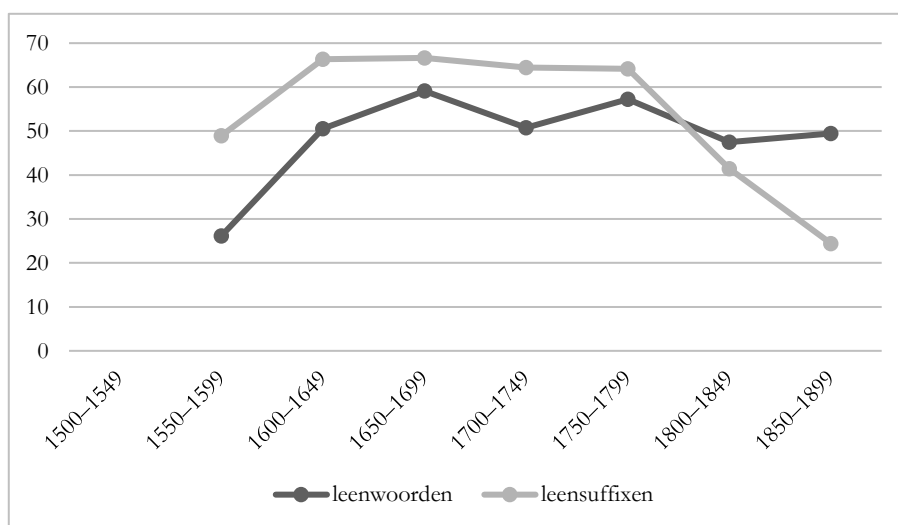
In figuur 7.10 zijn de resultaten van de vorige twee paragrafen gecombineerd: per periode en per domein wordt de tokenfrequentie per 1.000 woorden weergegeven van de leenwoorden waarvan zeker is dat die uit het Frans komen.



Figuur 7.10. Aantal leenwoorden dat zeker uit het Frans komt (tokens) per 1.000 woorden per periode per domein in het *LOL*-corpus.

Figuur 7.10 maakt duidelijk dat de afzonderlijke domeinen zich niet allemaal gedragen volgens het algemene beeld uit figuur 7.4 in paragraaf 7.4.1 hierboven. Wanneer alle domeinen bij elkaar worden genomen, zit de piek van het gebruik van Franse leenwoorden in de periode 1700–1749, maar dit is niet altijd het geval voor de individuele domeinen. Liefdadigheid, Publieke Opinie en Privé halen hun hoogste tokenfrequenties wel in 1700–1749, maar de andere domeinen niet. Het beeld is ook heel anders dan bij de leensuffixen uit het vorige hoofdstuk, waar de tokenfrequenties van Academie en Liefdadigheid in vrijwel alle perioden ver boven de rest uitstaken en waar die twee domeinen ook de enige waren waar een duidelijke op- en neerwaartse trend van het aantal leensuffixen door de eeuwen heen te zien was, terwijl de rest van de domeinen op een stabiel laag aantal leensuffixen bleef. Als naar de resultaten van de leenwoorden in figuur 7.10 wordt gekeken, dan lopen de lijnen van de verschillende domeinen veel meer door elkaar, waarbij een grove indeling gemaakt kan worden tussen Literatuur, Privé en Religie aan de ene kant als domeinen met een lage tokenfrequentie gedurende de vier eeuwen die het corpus beslaat, en Academie, Economie, Liefdadigheid en Publieke Opinie aan de andere kant met hogere tokenfrequenties.

In de rest van deze paragraaf worden de resultaten van de zeven domeinen elk afzonderlijk besproken in alfabetische volgorde aan de hand van figuren waarin zowel de grafiek is opgenomen van het betreffende domein uit figuur 7.10 als de grafiek van het domein uit het leensuffixonderzoek in hoofdstuk 5. Allereerst zijn in figuur 7.11 de resultaten voor het domein Academie weergegeven.



