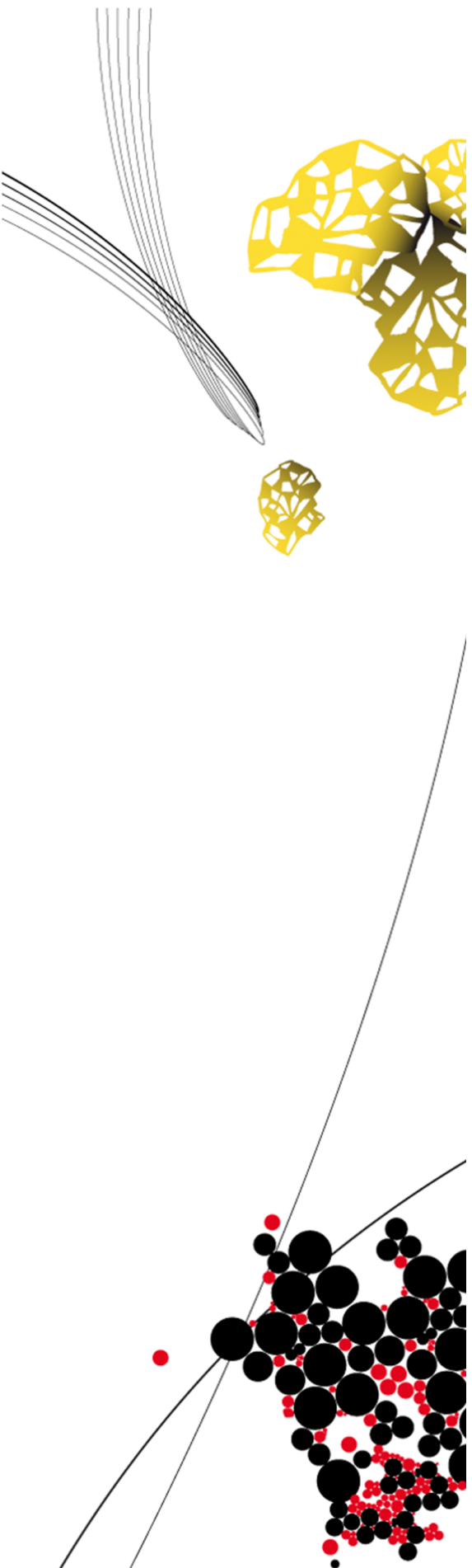




UNIVERSITY OF TWENTE.

Master Educatie en Communicatie in de Bètawetenschappen

Computationeel denken bij vwo wiskunde A

Verslag van Onderzoek van Onderwijs

J.M. Neeft

30 September 2019

Supervisors:

Mark Timmer

Tom Coenen

Department of Teacher Development,
Faculty of Behavioural, Management
and Social sciences.

Samenvatting

In dit ontwerponderzoek is onderzocht hoe computationele denkvaardigheden bij leerlingen die het vak wiskunde A volgen ontwikkeld kunnen worden. Verschillende aspecten spelen hierbij een belangrijke rol, waaronder modelleren, simuleren, probleemoplossen en het gebruik van data. Het onderzoek heeft geleid tot een lijst geschikte onderwerpen die gebruikt kunnen worden om computationeel denken te bevorderen, waarna een praktische opdracht over het K-means-algoritme is ontworpen. In deze opdracht staat het maken van een aanbevelingssysteem voor Netflix centraal. Elke opgave is door middel van een bestaand classificatiemodel gekoppeld aan computationele denkvaardigheden. Er heeft een try-out van de opdracht plaatsgevonden, waarna uitwerkingen van leerlingen en interviews zijn gebruikt om de praktische opdracht te evalueren. Geconcludeerd kan worden dat de context ondersteunend kan werken en dat het gebruik van technologie een meerwaarde heeft. Vragen die betrekking hadden op enkele complexe wiskundige begrippen werden nog niet juist beantwoord. Naar aanleiding van de try-out zijn suggesties gedaan voor vervolgonderzoek en ter verbetering van de opdracht.

Sleutelwoorden: computationeel denken, wiskundig denken, K-means algoritme, praktische opdracht, vwo wiskunde A

Voorwoord

“The world is full of magical things patiently waiting for our wits to grow sharper.”

Bertrand Russell

‘Waarom zou ik nu wiskunde moeten leren?’ is een vraag die door menig leerling weleens gesteld zal worden. Mijn overtuiging is dat wiskundeonderwijs leerlingen in ieder geval zou moeten uitdagen om kritisch te denken en te redeneren. Het gaat erom te kunnen bedenken welk wiskundig gereedschap bruikbaar is om een probleem aan te pakken. Hierin sta ik niet alleen, denk maar eens aan de veelgeciteerde uitspraak van Polya [25] over wiskundeonderwijs: “First and foremost, it should teach young people how to THINK.” Ik hoop dat dit onderzoek een bijdrage kan leveren aan de verdere ontwikkeling van wiskundig denken op de middelbare school.

Ik wil allereerst Mark Timmer bedanken voor zijn uitstekende begeleiding. Mark, je bent heel enthousiast en betrokken. Ik heb veel van je geleerd bij de colleges Vakdidactiek en tijdens dit afsluitende onderzoek heb je mijn concepten en ideeën altijd snel van goede feedback kunnen voorzien. Jouw hulp, en lichte druk om af en toe met een nieuwe versie te komen, heeft veel bijdragen aan de goede afloop van dit onderzoek. Ik wil ook Tom Coenen bedanken voor het overleg over mijn onderzoeksopzet, het lezen van mijn conceptversie en het zijn van tweede beoordelaar. Tot slot wil ik Jos Tolboom bedanken voor het gesprek over computationeel denken bij SLO.

Dit onderzoek is niet zonder slag of stoot tot stand gekomen. Naast mijn master ben ik de uitdaging aangegaan om 4 dagen per week te werken als docent wiskunde op een middelbare school. Dat was in het begin zeker niet altijd even makkelijk te combineren. Zonder enige ervaring lesgeven aan groepen van 30 pubers is een uitdaging. Ik ben best trots hoe ik dat gedaan heb. Ik hoop anderen te kunnen inspireren om mijn voorbeeld te volgen: het docentschap is mooi en waardevol!

Jelle Neeft
– 18 september 2019

Inhoudsopgave

1	Inleiding	4
1.1	Aanleiding	4
1.2	Onderzoeksvragen	5
1.3	Leeswijzer	6
2	Methode	8
2.1	Procedure	8
2.2	Respondenten	9
2.3	Dataverzameling	9
3	Theoretisch en conceptueel kader	11
3.1	Computationeel denken	11
3.2	Mogelijkheden voor computationeel denken bij vwo wiskunde A	15
3.2.1	Onderwerpen	15
3.2.2	Statistisch leren	18
3.3	K-means algoritme	19
3.4	Didactische achtergrond	20
4	Ontwerpeisen	22
4.1	Doelgroep en organisatie	22
4.2	Inhoudelijke leerdoelen	23
5	Resultaten	24
5.1	Ontwerp	24
5.2	Try-out	28
5.2.1	Voorwaarden	28

5.2.2	Leerdoelen	30
5.3	Interpretatie	35
6	Conclusie en discussie	36
6.1	Conclusie	36
6.2	Discussie	37
6.2.1	Limitaties	37
6.2.2	Implicaties en aanbevelingen	38
	Bibliografie	41
	Bijlagen	44
	Bijlage A: Interviewleidraad	44
	Bijlage B: Opdracht	45
	Bijlage C: Correctievoorschrift	59
	Bijlage D: Interviews	65

Hoofdstuk 1

Inleiding

De wereld om ons heen is aan verandering onderhevig. Eén van de ontwikkelingen is het toenemende gebruik van digitale technologie, zowel in het dagelijks leven als in de wetenschap. Het wiskundeonderwijs zal zich hierop aan moeten passen, om zo ook nu en in de toekomst als gereedschap te kunnen dienen voor het dagelijks leven en als taal der wetenschap. Dit onderzoek sluit aan bij het toenemende gebruik van digitale technologie in zowel de maatschappij als de wetenschap.

1.1 Aanleiding

Directe aanleiding van dit onderzoek is de vraag vanuit het SLO om lessen te ontwerpen waarbij wiskundig en computationeel denken centraal staan (computationeel denken is het zodanig formuleren van problemen dat ze met behulp van computertechnologie opgelost kunnen worden). Hier is reden toe, gezien de volgende twee recente ontwikkelingen in het Nederlandse onderwijs. Ten eerste is bij de invoering van de nieuwe curricula op havo en vwo in 2015 wiskundig denken een belangrijk aandachtspunt geworden [10]. Ten tweede pleit het curriculumvernieuwingsplatform voor wiskunde, curriculum.nl, en ook het platform voor informatica, voor toenemende aandacht voor digitale geletterdheid en computationeel denken [3]. Dit onderzoek sluit aan bij het op 01-12-2018 gestartte onderzoeksproject 'Computationeel denken en wiskundig denken: digitale geletterdheid in wiskundecurricula' van het NRO waarin vijf scholen, twee universiteiten en het SLO samen een theoriegedreven ontwerpstudie opzetten waarvan één van de doelen is om empirisch gevalideerde leeractiviteiten te ontwikkelen waarin kernaspecten van computationeel en wiskundig denken voorkomen.

Drijvers, Van Streun en Zwaneveld [11] onderscheiden drie didactische functies van ICT: ICT als gereedschap, ICT als oefeninstrument en ICT als ondersteuner voor begripvorming. In het wiskundeonderwijs wordt al steeds meer gebruik gemaakt van digitale tools, bijvoorbeeld voor het presenteren van kennis, digitaal lesmateriaal en digitale toetsen [13]. Zo speelt ICT in de vernieuwde statistiekprogramma's van wiskunde A en C een belangrijke rol [10]. Bovendien wordt in vervolgopleidingen en toepassingen ICT gebruikt als gereedschap om gegevens in grote datasets te verwerken en om bepaalde onderzoeksvra-

gen te kunnen beantwoorden. Om goed gebruik te kunnen maken van de mogelijkheden die ICT biedt moet naast technische vaardigheid (zoals programmeren en kennis van statistische pakketten als Excel en R), vooraf goed worden nagedacht hoe een probleem zo geformuleerd en geautomatiseerd kan worden dat deze kan worden opgelost (met behulp van technologische middelen). Naar dit laatste wordt vaak gerefereerd met de term computationeel denken [31]. Traditioneel gezien wordt dit geschaard onder het vakgebied informatica, maar aangezien niet alle leerlingen informatica op de middelbare school hebben en computationeel denken goed past binnen het thema 'wiskundige denkactiviteiten', lijkt het schoolvak wiskunde een prima plek om computationeel denken te bevorderen.

Een andere belangrijke motivatie voor het introduceren van computationeel denken in de wiskundeles is de snel veranderende werkelijkheid in de wetenschap en het bedrijfsleven [2]. In de laatste 20 jaar heeft ruwweg elke discipline gerelateerd aan natuurwetenschap en wiskunde een toename gezien in het gebruik van computationele elementen. Door computationeel denken en het gebruik van high-tech tools in de klas te brengen krijgen leerlingen een meer realistisch beeld van welke disciplines er zijn. Daarmee bereidt het leerlingen voor op succesvolle carrières [1]. Vanuit een didactisch perspectief kan computationeel denken en het gebruik van geschikte digitale tools het begrip van wiskundige concepten verdiepen [12]. Het tegengestelde is ook waar: wiskunde kan een goede context zijn om computationeel denken op toe te passen [38]. Computationeel denken past goed bij het wiskundig denken uit de vernieuwde wiskundeprogramma's waarin aspecten als probleemoplossen, modelleren, abstraheren en analytisch denken van belang zijn.

De wederkerige relatie tussen wiskunde en computationeel denken – computationeel denken om het wiskundecurriculum te verrijken en het gebruiken van het vak wiskunde om computationeel denken te promoten – vormt de basis van de argumentatie om computationeel denken en wiskunde samen te brengen.

1.2 Onderzoeksvragen

Het onderzoek tracht de volgende hoofdvraag te beantwoorden:

Hoe kan een werkmiddag vwo-leerlingen bij het vak wiskunde A ondersteunen bij het ontwikkelen van computationele denkvaardigheden?

Er is gekozen voor deze specifieke afbakening, namelijk vwo-leerlingen bij het vak wiskunde A, vanwege de vraag naar materiaal voor het domein Kansrekening en Statistiek (dat behoort tot het curriculum voor wiskunde A) en het praktische aspect dat er een groep leerlingen beschikbaar was om het onderwijsontwerp bij uit te testen.

Op basis van de hoofdvraag zijn meerdere deelvragen geformuleerd. Deze zijn opgesplitst in twee delen [33]: deelvragen voor de eerste fase waarin gezocht zal worden naar ontwerpeisen en deelvragen voor de tweede fase waarin het ontwerp gemaakt en getest zal worden.

Deelvragen onderzoeksfase 1 (ontwerpeisen):

- Wat wordt er precies bedoeld met computationeel denken?

- Welke vaardigheden hebben betrekking op computationeel denken?
- Welke lesontwerpen om computationeel denken te bevorderen bestaan al?
- Welke didactische aspecten zijn belangrijk bij het aanleren van computationeel denken?

Deelvragen onderzoeksfase 2 (ontwerp):

- Hoe zou een ontwerp voor het ontwikkelen van computationeel denken bij vwo-leerlingen die wiskunde A volgen eruit kunnen zien?
- In hoeverre voldoet het ontwerp aan de ontwerpeisen?
- In hoeverre voldoet het ontwerp in de praktijk?
- Op welke wijze moet het ontwerp (na het testen) worden aangepast?

1.3 Leeswijzer

De structuur van het verslag wordt nu beschreven. In hoofdstuk 2 wordt allereerst de onderzoeksmethode beschreven. Hoewel het wellicht gebruikelijker is om te starten met het literatuuronderzoek, is het hier wenselijk om dat niet te doen, aangezien het literatuuronderzoek een onderdeel is van de onderzoeksmethode. In de eerste paragraaf wordt de onderzoeksprocedure besproken, waarna een beschrijving van de respondenten gegeven wordt. Het hoofdstuk sluit af met een beschrijving van de manieren waarop data verzameld zal worden.

Hoofdstuk 3, dat het theoretisch kader bevat, vormt de opbouw naar de ontwerpeisen. In dit hoofdstuk worden de deelvragen uit de eerste onderzoeksfase beantwoord, waarna de ontwerpeisen in hoofdstuk 4 opgesteld kunnen worden. Hoofdstuk 3 start met een beschrijving van de term computationeel denken, waarna het model Weintrop wordt geïntroduceerd. Dit model wordt in het vervolg van dit onderzoek gebruikt om het ontwerp mee te toetsen. De tweede paragraaf beschrijft eerdere ontwerpen voor computationeel denken en mogelijkheden die in de literatuur worden aangedragen. Vervolgens geven we in de derde paragraaf achtergrond over het door ons gekozen onderwerp van het ontwerp: het K-means algoritme. Het tweede hoofdstuk wordt afgesloten met didactische theoriën die van belang zijn in de ontwerpfasen.

In hoofdstuk 4 staan de ontwerpeisen die op zijn gesteld naar aanleiding van de uitkomsten van het theoretisch en conceptuele onderzoek. We delen de ontwerpeisen in twee delen op: organisatorische eisen en kenmerken van de doelgroep enerzijds en (vak)didactische en inhoudelijke eisen anderzijds. Hoofdstuk 5 beschrijft vervolgens de resultaten, welke bestaan uit zowel het ontwerp zelf als een analyse van de try-out. Eerst komt het ontwerp aan bod. De gehele opdracht wordt opgedeeld in opgedragen delen. Daarnaast wordt, aan de hand van het model van Weintrop, toegelicht waarom deze opdracht computationeel denken bevordert. Hierna volgt een analyse van de try-out, waarbij aan de hand

van interviews en leerlinguitwerkingen wordt beoordeeld of de opdracht voldoet aan de ontwerpeisen. In de laatste paragraaf interpreteren we enkele resultaten.

Hoofdstuk 6 concludeert onze bevindingen en bediscussieert dit onderzoek. Er worden limitaties en implicaties voor de praktijk gegeven. Tot slot zijn er aanbevelingen, zowel ter verbetering van de opdracht als voor vervolgonderzoek.

Als laatste kan in bijlage A de interviewleidraad gevonden worden. In bijlage B staat de gehele opdracht, zoals uitgevoerd bij de try-out. Bijlage C bevat een modeluitwerking van de opdracht. Dit verslag sluit af met bijlage D, waarin de uitwerkingen van de leerlinginterviews staan.

Hoofdstuk 2

Methode

In dit hoofdstuk presenteren we onze methode. We starten met het beschrijven van de procedure, waarna we de respondenten beschrijven die aan het onderzoek hebben meegedaan. Tot slot beschrijven we hoe de dataverzameling heeft plaatsgevonden en waaruit deze bestaat.

2.1 Procedure

Dit ontwerponderzoek is uitgevoerd op een manier die zijn oorsprong vindt in de Twentse onderwijskunde. Ontwerponderzoek kent in deze traditie twee fasen: vooronderzoek en ontwerp. In de eerste fase staan de verkenning van het onderwijsprobleem centraal en worden criteria en randvoorwaarden genoemd. De ontwerpfase is een cyclisch proces van evalueren, overleggen en bijstellen [32].

Het vooronderzoek bestaat uit het vormen van een theoretische en conceptueel kader (zie hiervoor hoofdstuk 3). Dit kader is ontstaan door literatuuronderzoek te doen naar wat computationeel denken is en welke vaardigheden daarop betrekking hebben. Er is een verkenning van mogelijke onderwerpen en opdrachten gedaan die geschikt zijn om computationeel denken te bevorderen. Hierbij is gebruikt gemaakt van beschreven ervaringen en suggesties in de literatuur. Tot slot zijn didactische aspecten geïnventariseerd die een uitgangspunt kunnen vormen bij het ontwerpen van de opdracht. Deze fase heeft als eindresultaat een lijst van ontwerpeisen die gepresenteerd wordt in hoofdstuk 4.

In de voorverkenning is een inventarisatie gedaan naar mogelijke onderwerpen die geschikt zijn voor het bevorderen van computationeel denken. Van het meest veelbelovende idee is een eerste ontwerp gemaakt waarna deze iteratief werd besproken met een expert en vervolgens verbeterd. Het resultaat hiervan is een opdracht die in een try-out aan een realistische praktijktest onderworpen is bij de doelgroep. De try-out startte met een korte introductie door de docent, waarin het doel (computationeel denken) verteld werd en duidelijk werd gemaakt dat deze opdracht gemaakt is ter afsluiting van de onderwijsmaster van de heer Neeft. Er is verder geen inhoudelijke klassikale toelichting op de opdracht geweest. Tijdens de try-out is bekeken of het ontwerp aan de opgestelde voorwaarden

voldeed en of de leerdoelen werden behaald. Op woensdag 12 juni 2019 hebben daartoe 16 leerlingen uit vwo 5 van een middelbare school in Zeist in duo's de opdracht doorlopen. De opdrachten zijn gecorrigeerd en er zijn, binnen enkele dagen na afloop, een aantal leerlingen geïnterviewd over het ontwerp. Zinssneden uit de interviews zijn gebruikt om te bepalen in hoeverre leerdoelen zijn bereikt. Daarnaast zijn voor de bepaling in hoeverre leerdoelen behaald zijn de antwoorden die leerlingen gaven op de opdrachten gebruikt. De conclusies uit deze analyses zijn input voor verbeteringen in het ontwerp. Voorts hebben de leerlingen een cijfer gekregen voor hun opdracht, welke ze als bonus konden gebruiken voor de laatste reguliere schriftelijke toets van het schooljaar. Het cijfer maakt geen deel uit van het schoolexamen, omdat het programma daarvoor al eerder was vastgelegd.

De onderzoeksopzet is op 29-05-2019 goedgekeurd door de BMS Ethics Committee van de Universiteit Twente onder het nummer 190912.

2.2 Respondenten

Twee typen respondenten spelen een rol bij dit onderzoek. Allereerst de expert die heeft geholpen in het ontwerpproces door feedback te geven op de verschillende versies van het onderwijsontwerp, Mark Timmer. Hij is wiskundedocent op een middelbare school, gepromoveerd theoretisch informaticus, universitair vadicus en heeft veel ervaring met het ontwerpen en beoordelen van les- en leermateriaal. Ten tweede zijn er leerlingen die meegedaan hebben met de try-out van ons ontwerp in de praktijk. Een deel van deze leerlingen heeft ook deelgenomen aan de leerlinginterviews.

De try-out heeft plaatsgevonden op 12 juni 2019 op een middelbare school in Zeist van 12.40 uur tot 15.45 uur (met een uitloop tot 16.00 uur). In totaal hebben 16 leerlingen uit een vwo 5 wiskunde A klas het gemaakte onderwijsontwerp doorgewerkt. Van deze 16 leerlingen waren er 6 vrouw en 10 man. Er is in duo's gewerkt. Eén leerling moest door een medische afspraak eerder weg; de opdracht is door de andere persoon van het groepje afgemaakt.

Voor de interviews zijn 4 leerlingen geselecteerd op basis van hun voortschrijdende gemiddelde voor wiskunde A. Om een zo goed mogelijke afspiegeling te krijgen van de populatie zijn de volgende leerlingen geselecteerd op basis van het percentiel. De 4 leerlingen waarvan het voortschrijdende gemiddelde het dichtst bij het k^e percentiel ligt, met $k = 20\%$, 40% , 60% , 80% zijn geselecteerd voor het interview. De leerlingen kregen dit niet te horen totdat de opdracht voltooid was, om zo de resultaten niet te beïnvloeden. Er zijn zowel vrouwelijke als mannelijke respondenten geselecteerd. De interviews zijn allen binnen 2 dagen na het uittesten van het ontwerp afgenomen, zodat de leerlingen de opdrachten nog vers in hun geheugen hadden.

2.3 Dataverzameling

Op twee momenten heeft er dataverzameling plaatsgevonden: bij de try-out van het ontwerp en bij het afnemen van de leerlinginterviews.

Ten eerste is er de try-out die plaatsgevonden heeft op een middelbare school in Zeist. De leerlingen hebben de uitwerkingen van de opdracht in Word verwerkt en na afloop van de middag verzonden naar de onderzoeker. Opgave 2 is op verzoek van de onderzoeker overgeslagen, omwille van de tijd en technologische beperkingen op de school. Tijdens het uitvoeren van de opdrachten hebben wij bijzonderheden die voor de evaluatie van de opdracht van belang zijn genoteerd. Vervolgens zijn de uitwerkingen door de onderzoeker aan de hand van een antwoordmodel nagekeken. Voor elke vraag is de zogenoemde p' -waarde berekend, welke berekend wordt door de gemiddelde score op een opgave te delen door de maximaal haalbare score op die opgave. Het kan gezien worden als een maat voor de moeilijkheid van een vraag (daargelaten dat het ook kan dat een vraag een lagere p' -waarde heeft door tijdgebrek; dit moet in de analyse worden meegenomen). De p' -waarden gebruiken we in de analyse van ons ontwerp in hoofdstuk 5.