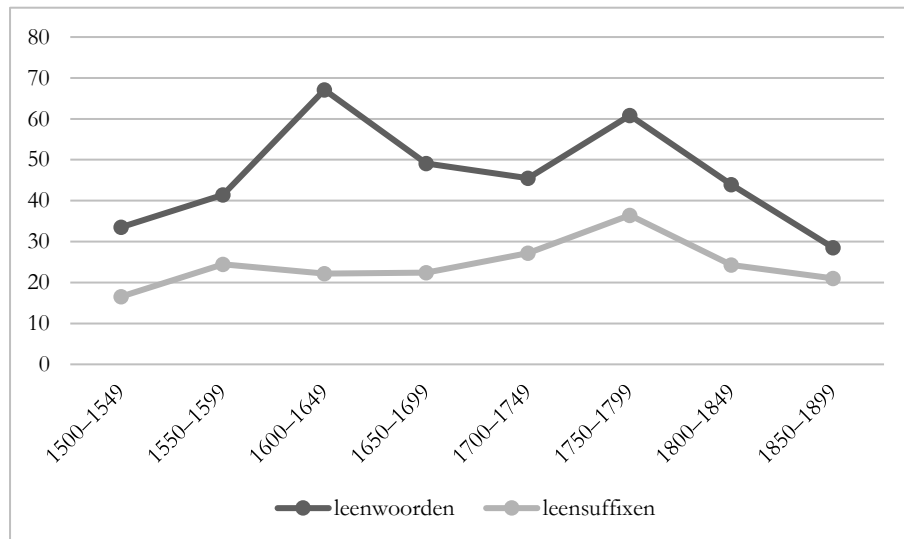
Figuur 7.11. Aantal leenwoorden dat zeker uit het Frans komt (tokens) en aantal leensuffixen (tokens) per 1.000 woorden per periode in het *LOL*-corpus voor het domein Academie.

Bij het domein Academie is het allereerst opvallend dat het aantal leensuffixen bijna altijd groter is dan het aantal leenwoorden in dezelfde periode. Dat het aantal leensuffixen bij Academie in totaal ook hoger is dan het aantal leenwoorden, was ook al geconcludeerd bij de bespreking van figuur 7.9. Alleen de twee laatste perioden wijken af van dit patroon: in zowel 1800–1849 als in 1850–1899 zijn de frequenties omgedraaid en zijn er meer tokens van Franse leenwoorden dan van Franse leensuffixen. In vergelijking met de eerste helft van de negentiende eeuw ontwikkelt de tokenfrequentie van de leenwoorden zich ook anders dan die van de leensuffixen: de eerste stijgt licht ten opzichte van 1800–1849, terwijl de andere juist daalt. De daling van het aantal leensuffixen in de periode 1850–1899 is met name toe te schrijven aan werkwoorden op *-eren* en zelfstandige naamwoorden op *-atie*, die tot en met de periode 1800–1849 veelvuldig aanwezig zijn in het corpus maar in de periode

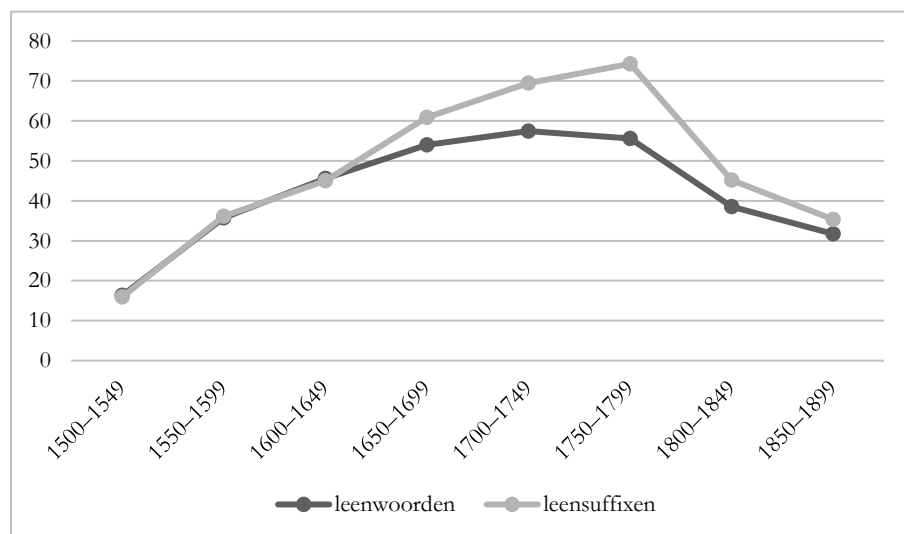
1850–1899 nauwelijks meer voorkomen in dit domein. De stijging van de leenwoorden in dezelfde periode heeft twee oorzaken: enerzijds komen er meer leenwoorden zonder leensuffix voor in de periode 1850–1899 dan in de voorgaande periode (onder andere door de woorden *adres* en *florijn*, zoals al is besproken in paragraaf 7.4.1) en anderzijds zijn de meeste leenwoorden met leensuffix in 1850–1899 ook met zekerheid Franse leenwoorden, die dan ook bijna allemaal zijn meegeteld bij bovenstaande resultaten (18 van de 104 tokens zijn in de categorie ‘mogelijk uit het Frans’ geplaatst). Bij de periode 1800–1849 is dit anders: daar is meer dan een derde van alle leenwoorden met leensuffix een mogelijke en geen zekere ontlening uit het Frans (61 van de 173), waardoor deze tokens niet meetellen in bovenstaande grafiek. Ook hier gaat het om diverse woorden op *-eren* en *-atie* en om het woord *universiteit*, dat nog maar zelden voorkomt in de periode 1850–1899.