De dataverzameling is uitgevoerd conform de ethische richtlijnen voor wetenschappelijk onderzoek. De leerlingen hebben een week voor uitvoering van de werkmiddag een brief ondertekend waarin uitgelegd werd hoe de dataverzameling plaats zou vinden. Door het ondertekenen gaven de leerling toestemming om geïnterviewd te worden betreffende hun ervaringen met de werkmiddag. Tevens zijn zij door het ondertekenen akkoord gegaan met het gebruik van hun verslag voor dit onderzoek.

Ten tweede hebben wij leerlingenterviews afgenomen na afloop van de try-out. Interviews kunnen op een continue schaal met aan het ene uiterste gestructureerd en aan het andere uiterste ongestructureerd geplaatst worden. Het gestructureerde interview kan als kwantitatief beschouwd worden en bevat veel gesloten vragen. Het semi-gestructureerde en het ongestructureerde interview laten zich karakteriseren door toenemende flexibiliteit en een afname in structuur. In tegenstelling tot het gestructureerde interview biedt het semi-gestructureerde interview de mogelijkheid om dieper door te vragen en antwoorden die gegeven worden beter te grijpen [20]. In een semi-gestructureerd interview wordt er vooraf wel een lijst met onderwerpen en vragen opgesteld door de interviewer, maar heeft de interviewer tijdens het interview de mogelijkheid om dieper op bepaalde zaken in te gaan. Een aantal interviewvragen zijn taakgericht, zodanig dat de conceptuele kennis van leerlingen beoordeeld kan worden, maar ook zodanig dat er mogelijkheid is om het begrip te vergroten. Een nadeel van deze techniek is dat de uitkomst afhankelijk is van de vaardigheid van de interviewer, echter, de interviews zullen wel allemaal afgenomen worden door dezelfde interviewer. In Bijlage A bij dit verslag is de ontworpen interviewleidraad voor het semi-gestructureerde interview opgenomen.

Hoofdstuk 3

Theoretisch en conceptueel kader

Dit hoofdstuk vormt de opbouw naar de ontwerpeisen. We beantwoorden de deelvragen van onderzoeksfase 1. Om goede ontwerpeisen op te kunnen stellen is het allereerst van belang om te weten wat computationeel denken nu precies is en welke vaardigheden daarbij horen. Daarnaast onderzoeken we in dit hoofdstuk welk materiaal er al beschikbaar is voor het bevorderen van computationeel denken bij vwo wiskunde A leerlingen. Tot slot noemen we didactische aspecten die van belang zijn.

3.1 Computationeel denken

Over de definitie van de term computationeel denken heerst in de literatuur weinig consensus, met name wanneer computationeel denken bedoeld wordt om te gebruiken in een wiskundige context [38]. Bekende definities zijn die van Wing [39] en Lu & Fletcher [19]. Wing [39] stelt: ‘computationeel denken bevat probleemoplossen, het ontwerpen van systemen en het begrijpen van menselijk gedrag, door gebruik te maken van fundamentele concepten uit de informatica’. Volgens Lu&Fletcher [19] bevat computationeel denken de gedachteprocessen die nodig zijn om problemen op zo een manier te formuleren en begrijpen dat ze op kunnen worden gelost met behulp van berekeningen.

Uit de review die Weintrop et al. [38] deed naar in de literatuur veelgenoemde vaardigheden en aspecten behorende bij computationeel denken, volgt de volgende lijst:

1. Vaardigheid om met problemen om te gaan die een open-einde hebben.
2. Aspecten van problemen kunnen abstraheren
3. Doorzettingsvermogen betreffende uitdagende problemen.
4. Problemen kunnen herformuleren in een herkenbaar probleem.
5. Vertrouwd zijn met complexiteit.
6. Sterke en zwakke aspecten kunnen benoemen van representaties en datasystemen.
7. Ideeën representeren zodat ze computationeel waardevol zijn.
8. Algoritmische oplosmethodes genereren.
9. Een probleem opdelen in kleinere problemen.
10. Herkenning wanneer een algoritme dubbelzinnig is.

Deze vaardigheden zijn erg breed geformuleerd en vrij moeilijk meetbaar. In dit onderzoek maken we daarom gebruik van de classificatie die Weintrop ontwierp. Deze classificatie gebruiken we enerzijds om ons onderwijsontwerp vorm te geven en anderzijds om te beoordelen in hoeverre ons ontwerp computationeel denken bevordert.

De classificatie van Weintrop bestaat uit vier categorieën: data, modelleren & simuleren, computationeel probleemoplossen en systeembenken. Elk van deze categorieën bestaat uit vijf tot zeven veelgebruikte praktijken, welke te zien zijn in figuur 3.1. Deze praktijken verschillen van de lijst van tien aspecten hierboven in die zin dat ze niet slechts de vaardigheid en het concept benadrukken, maar ook de specifieke kennis behorende bij elke praktijk. In de taxonomie worden de categorieën als niet-overlappend gepresenteerd, maar in de praktijk zijn de onderdelen uit verschillende categorieën met elkaar verweven en afhankelijk; ze worden vaak gecombineerd om specifieke (wiskundige) doelen te bereiken.

Data Practices	Modeling & Simulation Practices	Computational Problem Solving Practices	Systems Thinking Practices
Collecting Data	Using Computational Models to Understand a Concept	Preparing Problems for Computational Solutions	Investigating a Complex System as a Whole
Creating Data	Using Computational Models to Find and Test Solutions	Programming	Understanding the Relationships within a System
Manipulating Data	Assessing Computational Models	Choosing Effective Computational Tools	Thinking in Levels
Analyzing Data	Designing Computational Models	Assessing Different Approaches/Solutions to a Problem	Communicating Information about a System
Visualizing Data	Constructing Computational Models	Developing Modular Computational Solutions	Defining Systems and Managing Complexity
		Creating Computational Abstractions	
		Troubleshooting and Debugging	

Figuur 3.1: Taxonomie van Weintrop[38]

In de volgende opsomming bespreken we de vier categorieën, met als doel om te verduidelijken wat elke praktijk inhoudt. Deze beschrijvingen gebruiken we om opgaven te ontwerpen die computationeel denken bevorden. In hoofdstuk 5 koppelen we de opdrachten aan de verschillende praktijken die hieronder beschreven staan om te laten zien dat onze opgaven computationeel denken bevorderen.

1. Categorie Data

1.1. Data verzamelen

Computationele instrumenten kunnen gebruikt worden in diverse fases van het dataverzamelingsproces, zoals het doen van metingen en het opslaan van gegevens.

1.2. *Data creëren*

Soms zijn fenomenen niet meetbaar of niet te observeren. Computersimulaties kunnen dan behulpzaam zijn in het creëren van data op een schaal die anders onmogelijk was.

1.3. *Data manipuleren*

Om grote datasets te kunnen analyseren en interpreteren is het essentieel om snel en betrouwbaar te kunnen filteren, opschonen, sorteren, normaliseren, samenvoegen en splitsen.

1.4. *Data analyseren*

Computationele instrumenten maken het mogelijk om grote hoeveelheden data op een effectieve en betrouwbare manier te analyseren. Analyse-strategieën die hieronder vallen zijn bijvoorbeeld het zoeken naar patronen, regels vinden voor het categoriseren van data én trends en correlaties identificeren.

1.5. *Data visualiseren*

Het communiceren van resultaten is een belangrijke component van elk data-onderzoek; visualisering is hierbij een krachtige strategie. Computationele instrumenten kunnen hierbij van dienst zijn.

2. **Categorie Modelleren en Simuleren**

2.1. *Een concept begrijpen*

Modellen kunnen als krachtig leerinstrument gebruikt worden om (wiskundige) concepten beter te begrijpen. Ze geven leerlingen meer eigenaarschap.

2.2. *Oplossingen vinden en testen*

Modellen kunnen gebruikt worden om hypothesen te testen en oplossingen van problemen te vinden. De computationele modellen maken het mogelijk om snel en goedkoop meerdere oplossingen te testen, wat erg behulpzaam kan zijn wanneer er voor parameters bepaalde waarden gekozen moeten worden.

2.3. *Modellen beoordelen*

Een belangrijk aspect van modelleren is om te begrijpen in hoeverre het model de werkelijkheid representeert en hoe de gemaakte aannames van invloed zijn op de uitkomst van het model.

2.4. *Modellen ontwerpen*

Het ontwerpen van een model vereist het maken van technische, methodologische en conceptuele keuzes.

2.5. *Modellen maken*

Voor het maken van een model moet men de modelonderdelen zo kunnen formuleren dat een computer het kan lezen. Soms zal dit programmeren vereisen; soms zijn er instrumenten beschikbaar die kunnen ondersteunen.

3. **Categorie Computatieel probleemoplossen**

3.1. *Herformuleren*

Een probleem zodanig formuleren dat de computationele instrumenten gebruikt kunnen worden. Veelgebruikte strategieën hiervoor zijn opdelen in kleinere problemen, versimpelen en koppelen aan problemen waarvoor al een computationeel instrument bekend is.

3.2. *Programmeren*

Dit onderdeel bestaat uit het begrijpen en aanpassen van programma's gemaakt door anderen en het schrijven van nieuwe code vanaf het begin.

3.3. *Instrumenten beoordelen en kiezen*

Een probleem kan vaak met behulp van meerdere computationele instrumenten worden opgelost. Het gaat hier om het identificeren van zwakke en sterke punten van verschillende instrumenten in relatie met het probleem. De keuze voor het instrument kan gemaakt worden op basis van bijvoorbeeld functionaliteit, omvang en flexibiliteit, datatype en aanwezige kennis.

3.4. *Oplossingscomponenten ontwerpen*

Het gaat hier om het ontwerpen van kleine, modulaire, herbruikbare componenten. Het voordeel hiervan is dat de analyse van de oplossing op een gedetailleerd, modulair niveau mogelijk is en dat andere problemen wellicht ook van de modulaire componenten gebruik kunnen maken.

3.5. *Computationele abstractie*

Het vermogen tot abstraheren is noodzakelijk om problemen die structureel hetzelfde zijn (maar op detailniveau verschillen) op te kunnen lossen. Het gaat hierbij om de vaardigheid om bij een probleem de belangrijke aspecten naar de voorgrond te kunnen halen en zo gelijkenissen met andere problemen te kunnen maken.

3.6. *Foutopsporing*

Foutopsporing gaat over het uitvinden van waarom iets niet werkt of zich niet gedraagt zoals verwacht. Een belangrijke aanpak die hierbij hoort is het systematisch testen van het systeem om de bron van de fout te isoleren (modulaire componenten kunnen hierbij van nut zijn).

4. Categorie Systeendenken

4.1. *Complex systeem als geheel onderzoeken*

Soms is het nodig om te focussen op het geheel om de eigenschappen van het systeem onderzoeken. Het gaat hier dan om de vaardigheid om de input en output van het systeem te bekijken en te meten, waarbij de details van de onderliggende processen als een black box worden gezien.

4.2. *Relaties begrijpen*

Waar sommige vragen beantwoord kunnen worden door naar een complex systeem als geheel te kijken, kunnen andere vragen juist beantwoord worden door te begrijpen wat de wisselwerking is tussen de verschillende onderdelen van een systeem.

4.3. *Wisselen van schaalniveau*

Systemen kunnen geanalyseerd worden van micro-niveau (waarbij gekeken wordt naar de kleinste elementen) tot macro-niveau (waarbij het systeem als geheel bekeken wordt). Elk niveau, en ook een combinatie of iteratie, kan tot nieuwe inzichten leiden.

4.4. *Informatie communiceren*

Het kan lastig zijn om de kennis over een groot en complex systeem over te dragen. Om dit wel succesvol te doen is het belangrijk om duidelijke visualisaties te maken van de belangrijkste aspecten van wat geleerd is en om goede keuzes te maken in wat weggelaten kan worden.

4.5. *Complexiteit managen*

Een effectief systeem is in staat om het beoogde doel te bereiken, terwijl het zijn grootte en complexiteit beperkt.

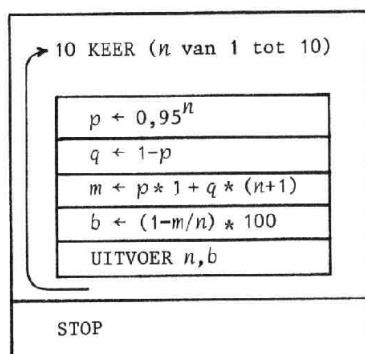
3.2 Mogelijkheden voor computationeel denken bij vwo wiskunde A

3.2.1 Onderwerpen

Van der Meulen [34] ontwierp eerder een werkmiddag in het kader van computationeel denken voor 5 vwo wiskunde B. Hij deed dit rondom het algoritme-ontwerpparadigma Branch-and-Bound. Tijdens de werkmiddag gingen leerlingen aan de slag met opdrachten over het optimaliseren van een taakplanning en maakten daarbij kennis met algoritmen, heuristieken en computationele complexiteit. Hierbij zijn de didactische principes van geleid heruitvinden en relationeel begrip gebruikt. Geconcludeerd wordt dat het introduceren van puzzelachtige problemen en een goed gekozen algoritme-ontwerpparadigma in staat zijn om aspecten van computationeel denken aan te leren, maar dat daarbij wel

voldoende sturing van de docent en/of de opdracht nodig is. Het lopende praktijkgerichte onderwijsonderzoek onder leiding van Drijvers [23] stelt voor wiskunde B voor om materiaal te ontwikkelen dat calculusproblemen oplost met behulp van Geogebra. Dit sluit aan bij het curriculum en het gebruik van software in de klas. Benakli [5] ontwierp vier computationele projecten bij het onderwerp Calculus, gebruik makende van de gratis open-source programmeertaal R. Zijn projecten gaan onder andere over het visualiseren van niet-differentieerbare functies en het benaderen van integralen met behulp van simulatie.

In een uitgave uit de Hewet-reeks (uit 1983) van het Freudenthal Instituut [14], uitgegeven ter vernieuwing van de examenprogramma's wiskunde destijds, zijn elementen te vinden die computationeel denken bevorderen. Zo is in de uitgave over kansverdelingen meermaals een structuurdiagram te vinden. Deze diagrammen zijn eigenlijk niets anders dan algoritmes met een hoog abstractiegehalte. Een voorbeeld van een dergelijk diagram is te vinden in onderstaand figuur 3.2. Leerlingen moeten dergelijke diagrammen kunnen lezen, gebruiken en "vertalen in een programma en laten verwerken door een computer". Deze diagrammen en opdrachten zijn in hedendaagse wiskundemethodes niet meer terug te vinden.

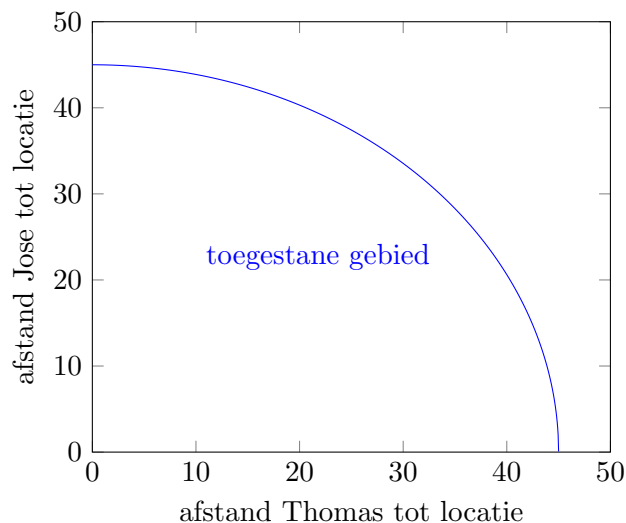


Figuur 3.2: Structuurdiagram uit HEWET [14]

In het hedendaagse examenprogramma voor vwo wiskunde A is het volgende te vinden: "de kandidaat beheerst statistisch ICT-gebruik (...) om grote datasets te interpreteren en te analyseren" [8]. Het materiaal dat hiervoor ontwikkeld is (zie bijvoorbeeld de dataset en opgaven gemaakt door CITO [7]) focust zich op het doorlopen van de statistische cyclus, met behulp van Excel of Vu-Stat. Dit materiaal bevordert het computationele denken echter nauwelijks, er is alleen aandacht voor de data-categorie uit de taxonomie van Weintrop. In dit onderzoek gaan wij expliciet op zoek naar een opgavenset die wel aansluit bij het programma van wiskunde A, maar waarvoor meer nodig is dan het uitrekenen van statistische maten die in één stap te berekenen zijn.

Benakli et al. [5] omschreef naast vier computationele projecten bij calculus ook meerdere projecten voor kansrekening en data-analyse. Deze overstijgen het niveau van de middelbare school, maar zijn mogelijk wel geschikt na aanpassing. Zo luidt één van de opdrachten onder het kopje kansrekening als volgt: Jose en Thomas kunnen met elkaar communiceren als ze minder dan 45 km van elkaar verwijderd zijn. Jose en Thomas zijn beiden naar

dezelfde locatie onderweg, waarbij Jose uit het noorden komt en Thomas uit het westen. Beiden zijn momenteel binnen 50 km van de locatie. Bereken door te simuleren de kans dat Jose en Thomas met elkaar kunnen communiceren. Zie voor een visualisatie figuur 3.3. Natuurlijk kan dit probleem vrij eenvoudig exact worden opgelost (met integraalrekening bijvoorbeeld), maar het kan ook gebruikt worden om het idee achter Monte Carlo simulaties uit te leggen. De kracht van deze methode is dat het ook bruikbaar is wanneer geen exacte antwoorden gevonden kunnen worden (wat mooi aansluit bij *computationele abstractie*). Benakli noemt expliciet dat studenten worden uitgedaagd om wiskundig en computationeel te denken. Enkele praktijken van de taxonomie van Weintrop die gebezigd worden zijn het *visualiseren* van het probleem met een diagram, *herformuleren*, *modellen beoordelen*, *modellen ontwerpen en maken*, *programmeren* (met R) en *informatie communiceren* (het schrijven van een rapport).



Figuur 3.3: Diagram bij het beschreven computationele project van Benakli [5]

Een voorbeeld van een uitgewerkte opdracht die computationeel denken bevordert is de opdracht 'Vier in een rij' ontworpen onder leiding van het Freudenthal instituut voor de Wiskunde A-lympiade [15]. In deze opdracht onderzoeken vwo wiskunde A leerlingen een kassasysteem bij een supermarkt. De opdracht start met een eenvoudig systeem van één kassa en vooraf bekende aankomsttijden van klanten. Gedurende de opdracht neemt de complexiteit toe: er komen meer kassa's, kassa's kunnen sluiten en openen en de kassastroom wordt stochastisch. Het dynamische karakter van het systeem maakt dat er veel is om over na te denken. Om de opdracht tot een goed einde te brengen doen leerlingen onder andere het volgende: ze onderzoeken data, visualiseren data, abstraheren het probleem, vertalen het probleem naar een programma dat met de computer kan worden opgelost, bouwen een model, testen hun model, passen hun model aan op basis van hun bevindingen, proberen verschillende modellen uit en onderzoeken relaties in het systeem.

Naast simulatie (zoals het voorbeeld van Benkali) en modelleren (zoals de opdracht van de Wiskunde A-lympiade) zijn technieken uit het veld *statistisch leren* geschikt voor het bevorderen van computationeel denken [16]. Onder statistisch leren verstaan wij een set

van technieken die het doel heeft om patronen te vinden in de data [18]. Gesuperviseerd statistisch leren gaat over het bouwen van een statistisch model voor het voorspellen, of schatten, van een uitkomst op basis van input.

Gavalda [16] noemt de mogelijkheid om Machine Learning (wij gebruiken liever het bredere begrip statistisch leren) te gebruiken als onderwerp voor lesmateriaal voor leerlingen van de middelbare school. Hij stelt dat statistisch leren een uniek onderdeel van informatica is met de volgende drie eigenschappen:

1. De basistechnieken kunnen uitgelegd worden met wiskunde die op de middelbare school behandeld is.
2. Zonder programmeerervaring kan er toch direct aan de slag gegaan worden en kunnen er resultaten gegenereerd worden die leerlingen kunnen interpreteren.
3. Nuttige, realistische toepassingen en uitdagingen in de wetenschap kunnen worden beschreven. Deze kunnen leerlingen uitnodigen om meer geavanceerde technieken te willen leren.

Dit lijken goede eigenschappen voor het succesvol implementeren van een opdracht over statistisch leren in een vwo 5 wiskunde A klas en sluit goed aan bij de opmerking uit de inleiding waarin gesteld wordt dat één van de doelen van computationeel denken en het gebruik van high-tech tools in de klas is dat leerlingen een realistisch beeld krijgen van welke disciplines er zijn en hoe wiskunde gebruikt wordt in de wetenschap en het bedrijfsleven. Door aan te sluiten bij de voorkennis van leerlingen van bepaalde wiskundige concepten kan bereikt worden dat het begrip hiervan verder verdiept wordt.

Wij kiezen statistisch leren als onderwerp van de werkmiddag.

3.2.2 Statistisch leren

Gavalda [16] noemt een aantal basistechnieken die hij geschikt acht, waaronder:

- Eenvoudige lineaire classificatiemethoden
Een voorbeeld daarvan is de *perceptron*. Dit is een algoritme om een binaire classificatie uit te kunnen voeren aan de hand van een lineaire functie.
- K-nearest-neighbours classificatie
Dit is een methode om een punt te classificeren gebaseerd op de classificatie van andere punten. De k staat voor het aantal burens dat de classificatie van het nieuwe punt bepaalt.
- Naive Bayes
Een naïeve bayes classificatie is een techniek gebaseerd op de regel van Bayes. De term 'naïeve' komt van de naïeve assumptie dat de inputvariabelen onafhankelijk van elkaar zijn.
- Frequent set mining
Deze techniek vindt frequent voorkomende patronen in een dataset.

- Eenvoudige clusteringalgoritmes

Het doel van deze algoritmes is om de dataset in groepen op te delen, zodanig dat objecten in een groep meer overeenkomstig met elkaar zijn dan met objecten uit andere groepen. Het ontbreken van een eenduidige definitie van clustering is een reden waarom er veel verschillende algoritmes zijn, zoals centrum-gebaseerde technieken (zoals K-means-clustering) en agglomeratieve technieken (zoals hierarchische clustering).