Afgezien van de laatste perioden lijkt het verloop van beide grafieken wel op elkaar, met een stijging tussen de tweede helft van de zestiende eeuw en de eerste helft van de zeventiende eeuw, een min of meer stabiele zeventiende en achttiende eeuw en een daling in de eerste helft van de negentiende eeuw.

In figuur 7.12 op de volgende pagina zijn de resultaten opgenomen voor het domein Economie. Zoals blijkt uit deze figuur heeft het domein Economie, anders dan bij Academie het geval was, gedurende alle acht de onderzochte perioden meer leenwoorden dan leensuffixen in de teksten zitten. Daarbij volgt de tokenfrequentie van de leenwoorden niet altijd het patroon van de leensuffixen. Het opvallendste verschil is de piek bij de leenwoorden in de periode 1600–1649, waar de leensuffixen juist stabiel blijven. De belangrijkste oorzaak hiervan is al kort genoemd in paragraaf 7.4.1: in het domein Economie komen veel weverstermen voor als *baai*, *grein* en *fustein*, die geen leensuffix bevatten maar waarvan wel met zekerheid vastgesteld is dat ze uit het Frans afkomstig zijn. Met name in de periode 1600–1649 komen deze woorden zeer vaak voor. De andere piek – in de periode 1750–1799 – komt wel overeen met de piek bij de leensuffixen. Wanneer naar de typefrequentie in plaats van de tokenfrequentie bij het domein Economie wordt gekeken, dan blijkt ook dat de pieken in 1600–1649 en 1750–1799 slechts door een beperkt aantal leenwoordtypen worden veroorzaakt die heel frequent zijn in deze perioden. De piek in het aantal leenwoordtypen is te vinden in 1700–1749 met ongeveer 23 leenwoordtypen per 1.000 woorden, terwijl deze frequentie zowel in 1600–1649 als in 1750–1799 rond de 20,5 per 1.000 woorden ligt.



Figuur 7.12. Aantal leenwoorden dat zeker uit het Frans komt (tokens) en aantal leensuffixen (tokens) per 1.000 woorden per periode in het *LOL*-corpus voor het domein *Economie*.

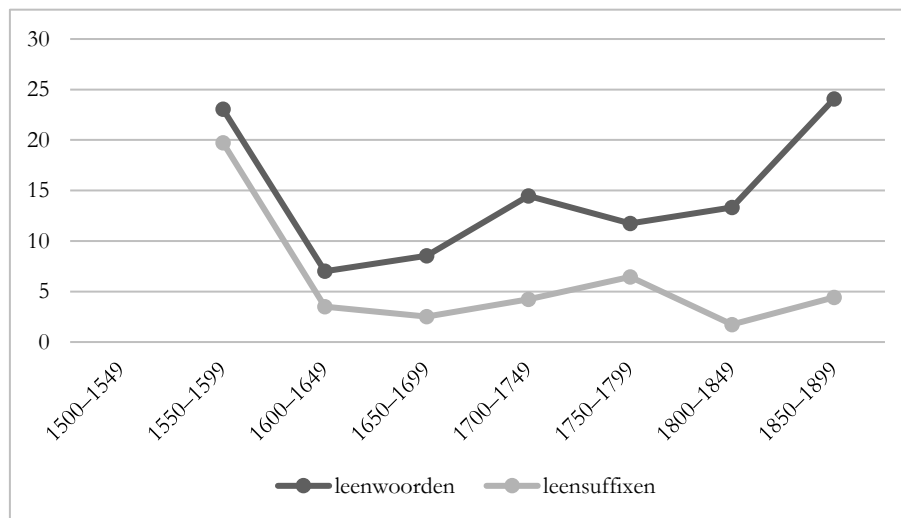


Figuur 7.13. Aantal leenwoorden dat zeker uit het Frans komt (tokens) en aantal leensuffixen (tokens) per 1.000 woorden per periode in het *LOL*-corpus voor het domein *Liefdadigheid*.

Figuur 7.13 geeft de leenwoord- en leensuffixfrequentie voor het domein *Liefdadigheid*. In de eerste drie perioden van het corpus zijn beide

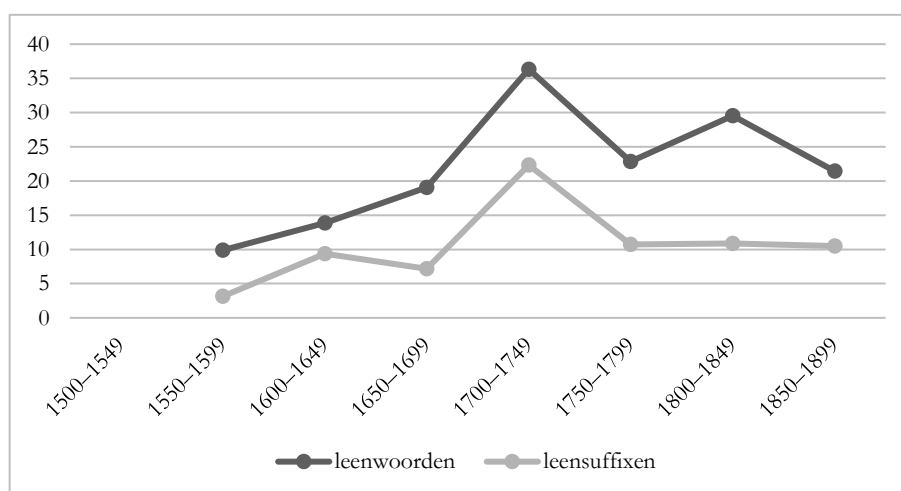
tokenfrequenties opvallend gelijk aan elkaar: er worden in deze perioden dus ongeveer evenveel leenwoorden als leensuffixen gebruikt per 1.000 woorden. Vanaf de tweede helft van de zeventiende eeuw verandert dit en hebben de leensuffixen een hogere frequentie dan de leenwoorden, maar beide grafieken blijven min of meer hetzelfde patroon volgen. De hoogste tokenfrequentie ligt voor de leenwoorden echter iets eerder dan voor de leensuffixen: in de periode 1700–1749, en bij de leensuffixen in de periode 1750–1799. Beide frequenties nemen in de negentiende eeuw af.

Een verklaring voor de hogere frequentie van leensuffixen dan leenwoorden in dit domein is al kort besproken bij figuur 7.9: er komen veel frequente woorden voor in het domein Liefdadigheid die een Frans leensuffix bevatten maar waarvan de Franse oorsprong niet zeker is, zoals *testament* en *compareren*, en daarnaast zijn er ook veel termen die wel een Frans leensuffix hebben, maar waarvan zeker is dat ze niet uit het Frans komen, zoals *admitteren*, dat uit het Latijn komt, en *legateren*, dat een Nederlandse nieuwvorming is van het Latijnse leenwoord *legaat* en het Franse leensuffix *-eren*. Deze woorden zijn uiteraard niet meegeteld bij de leenwoorden in figuur 7.13, maar wel bij de leensuffixen, waardoor de resultaten voor deze laatste variabele relatief hoog uitvallen.



Figuur 7.14. Aantal leenwoorden dat zeker uit het Frans komt (tokens) en aantal leensuffixen (tokens) per 1.000 woorden per periode in het *LOL*-corpus voor het domein Literatuur.