In de ideale situatie wordt de methode uitgelegd, waarna leerlingen dit in een context toe kunnen passen. Gavalda stelt dat het een lastig probleem is om met de juiste voorbeelden te komen: enerzijds moet het realistisch zijn om interessant te zijn, anderzijds moet het niet te moeilijk zijn zodat er binnen niet al te lange tijd redelijke resultaten behaald worden. Tot op heden lijkt dergelijk materiaal (op het gebied van statistisch leren) niet te zijn ontworpen.

3.3 K-means algoritme

Het organiseren van data in groepen is één van de meest fundamentele methodes voor begrip en leren. Clusteranalyse is de formele naam voor de studie naar algoritmes en methodes om objecten in groepen op te delen aan de hand van gemeten of ervaren eigenschappen [17]. Het K-means algoritme is het bekendste, populairste en meest simpele data-clustering algoritme [17]. Vanwege de eenvoudigheid van het algoritme, de mogelijkheid die wij zien om realistische, aansprekende voorbeelden te maken, de eigenschappen die Gavalda benoemde en de vele aspecten van het algoritme die aangepast en geanalyseerd kunnen worden [6] - wat veel mogelijkheid geeft tot systeemdenken - denken wij dat dit onderwerp geschikt is voor het bevorderen van computationeel denken voor vwo-leerlingen bij wiskunde A.

Het ontwerpparadigma dat wij gekozen hebben is het K-means algoritme.

Het K-means algoritme is een aanpak om een dataset op te delen in K niet-overlappende clusters [18]. Voordat het algoritme gebruikt wordt moet bepaald worden hoeveel clusters er gewenst zijn; daarna zal het algoritme elke datapunt aan een cluster toewijzen. Het wiskundige idee achter het K-means algoritme is dat een goede clustering zorgt voor een zo klein mogelijke variantie binnen een cluster. De *binnen-clustervariantie* is een maat voor hoeveel de observaties binnen het cluster verschillen. Het doel is dan om de dataset zo op te delen in niet-overlappende clusters zondanig dat de som van alle *binnen-clustervarianties* zo klein mogelijk is. Om de *binnen-clustervariantie* te kunnen berekenen is een afstand-maat nodig: de meest gemaakte keus is de Euclidische afstand. Het optimalisatieprobleem dat is ontstaan (het minimaliseren van de som van de *binnen-clustervariantie*) is erg moeilijk op te lossen, tenminste als er gezocht wordt naar een globale oplossing. Het probleem is NP-hard [21]. Het kan worden aangetoond (zie bijvoorbeeld [18], p.388) dat het K-means algoritme een lokaal optimum geeft door het volgende te doen:

1. Wijs elk punt aan een cluster, $1, \dots, K$, toe.

2. Doe iteratief het volgende (totdat er geen verandering meer is):
 - a) Bereken van elk cluster het clustercentrum (door voor elke coördinaat het gemiddelde te nemen over alle datapunten uit het cluster)
 - b) Wijs elk punt toe aan het cluster bij het dichtbijzijnde clustercentrum (op basis van de Euclidische afstand).

Het lokale optimum kan theoretisch oneindig 'slecht' zijn, maar veel empirisch onderzoek toont aan dat, wanneer de data goed clusterbaar is, een zo-goed-als optimale oplossing gevonden kan worden [21].

Terwijl het algoritme bekend staat om zijn eenvoudigheid en praktische snelheid, is de bovengrens van de berekeningstijd exponentieel in het aantal punten [35]. Het algoritme is gevoelig voor de manier waarop de clusters worden geïnitieerd. Dit probleem wordt gedeeld met andere clusteringalgoritmes die een zogenaamde *hill-climbing strategy* gebruiken (na elke iteratie is de volgende oplossing minstens zo goed). Celebi [6] vergeleek acht initialisatiemethodes en vond dat veelgebruikte initialisatiemethodes vaak slecht presteren, maar dat er goede alternatieven zijn. Hoewel er geen garantie is voor het bereiken van een globaal optimum, is convergentie wel gegarandeerd. Naast het probleem van initialisatie moet het aantal clusters vooraf bepaald worden, terwijl het in veel situaties niet duidelijk is hoeveel clusters er zouden moeten zijn. Dit probleem kan deels worden ondervangen door meerdere initiële waarden voor K te kiezen en dan de 'beste' te gebruiken [24].

3.4 Didactische achtergrond

In deze sectie bespreken we twee didactische principes die ten grondslag liggen aan het ontwerp van onze opdracht: het abstractieproces en de constructivistische leertheorie.

Het meest belangrijke en high-level gedachteproces bij computationeel denken is het abstractieproces [39]. Deze abstractie wordt gebruikt in het herkennen en definiëren van patronen, generalisaties van specifieke gebeurtenissen en parametrisaties. Door abstractie kunnen gemeenschappelijke eigenschappen van objecten gevonden worden, terwijl irrelevante details achterwege worden gelaten. Zo is een algoritme het abstracte proces dat invoer neemt en via een serie aan stappen, een eindresultaat produceert dat aan bepaalde eisen voldoet. Skemp [28] schreef dat abstraheren een activiteit is waarbij we ons bewust worden van overeenkomsten. Tall [30] definieert twee soorten van abstractie waarin objecten mentale eenheden geworden zijn: *operationele abstractie* waarin de acties op objecten centraal staan en *structurele abstractie* waarin de eigenschappen van objecten centraal staan. Deze tweedeling is vergelijkbaar met Sfards [26] opdeling van concepten in operationeel en structureel. Uitdrukkingen die zowel geïntrepreteerd kunnen worden als een proces van evaluatie en een concept waarop zelf operaties uitgevoerd kunnen worden worden procepten genaamd [30]. Deze proces-object dualiteit gaat over het idee dat een wiskundig begrip zowel als proces als object gezien kan worden. Leerlingen leren vaak eerst de proceskant, waarna het mogelijk is om de objectkant te ontwikkelen. Deze ontwikkeling staat bekend als reïficatie [26]. In ons ontwerp willen we hiermee rekening houden door leerlingen eerst aan de slag te laten gaan met de proceskant van het K-means algoritme;

leerlingen voeren het algoritme uit. Nadat zij dit onder de knie hebben zal een hoger abstractieniveau worden aangenomen door te kijken naar de verschillende onderdelen en eigenschappen van het algoritme. Zo zal bijvoorbeeld begrepen moeten worden wat de invloed is van bepaalde input en parameters op de uitkomst. Deze structurele abstractie kan gekoppeld worden aan het systeemdenken uit de taxonomie van Weintrop.

In de constructivistische leertheorie wordt kennis verworven door actieve deelname aan een leeractiviteit, niet slechts door het overbrengen van informatie van de docent naar de leerling [4]. Piaget en Vygotsky zijn de grondleggers van de theorie. Vygotsky's werk (zie [36]) concentreert zich rond 3 concepten: de zone van naaste ontwikkeling, scaffolding en de sociale-culturele natuur van leren. De zone van naaste ontwikkeling is een status waarin een leerling moeite heeft om een probleem op te lossen, maar daartoe wel in staat zou zijn met enige hulp van een docent. Scaffolding gaat over het proces van hulp bieden aan de lerende, waarbij de hoeveelheid hulp afbouwt naarmate de leerling in staat is om de taak meer zelf te volbrengen. Vygotskys werk liet zien dat leren niet een geïsoleerde activiteit is, maar dat dit plaats vindt in een sociaal-culturele context. Technologie staat bekend als een krachtig instrument om op effectieve wijze volgens constructivistische wijze kennis over te dragen, omdat leerlingen in een technologische setting in staat zijn om kennis en vaardigheden toe te passen in een betekenisvolle context [4].

Hoofdstuk 4

Ontwerpeisen

In dit hoofdstuk worden de ontwerpeisen gepresenteerd die volgen uit het theoretisch en conceptueel kader. Ontwerpeisen zijn richtlijnen voor de eisen waaraan het uiteindelijke ontwerp moet voldoen [33]. We delen de ontwerpeisen in twee delen op: organisatorische eisen en kenmerken van de doelgroep enerzijds en (vak)didactische- en inhoudelijke eisen anderzijds.

4.1 Doelgroep en organisatie

Voorwaarde 1: De praktische opdracht is in ongeveer drie uur af te ronden.

Het doel is om een werkmiddag te ontwerpen die in drie uur gemaakt kan worden. Idealiter zijn er weinig leerlingen die sneller klaar zijn en weinig die veel meer tijd nodig hebben, maar feit is dat er vaak grote verschillen tussen leerlingen zijn. Van belang is dat alle leerlingen in elk geval het K-means algoritme zowel met de hand als met behulp van ICT hebben uitgevoerd, want anders is er weinig te analyseren na afloop.

Voorwaarde 2: De praktische opdracht sluit aan bij de belevingswereld van de leerlingen.

Wang en Eccles [37] toonden aan dat hoe meer bij de belevingswereld werd aangesloten, hoe hoger de intrinsieke motivatie en de verwachting over het eigen kunnen. Gavalda [16] noemde het vinden van goede voorbeelden en een juiste context de grootste uitdaging.

Voorwaarde 3: De opdracht sluit aan bij het niveau en de voorkennis van de leerlingen.

Ons ontwerp sluit aan bij het domein Kansrekening en Statistiek uit het vwo wiskunde A programma. We gaan ervan uit dat de leerlingen basisvaardigheden in Excel hebben, maar weinig tot geen ervaring met andere statistische software. Vanwege voorwaarde 1 (de tijdsbeperking) willen we niet te veel tijd besteden aan het leren van een softwareprogramma. We steken de tijd die we hebben liever in het behalen van inhoudelijke doelen die computationeel denken bevorderen. Dit sluit aan bij wat de vernieuwingscommissie cTWO in haar eindrapport in 2008 schreef over het gebruik van ICT bij wiskunde [9]: ICT moet geen doel zijn, maar een middel. Het is 'use to learn' in plaats van 'learn to use'.

4.2 Inhoudelijke leerdoelen

Om ervoor te zorgen dat verschillende aspecten van computationeel denken in ons ontwerp vertegenwoordigd zijn, zijn vier leerdoelen opgesteld. Om de leerdoelen te bereiken moet een leerling zowel operationele abstractie als structurele abstractie ontwikkelen. De leerdoelen zijn zowel input voor het ontwerp als een framework om te testen of de beoogde resultaten zijn behaald aan het einde van de opdracht. De vier leerdoelen zijn als volgt:

1. *De leerling kan omschrijven wat een algoritme is en kan zowel in praktijksituaties als uit de wiskundeles voorbeelden benoemen.*

Om het K-means algoritme te begrijpen en abstractie te ontwikkelen is het van belang dat de leerling allereerst begrijpt wat een algoritme is. In de opdracht zal de leerling aan de slag gaan met zowel een praktisch voorbeeld van een algoritme als meer wiskundige voorbeelden. Opnieuw liggen de voorbeelden zo dicht mogelijk tegen de belevingswereld van de leerlingen.

2. *De leerling kan omschrijven waarvoor het K-means algoritme dient.*

De minimale beheersing is dat de leerling kan uitleggen wat het algoritme bereikt, zowel zonder als met context. De leerling kan daarbij uitleggen wat een cluster is. Wanneer de leerling een hoger operationeel abstractieniveau heeft bereikt kan hij/zij ook in meer algemene termen, zonder context, uitleggen wat het algoritme bereikt. De leerling beheerst dit leerdoel nog beter wanneer hij/zij in staat is om voorbeelden te noemen uit de dagelijkse praktijk waarbij het K-means algoritme gebruikt kan worden (transfer).

3. *De leerling kan het K-means algoritme uitvoeren, zowel handmatig als met behulp van technologische hulpmiddelen.*

De leerlingen kunnen stap-voor-stap het K-means algoritme uitvoeren, wat doelt op het begrijpen van de operationele kant van het object. Leerlingen kunnen aan de hand van een grafische weergave de stappen van het algoritme doorlopen. Daarna kunnen zij een computationeel voorbeeld doorekenen. Tot slot kunnen leerlingen, op een geleide manier, het K-means algoritme in Excel implementeren. Het doel is dat leerlingen begrijpen dat het berekenen met de hand tijdrovend is en niet haalbaar voor grote datasets. Zij zien hier het nut en de kracht van ICT als ondersteuning.

4. *De leerlingen kan kenmerkende aspecten van algoritmiek, zoals initialisatie, convergentie en uniciteit beschrijven en kan deze begrippen toepassen bij het analyseren van het K-means-algoritme.*

Het doel is om structurele abstractie te ontwikkelen door te kijken naar eigenschappen van objecten. Er moet zowel naar het systeem als geheel worden gekeken als naar specifieke delen van het algoritme. Het doel is dat leerlingen inzien dat het K-means algoritme tot een lokaal optimum zal leiden en dat de uitkomst afhankelijk is van de initialisatie (zoals de positie van de clustercentrums en de hoeveelheid clusters).

Hoofdstuk 5

Resultaten

Dit onderzoek heeft een praktische opdracht ter bevordering van computationeel denken opgeleverd. Deze gehele opdracht, zoals deze gebruikt is bij de try-out, is te vinden in Bijlage B. In Bijlage C is een correctievoorschrift te vinden. In dit hoofdstuk zal eerst de praktische opdracht doorgenomen worden en toegelicht worden waarom deze computationeel denken bevordert. Vervolgens wordt een analyse van de try-out gepresenteerd. Tot slot interpreteren we de resultaten.

5.1 Ontwerp

Het doel van ons onderwijsontwerp is het bevorderen en ontwikkelen van computationeel denken bij leerlingen. In deze paragraaf wordt eerst de context waarin deze opdracht is gemaakt toegelicht waarna de verscheidene onderdelen uit het ontwerp gekoppeld worden aan de vier categorieën uit de taxonomie van Weintrop (zie hoofdstuk 3). Door per leeractiviteit deze koppeling te maken trachten we aan te tonen dat de leeractiviteiten kunnen bijdragen aan het doel om computationeel denken te ontwikkelen en te bevorderen. Indien wenselijk geacht door de auteur worden daarnaast de keuzes voor leeractiviteiten en de relevantie ervan verder toegelicht door gebruik te maken van didactische uitgangspunten.

Context

Bij het kiezen van de context speelden er een aantal aspecten: de context moest er één zijn waarmee jongeren zich kunnen indentificeren (om aan te sluiten bij de bevelingswereld), het moest geschikt zijn om het K-means algoritme op toe te passen (oftewel: er moet een betekenisvolle uitkomst zijn) en er moet een dataset beschikbaar zijn waarop het K-means algoritme uitgevoerd kan worden. Een goede context lijkt ons het ontwerpen van **een aanbevelingssysteem voor Netflix**. Netflix heeft enerzijds verschillende datasets beschikbaar gesteld, welke bekend staan als de 'MovieLens user rating datasets'. Er zijn twee datasets beschikbaar voor educatieve doeleinden: de eerste bevat 100.000 ratings en de tweede 27.000.000 ratings. In de dataset zijn de volgende variabelen opgenomen: gebruikersnummer, filmnummer, beoordeling (schaal 0-5) en genre. Naast het bestaan van een geschikte dataset verwachten wij dat het aanbevelen van films een interessante context voor jongeren is waarmee zij zich indentificeren. Tot slot blijkt uit het onderzoek

van software-developer Shea [27] dat de combinatie van een Netflix-dataset en een K-means algoritme betekenisvolle resultaten oplevert. Shea vergelijkt eerst de gemiddelde beoordeling van individuele gebruikers op de genres romantiek en science-fiction met elkaar en probeert deze te clusteren. Hij bediscussieert het kiezen van het aantal clusters (de K uit het K-means algoritme) en maakt het probleem daarna complexer door meer genres toe te voegen (waardoor visualisatie steeds ingewikkelder wordt). Tot slot voert Shea het K-means algoritme uit op de individuele beoordelingen van films. De resulterende clusters bestaan dan uit individuen met een vergelijkbare filmsmaak, welke gevisualiseerd kunnen worden met behulp van een heatmap. Het advies dat een gebruiker gegeven kan worden is dan, uit het cluster waarin de gebruiker zich bevindt, de nog niet bekeken film met het hoogste gemiddelde.

We kijken nu naar de verschillende onderdelen van de ontworpen opdracht (zie bijlage B). De volgende leeractiviteiten passen op onderstaande wijze in de vier categorieën van de taxonomie van Weintrop:

	Data					Modelleren & Simuleren					Computationeel Probleemoplossen						Systeemdenken				
	1	2	3	4	5	1	2	3	4	5	1	2	3	4	5	6	1	2	3	4	5
1	X																				
2				X	X																
3				X																	
4													X		X						
5															X						
6								X								X					
7						X															
8							X														
9												X		X							
10																		X			
11																	X				
12																	X				
13													X								
14				X																X	

Tabel 5.1: Koppeling ontwerp aan taxonomie van Weintrop. Verticaal staan de opdracht-nummers, horizontaal staat de nummering uit de taxonomie van Weintrop uit hoofdstuk 3.1.

Uit de tabel blijkt dat alle vier de categorieën in de opdracht vertegenwoordigd zijn. Hier is bij het ontwerpen ook naar gestreefd. De categorie Data heeft 5 kruisjes, de categorie Modelleren & simuleren heeft er 3, de categorie Probleemoplossen 7 en systeemdenken 5. In de opdracht verschuift chronologisch de aandacht van data, via modelleren & simuleren en computationeel probleemoplossen, naar systeemdenken. Niet elke subcategorie heeft zijn plaats gevonden in de opdracht. De volgende subcategorieën ontbreken: data creëren (1.2), data manipuleren (1.3), modellen maken (2.4), modellen ontwerpen (2.5), herformuleren (3.1), wisselen van schaalniveau (4.3) en complexiteit managen (4.5). Enkele

subcategorieën komen meerdere malen terug: *analyseren van data* (1.4) komt drie maal voor, *computationele abstractie* (3.5) twee maal en *het systeem als geheel onderzoeken* (4.1) ook twee maal. De wenselijkheid hiervan wordt in paragraaf 6.2 besproken.

We lopen nu het ontwerp door:

Introductie (p.3)

Het doel van de introductie is om de interesse van de leerling te wekken door een probleem uit hun belevingswereld te gebruiken als doelstelling van de werkmiddag. Wang en Eccles [37] toonden aan dat hoe meer bij de belevingswereld werd aangesloten, hoe hoger de intrinsieke motivatie en de verwachting over het eigen kunnen. Er is gekozen voor een korte introductie, zodat leerlingen snel aan de slag kunnen.

Opdracht 1: data verzamelen (p.3)

Leerlingen bedenken welke gegevens nodig zijn om een aanbevelingssysteem op te kunnen zetten. Vanuit constructivistisch perspectief is dit beter dan wanneer deze gegevens direct gegeven worden, omdat leren gebeurt door actieve constructie door de lerende en niet door passieve opname [4].

Opdracht 2: data analyseren en visualiseren (p.3)

Leerlingen maken kennis met de data door wat beschrijvende statistiek toe te passen. Het visualiseren van data is in de wiskunde, en in de betawetenschappen in het algemeen, een krachtig instrument voor zowel het analyseren als het delen van data.

(p.4)

Het begrip ‘clusteren’ wordt geïntroduceerd; ook dit wordt in de context van Netflix gedaan. Een probleem dat hierbij kan optreden is dat clusteren verweven raakt met de context dat de leerlingen het daar niet meer los van kunnen zien. Er is dan een probleem van transfer [22]. Ons bewust zijnde van deze beperking verantwoorden wij hier de keus om het begrip toch in de context te behouden door de tijdsbeperkingen en de hogere abstractie van een algemene formule die bij voorkeur vermeden wordt.

Opdracht 3: data analyseren

Leerlingen passen het begrip cluster toe op een voorbeeld, zodat zij direct geactiveerd worden om met het nieuwe begrip aan de slag te gaan. Er zijn verschillende strategieën die ingezet kunnen worden om data te analyseren. Eén daarvan is het definiëren van regels om data te categoriseren. De leerlingen denken na over zowel de betekenis van een cluster in dit specifieke voorbeeld en wat het gevolg is van de keuze van het aantal clusters. Er wordt verwacht dat leerlingen erachter komen dat ze hier graag een computationele methode willen, omdat het betrouwbaarder is dan het vinden van patronen op het blote oog. Dit is een goede opstap voor het computationele vervolg in latere delen.

Opdracht 4: instrumenten beoordelen en kiezen, computationele abstractie

Leerlingen onderzoeken twee afstandsmaten: de Manhattan afstand en de Euclidische afstand. Verwacht wordt dat opgave 4c voor veel (wiskunde A) leerlingen lastig is, omdat ze niet gewend zijn om een ingewikkelde formule in algemene termen te beschrijven. Aangezien absoluutstrepn of gebroken functievoorschriften niet in het programma van wiskunde A zitten, laten we de algemene formule voor de Manhattan-afstand achterwege.

Opdracht 5: computationele abstractie

In deze opdracht zetten we de stap van 2D naar 3D. Het doel is dat leerlingen inzien dat de afstandsmaten ook in een hogere dimensie gebruikt kunnen worden zonder ingewikkelde notatie te introduceren waarbij ze de Euclidische afstandsformule veralgemeniseren naar hogere dimensies.

Opdracht 6: modellen beoordelen en foutopsporing

Het begrip algoritme wordt gekoppeld aan de geleerde stof in de wiskundeles, maar ook aan een praktijkvoorbeeld door gebruik te maken van de ‘Math with bad drawings’ strips. Zo moet duidelijk worden dat algoritmes overal zijn. Leerlingen beoordelen de gegeven algoritmes en zoeken naar fouten.

Opdracht 7: concept begrijpen

Door te experimenteren met het K-means algoritme aan de hand van een online tool krijgen leerlingen zowel een beeld bij de werking van het algoritme als een beter begrip van het concept clustering en clustercentrum. Ook hier gebeurt leren door actieve constructie.

Opdracht 8: oplossingen vinden en testen

In deze opdracht werken leerlingen het numerieke voorbeeld dat in de tekst boven de opgaven geïntroduceerd is verder uit. Ze hebben nu zowel verbaal, grafisch als numeriek de werking van het K-means algoritme gezien. Door het gebruik van al deze verschillende representaties (verbaal, numeriek en grafisch) kan een rijk cognitief schema ontstaan [11].