Het domein Literatuur in figuur 7.14 op de vorige pagina heeft twee perioden waarin het leenwoordgebruik duidelijk hoger ligt dan in de andere perioden: 1550–1599 en 1850–1899. Bij de leensuffixen is in de periode 1550–1599 eenzelfde piek te zien in de tokenfrequentie, maar niet in de periode 1850–1899. Het hoge leenwoord- en leensuffixgebruik in de periode 1550–1599 kan verklaard worden door het feit dat de teksten uit deze periode geschreven zijn door rederijders. De rederijders stonden bekend om hun veelvuldige gebruik van Franse leenwoorden (zie bijv. Keßler 2012: 56). Bij de periode 1850–1899 lijkt het hoge leenwoordgebruik vooral aan het toneelstuk *Magdalena Moons* te liggen. In dit toneelspel – waarin de gebeurtenissen rondom het Leidens Ontzet centraal staan – komen enkele woorden voor die te maken hebben met oorlogsvoering als *plan* (3 keer), *kapitein* (10 keer) en *soldaat* (15 keer). Aangezien de algehele frequentie van leenwoorden in het domein Literatuur erg laag is in vergelijking met de andere domeinen, zorgen deze woorden al snel voor een relatief hoge frequentie in de tweede helft van de negentiende eeuw.

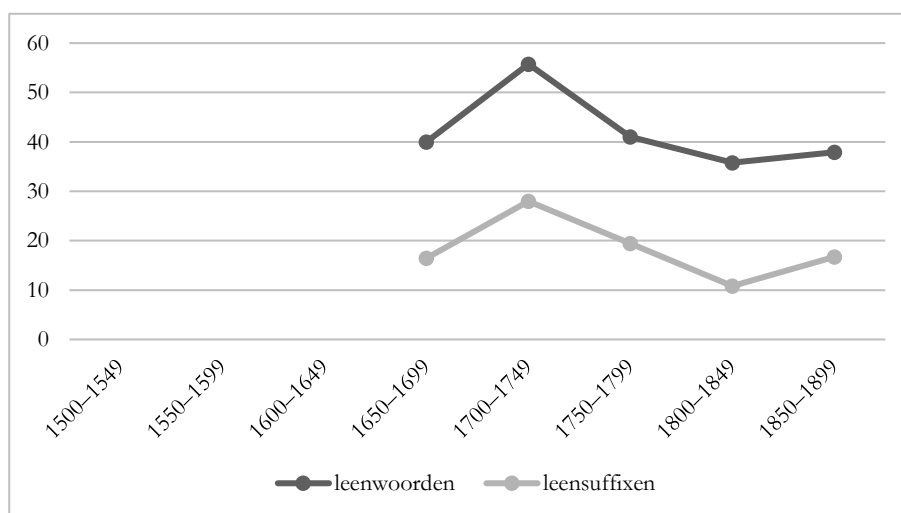


Figuur 7.15. Aantal leenwoorden dat zeker uit het Frans komt (tokens) en aantal leensuffixen (tokens) per 1.000 woorden per periode in het *LOL*-corpus voor het domein Privé.

In figuur 7.15 zijn de resultaten weergegeven voor het Privédomein. Ook bij dit domein is het aantal leenwoorden standaard hoger dan het aantal leensuffixen, maar daarnaast vertonen beide lijnen een bijna identiek verloop, met een stijging in de zestiende en zeventiende eeuw en een

duidelijke piek in de eerste helft van de achttiende eeuw. Hoewel de laatste drie perioden bij de leensuffixen een redelijk stabiel beeld laten zien, is dit bij de leenwoorden niet helemaal het geval: de perioden 1750–1799 en 1850–1899 zijn qua aantal leenwoorden redelijk gelijk aan elkaar, maar het aantal leenwoorden in de eerste helft van de negentiende eeuw is wel wat hoger. Met name het veelvuldige gebruik van *papa* en *mama* valt op in deze periode (9 respectievelijk 11 keer, tegenover 2 respectievelijk 4 keer in de periode 1750–1799).