Opdracht 9: programmeren en oplossingscomponenten ontwerpen

Het geheel uit het niets programmeren van het K-means algoritme leek ons binnen het tijdsbestek van deze opdracht niet mogelijk voor leerlingen zonder programmeerervaring. Om toch wat ervaring op te doen met het implementeren van het algoritme hebben we ervoor gekozen om het k-means algoritme in Excel te programmeren met behulp van ‘Visual Basic for Applications’. Op een aantal plekken zijn de formules weggehaald en de leerling moet deze zelf weer aanvullen. Hiervoor moet de leerlingen goed begrijpen hoe het algoritme precies werkt, bijvoorbeeld dat een punt toegewezen wordt aan het cluster waarvan het clustercentrum het dichtste bij ligt. Om deze opdracht uit te kunnen voeren is basisvaardigheid van Excel nodig, wat de meeste leerlingen in de onderbouw van de middelbare school geleerd (zouden moeten) hebben.

Opdracht 10: relaties begrijpen

In de drie deelopdrachten worden leerlingen uitgedaagd om na te denken over de garantie tot convergentie, de invloed van de initialisatie en factoren die van invloed zijn op de snelheid waarmee een oplossing gevonden wordt. De opdrachten vragen duidelijk een hoger abstractieniveau dan de voorgaande vragen die procesgericht waren. De leerling moet nu de stap zetten naar *structurele abstractie*. Om de opdracht tot een goed einde te kunnen brengen is het van belang dat leerlingen de relaties tussen de verschillende onderdelen van het algoritme begrijpen. De leerling wordt uitgedaagd een onderzoekende houding aan te nemen en wordt daarin gefaciliteerd door een website die het makkelijk maakt verschillende initialisaties en datasets te gebruiken. Sfard geeft aan dat visualisaties vaak behulpzaam zijn voor het verkrijgen van structurele abstractie [26].

Opdracht 11 en 12 : complex systeem als geheel begrijpen

Deze opdrachten gaan verder in op de keuze van de clustercentrums en de invloed daarvan op de uitkomst. Er dient hier gekeken te worden naar de input (clustercentrums en dataset)

en de output (de clustering). Het tussenliggende proces kan als black-box worden gezien. Opnieuw is een online computationele tool beschikbaar om snel analyses te kunnen doen, zodat er vooral tijd is voor het denkwerk.

Opdracht 13: instrumenten beoordelen en kiezen

Het K-means algoritme is niet altijd het geschikte instrument om een clustering uit te voeren. In deze opdracht bekijkt de leerling een voorbeeld waarbij het K-means algoritme geen 'logische' clustering oplevert.

Opdracht 14: data analyseren en informatie communiceren

Tot slot grijpen we terug naar de context waarmee deze opdracht begon: het opzetten van een aanbevelingssysteem voor Netflix. Omwille van de tijd en de moeilijkheidsgraad laten we leerlingen zelf niets ontwerpen, maar gebruiken we de K-means clustering van Rhys Shea op de MovieLens dataset (met honderduizend beoordelingen). We gebruiken enkele van zijn kleurrijke figuren om leerlingen de uitkomsten van een clustering op de zeer grote dataset te laten interpreteren. Zo tonen we hoe je meerdimensionale data goed kunt weergeven en geven we een beeld van wat je met het K-means algoritme kunt bereiken bij een aansprekend praktijkvoorbeeld. We geven tot slot mee dat modellen gebruikt kunnen worden om voorspellingen te doen.

5.2 Try-out

In deze sectie bespreken we de resultaten van de try-out. De analyse maakt gebruik van de voorgestelde instrumenten: de leerlinguitwerkingen van de opdrachten en de 4 afgenomen interviews. De leerlinguitwerkingen zijn aan de hand van het correctievoorschrift nagekeken, waarna van elke opgave de p'-waarde is berekend. De volledige uitwerkingen van de 4 interviews zijn te vinden in Bijlage D. In deze sectie maken we gebruik van tabellen waarin de antwoorden van leerlingen op enkele vragen zijn samengevat. We bespreken nu eerst in hoeverre ons ontwerp aan de ontwerpvoorwaarden voldoet en daarna of de leerdoelen zijn bereikt. Zie voor de ontwerpvoorwaarden en leerdoelen de beschrijving in hoofdstuk 4.

5.2.1 Voorwaarden

Voorwaarde 1: De praktische opdracht is in ongeveer drie uur af te ronden.

Tijdens de uitvoering van de praktische opdracht hebben wij gemerkt dat sommige leerlingen meer tijd nodig hadden voor het uitvoeren van de opdrachten dan de drie uur die beschikbaar was. Van de acht duo's waren er drie die na 2 uur en 45 minuten alle opdrachten hadden afgerond. De andere duo's waren bij het inlevermoment, drie uur na aanvang van de opdracht, nog bezig: vier duo's kwamen tot en met opdracht 9, één duo kwam tot en met opgave 11.

Enkele leerlingen geven in het interview aan dat zij sneller door de opdrachten heen zouden kunnen gaan wanneer er een klassikale instructie zou zijn over de werking van het K-means algoritme, zie tabel 5.2. Minder dan de helft van de leerlingen heeft de opdracht binnen

Leerling 1	We hebben het niet afgekregen omdat we te weinig tijd hadden, maar ook omdat we het gewoon niet helemaal begrepen. Het zou fijn zijn als het stappenplan een keer helemaal klassikaal voorgedaan zou worden.
Leerling 2	Nee, we zijn tot en met opgave 9c gekomen. Het zou tijd schelen als het algoritme stap-voor-stap door de docent voorgedaan zou worden.
Leerling 3	Ja, we waren al eerder klaar.
Leerling 4	Ik heb het niet afgekregen, maar dat lag niet alleen de lengte hoor, ik was ook niet zo gemotiveerd.

Tabel 5.2: Antwoorden op de vraag of de leerling genoeg tijd had om de opgaven te maken

de 3 uur afgerond. Daarmee is in de try-out niet aan de ontwerpeis voldaan.

Voorwaarde 2: De praktische opdracht sluit aan bij de belevingswereld van de leerlingen.

Het is aan de hand van de leerlingenuitwerkingen en interviews lastig te beoordelen of aan deze voorwaarde voldaan is. In de interviews is niet gevraagd naar hoe de leerlingen de context van deze opdracht vonden en hoe de opdracht aansloot bij hun interesse. Uit de observaties blijkt wel dat de meeste leerlingen enthousiast met de opdracht aan de slag gingen, maar het is lastig na te gaan waar dat aan ligt. In de interviews wordt genoemd dat het leuk is dat de opdrachten anders zijn dan in de les. Veel vragen uit de interviews worden beantwoord in de Netflix-context, terwijl hier niet expliciet om gevraagd is. Dit duidt er in elk geval op dat de context ondersteunend is voor het gebruik van het algoritme.

Voorwaarde 3: De opdracht sluit aan bij het niveau en de voorkennis van de leerlingen.

Om te testen of de opdrachten van het juiste niveau zijn, hebben we, zoals aangekondigd in hoofdstuk 2, van elke opgave de p'-waarde berekend. Deze zijn te vinden in tabel 5.3. Het doel van het rapporteren van de p'-waarden is om de moeilijkheid van de opgaven te bepalen en aan de hand daarvan eventuele suggesties te doen om de opdracht te verbeteren of verduidelijken. Om deze reden zijn van opdracht 10 t/m 14 alleen de uitwerkingen van leerlingen meegenomen die de opdracht gemaakt hebben.

Opgave	1	3a	3b	3c	3d	4a	4b	4c	5a	5b	5c	6a	6b	8
p'-waarde	1	0.75	0.66	0.31	0.75	1	0.75	0.44	0.67	0.67	0.50	0.50	0.56	0.29
Opgave	9a	9b	9c	10a*	10b*	10c*	11*	12*	13a*	13b*	14a*	14b*	14c*	14d*
p'-waarde	0.71	0.57	0.43	0	0.25	1	0.75	0.83	1	0	1	1	0	0.67

Tabel 5.3: p'-waarde per opgave bij de try-out (bij de opgaven met * zijn de leerlingen die de opgaven niet hebben gemaakt niet meegenomen)

De opgaven uit het introducerende eerste deel van de opdracht, te weten opgaven 1 tot en met 6, zijn door de leerlingen veelal juist beantwoord. De p'-waarde van deze opgaven is in de meeste gevallen ruim hoger dan 0.5, wat betekent dat de leerlingen gemiddeld ruim meer dan de helft van de punten op deze opgaven behaalden. Opgave 3c en 4c,

waarin gevraagd werd naar een nadeel van een groot aantal clusters en een algemene formule voor de Euclidische afstand, werden relatief lastig gevonden met p' -waarde van respectievelijk 0.31 en 0.44. Opgave 8, waarbij de leerlingen het K-means algoritme uit moesten voeren op een voorbeeld heeft een lage p' -waarde van 0.29. De meeste leerlingen konden enkele stappen van één iteratie uitvoeren, maar gingen aan het iteratieve karakter voorbij en trokken geen onderbouwde conclusie. Bij deze opgave werden daarnaast veel kleine slordigheidsfouten gemaakt. Bij opdracht 9, waarvan de deelopgaven p' -waarden van rond 0.5 hebben, liepen enkele leerlingen vast bij het gebruik van Excel. Geen van de leerlingen maakte gebruik van de tip om met dollar-tekens celverwijzingen vast te zetten. De opgaven uit het laatste deel van de opdracht, te weten opgave 10 tot en met 14, hebben wisselende moeilijkheidsgraad. Zo zijn opgaven 10c, 13a, 14a en 14b door alle leerlingen juist beantwoord, maar wist geen enkele leerling opgave 10a, 13b of 14c correct te beantwoorden. Met name bij opgave 10a, een abstract wiskundige vraag waarin gevraagd wordt naar een argument waarom er altijd convergentie op zal treden, kwam geen enkel antwoord in de buurt. De gegeven antwoorden hebben allen betrekking op het feit dat de punten niet meer verschuiven als er convergentie is opgetreden, maar er worden geen argumenten gegeven waarom dat altijd het geval is. Bij opgave 13b waren de antwoorden op het verkeerde schaalniveau; er werd niet begrepen dat om een onderdeel uit het algoritme werd gevraagd. Bij opgave 14c gaven leerlingen de definitie van de kleuren, maar werd er geen praktische betekenis gegeven aan het verschil tussen beide clusters.

Leerling 1	Het K-means algoritme in de praktijk toepassen en alle stappen netjes opschrijven.
Leerling 2	De theorie was lastig, vooral al die berekeningen en hoe het algoritme in elkaar zat.
Leerling 3	Opgave 9! Het was een beetje gezakt hoe je alles in moest voeren.
Leerling 4	Ik snapte eerst het algoritme bij opgave 7 niet. Bij opgave 9 vond ik de formule lastig.

Tabel 5.4: Antwoorden op de vraag welke onderdelen lastig waren

In de leeringinterviews is gevraagd naar welke onderdelen de leerling lastig vond. De antwoorden op deze vraag staan samengevat in tabel 5.4. De leerlingen geven aan dat ze het lastig vinden om het algoritme stap-voor-stap te doorlopen en alle noodzakelijke berekeningen uit te voeren. Het complexe model als geheel wordt door twee van de vier geïnterviewde leerlingen als lastig beschouwd. Het computationele aspect ('de berekeningen') wordt door drie leerlingen als een moeilijkheid beschouwd, waarbij twee leerlingen dit relateren aan het gebruik van technologie in opgave 9.

Concluderend voldoet de opdracht in de try-out niet aan voorwaarde 1. Het is voor voorwaarde 2 lastig te beoordelen of hieraan voldaan is. Wat betreft voorwaarde 3 worden de eerste opgaven door de leerlingen als eenvoudig ervaren, waarna de moeilijkheidsgraad toeneemt (de p' -waarden vormen een gemengd beeld). Leerlingen kunnen een tekort aan vaardigheden hebben voor het effectief gebruik van Excel (onvoldoende voorkennis).

5.2.2 Leerdoelen

1. *De leerling kan omschrijven wat een algoritme is en kan zowel in praktijksituaties als*

uit de wiskundeles voorbeelden benoemen.

We hebben getest of dit leerdoel behaald is door in de interviews te vragen naar een omschrijving van een algoritme en de leerlingen voorbeelden te laten geven van algoritmes uit de wiskundeles en/of de praktijk. Tabel 5.5 geeft een samenvatting weer van wat leerlingen hebben geantwoord op deze vragen.

Leerling 1	Een stappenplan dat elke keer opnieuw wordt toegepast om een bepaald resultaat te behalen. Voorbeeld: Formules en een afwasmachine.
Leerling 2	Een stappenplan die je iets geeft. Voorbeeld: hypothesetoetsing.
Leerling 3	Een cirkelredenering die je elke keer volgt, eigenlijk een stappenplan. Voorbeeld: stelling van Pythagoras.
Leerling 4	Iets met elke keer dezelfde stappen dat een uitkomst geeft. Voorbeelden: normaalverdeling en opstaan.

Tabel 5.5: Antwoorden op de vraag wat een algoritme is en voorbeelden van een algoritme.

Uit bovenstaande tabel blijkt dat alle leerlingen weten dat een algoritme een stappenplan is. De twee voorbeelden uit de praktijk die leerlingen geven zijn goed, maar het blijkt nog lastig te zijn om goede voorbeelden uit de wiskundeles te benoemen (gezien de onjuiste antwoorden formules en normaalverdeling).

Opgave 6 sluit aan bij dit leerdoel. Uit tabel 5.3 volgt dat er de p'-waarden 0.50 en 0.56 zijn.

2. De leerling kan omschrijven waarvoor het K-means algoritme dient.

Wij hebben getest of dit leerdoel behaald is door in de interviews te vragen of de leerling kan uitleggen wat het doel is van het K-means algoritme. Tabel 5.6 geeft een samenvatting weer van wat leerlingen hebben geantwoord op deze vraag. Uit deze tabel blijkt dat alle vier leerlingen het begrip cluster weten te noemen (waarbij aangetekend moet worden dat niet altijd goed is doorgevraagd of leerlingen kunnen uitleggen wat nu een cluster is). Wat opvalt is dat drie leerlingen in hun antwoord gebruik maken van de context Netflix.

Leerling 1	[Legt vooral uit hoe het werkt, maar na lange uitleg:] Het aanraden van films op Netflix. Algemener maak je clusters.
Leerling 2	De opdrachten gingen over wat de verdeling is van de genres die mensen kijken. Je krijgt clusters: groepen van mensen die heeft erg dicht in de buurt liggen qua interesse.
Leerling 3	Om zo goed mogelijk clusters, van bepaalde voorkeuren, te maken. Je had verschillende punten, van voorkeur van films bijvoorbeeld. Dan wil je zo nauwkeurig mogelijk een cluster vormen bij mensen die dezelfde interesses hebben.
Leerling 4	Het bepalen van het centrum van de clusters en welke punten erbij horen.

Tabel 5.6: Antwoorden op de vraag waarvoor het K-means algoritme dient

3. *De leerling kan het K-means algoritme uitvoeren, zowel handmatig als met behulp van technologische hulpmiddelen.*

Het behalen van het derde leerdoel hebben we getest door in de interviews leerlingen een testvraag te laten maken. In de testvraag lossen leerlingen een probleem op vergelijkbaar aan het probleem uit opdracht 8, er zijn twee clustercentrums gegeven en vier punten. De leerlingen moeten laten zien hoe de eerste iteratie verloopt en volgens aangeven wat het stopcriterium is. De leerling is vrij om een grafische weergave te maken, maar er wordt vereist dat er berekeningen gemaakt worden (gezien het computationele karakter van deze opdracht). Zo moeten afstanden gekwantificeerd worden. Alle vier de leerlingen waren in staat om uit te leggen hoe het K-means-algoritme gebruikt kon worden om het probleem op te lossen. Twee leerlingen liepen wel vast op het computationele gedeelte: één leerling kon de afstand tussen twee punten niet goed berekenen en één leerling wist niet hoe de nieuwe positie van het clustercentrum berekend moest worden. Deze twee leerlingen startten met maken van een schets en probeerden van daaruit het probleem op te lossen. Hun computationele kennis schoot tekort. De twee andere leerlingen lossen het probleem in zijn geheel correct op door direct te rekenen. Wat opvalt is dat zij later, wanneer ze toch wat vastlopen, opmerken dat een schets behulpzaam is. Zij hebben een rijker schema opgebouwd en kunnen schakelen tussen verschillende representaties.

In tabel 5.7 staan de bevindingen van de testvraag samengevat.

Leerling 1	- Begint direct met rekenen, maar komt er zelf achter dat een schets behulpzaam is. - Vergeet eerst de stelling van Pythagoras toe te passen, maar komt daar zelf op terug. - Juist gebruik van begrippen als cluster en clustercentrum
Leerling 2	- Weet eerst niet goed hoe te beginnen, na de tip om een schets te maken lukt het wel - Kan afstanden berekenen en punten correct aan clusters toewijzen - Weet niet meer hoe je de nieuwe positie van het clustercentrum kunt berekenen.
Leerling 3	- Begint direct met rekenen, maar merkt op dat een schets maken misschien toch handig was - Rekent alles correct uit.
Leerling 4	- Maakt direct een schets - Weet niet meer hoe je afstanden tussen punten uit kunt rekenen. - Zet de juiste iteratieve stappen, maar maakt fouten in de berekeningen

Tabel 5.7: Interessante aspecten uit interviews bij het maken van de testvraag

Opgave 8 testte de vaardigheid om het K-means algoritme handmatig uit te voeren, terwijl in opgave 9 gebruik gemaakt moest worden van Excel. Op opgave 8 is laag gescoord met een p'-waarde van 0.29. Bij opgave 9, bestaand uit drie deelopgaven, zijn de p'-waarden achtereenvolgens 0.71, 0.57 en 0.43. De leerlingen scoren beter op de opgave waarbij gebruik gemaakt werd van technologische hulpmiddelen.

4. *De leerling kan kenmerkende aspecten van algoritmiek, zoals initialisatie, convergentie en uniciteit beschrijven en kan deze begrippen toepassen bij het analyseren van het K-means algoritme.*

In de interviews zijn twee vragen gesteld die aansluiten bij dit leerdoel. De eerste vraag is om uit te leggen wat de oorzaak kan zijn van het feit dat het algoritme de ene keer heel snel tot een oplossing komt en het in andere gevallen veel langer kan duren. In tabel 5.8 staan de antwoorden van leerlingen op deze vraag samengevat. Alle vier de leerlingen hebben begrepen dat de spreiding van de punten van belang is voor de convergentiesnelheid, de formulering hoe dat precies werkt is niet altijd precies. De initialisatie van de clustercentrums wordt niet genoemd, één leerling noemt wel het aantal clustercentrums.

Leerling 1	Als de punten dichtbij één punt liggen en de afstand tot de andere punten groot is, dan zal het snel opgelost zijn.
Leerling 2	Het kan aan het aantal clustercentrums liggen en aan de spreiding van de punten.
Leerling 3	Als de punten veel meer verspreid zijn zal het langer duren. Dan gaan de clustercentrums veel meer bewegen.
Leerling 4	Als de punten dicht bij elkaar liggen, dan gaat het de hele tijd heen en weer. Als de punten ver uit elkaar liggen en de clustercentrums zijn eenmaal verdeeld, dan gaan de punten niet meer naar andere clusters.

Tabel 5.8: Antwoorden op de vraag waarom er soms veel iteraties nodig zijn en waarom soms weinig

De tweede vraag uit de interviews die bij dit leerdoel past is de vraag waarom het algoritme, bij dezelfde dataset, niet altijd tot dezelfde oplossing komt. Leerlingen hebben dit ervaren bij vraag 11 tot en met 13. De antwoorden van de leerlingen op deze vraag zijn in tabel 5.9 samengevat. Drie van de vier leerlingen geven in hun antwoord aan dat de oplossing afhangt van de initialisatie van de clustercentrums (wat juist is), waarbij één leerling expliciet de onjuiste conclusie trekt dat een andere initialisatie altijd een andere oplossing geeft.

Leerling 1	Als ik de clustercentrums iets verplaats, dan komen de punten in een ander cluster. Dat hoeft op termijn niet meer goed te komen, waardoor je een andere oplossing krijgt.
Leerling 2	Ik zou het niet weten.
Leerling 3	Het hangt er vanaf waar je de clustercentrums plaatst. Als je alle drie de clusters bij elkaar laat starten, dan is het logisch dat ze niet allemaal verspreiden.
Leerling 4	Als het clustercentrum op een andere plek begint krijg je altijd een andere oplossing.

Tabel 5.9: Antwoorden op de vraag waarom het algoritme niet altijd tot dezelfde oplossing komt

Opgave 10c, 11 en 12 hebben betrekking op de initialisatie van het K-means algoritme en hebben hoge p'-waarden. Op opgaven 10a en 10b over de convergentie van het K-means algoritme en de uniciteit van de oplossing is slecht gescoord; de p'-waarden zijn respectievelijk 0 en 0.25. Dit zijn wat ingewikkeldere wiskundige concepten die in dit tijdsbestek nog niet begrepen zijn.

5.3 Interpretatie

In het ontwerp hebben we opdrachten ontworpen passende bij alle vier de categorieën uit de taxonomie van Weintrop. In de resultatensectie blijkt (zie tabel 5.1) dat de opdrachten niet helemaal evenwichtig verdeeld zijn over de verschillende categorieën; de categorie computationeel denken lijkt veel meer vertegenwoordigd dan de categorie modelleren & simuleren. Is dat erg? Wij denken van niet, ten eerste aangezien een kruisje in de tabel niet vertelt hoe veelomvattend de opdracht is. Ten tweede hoeft het geen doel te zijn om alle categorieën in exact dezelfde mate terug te laten komen. Het doel van ons ontwerp is het bevorderen van computationeel denken. Het is daarbij wenselijk om verschillende aspecten, die worden beschreven in de vier categorieën, in de opdrachten voorbij te laten komen, om zo de verscheidende vaardigheden aan te leren. In de opzet van ons ontwerp is er gekozen voor een al bestaand model (het K-means algoritme), waardoor *model maken* en *model ontwerpen* niet aan de orde is gekomen. Deze keuze is gemaakt vanwege 1) de beperkte tijdsduur van 3 uur voor de try-out van de opdracht, 2) het ontbreken van eerdere ervaring met computationeel denken door leerlingen en 3) de beperkte ICT-vaardigheid van leerlingen.