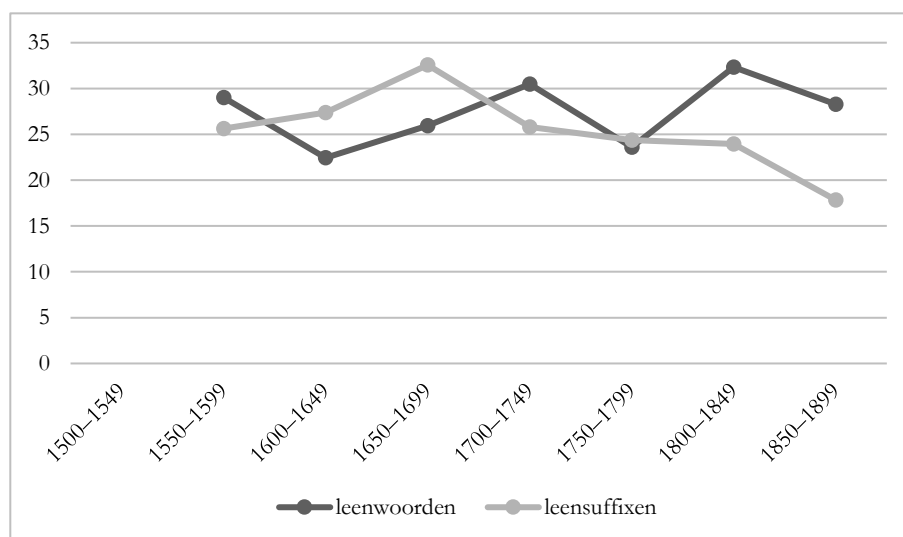
De resultaten van de Publieke Opinie in figuur 7.16 hieronder vertonen vrijwel precies dezelfde ontwikkeling in het gebruik van de leenwoorden als in het gebruik van de leensuffixen, met een piek in de eerste helft van de achttiende eeuw en daarna een relatief stabiel gebruik in de laatste drie perioden.



Figuur 7.16. Aantal leenwoorden dat zeker uit het Frans komt (tokens) en aantal leensuffixen (tokens) per 1.000 woorden per periode in het *LOL*-corpus voor het domein Publieke Opinie.

Waar de relatie tussen de leenwoordfrequentie en de leensuffixfrequentie bij de meeste besproken domeinen duidelijk is, is dat bij het domein Religie niet het geval. De resultaten staan in figuur 7.17 op de volgende pagina. Uit deze figuur blijkt dat het niet zo is dat ofwel de leenwoorden ofwel de leensuffixen duidelijk frequenter zijn bij het domein Religie: in de perioden 1550–1599, 1700–1749, 1800–1849 en 1850–1899 zijn de leenwoorden frequenter, terwijl de leensuffixen vaker voorkomen dan de

leenwoorden in de overige perioden. Met name de daling van het aantal leenwoorden in de periode 1600–1649 in vergelijking met de periode 1550–1599 en de stijging van het gebruik van Franse leenwoorden in de periode 1800–1849 in vergelijking met de periode 1750–1799 is opvallend, aangezien deze ontwikkelingen zich in de andere domeinen niet op deze manier voordoen. Er blijkt echter geen duidelijke verklaring te zijn voor dit patroon: er zijn geen grote verschillen in het gebruik van specifieke woorden of woordfrequenties tussen de perioden die dit verloop kunnen verklaren. Het gebruik van Franse leenwoorden in het domein Religie blijkt dus wat te fluctueren door de tijd heen tussen de 22,4 en 32,3 leenwoorden per 1.000 woorden.



Figuur 7.17. Aantal leenwoorden dat zeker uit het Frans komt (tokens) en aantal leensuffixen (tokens) per 1.000 woorden per periode in het *LOL*-corpus voor het domein Religie.

Concluderend heeft de bespreking van de leenwoordfrequenties per periode per domein en de vergelijking met de leensuffixfrequenties per periode per domein nogmaals laten zien dat er grote verschillen zitten in het gebruik en in de diachrone ontwikkeling in het gebruik van leenwoorden tussen de verschillende domeinen. In de domeinen Literatuur, Privé en Religie worden relatief lage aantallen leenwoorden gevonden, terwijl in de andere domeinen de frequenties hoger liggen. Daarnaast zijn er verschillen tussen de domeinen in hoeverre het verloop van de

leenwoordfrequenties overeenkomt met die van de leensuffixfrequenties in hetzelfde domein. Bij sommige domeinen (Academie, Liefdadigheid, Religie) zit de leensuffixfrequentie soms boven de leenwoordfrequentie, terwijl het aantal leenwoorden bij de andere domeinen consequent boven dat van de leensuffixen zit. Ook vertonen niet alle domeinen hetzelfde patroon bij de leenwoorden en de leensuffixen: bij Religie en in mindere mate bij Academie en Economie wijkt het verloop van de leenwoordfrequenties enigszins af van dat van de leensuffixen. Over het algemeen lopen de patronen van leenwoorden en leensuffixen bij de verschillende domeinen echter wel gelijk. Wat de invloed van het Frans op het Nederlandse lexicon betreft is in elk geval duidelijk dat het afhankelijk is van het domein hoe groot die invloed is geweest en in welke periode die invloed het grootst was.

7.5 Besluit

De uitkomsten van het leenwoordenonderzoek die in dit hoofdstuk besproken zijn, waarbij de resultaten per periode, per domein en per periode per domein zijn gepresenteerd, hebben laten zien dat het gebruik van leenwoorden in het *LOL*-corpus in hoge mate overeenkomt met het gebruik van leensuffixen in het corpus zoals besproken in hoofdstuk 5. Als naar het algemene patroon per periode wordt gekeken, dan laten zowel de Franse leenwoorden als de leensuffixen een toename zien in het gebruik van de zestiende tot en met de eerste helft van de achttiende eeuw, waarna het aandeel Franse elementen in de tweede helft van de achttiende en in de negentiende eeuw weer afneemt. Het maakt hierbij weinig uit naar welke categorie van Franse elementen je kijkt: de leensuffixen, de leenwoorden die zeker uit het Frans komen, alle leenwoorden (inclusief de woorden die mogelijk uit het Frans komen), of een vergelijking waarbij de leenwoorden die zeker uit het Frans komen met een leensuffix en leenwoorden die zeker uit het Frans komen zonder leensuffix afzonderlijk worden bekeken; het patroon is telkens min of meer hetzelfde. Dit geldt ook voor de resultaten per domein en per periode per domein. Religie is de belangrijkste uitzondering: dit domein heeft relatief veel leensuffixen en weinig leenwoorden in vergelijking met de andere domeinen en ook vertoont het verloop van de leensuffixen en leenwoorden over de verschillende perioden niet hetzelfde patroon.

Nu het algemene beeld van zowel de leensuffixen als de leenwoorden uit het Frans duidelijk is, dat wil zeggen, het algemene beeld

van de geleende elementen waarvan vaststaat dat die uit het Frans afkomstig zijn, kan een beeld worden geschetst van de domeinen en tijdvakken waarin deze door het Frans beïnvloede elementen voorkomen. Het leensuffixenonderzoek en het leenwoordenonderzoek samen laten zien dat de zestiende, zeventiende en de eerste helft van de achttiende eeuw de perioden zijn van toenemende invloed van het Frans op het Nederlands: in deze perioden stijgt het gebruik van uit het Frans afkomstige elementen, terwijl dit in de laatste drie perioden (de tweede helft van de achttiende en de negentiende eeuw) juist afneemt. De eerste helft van de achttiende eeuw is de periode waarin de invloed van het Frans het hoogtepunt bereikte.

Wat betreft de domeinen lijken de domeinen Academie en Liefdadigheid over alle perioden van het corpus bezien het meest beïnvloed door het Frans: zowel bij de leensuffixen als bij de leenwoorden lieten beide domeinen hoge frequenties zien. Economie, Publieke Opinie en Religie nemen een middenpositie in, waarbij Economie en Publieke Opinie relatief veel leenwoorden bevatten maar weinig leensuffixen en Religie juist wel een hoge leensuffixfrequentie maar weer een relatief lage leenwoordfrequentie heeft. De domeinen Literatuur en Privé blijken in vergelijking met de andere domeinen weinig Franse elementen te bevatten, aangezien die in beide onderzoeken de laagste gebruiksfrequenties hadden. In deze domeinen is de invloed van het Frans dus beperkter geweest dan in de andere domeinen.

Met het hier geschetste patroon van het gebruik van elementen waarvan is vastgesteld dat die uit het Frans komen, kan worden bekeken of het gebruik van elementen waarvan wordt geclaimd dat die door het Frans zijn beïnvloed hetzelfde beeld vertonen als dat van het hier geschetste profiel. Eén van deze geclaimde Franse elementen is het gebruik van de relativa *dewelke* en *hetwelk*. Wanneer deze elementen frequent voorkomen in dezelfde tijdvakken en domeinen als de tijdvakken en domeinen waarin de leenwoorden en leensuffixen het frequentst waren, dan zou dit een indicatie zijn dat deze geclaimde elementen daadwerkelijk onder invloed van het Frans hebben gestaan. In hoofdstuk 10 wordt onderzocht of het patroon dat is gevonden voor de Franse leensuffixen en leenwoorden ook opgaat voor de relativa *dewelke* en *hetwelk*, maar eerst wordt in de volgende twee hoofdstukken dieper ingegaan op enkele aspecten van de leenwoorden uit het *LOL*-corpus om een beter beeld te krijgen van de invloed van het Frans op het Nederlandse lexicon.