Wat aan de hand van tabel 5.1 opvalt is dat de eerste opdrachten vooral in de categorie *data* vallen, waarna opdrachten volgen die te classificeren zijn in de categorie *modelleren & simuleren* en *computationeel denken*. De laatste opgaven (10 t/m 14) vallen grotendeels in de categorie *systeemdenken*. De opbouw van deze opdracht kan gekoppeld worden aan het abstractieproces dat beschreven is in paragraaf 3.4. De opdracht start met het verkennen van de data, wat nog vooral beschrijvend is en matig denktensief. Wanneer we kijken naar het concept 'K-means algoritme' wordt eerst gewerkt aan het algoritme als een *proces van evaluatie*. De leerlingen verkrijgen *operationele abstractie* door acties uit te voeren. Dit gebeurt grotendeels in de opdrachten die vallen onder de categorie *modelleren & simuleren*, in onze opdracht zijn dat opgaven 7 t/m 9. In het laatste deel van de opdracht staat het ontwikkelen van *structurele abstractie* centraal, waarbij leerlingen de eigenschappen van het object verkennen. De ontwikkeling waarbij eerst het concept als proces wordt bekeken en pas daarna als object noemen we reïficatie [26].

In de try-out bleek dat meer dan de helft van de leerlingen de opdracht niet binnen de tijd af heeft gekregen. Dit kan enerzijds komen doordat de opdracht daadwerkelijk te lang is, maar het kan ook komen doordat leerlingen vast zijn gelopen op een aantal opgaven of door een gebrek aan motivatie. Uit de observatie van de try-out blijkt dat leerlingen veel tijd hebben besteed aan opdracht 8 en 9. Hoewel leerlingen over het algemeen hard gewerkt hebben, kan het feit dat de opdracht niet meetelde voor het schoolexamen een rol hebben gespeeld in de snelheid waarmee leerlingen de opdrachten hebben gemaakt. Eén leerling geeft ook in het interview aan dat de motivatie niet zo hoog was.

Hoofdstuk 6

Conclusie en discussie

In dit hoofdstuk trekken we eerst onze conclusies. Daarna bespreken we limitaties van dit onderzoek en geven we aanbevelingen voor toekomstig onderzoek en verbeteringen in het ontwerp.

6.1 Conclusie

In dit onderzoek hebben we de volgende hoofdvraag beantwoord: Hoe kan een onderwijs-ontwerp vwo-leerlingen wiskunde A ondersteunen bij het ontwikkelen van computationele denkvaardigheden? Deze onderzoeksvraag is beantwoord door het ontwikkelen van een praktische opdracht met een lengte van circa 3 uur waarin verschillende aspecten van computationeel denken aan bod komen.

Een veelgebruikte definitie van wat computationeel denken is, is de definitie van Wing [39]: ‘computationeel denken bevat probleemoplossen, het ontwerpen van systemen en het begrijpen van menselijk gedrag, door gebruik te maken van fundamentele concepten uit de informatica’. Op basis van deze definitie en de in de literatuur genoemde breed geformuleerde vaardigheden en aspecten die horen bij computationeel denken, is het moeilijk te bepalen in hoeverre een opdracht met computationeel denken te maken heeft. Een werkbaar instrument lijkt de taxonomie van Weintrop et al [38], welke vier categorieën voor computationeel denken definieert: data, modelleren & simuleren, computationeel probleemoplossen en systeemdenken. Elke categorie bestaat uit meerdere praktijken, met bijbehorende kennis, concepten en vaardigheden.

Uit de literatuurstudie blijkt dat verschillende onderwerpen geschikt zijn voor een opdracht om computationeel denken te bevorden, zoals simulatie, modelleren en technieken van statistisch leren. We kiezen voor statistisch leren, omdat de basistechnieken uitgelegd kunnen worden met wiskunde die op de middelbare school behandeld is, er zonder programmeerervaring direct aan de slag gegaan kan worden en er resultaten gegenereerd kunnen worden die leerlingen kunnen interpreteren. Tot slot is het mogelijk nuttige, realistische toepassingen en uitdagingen in de wetenschap te beschrijven. Door aan te sluiten bij de voorkennis van leerlingen van bepaalde wiskundige concepten kan bereikt worden

dat het begrip hiervan verder verdiept wordt.

Het onderwerp van deze praktische opdracht is het K-means algoritme, welke toegepast wordt bij het ontwerpen van een aanbevelingssysteem voor Netflix. Het K-means algoritme is het bekendste, populairste en meest simpele data-clustering algoritme [29]. Vanwege de eenvoudigheid van het algoritme, de mogelijkheid die wij zien om realistische, aansprekende voorbeelden te maken en de vele aspecten van het algoritme die aangepast en geanalyseerd kunnen worden [6] - wat veel mogelijkheid geeft tot systeemdenken - denken wij dat dit onderwerp geschikt is voor het bevorderen van computationeel denken bij vwo-leerlingen die wiskunde A volgen. In het ontwerp is rekening gehouden met twee didactische principes die passen bij het ontwikkelen van computationeel denken: de ontwikkeling van operationele abstractie (procesgericht) naar structurele abstractie (objectgericht) en de constructivistische leertheorie. Er is laten zien dat het K-means algoritme zich goed leent om een breed scala aan aspecten van computationeel denken aan bod te laten komen. Dit is gedaan door de opgaven te koppelen aan de vier categoriën uit de classificatie van Weintrop voor computationeel denken, welke alle vier aan bod komen.

De praktische opdracht is uitgetoetst door 16 vwo 5 wiskunde A leerlingen op een middelbare school in Zeist. Uit de analyse van de leerlinguitwerkingen en de vier interviews blijkt dat de opdracht door ruim de helft van de leerlingen niet binnen de tijd kon worden afgerond. De p'-waarden, die de moeilijkheid van de opgaven aangeven, laten een gemengd beeld zien. Uit de antwoorden van de 4 interviews blijkt dat de leerlingen een beeld hebben van wat een algoritme is en het achterliggende idee van het K-means algoritme, clustering, begrijpen. De leerlingen betrekken veelal de context van de opdracht in hun antwoorden: het opzetten van een adviesstelsel voor Netflix. Bij toepassen van het K-means algoritme op een casus, behalen de leerlingen een beter resultaat wanneer gebruik gemaakt kan worden van Excel dan wanneer alle stappen en berekeningen handmatig gedaan moeten worden. Tot slot weten leerlingen vragen over initialisatie juist te beantwoorden, maar worden vragen omtrent de begrippen convergentie en uniciteit nog niet correct beantwoord.

6.2 Discussie

6.2.1 Limitaties

We bespreken in deze paragraaf de limitaties van de procedure, de respondenten en de instrumenten.

In dit onderzoek is een werkmiddag gemaakt in het kader van computationeel denken. Het opzetten van de studie en het ontwerpen van het materiaal, inclusief meerdere malen consultatie van een expert en aanpassingen van het ontwerp, kost veel tijd. Het ontwerp had vooraf nog voorgelegd kunnen worden aan een medestudent of belanghebbende om het ontwerp aan te scherpen. Door de beperkte tijd is het niet mogelijk geweest het materiaal na de try-out verder te herontwerpen en nogmaals te testen op grotere schaal. Dit paste ook niet in de reikwijdte van dit onderzoek. Desalniettemin zou het de kwaliteit van het ontwerp wel ten goede komen om het ontwerpproces nog verder te doorlopen.

Tevens is er voor één geschikt onderwerp (om computationeel denken te bevorderen voor

vwo wiskunde A) een ontwerp gemaakt. Dit limiteert de mogelijkheid om het ontwerp en de resultaten van de try-out te vergelijken met andere ontwerpen ter bevordering van computationeel denken. Daarnaast ontbreekt een objectieve manier om te meten in hoeverre leerlingen computationeel denken hebben ontwikkeld. Dit had gedaan kunnen worden door een meting te doen voorafgaand aan de opdracht en dit te vergelijken met een meting achteraf. Hiervoor moet dan een geschikt instrument worden ontworpen.

Het aantal respondenten, te weten 16 leerlingen, is klein. Voor een try-out kan het aardige inzichten geven (om een eerste indruk te krijgen van hoe de opdrachten ervaren worden en waar mogelijkheden zijn tot verbetering), maar voor het systematisch testen van het materiaal (en om conclusies te kunnen trekken over het bevorderen van computationeel denken) is een grotere groep respondenten nodig. Een andere beperking van dit onderzoek is dat deze opdracht is uitgevoerd in de middag nadat de leerlingen eerder al andere lesactiviteiten hebben uitgevoerd, hierdoor zou enige 'frisheid' kunnen ontbreken. Wellicht een groter probleem voor de motivatie was het ontbreken van een beoordeling in de vorm van een cijfer dat meetelt voor het examendossier. Dit zou de (extrinsieke) motivatie kunnen verhogen.

De interviews zijn niet direct na afloop van de try-out uitgevoerd, maar binnen twee dagen. Hierdoor is er het gevaar dat enige kennis weg is gezakt. Een voordeel hiervan is echter dat bekeken kan worden of de kennis en vaardigheden iets langer zijn blijven beklijven dan direct na afloop van de opdracht. Het afnemen van interviews ter beoordeling van het behalen van leerdoelen kan subjectief gevonden worden. De kwaliteit van het interview is daarnaast afhankelijk van de ervaring en kunde van de interviewer.

Tot slot zijn de leerlingantwoorden geanalyseerd aan de hand van de p'-waarde. Een beperking van deze p'-waarde is dat het alleen aangeeft welke proportie van het totaal aantal punten op een vraag is behaald en geen antwoord geeft op de vraag waardoor deze waarde tot stand is gekomen. Er heeft er geen diepte-analyse plaatsgevonden van de individuele antwoorden. Uit deze antwoorden zou nuttige informatie over de opdracht gehaald kunnen worden, wat wellicht tot nieuwe (of sterkere) conclusies zou kunnen leiden.

6.2.2 Implicaties en aanbevelingen

Docenten kunnen het ontworpen materiaal gebruiken als praktische opdracht, bijvoorbeeld als onderdeel van de vrije keuzeruimte of als lesaanvulling voor leerlingen die sneller dan gebruikelijk door het reguliere programma gaan. Deze opdracht lijkt ook geschikt om te gebruiken voor wiskunde B, gezien het beperkte beroep op kennis van kansrekening en statistiek (wat niet in het programma voor wiskunde B zit). Het materiaal leent zich gezien de aard ook voor het gebruik bij de schoolvakken informatica en Natuur, Leven en Technologie (NLT). Bij NLT zou het bijvoorbeeld goed aansluiten bij het doel om 'wiskunde als gereedschap' te gebruiken. Gezien de complexiteit en de hoge mate van abstractie raden we niet aan om de opdracht te gebruiken in de onderbouw.

Wij hopen dat dit onderzoek laat zien dat computationeel denken belangrijk is, actueel en leuk is en dat er vele mogelijkheden zijn om lesmateriaal te maken die computationeel denken bevorderen. We hopen dat SLO, onderzoekers, docenten en andere belanghebben-

den dit onderzoek kunnen gebruiken voor het verder ontwerpen en testen van materiaal omtrent computationeel denken. De taxonomie van Weintrop kan, zoals in dit onderzoek is laten zien, een waardevol instrument zijn bij het ontwerpen van een opdracht omtrent computationeel denken. We hopen dat computationeel denken een volwaardige plaats zal krijgen in het (nieuwe) curriculum op de middelbare school. Of dat nu bij wiskunde is of bij één van de andere vakken.

Enkele aanbevelingen voor vervolgonderzoek in het kader van computationeel denken zijn:

- Test, na doorontwikkeling van de opdracht aan de hand van de aanbevelingen hieronder, het ontworpen lesmateriaal op grotere schaal uit.
- Ontwikkel een instrument waarbij valide gemeten kan worden in hoeverre een activiteit computationeel denken bevordert (zodat beter gemeten kan worden of leerdoelen behaald worden).
- Ontwikkel met gebruik van de taxonomie van Weintrop andere opdrachten in het kader van computationeel denken. Enkele suggesties voor onderwerpen zijn Monte Carlo simulaties, modelleren van praktijkproblemen en (andere) technieken van statistisch leren.
- Onderzoek hoe computationeel denken in een nieuw curriculum zijn plaats kan vinden en wat de rol van opdrachten en onderwerpen, zoals die vergelijkbaar aan dit onderzoek, daarin is.

We sluiten dit onderzoek af met enkele aanbevelingen voor het uitvoeren en doorontwikkelen van deze opdracht:

- Maak een bewuste afweging over het wel of niet geven van klassikale uitleg. Uit de resultatensectie blijkt dat enkele leerlingen aangeven dat zij een klassikale toelichting over de werking van het K-means algoritme wenselijk vinden. Deze hulp van de docent kan gezien worden als scaffolding, waarbij de hulp van een 'more knowledgeable' de leerling vooruit kan helpen waarneer de gestelde vragen en geboden hulp op het juiste niveau zijn. Zie hiervoor de theorie van Vygostky uit het Theoretisch kader (hoofdstuk 2.4). Wanneer hiervoor gekozen wordt denken wij dat leerlingen wellicht sneller in staat zijn om het proces te begrijpen, echter de opdracht verliest dan wel iets van zijn onderzoekende karakter. Een voordeel van het klassikaal bespreken kan zijn dat leerlingen minder tijd hoeven te besteden aan de opdrachten over het algoritme als proces (de categorie *modelleren en simuleren*), waardoor er meer tijd over blijft voor de opdrachten die structurele abstractie ontwikkelen. Dit kan ervoor zorgen dat leerlingen de opdracht wel afkrijgen binnen de tijd, wat nu vaak niet het geval was.
- Onderzoek of het wenselijk is om deze opdracht in zijn geheel achter elkaar uit te voeren of te verdelen over meerdere lessen. Enige tijd om de stof te laten bezinken zou een beter resultaat op kunnen leveren.

- Maak een werkblad bij opgave 8 (als bijlage) dat leerlingen in kunnen vullen. Door dit werkblad is het duidelijker dat een computationeel antwoord van de leerlingen verwacht wordt. Maak expliciet duidelijk dat de leerlingen het hele proces moeten doorlopen totdat convergentie op zal treden.
- Geef bij de opdrachten over convergentie en uniciteit, zoals opgaven 10a en 10b, aan wat voor soort antwoord er van de leerling verwacht wordt of maak een voorbeeld dat leerlingen op weg helpt in de juiste richting te denken. Leerlingen zijn, bij wiskunde A in elk geval, niet gewend om dit soort abstracte vragen te beantwoorden waarbij abstract wiskundig redeneren vereist is. Uit de try-out blijkt dat leerlingen deze vragen lastig vinden om te beantwoorden. Zo had opdracht 10a een p'-waarde van 0; geen van de leerlingen gaf het juiste antwoord.
- Als er meer tijd beschikbaar is om deze opdracht uit te voeren, of de leerlingen al enige programmeerervaring hebben, dan zou het zelf laten programmeren van het K-means algoritme een mooie uitbreiding zijn. Bij leerlingen die het vak informatie volgen is wellicht al genoeg voorkennis aanwezig om het algoritme te laten programmeren.
- De leerlingen vonden de eerste opdrachten goed te doen, maar hadden veelal moeite met de eerste opgaven over het K-means algoritme. De opdracht kan verbeterd worden door deze sprong in moeilijkheidsgraad te verkleinen en een aantal instapopdrachten over het K-means algoritme toe te voegen. Elke van deze opgaven zou bijvoorbeeld in kunnen zoomen op één stap uit het algoritme. Een mogelijkheid is om hierin te differentiëren en deze opgaven alleen te laten maken door leerlingen die niet direct opgaven 8 kunnen maken, maar tussenstappen nodig hebben.
- Verbeter, vervang of verwijder opdracht 2. Het doel van deze opdracht is het visualiseren van data. Deze opdracht is niet getest in de try-out, omdat er verwacht werd dat dit te veel tijd zou kosten. De opdracht moet zodanig zijn dat de voorkennis en ICT-vaardigheden van leerlingen aansluiten bij wat er gevraagd wordt.

Bibliografie

- [1] Augustine, M. (2007). *Rising Above the Gathering Storm*. National Academies Press, Washington, D.C.
- [2] Bailey, D. and Borwein, J. (2011). High-precision numerical integration: Progress and challenges. *Journal of Symbolic Computation*, 46(7):741–754.
- [3] Barendsen, E., Grgurina, N., and Tolboom, J. (2016). A new informatics curriculum for secondary education in The Netherlands. *Lecture Notes in Computer Science (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)*, 9973:105–117.
- [4] Becker, K. (2002). Constructivism and the Use of Technology. *Technology Teacher*.
- [5] Benakli, N., Kostadinov, B., Satyanarayana, A., and Singh, S. (2017). Introducing computational thinking through hands-on projects using R with applications to calculus, probability and data analysis. *International Journal of Mathematical Education in Science and Technology*, 48(3):393–427.
- [6] Celebi, M. E., Kingravi, H. A., and Vela, P. A. (2013). A comparative study of efficient initialization methods for the k-means clustering algorithm. *Expert Systems with Applications*, 40(1):200–210.
- [7] CITO (2019). Datasets en bijbehorende opgaven. Retrieved via <https://www.uu.nl/files/datasets-en-bijbehorende-opgaven-cito-nwd-2019pdf>.
- [8] College voor Examens (2017). VWO WISKUNDE A: Syllabus Centraal Examen 2019. Technical report.
- [9] CTWO (2008). Use to learn: Naar een zinvolle integratie van ICT in het wiskundeonderwijs. Technical report.
- [10] CTWO (2015). Wiskunde op havo en vwo per 2015. Technical report.
- [11] Drijvers, P., van Streun, A., and Zwaneveld, G. (2015). *Handboek vakdidactiek wiskunde*.
- [12] Eisenberg, M., Eisenberg, M., Eisenberg, A., Eisenberg, A., Gross, M., Gross, M., Kaowthumrong, K., Kaowthumrong, K., Lee, N., Lee, N., Lovett, W., and Lovett, W. (2002). Computationally-Enhanced Construction Kits for Children: Prototype and

- Principles. *The fifth International Conference of the Learning Sciences (ICLS)*, pages 79–85.
- [13] Eshuis, P., Houtenbos, R., and Boertjens, J. (2016). Digitalisering van wiskunde in het voortgezet onderwijs. *Nieuw Archief voor Wiskunde*, 17(1):10–14.
 - [14] Freudenthal Instituut (1983). Kansverdelingen (Hewet-reeks). Retrieved from <http://www.fisme.science.uu.nl/toepassingen/00554/>.
 - [15] Freudenthal Instituut (2012). Wiskunde A-lympiade (voorrunde): Vier in een Rij. Retrieved from <http://www.fisme.science.uu.nl/toepassingen/28161/>.
 - [16] Gavaldà, R. (2008). Machine Learning in Secondary Education? In *Proceedings TML*.
 - [17] Jain, A. K. (2010). Data clustering: 50 years beyond K-means. *Pattern Recognition Letters*, 31(8):651–666.
 - [18] James, G., Witten, D., Hastie, T., and Tibshirani, R. (2013). *Springer Texts in Statistics An Introduction to Statistical Learning - with Applications in R*.
 - [19] Lu, J. and Fletcher, H. (2009). Thinking About Computational Thinking. *ACM SIGCSE Bulletin*, 41(1):260–264.
 - [20] Margaret C. Harrell; Melissa A. Bradley (2009). *Data Collection Methods Semi-Structured Interviews and Focus Groups*. Rand National Defense Research Institute.
 - [21] Meilă, M. (2006). The uniqueness of a good optimum for K-means. In *Proceedings of the 23rd international conference on Machine learning*, pages 625–632.
 - [22] Nelissen, J. (2012). Context, motivatie en betekenis. *Panama-Post - Reken-wiskundeonderwijs: onderzoek, ontwikkeling, praktijk*, 31(3):16–22.
 - [23] NRO (2018). Computatieel denken en wiskundig denken: digitale geletterdheid in wiskundecurricula. Retrieved from <https://www.nro.nl/onderzoeksprojecten-vinden/?projectid=40-5-18540-130-computatieel-denken-en-wiskundig-denken-digitale-geletterdheid-in-wiskundecurricula>.
 - [24] Peña, J., Lozano, J., and Larrañaga, P. (1999). An empirical comparison of four initialization methods for the K-Means algorithm. *Pattern Recognition Letters*, 20(10):1027–1040.
 - [25] Polya, G. (2006). On Learning, Teaching, and Learning Teaching. *The American Mathematical Monthly*, 70(6):605.
 - [26] Sfard, A. (1991). On the dual nature of mathematical conceptions: Reflections on processes and objects as different sides of the same coin. *Educational studies in mathematics*, 22(1):1–36.
 - [27] Shea, R. (2018). K-means clustering of movie ratings. Retrieved from https://programming.rhysshea.com/K-means_movie_ratings/.

- [28] Skemp, R. R. (1976). Relational Understanding and Instrumental Understanding. *Mathematics Teaching*, 77(1):20–26.
- [29] Steinhaus, H. (1956). Sur la division des corps materiels en parties. *Bulletin de l'academie*, 4(12).
- [30] Tall, D. (2013). *How humans learn to think mathematically*. Cambridge University Press.
- [31] Timmer, M. and Tolboom, J. (2019). Computational thinking-een vooruitblik op het wiskundeonderwijs van de toekomst. *Nieuw Archief voor Wiskunde*, 20(1):42–45.
- [32] van den Berg, E. and Kouwenhoven, W. (2008). Ontwerponderzoek in vogelvlucht. *Tijdschrift voor lerarenopleiders*, 29(4):20–26.
- [33] van der Donk, C. and van Lanen, B. (2009). *Praktijkonderzoek in de school*.
- [34] van der Meulen, J. and Timmer, M. (2018). Computational Thinking in vwo 5; Een werkmiddag over 'branch en bound'. *Euclides*, 94(3):14–17.
- [35] Vattani, A. (2011). k-means Requires Exponentially Many Iterations Even in the Plane. *Discrete & Computational Geometry*, 45(4):596–616.
- [36] Vygotsky, L. (1978). *Mind in society: The development of higher psychological processes*. Cambridge. Harvard University Press.
- [37] Wang, M.-T. and Eccles, J. S. (2013). School context, achievement motivation, and academic engagement: A longitudinal study of school engagement using a multidimensional perspective. *Learning and Instruction*, 28:12–23.
- [38] Weintrop, D., Beheshti, E., Horn, M., Orton, K., Jona, K., Trouille, L., and Wilensky, U. (2016). Defining Computational Thinking for Mathematics and Science Classrooms. *Journal of Science Education and Technology*, 25(1):127–147.
- [39] Wing, J. (2006). Computational Thinking. *Communication of the ACM*, 49(3):33–35.

Bijlagen

Bijlage A: Interviewleidraad

Doelen:

- Mening van de leerlingen over de lengte, moeilijkheid en motivatie voor de opdracht.
- Testen of de leerdoelen zijn behaald.

Inleiding

- Hoe vond je het om aan deze opdracht te werken?
- Had je genoeg tijd om alle opdrachten te maken?
- Waren er opdrachten waarbij je niet begreep wat er gevraagd werd? Welke?
- Wat was de moeilijkste opdracht?

Kern

- Wat is een algoritme? Kun je een voorbeeld geven van een algoritme?
- Waarvoor kun je het K-means algoritme gebruiken? Wat doet het K-means algoritme precies?
- Je zou ook kunnen clusteren door op het oog punten in clusters in te delen. Kun je daar iets tegenin brengen? Waarom zou Netflix dit willen doen? Hint: drie-dimensionaal.
- Hoe weet je nu hoeveel clusters er zijn?
- Het K-means algoritme komt niet altijd tot dezelfde oplossing. Waar kan dat aan liggen?
- Je zag dat het k-means-algoritme soms snel tot een oplossing komt en dat het soms lang duurt. Kun je uitleggen waardoor dat is?

- Testvraag: Gegeven zijn de volgende punten (1,4), (2,4), (4,2) en (5,1). Neem de clustercentrums (1,4) en (2,4). Voer het K-means algoritme uit.

Hint: schets maken

Slot

- Aanvullende opmerkingen leerling.

Bijlage B: Opdracht

Inleiding

Welkom op deze werkmiddag. Je gaat vandaag van 12:40 uur tot 16:00 uur aan de slag met het k-means algoritme. Dit onderwerp ligt op het raakvlak van wiskunde en informatica. Je gaat het k-means algoritme toepassen op een aanbevelingssysteem voor Netflix-films. Wat dit algoritme precies inhoudt en hoe je dit gaat toepassen, zal in de loop van de opdracht duidelijk worden.

De opdracht maak je in tweetallen. Het is niet toegestaan om met andere groepjes over de opdracht te overleggen. Als jullie ergens echt niet uitkomen, mag je uiteraard wel de docent om extra uitleg vragen.

Jullie eindproduct is een getypt verslag waarin de opgaven uit dit boekje zijn uitgewerkt. Het verslag leveren jullie vandaag tussen 15:45 en 16:00 in bij de docent.

Deze opdracht bestaat uit 7 hoofdstukken. Hoofdstuk 1 t/m 4 vormen het voorwerk. Besteed hier niet meer dan 45 minuten aan. In hoofdstuk 5 leer je het k-means algoritme. Dit is de kern van deze werkmiddag. Besteed hier voldoende tijd aan, zo'n 60 tot 90 minuten. Mochten jullie niet uit een bepaalde opgave komen uit hoofdstuk 1 t/m 4, maak je dan niet druk. Kom je niet uit een vraag in hoofdstuk 5, vraag dan om hulp. Probeer wel, zolang de tijd dit toelaat, alle opgaven en schrijf ook over elke opgave iets in jullie verslag (naast het antwoord kan dit bijv. zijn wat jullie geprobeerd hebben, waarom jullie er niet uitkomen, wat voor soort antwoord jullie verwachten dat er uitkomt, etc.).

We raden jullie aan om de opdrachten niet op te delen, dat gaat niet lukken! Samen erdoorheen en telkens overleggen en samen nadenken zal tot de beste antwoorden leiden.

Veel succes met deze opdracht!

1. Introductie

Netflix is een online streamingdienst waarbij consumenten films, series en documentaires kunnen bekijken. Er zijn ruim 100 miljoen gebruikers wereldwijd en het bedrijf maakt elk jaar miljarden dollar omzet. Een belangrijk aspect is dat Netflix in staat is om gebruikers films te adviseren op basis van hun voorkeuren en kijkgeschiedenis.

Vandaag kruip je in de rol van data-analist bij Netflix. Je gaat verschillen en overeenkomsten in de filmvoorkeuren van Netflix-gebruikers onderzoeken. Dit doe je aan de hand van de beoordeling die gebruikers aan een film geven. Kunnen deze beoordelingen gebruikt worden om een aanbevelingssysteem op te zetten?

In 2006 schreef Netflix een prijsvraag uit waarin zij het publiek vroeg om een zo goed mogelijke manier te vinden om de beoordelingen van gebruikers te voorspellen. Hiervoor was er een dataset beschikbaar met maar liefst 100.000 beoordelingen van films. Degene die meer dan 10% verbetering wist te bereiken ten opzichte van de voorspelling van Netflix kon maar liefst 1.000.000 dollar winnen.

Opdracht 1

Welke data denk je dat er gebruikt kan worden om een aanbevelingssysteem op te zetten? Noem ten minste drie variabelen. ■

Elke film in de database van Netflix heeft beoordelingen van gebruikers die de film bekeken hebben. Een film heeft een beoordeling van 0 tot en met 5, waarbij 5 de hoogste beoordeling is. Verder is er informatie beschikbaar over het genre van elke film.

Films in de dataset behoren tot één of meerdere van de volgende genres:

- Action
- Adventure
- Animation
- Comedy
- Children
- Crime
- Documentary
-

Opdracht 2

Ga naar <https://www.vustat.eu/apps/stat/index.html>.

Open het bestand (...).

Bekijk alle beoordelingen voor de film Star Track (2009).

a) Wat is de gemiddelde beoordeling van deze film?

In het bestand (...) staan voor de films die gelabeld zijn met het genre Romantiek en Sciencefiction de gemiddelde beoordelingen die de gebruikers geven aan films met dit genre.

b) Geef deze gemiddelde beoordelingen op een overzichtelijke manier weer in één diagram.

Er wordt het volgende beweerd: een gebruiker met een gemiddelde beoordeling van 5 voor Sciencefiction films is een liefhebber van dit soort films.

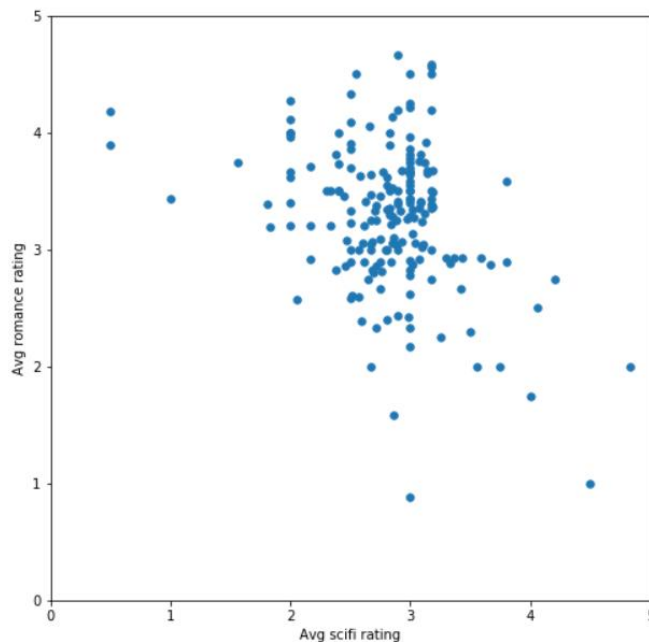
c) Geef een argument waarom dit niet juist hoeft te zijn. ■

2. Clustering

Het doel van het clusteren van data is het opdelen van de dataset in groepen bestaande uit soortgelijke items. Deze groepen worden clusters genoemd. In de context van Netflix willen we bijvoorbeeld groepen maken met gebruikers die een gelijke smaak hebben zodat we zo op basis van data uit deze groep films kunnen adviseren aan de gebruikers.

Ter illustratie van het begrip cluster bekijken we onderstaande puntenwolk bestaande uit 200 punten. Elk punt stelt een gebruiker voor en bestaat uit twee coördinaten: een horizontale coördinaat voor de gemiddelde romantiekbeoordeling en een verticale coördinaat voor de gemiddelde sci-fi beoordeling.

Voorbeeld: een gebruiker heeft films met het label romantiek beoordeeld met een gemiddelde van 3 en sci-fi met een gemiddelde van 3,5. Dit levert het punt (3 ; 3,5) op.



Opdracht 3

Stel we deze puntenwolk in twee delen willen splitsen, door bijvoorbeeld een deel van de punten rood te kleuren en een deel groen. Het doel is om het zo op te delen, dat de personen in een groep een enigszins overeenkomstige filmsmaak hebben. Elk punt uit de puntenwolk moet tot één van de twee delen behoren. Er zijn vele manieren om deze puntenwolk in twee delen op te delen.

a) Geef grafisch weer welke twee groepen jij zou maken en licht toe hoe je elk cluster zou kunnen karakteriseren.

Je kunt ook drie clusters maken.

b) Geef grafisch weer welke drie clusters jij zou maken en licht toe hoe je elk cluster zou kunnen omschrijven.

Wanneer het aantal clusters dat je maakt toeneemt, zal elk cluster gemiddeld uit minder punten bestaan.

c) Noem een voordeel en een nadeel van een groot aantal clusters.

Neem aan dat er geen lege cluster zijn.

d) Wat is het maximaal aantal clusters dat je zou kunnen maken? ■

3. Afstandsmaten

We vragen ons nu af hoeveel de ene film op de andere film lijkt. Om de mate van overeenkomst te bepalen bekijken we de afstand tussen twee punten. We gaan twee verschillende *afstandsmaten* bekijken: de Manhattan-afstand en de Euclidische afstand.

Manhattan afstand:

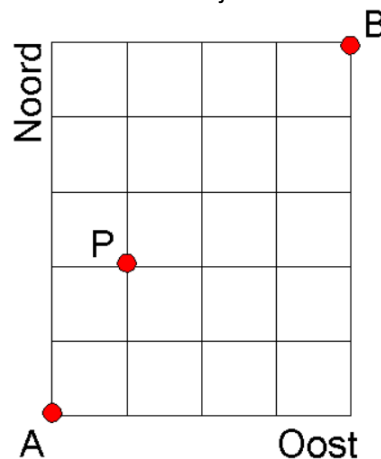
De naam van deze afstandsmaat verwijst naar de roostervormige opzet van de meeste straten op het eiland Manhattan in New York. Een route tussen twee punten is altijd de kortst mogelijk weg tussen twee punten, je loopt dus niet 'om'. Verder is schuin oversteken niet toegestaan, een route gaat altijd over de lijnen in het rooster. De Manhattan-afstand tussen twee punten is nu de lengte van de route (dat is hetzelfde als de optelling van de verticale en horizontale afstand).

Euclidische afstand:

De Euclidische afstand tussen twee punten is de lengte van de kortste verbindingsweg, waarbij nu niet langer over het rooster gelopen hoeft te worden. Er hoeft dus geen gebruik gemaakt te worden van een rooster. De afstand kan berekend worden met behulp van de stelling van Pythagoras.

Opdracht 4

- a) Teken in onderstaand rooster een route tussen punt A en punt B waarmee je de Manhattan afstand kunt berekenen en een route waarmee je de Euclidische afstand kunt berekenen.



- b) Bereken de Manhattan afstand en de Euclidische afstand tussen punt A en B.

Stel punt C heeft coördinaten (x_C, y_C) en punt D heeft coördinaten (x_D, y_D)

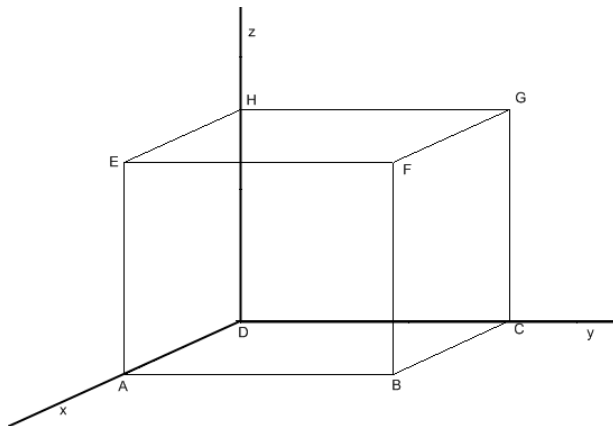
- c) Stel een formule op waarmee je de Euclidische afstand kunt berekenen tussen $C(x_C, y_C)$ en $D(x_D, y_D)$. ■

Je hebt zojuist twee verschillende afstanden berekend tussen twee punten die beide twee coördinaten hebben. Je kunt ook de afstand berekenen tussen twee punten met meer dan twee coördinaten. Je zou bijvoorbeeld kunnen kijken naar de coördinaten die behoren bij een film. Elk

coördinaat geeft aan of de film tot een bepaald genre behoort (met de afspraak dat een 1 aangeeft dat de film wel tot dit genre behoort en een 0 niet).

We kijken nu een situatie met drie coördinaten. De eerste coördinaat geeft aan of de film tot het genre Animatie behoort, het tweede coördinaat of een film tot het genre Comedy behoort en het derde coördinaat de de film tot het genre Children behoort. De film Toy Story heeft in dit geval de coördinaten (1,1,1).

Drie dimensionale punten kunnen weergegeven worden in een xyz-assenstelsel. In onderstaand figuur zie je dit assenstelsel. In het assenstelsel is de kubus ABCD.EFGH getekend. De ribben van deze kubus hebben lengte 1.



Opdracht 5

Punt D ligt in de oorsprong.

- Geef een zelfbedacht voorbeeld van een film waarvan de coördinaten overeenkomen met punt D.
- Welk van de acht hoekpunten correspondeert met de film Toy Story?
- Bereken de Euclidische afstand tussen punt D en het punt dat hoort bij Toy Story. ■

4. Algoritme

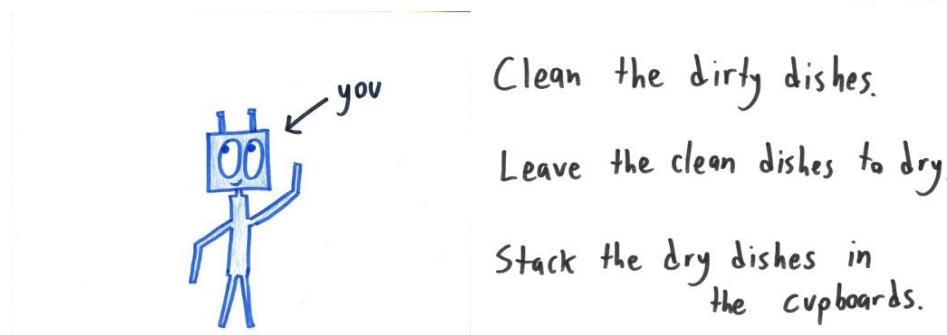
Een algoritme is een vaste reeks instructies of stappen die je van een startpunt naar een eindpunt leiden. Voor veel dingen die je doet maak je eigenlijk gebruik van een algoritme. Denk maar eens aan het koken aan de hand van een recept of het doen van de afwas. Ook in de wiskundeles maak je veel gebruik van algoritmes. Een voorbeeld is het opstellen van een raaklijn aan een grafiek. Om dat te doen volg je altijd hetzelfde stappenplan.

Opdracht 6

- a) Beschrijf een ander algoritme dat je gebruikt hebt in de wiskundeles.

Het is belangrijk dat de stappen in het algoritme duidelijk zijn.

Stel je voor dat je een huishoudrobot bent. Jouw taak is om de afwas te doen. Hierbij voer je het algoritme uit dat jou gegeven is, zonder daarbij na te denken. In onderstaand figuur vind je het algoritme dat bestaat uit drie stappen.



Er kunnen nogal veel dingen misgaan. In de bijlage staan vier afbeeldingen op chronologische volgorde. In elke afbeelding is het algoritme voor de afwasrobot verbeterd, nadat duidelijk werd dat de robot zijn taak niet zoals verwacht uitvoerde.

- b) Geef bij elk afbeelding aan wat de robot fout gedaan zou kunnen hebben, voordat het algoritme aangepast moest worden. ■

5. K-means clustering

Nu je weet wat een algoritme is gaan we één van de bekendste clusteringalgoritmes bestuderen: K-means clustering. Het is een methode die N punten verdeelt over k clusters. Het doel is om de N punten op te delen in k clusters, zo dat elk cluster een betekenisvolle groep (films die op elkaar lijken) representeert die we van een goed advies over hun filmkeuze kunnen voorzien. Belangrijk is dat vooraf het aantal clusters bepaald moet worden. Het k-means algoritme is een iteratieve methode, wat aangeeft dat het gaat om een zich herhalend proces. Het algoritme gaat door tot de clusters stabiel zijn, dat wil zeggen dat de inhoud ervan niet meer verandert. Een iteratie is het doorlopen van één cyclus van het proces (stap 3 t/m 5).

Het algoritme werkt met behulp van clustercentrums. Deze clustercentrum worden eerst gekozen. Hoe deze gekozen worden, dat bepaal je zelf of laat je willekeurig bepalen. Het clustercentrum mag een datapunt zijn, maar ook elk ander willekeurig punt. In elke iteratie worden de clustercentrums opnieuw bepaald.

Het *K-means* algoritme werkt als volgt:

Voorbereiding (ook wel initialisatie)

Stap 1. Kies k , het aantal clusters dat je wilt maken

Stap 2. Kies de positie van de clustercentrums.

Toewijzing

Stap 3: Wijs elk punt toe aan het cluster waarvan het centrum het dichtst bij is.

Updaten clustercentrums

Stap 4. Bepaal van elk cluster het nieuwe clustercentrum door voor elke coördinaat het gemiddelde te nemen van alle punten in het cluster.

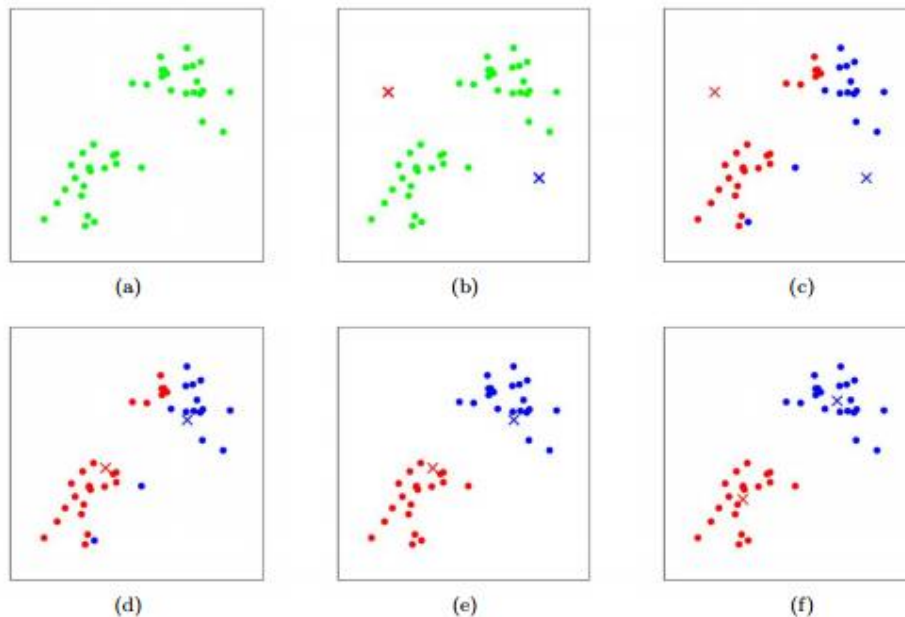
Controle

Stap 5. Controleer of er clustercentrums zijn die minstens één nieuwe coördinaat hebben.

- a. Zo ja; ga terug naar stap 3.
- b. Zo nee; je bent klaar.

Voorbeeld

Stap 1 In onderstaand figuur zie je een grafisch voorbeeld van het K-means algoritme met 2 clustercentrums ($k = 2$). In figuur (a) zijn alle datapunten gegeven.



Stap 2. In figuur (b) worden twee clustercentrums gekozen die gegeven worden door een rood kruisje en een blauw kruisje. In dit geval zijn de clustercentrums willekeurig gekozen.

Stap 3. In figuur (c) wordt elk punt toegewezen aan het dichtstbijzijnde clustercentrum. De punten die het dichtst bij het rode clustercentrum zijn worden rood gemaakt, de punten die het dichtst bij het blauwe clustercentrum worden blauw gekleurd.

Stap 4. In figuur (d) zijn de clustercentrums opnieuw berekend.

Stap 5. De locatie van de clustercentrums is veranderd, dus ga weer verder met stap 3.

Stap 3. In figuur (e) worden punten opnieuw toegewezen aan de clustercentrums.

Het algoritme stopt nadat de clustercentrums onveranderd blijven. Dat is na figuur (f).

Opdracht 7

Zorg dat je begrijpt hoe het algoritme werkt door te experimenteren op

<http://stanford.edu/class/ee103/visualizations/kmeans/kmeans.html>. Je hoeft over deze opdracht niets in het verslag te schrijven. ■

We bekijken nu een getallenvoorbeeld. Je gaat dit getallenvoorbeeld in de volgende opdracht zelf afmaken.

Gegeven zijn de volgende punten: (1,1), (2,1), (4,3) en (5,4).

Stap 1 Kies $k=2$.

Stap 2 Kies (1,1) en (2,1) als clustercentrums.

Stap 3 Bereken voor elk punt de Euclidische afstand tot de clustercentrums en wijs elk punt toe aan het dichtstbijzijnde clustercentrum.

Punt	Afstand tot cluster 1	Afstand tot cluster 2	Toegewezen cluster
(1,1)	0	1	1
(2,1)	1	0	2
(4,3)	$\sqrt{13}$	$\sqrt{8}$	2
(5,4)	5	$\sqrt{18}$	2

Stap 4

De coördinaten van clustercentrum 1 zijn nu: (1,1)

De coördinaten van clustercentrum 2 zijn nu: $\left(\frac{2+4+5}{3}, \frac{1+3+4}{3}\right) = \left(3\frac{2}{3}, 2\frac{2}{3}\right)$

Stap 5 Ja, clustercentrum 2 heeft nieuwe coördinaten. Ga weer verder met stap 3.

Opdracht 8

Maak bovenstaand voorbeeld af door het K-means algoritme verder uit te voeren. Doe dit op dezelfde manier als in het voorbeeld. Schrijf alle stappen uitgewerkt op. ■

Opdracht 9

Open het Excel-bestand (). In de puntenwolk staan 17 punten: de blauwe punten zijn de datapunten, de andere kleuren stellen de clustercentrums voor. Er ontbreken nog een aantal formules voordat er een K-means clustering met $k = 3$ uitgevoerd kan worden.

- Vul in Cel C25, D25 en E25 de formule in waarmee je de afstand tussen punt 1 en het clustercentrum berekent. Trek de formules door naar beneden zodat van elk punt de afstand tot de clustercentra berekend wordt. Schrijf de formule die je in cel A25 hebt ingevoerd op in je verslag.

Hint: bij verwijzingen naar de cellen van het clustercentrum gebruik je \$-tekens, daarmee zet je de celverwijzing vast als je de formule doortrekt. Voorbeeld voor cel B20: \$B\$20.

- In cel H25 staat de volgende formule: =ALS(XXX(C25:E25)=C25;"Cluster 1";ALS(XXX(C25:E25)=D25;"Cluster 2";"Cluster 3")). Vul op de plaats van XXX de juiste Excel-functie in. Schrijf in je verslag welke functie je hebt gebruikt.
- Voer de K-means clustering uit met de initialisatie van de clustercentrums (4,7), (2,9) en (7,5). Neem het diagram dat ontstaan is na convergentie van het algoritme op in je verslag. ■

Opdracht 10

a) Het K-means algoritme stopt nadat de clustercentrums niet meer verschuiven. Je kunt er zeker van zijn dat dit na een eindig aantal stappen gebeurt. Geef een reden waarom je hier zeker van kunt zijn.

b) Onderzoek met behulp van de website

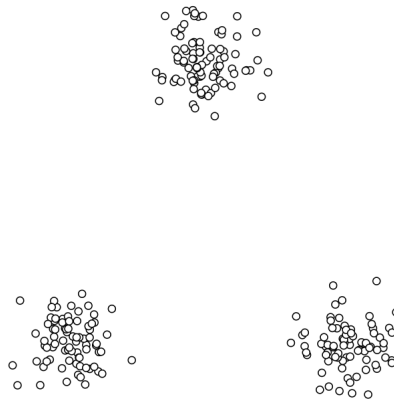
<http://stanford.edu/class/ee103/visualizations/kmeans/kmeans.html> of de clusters die resulteren na uitvoeren van het algoritme afhankelijk zijn van de keuze van de locatie van de clustercentrums in stap 2.

c) Het aantal iteraties voordat convergentie (de clustercentrums veranderen niet meer) optreedt is niet altijd hetzelfde. Noem ten minste drie factoren die van invloed zijn op het aantal iteraties dat nodig is. ■

6. Eigenschappen van K-means clustering

In de volgende reeks opdrachten bekijk je verschillende datasets.

We starten met de Gaussian mixture dataset in figuur X. Natuurlijk is het eigenlijk niet nodig om te clusteren bij dit voorbeeld, op het oog is het al overduidelijk wat de drie clusters zijn: de drie groepjes punten die bij elkaar liggen. Toch gaan we met dit voorbeeld aan de slag en beperken we ons niet tot $k = 3$, omdat deze dataset uitermate geschikt is om een aantal negatieve eigenschappen van het k-means algoritme te onderzoeken waar je rekening mee zou moeten houden als je het k-means algoritme gebruikt bij grotere datasets (waarbij je niet direct ziet hoeveel clusters er zijn en hoe die eruit zouden moeten zien).



Figuur X. Gaussian mixture dataset

Je gaat nu zelf op zoek naar initialisaties waarmee je verschillende eigenschappen van de k-means clustering kunt laten zien. Ga naar de volgende website: <https://www.naftaliharris.com/blog/visualizing-k-means-clustering/>. Kies in het eerste scherm de optie dat jij zelf de clustercentra wilt initialiseren ('I'll Choose') en kies vervolgens voor de 'Gaussian mixture'-verdeling. Je kunt nu zelf clustercentra kiezen door in het venster op een locatie te klikken. In de visualisatie neemt elk datapunt de kleur aan van het cluster waartoe het toegewezen is.

Opdracht 11: Een leeg cluster

Bedenk een initialisatie met $k = 4$ waarbij er na het uitvoeren van het k-means algoritme één van de clusters leeg is (het bevat geen datapunten). Maak een screenshot van het resultaat en zet dat in het verslag. ■

Opdracht 12: Verschillende uitkomsten bij hetzelfde aantal clusters.

Laat zien dat na initialisatie met $k = 3$ en het uitvoeren van het algoritme de volgende twee gevallen kunnen ontstaan:

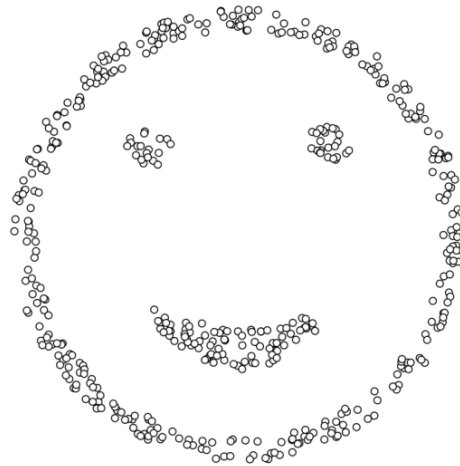
1. Elk cluster bestaat uit een groepje punten.
2. Eén cluster bestaat uit twee groepjes punten.

Maak van beide gevallen een screenshot. ■

Opdracht 13: Er is geen goede clustering mogelijk

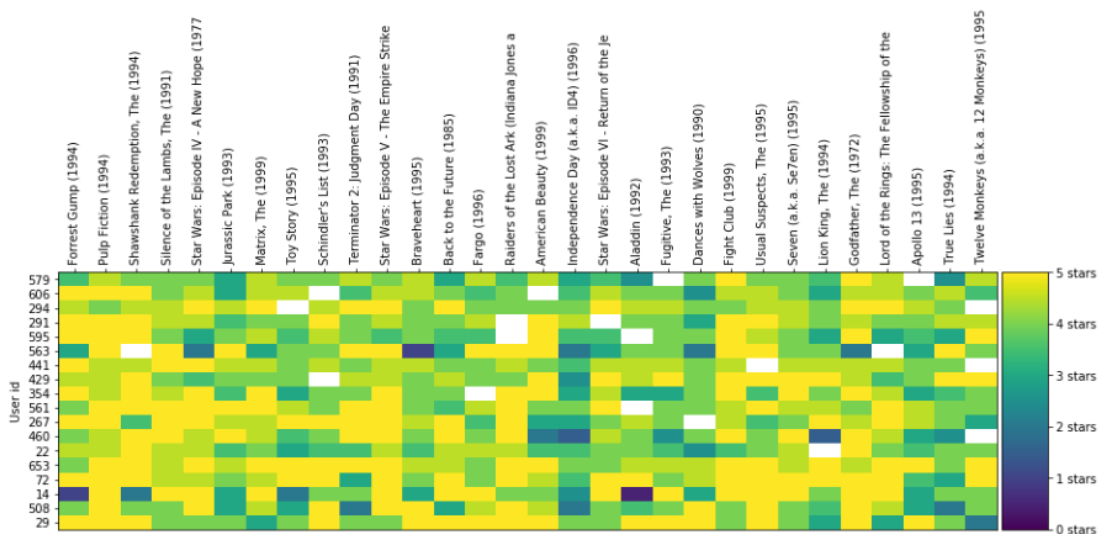
We bekijken nu de Smiley Face dataset. Zie de figuur hieronder.

- Welke clustering zou jij zonder gebruik te maken van het algoritme doen?
- Welk deel van het k-means algoritme zorgt ervoor dat er bij de Smiley Face dataset geen goede clustering zal ontstaan? ■



7. Film level clustering

Tot slot bekijken we de resultaten van een k-means clustering bij een grote dataset met gegevens van Netflix. In onderstaand figuur is een deel van de dataset grafisch weergegeven. Elke kolom is een film, elke rij een gebruiker. De kleur geeft de beoordeling die de gebruiker aan de film geeft aan. Zie hiervoor de kleurschaal naast het figuur.

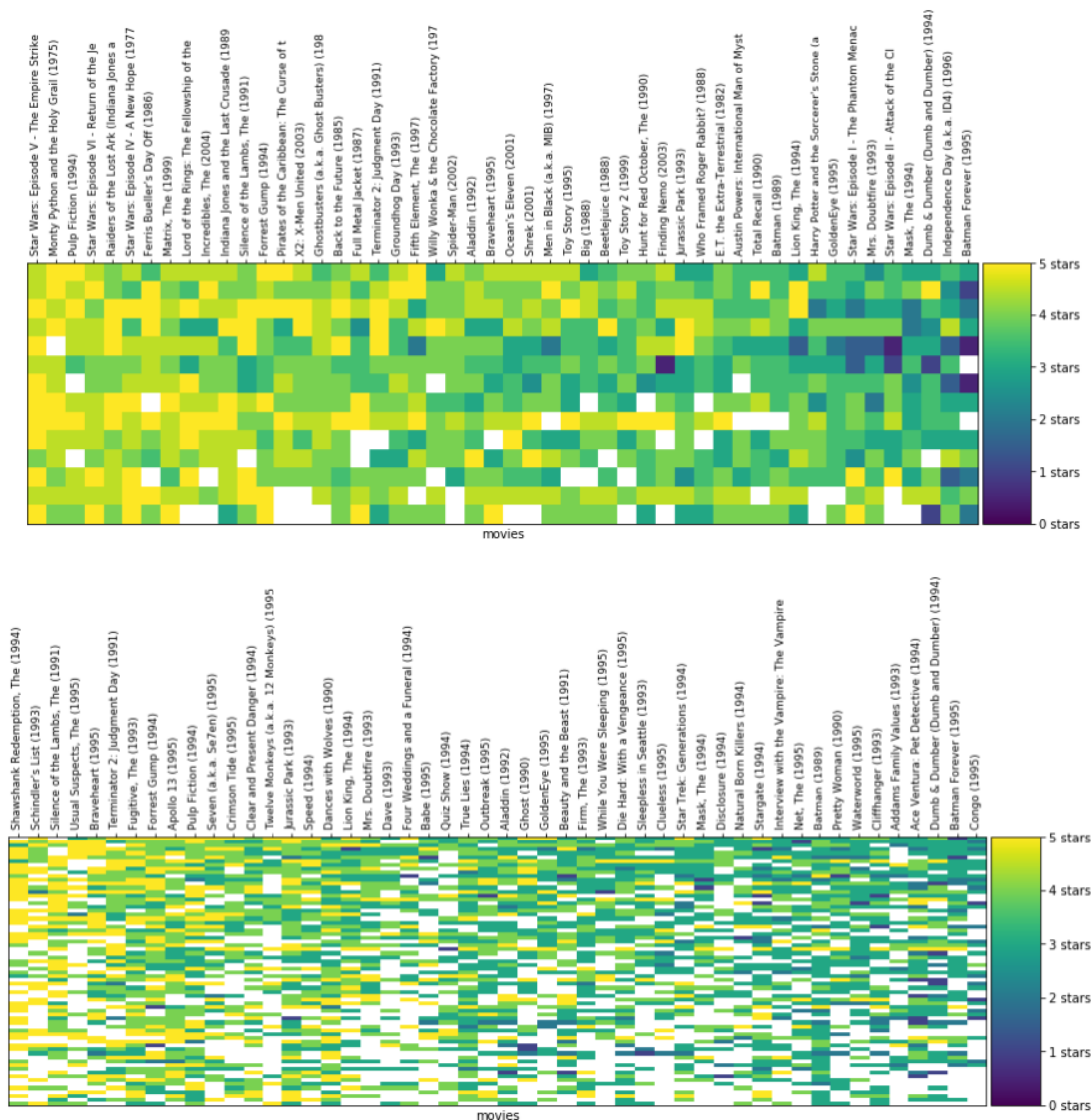


Opdracht 14

- Er zijn enkele cellen wit. Wat zou dit kunnen betekenen?

Het k-means algoritme wordt toegepast op deze dataset. Zo ontstaan 20 clusters van gebruikers met overeenkomstige filmsmaak. Deze clusters kunnen dan gebruikt worden om films te adviseren aan gebruikers.

In onderstaande twee figuren zie je twee verschillende clusters die zijn ontstaan na toepassing van het k-means algoritme.



- Streep in de volgende zin bij de onderstreepte woorden door wat niet van toepassing is: Hoe meer gelijk de smaak is van gebruikers in het cluster, hoe meer horizontale/verticale lijnen in hetzelfde/verschillende kleur er zichtbaar zijn.
- Sommige clusters zijn meer geel en andere clusters zijn meer groen/blauw. Wat is hiervan de praktische betekenis?

Er zijn nu clusters van gebruikers met gelijke smaak. Dit kun je gebruiken om films aan te bevelen die gebruikers nog niet hebben gezien.

- Beschrijf hoe je dit zou doen op basis van de beoordelingen van de films van andere gebruikers in dit cluster. ■

Clean the dirty dishes.

Leave the clean dishes to dry.

Stack the dry dishes in
the cupboards.*

*EXCEPT in the hour before a meal;
in that case, leave them

Clean the dirty dishes.
↳ but wait until all the large dirt particles are gone

Leave the clean dishes to dry.

Stack the dry dishes in
the cupboards.*

*EXCEPT in the hour before a meal;
in that case, leave them

Clean the dirty dishes.

↳ but wait until all the large dirt particles are gone*
*OR until removal of particles has stopped

Leave the clean dishes to dry.

Stack the dry dishes in
the cupboards.*

*EXCEPT in the hour before a meal;
in that case, leave them

Clean the dirty dishes.

↳ but wait until all the large dirt particles are gone*
*OR until removal of particles has stopped
FOR AT LEAST 5 MINUTES

Leave the clean dishes to dry.

↳ UNLESS it's right before a meal, and there aren't enough dishes, in which case manually dry them

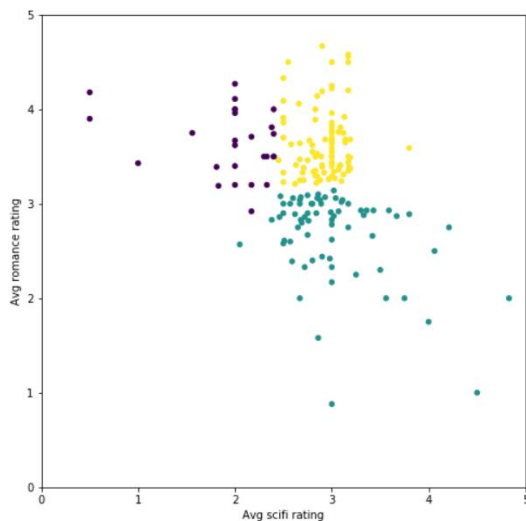
Stack the dry dishes in
the cupboards.*

*EXCEPT in the ^{45 minutes} hour before a meal;
in that case, leave them → assuming they're on the table or the counter, the dishes in the drying rack can go

Bijlage C: Correctievoorschrift

Voorbeeld uitwerkingen werkmiddag K-means

- 3p 1) Voorbeeld van juiste antwoorden: Filmnaam, filmgenre, gebruikersnummer, beoordeling.
Voor elk juist antwoordelement 1 punt.
- 2p 3a) Grafische weergave (1 punt)
Een toelichting waaruit blijkt dat de leerling in staat is om op een betekenisvolle manier de dataset in twee groepen op te delen (1 punt).
Bijvoorbeeld: Een cluster van gebruikers met een romance rating van lager dan 3 en een cluster van gebruikers met romance rating hoger dan 3. In de eerste groep zitten gebruikers die niet van romance houden, in de twee groep gebruikers die wel van romance houden.
- 2p 3b) Grafische weergave (1 punt)
Een toelichting waaruit blijkt dat de leerling in staat is om op een betekenisvolle manier de dataset in drie groepen op te delen (1 punt). Bijvoorbeeld: Een cluster met een hoge romance beoordeling en een lage scifi beoordeling. Een cluster met hoge scifi beoordeling en een hoge romance beoordeling. Een cluster met een lage romance beoordeling en een hoge beoordeling.



- 2p 3c) Voordeel: de clusters bestaan uit punten die soortgelijk zijn (1 punt)
Nadeel: veel rekentijd of de clusters hebben weinig praktische betekenis meer (1 punt).
- 1p 3d) Evenveel clusters als punten / 200 clusters
- 2p 4a) Diagonale route van A naar B voor Euclidische afstand (1 punt), een route van A naar B zonder omwegen voor Manhattan afstand (1 punt)
- 2p 4b) Manhattan: $4+5=9$
Euclidisch: $\sqrt{4^2 + 5^2} = \sqrt{41}$
- 2p 4c) $\sqrt{(x_D - x_C)^2 + (y_D - y_C)^2}$
Opmerking: wanneer als antwoord afstand² = $(x_D - x_C)^2 + (y_D - y_C)^2$ gegeven is, 1 punt toekennen
- 1p 5a) Een voorbeeld van een film die geen animatie is, geen comedy is en geen kinderfilm is.
- 1p 5b) punt F

- 2p 5c) $BD = \sqrt{2}$ (1 punt)
 $DF = \left(\sqrt{(\sqrt{2})^2 + 1^2} \right) \sqrt{3}$ (of 1,73) (1 punt)
- 1p 6a) Bijvoorbeeld: het algebraïsch berekenen van de extreme waarde of het vinden van de oplossingen van een kwadratische vergelijking van de vorm $ax^2 + bx + c = 0$ met de abc-formule.
- 4p 6b) Per afbeelding 1 punt:
- 1) de robot heeft de op tafel gezette borden al weer opgeruimd voordat de maaltijd begonnen is.
 - 2) direct wanneer het eten begint (en de borden 'vies' zijn) heeft de robot de vaat opgeruimd.
 - 3) de robot heeft de vaat laten staan, omdat er nog grove resten op waren blijven liggen.
 - 4) de robot heeft de vaat opgeruimd terwijl men nog niet klaar was met eten.
- 6p 8) Iteratie 2: bereken afstanden (2 punten), toewijzing clusters (1 punt), berekenen positie clustercentrums (1 punt)
 Toelichting dat er in de volgende iteratie niets meer gebeurt (eventueel met behulp van het doorrekenen van een volgende iteratie). (2 punten)

Iteratie 2:

Stap 3 Bereken voor elk punt de Euclidische afstand tot de clustercentrums en wijs elk punt toe aan het dichtstbijzijnde clustercentrum.

Punt	Afstand tot cluster 1	Afstand tot cluster 2	Toegewezen cluster
(1,1)	0	$\sqrt{9\frac{8}{9}} = 3,14$	1
(2,1)	1	$\sqrt{5\frac{5}{9}} = 2,36$	1
(4,3)	$\sqrt{13}$	$\sqrt{\frac{2}{9}} = 0,47$	2
(5,4)	5	$\sqrt{3\frac{5}{9}} = 1,89$	2

Stap 4

De coördinaten van clustercentrum 1 zijn nu: $(1\frac{1}{2}, 1)$

De coördinaten van clustercentrum 2 zijn nu: $(\frac{4+5}{2}, \frac{3+4}{2}) = (4\frac{1}{2}, 3\frac{1}{2})$

Iteratie 3:

Stap 3 Bereken voor elk punt de Euclidische afstand tot de clustercentrums en wijs elk punt toe aan het dichtstbijzijnde clustercentrum.

Punt	Afstand tot cluster 1	Afstand tot cluster 2	Toegewezen cluster
(1,1)	0,5	4,3	1
(2,1)	0,5	3,5	1
(4,3)	3,2	0,7	2
(5,4)	4,6	0,7	2

Stap 4

De coördinaten van clustercentrum 1 zijn nu: $(1\frac{1}{2}, 1)$

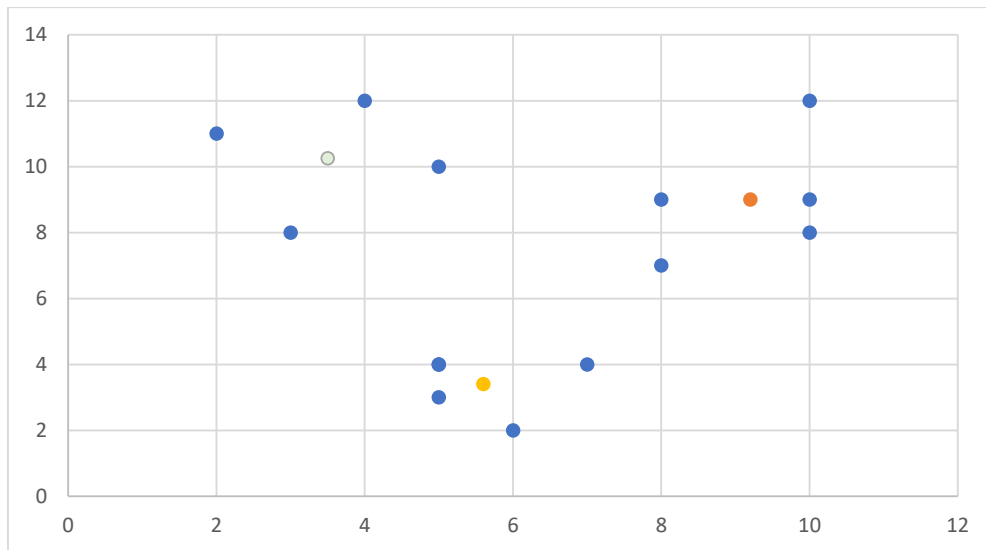
De coördinaten van clustercentrum 2 zijn nu: $(\frac{4+5}{2}, \frac{3+4}{2}) = (4\frac{1}{2}, 3\frac{1}{2})$

EINDE (want de clustercentrums zijn niet meer verschoven)

2p 9a) =WORTEL((F23-\$B\$20)^2+(G23-\$C\$20)^2)

1p 9b) De functie MIN

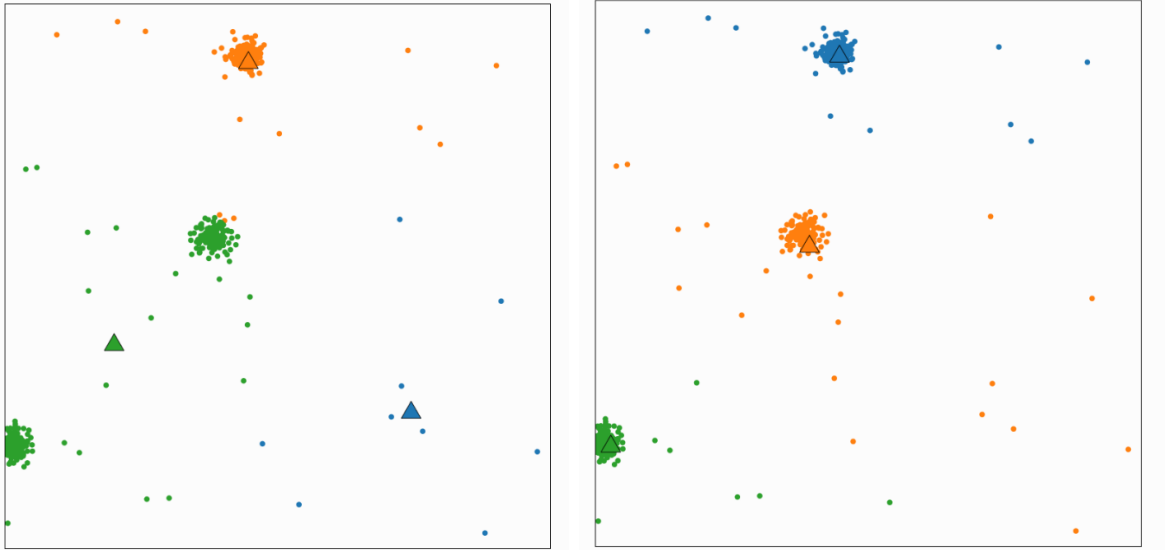
2p 9c)



1p 10a)

- In elke iteratie wordt de som van de afstand tussen de punten en de clustercentrums kleiner.
- Het aantal punten is eindig, dus er is maar een eindig aantal manier op te clusteren mogelijk.

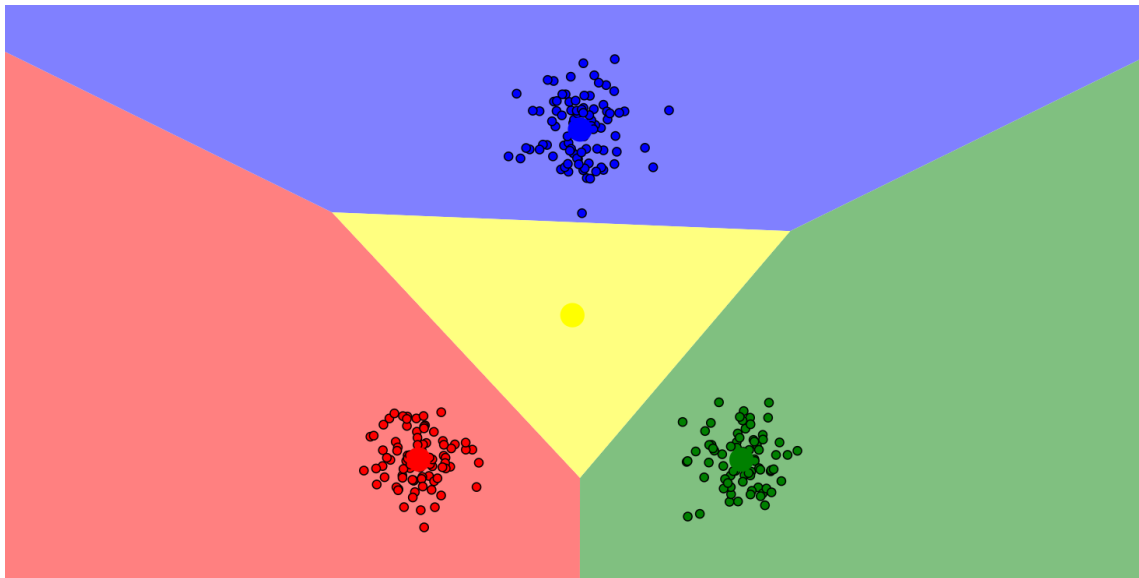
2p 10b) Ja, een goed tegenvoorbeeld waarbij een andere initialisatie (wel dezelfde data) tot een ander resultaat zal leiden:



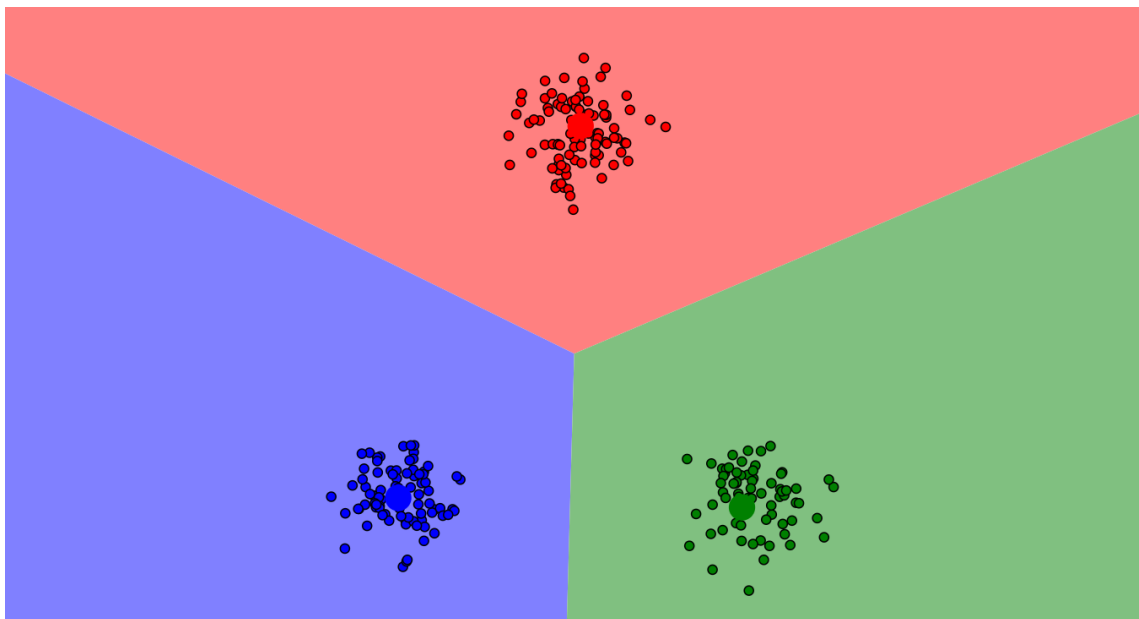
3p 10c) Drie van onderstaande redenen:

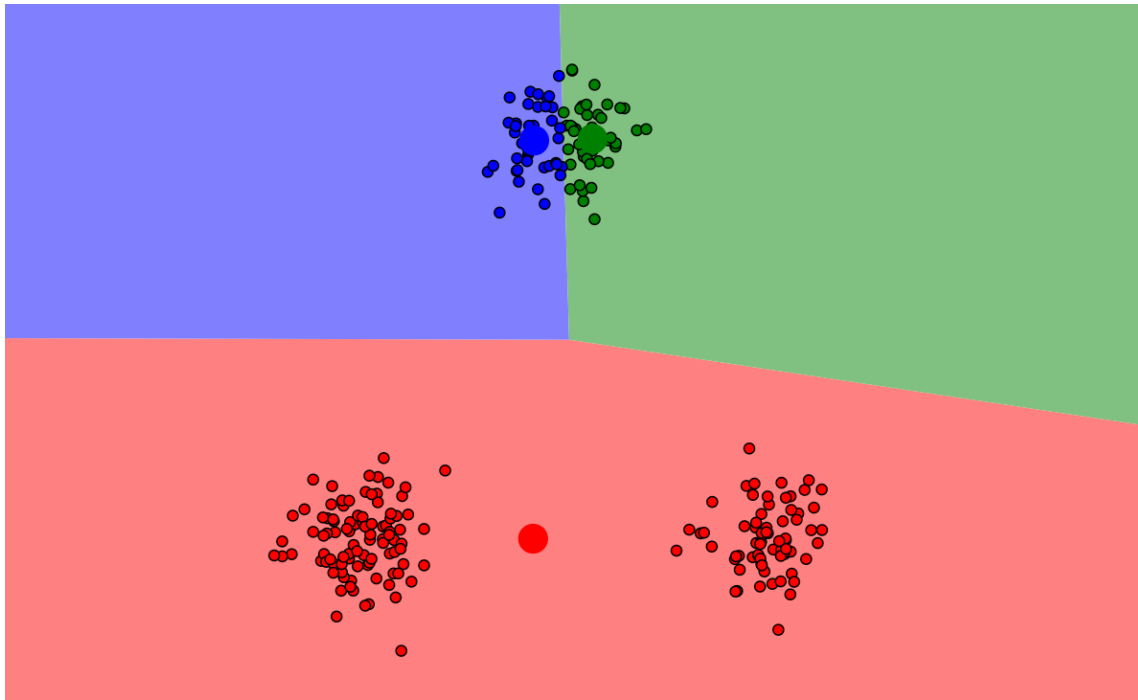
- de keuze van de clustercentrums
- het aantal clusters (k)
- de verdeling van de datapunten (duidelijke groepjes of niet)
- de hoeveelheid data

2p 11)



2p 12) Voor elk figuur 1 punt





- 1p 13a) Vier clusters, voor beide ogen elk één, voor de rand van het gezicht één en voor de mond één.
- 2p 13b) Het gebruik van de afstandsformule. Omdat het k-means algoritme met een afstandsformule werkt, zal het altijd punten die dichtbij elkaar liggen in hetzelfde cluster indelen. Hierdoor zullen de punten op de rand van het gezicht niet tot hetzelfde cluster behoren.
- 1p 14a) Dit zijn films die van een gebruiker geen beoordeling hebben gekregen.
- 2p 14b) Hoe meer gelijk de smaak is van gebruikers in het cluster, hoe meer verticale lijnen in gelijke kleur er zichtbaar zijn.
- 1p 14c) Clusters met veel gele vakjes bestaan uit gebruikers die een bepaalde groep films heel erg waarderen. In Clusters met veel blauwe/groene vakjes geven gebruikers een zelfde groep films een beoordeling van 2-3 sterren.
- 1p 14d) Per cluster zou ik de gemiddelde score per film berekenen. Ik zou dan de film met de hoogste gemiddelde score, die de gebruiker nog niet heeft gezien, adviseren.

Bijlage D: Interviews

Interview Leerling 1

Onderzoeker	Vond je het een leuke opdracht?
Leerling	De opzet van de opdracht was goed, maar ik vond het vooral een ingewikkelde opdracht.
Onderzoeker	Was was er ingewikkeld aan?
Leerling	Het was op een niveau dat wij niet eerder hebben gehad. Het was wel fijn dat er een opbouw was richting de moeilijkere opgaven. De eerste opdrachten waren niet zo moeilijk, zodat we eerst konden ontdekken hoe algoritmes nu werken. We werden daardoor niet meteen het diepe in gegooit.
Onderzoeker	Kun je aangeven welke opgaven je lastig vond?
Leerling	[Bladert door de opgaven]. De eerste vier opdrachten waren wel eenvoudig. Opdracht 5 kwam ik niet uit, dat was niet zo leuk. Opdracht 6 hebben we naar gekeken, maar kwamen we ook niet uit. Nouja, eigenlijk ben ik er trouwens wel een klein beetje uitgekomen. Ik heb wel een algoritme bedacht; ik kon wel bedenken wat er fout ging. Door de opdracht heb ik wel geleerd wat een algoritme is. Enne, hoe je dat kunt gebruiken.
Onderzoeker	En de opgaven erna over het K-means algoritme?
Leerling	Bij het K-means algoritme vond ik dat de opdracht ons wel goed leerde hoe een K-means algoritme in elkaar zit en hoe het werkt, maar het was moeilijk om dat in de praktijk te brengen. Ik vond het moeilijk om dat zo op te schrijven, met al die stappen naar het eindantwoord met punten en clusters die erin zitten.
Onderzoeker	Kun je iets bedenken waarmee je geholpen zou zijn, zodat je het beter in de praktijk kunt brengen?
Leerling	Eh, in principe was het stappenplan duidelijk, maar het zou nog fijner zijn als dat een keer helemaal klassikaal voorgedaan zou worden door jou.
Onderzoeker	Heb je het helemaal afgekregen?
Leerling	Nee, we zijn tot de Excel-opdracht gekomen. Dat we het niet afkregen komt deels omdat we te weinig tijd hadden, maar ook wel doordat we het gewoon niet helemaal begrepen. Wat ik al zei: het was een ingewikkelde opdracht.
Onderzoeker	Dat waren de inleidende opdrachten. Ik ga nu een paar vragen stellen over het onderwerp zelf. We hebben gekeken naar algoritmes, in het specifiek het K-means algoritme. Kun je mij vertellen wat een algoritme eigenlijk is?
Leerling	Een algoritme is een stappenplan die elke keer over en over wordt toegepast om een bepaald resultaat te halen, zodat iets, een apparaat of een wiskundeopgave, bij het juiste antwoord uitkomt.
Onderzoeker	Zou je een voorbeeld kunnen geven van een algoritme, anders dan het K-means algoritme?
Leerling	Bij de stelling van Pythagoras heb je een algoritme, namelijk een bepaald stappenplan hoe je dat berekent. [Denkt even] In principe zijn formules algoritmes.
Onderzoeker	Kun je een formule altijd als een algoritme zien?
Leerling	Nee, misschien dat niet. Geen goed voorbeeld. Maar dan een afwasmachine, die heeft het algoritme om eerst te wassen en dan te drogen. Niet andersom, dat werkt niet.
Onderzoeker	Dan over het K-means algoritme. Kun je vertellen wat dat algoritme doet?

Leerling	Nou, wat het doet: je hebt punten staan bij eh... bijvoorbeeld bij Netflix of mensen romantische films leuk vinden of science-fiction films. Dan werd je account als een puntje ingedeeld op die lijst op basis van je kijkgeschiedenis en op basis van je beoordelingen. En je hebt een aantal clusters, zoveel je wilt, laten we nu uitgaan van drie. Dan kijk je welke punten het dichtst bij welke clusters liggen. En als je die hebt gemarkeerd dan ga je het cluster verplaatsen naar het gemiddelde van al die punten. Dan ga je weer kijken of er andere punten dichterbij liggen. En dan verplaats je het cluster weer. Zo ga je door totdat er geen punten meer veranderen.
Onderzoeker	Heb je een idee wat het doel is van het algoritme?
Leerling	Dat kan bijvoorbeeld zijn om op Netflix aangeraden film te laten zien op je account.
Onderzoeker	Kun je dat nog iets algemener maken?
Leerling	Als er bijvoorbeeld iemand een nieuw account maakt en film gaat kijken, dan kun je dat punt erbij toveren en dan kunnen die clusters zich weer opnieuw gaan vormen.
Onderzoeker	Dus met dat K-means algoritme maak je clusters?
Leerling	Ja, dat klopt.
Onderzoeker	Heb je een idee hoe je kunt bepalen hoeveel clusters je zou willen maken?
Leerling	Dat ligt eraan. Weinig clusters heeft het voordeel dat je een grote groep hebt die het bepaald, zodat het meer naar het gemiddelde toekomt. Weinig clusters heeft het voordeel dat het specifiek aan jouw kijkwensen zal zijn. Het heeft allebei zijn voordelen en nadelen, maar je wilt natuurlijk niet heel veel of heel weinig clusters. Ik zou altijd rond de 3,4,5 doen; in dit geval dan.
Onderzoeker	Heb je een idee waar je dit nog meer zou kunnen toepassen naast Netflix?
Leerling	Eeeehhh.... Dat is oprecht een goede vraag. Daar moet ik even over nadenken. [Denkt na] Dat kun je toepassen op een beroepenwijzer, om te kijken in welk cluster je puntje komt op basis van de vragen die je hebt beantwoord. Dus om te bepalen welk beroep bij je past. [Lijkt heel tevreden over het antwoord]
Onderzoeker	Waarom je het K-means algoritme willen toepassen? Want soms kun je het ook gewoon zien, zoals in het eenvoudige voorbeeld in het boekje.
Leerling	Bij een algoritme kun je het nog preciezer maken en is het makkelijker om weer te geven op de computer. Als je het zelf bekijkt kun je niet heel precies met cijfers werken. Met wiskunde kun je dat veel preciezer bekijken.
Onderzoeker	Kun je dat nog koppelen aan dit figuur [wijst naar het 3D figuur]?
Leerling	Met die assen kun je het anders opdelen. Het zijn meer vierkantjes in de kubus. Je kijkt naar drie variabelen, bijvoorbeeld naar hoe leuk mensen romantische-, science-fiction en horrorfilms vinden. Dan kom je op basis van drie films ergens in de kubus.
Onderzoeker	Kan ik nog een as toevoegen?
Leerling	Is dat mogelijk 4D? Dat weet ik niet, ik heb er wel dingen over gehoord. Het zal misschien kunnen, maar ik kan het me niet voorstellen.
Onderzoeker	Dan nu een testvraagje over het K-means algoritme. Ik heb hier een aantal punten en twee clustercentrums. Zou je het K-means algoritme op dit voorbeeld kunnen toepassen?
Leerling	Jazeker. Ik maak gewoon twee kolommen, één met (1,3) en één met (2,4). Ik schrijf eronder welke bij welke horen. (1,4) ligt dichterbij (1,3). Het punt (3,4) ligt dichterbij (2,4).
Onderzoeker	Hoe bepaal je dat?
Leerling	Voor deze punten [wijst naar (3,4) en (2,4)]: het verschil tussen de y-as is 0 en het verschil met de x-as is 1. Dus de afstand is 1. Ik ga verder; het punt (4,2) is moeilijk. Die ligt denk ik dichterbij (3,4), want als ik 4-1 doe dan is dat 3 en 4-2 is 1. Dus dat zou 3 afstand moeten zijn. Is dat goed geredeneerd of niet? [Leerling is wat in verwarring]

Onderzoeker	Je hebt twee afstandsmaten geleerd.
Leerling	Oh wacht. Ik kan het tekenen. Ik maak nu twee assen. Ik heb het punt (4,2) en het punt (1,3) is hier en het punt (2,4) is hier. Dan ligt deze [wijst naar (2,4)] dichterbij.
Onderzoeker	Wat is nu je clustercentrum?
Leerling	Ik heb er twee. Dat zijn deze en deze [wijst de correcte punten aan en maakt er kruisjes van]
	Dan het punt (5,1). Dat is moeilijk. Je kon dit vast ook wel uitrekenen toch?
Onderzoeker	Ja, dat heb je geleerd.
Leerling	Oeps [lacht]. Ah, shit [brabbelt nog wat over verticale en horizontale verschillen]
	Ah wacht. Oh, de stelling van Pythagoras [tekent een driehoek]. De ene afstand is 4 en de andere 2. Dan wil je deze uitrekenen [wijst op de schuine zijde]. Dat is 16 en 4. Dus wortel 20. De andere afstand is wortel 13. Dat is minder groot. Dus hij ligt dichterbij (2,4).
Onderzoeker	Heel goed.
Leerling	Ja he. Dan kun je daarna gewoon het gemiddelde nemen. [Telt de x-coördinaten op en deelt het door drie, zelfde voor y]. Daarna kijken we of de punten verplaatsen. Punt (2,4) gaat veranderen. Daarna moeten we weer opnieuw de gemiddeldes uitrekenen. Daarna kunnen we stoppen.
Onderzoeker	Waarom kun je dan stoppen?
Leerling	Omdat de punten dan niet meer veranderen.
Onderzoeker	Dat was de testvraag. Dan wil ik tot slot nog even met je kijken naar de laatste vragen van de opdracht. Die heb je vanwege de tijd niet meer helemaal kunnen maken. Denk je dat dit algoritme bij dezelfde punten altijd dezelfde oplossing geeft?
Leerling	Nou nee, als ik dit clustercentrum iets verplaats, dan zouden andere punten bij dat ding horen, waardoor het gemiddelde ook veranderd. De punten komen dan in een ander cluster.
Onderzoeker	En komt dat 'op het einde' niet meer goed? Want je gaat door tot het stabiliseert.
Leerling	Als één van die punten verplaatst zal het niet hetzelfde zijn. Anders gezegd, als bij de eerste meting er een punt dichterbij een ander clustercentrum ligt, dan veranderd het. Als er bij de eerste meting geen verschil is, bijvoorbeeld als het clustercentrum één centimeter hoger ligt, dan maakt dat niet uit.
Onderzoeker	Soms kan het heel lang duren tot het algoritme klaar is, maar soms kan het ook heel snel gaan. Waar kan dat aan liggen?
Leerling	Dat kan liggen aan dat er veel punten op de grens van de clusters liggen. Dat ze net bij de één horen. Dat wanneer je de clusters verschuift, de kans groot is dat er een puntje bij een ander cluster gaat horen. Zodat de clusters dan weer veranderen. Terwijl als je twee grote bollen hebt met punten [tekent twee puntenwolken], dan is het al na één zet klaar.
Onderzoeker	Dus je zegt nu iets over de spreiding van de punten.
Leerling	Ja. Als de punten dichtbij één punt liggen en de afstand tot de andere punten groot is, dan zal het snel opgelost zijn.
Onderzoeker	Bedankt, dat waren de vragen. Heb je tot slot nog opmerkingen?
Leerling	Ik heb hier oprecht nog wel wat van geleerd.
Onderzoeker	Kun je nog specificeren wat je hebt geleerd?
Leerling	Wat een algoritme is en hoe dit K-means algoritme werkt.

Interview leerling 2

Onderzoeker	Wat vond je van de opdracht?
Leerling	Wel een leuke opdracht. Beetje hetzelfde idee als de wiskundedag opdracht. Echt wel heel anders dan de opgaven uit de les. Het was wel moeilijker.
Onderzoeker	Wat was er moeilijker?
Leerling	Nou, de theorie was ingewikkelder. Die berekeningen en hoe het in elkaar zat. Je moest er gewoon veel meer voor doen dan voor de vorige praktische opdracht.
Onderzoeker	Het was meer werk?
Leerling	Ja, dit was veel meer werk.
Onderzoeker	Kun je in het boekje aangeven welke opdrachten moeilijk waren?
Leerling	De opdrachten 1 tot en met 5 waren makkelijk. Vanaf opdracht 6 ging het minder. Het werd allemaal net iets te ingewikkeld.
Onderzoeker	Was de opdracht wel duidelijk?
Leerling	Ja dat wel, alleen het was gewoon moeilijk om het uit te voeren.
Onderzoeker	Zou er iets gedaan kunnen worden om het makkelijker te maken?
Leerling	Ja, als het klassikaal zou worden voor worden gedaan, hoe het moest enzo, stap-voor-stap, en wat het idee was, dan was het wel iets makkelijker geweest.
Onderzoeker	Tot hoe ver ben je gekomen?
Leerling	Tot en met 9c.
Onderzoeker	Kun je mij vertellen wat een algoritme is?
Leerling	Nou, dat is gewoon een stappenplan die je iets geeft, bijvoorbeeld die robot, om een bepaalde taak uit te voeren.
Onderzoeker	Kun je een voorbeeld noemen uit de les ?
Leerling	Bijvoorbeeld bij het onderwerp waarbij we nu bezig zijn met wiskunde (hypothesetoetsing, red). Die zes stappen die we altijd moeten zetten lijkt me een algoritme.
Onderzoeker	Wat levert het gebruik van het K-means algoritme op?
Leerling	Wat ik weet van de opdrachten die we hebben gedaan is dat je een overzicht krijgt van wat mensen veel kijken. De eerste opdrachten gingen over wat de verdeling is van mensen en welke genres ze kijken. Dan krijg je punten in een grafiek en dan kun je zien waar de clusters in een grafiek zijn, daar waar de mensen het dichtst bij elkaar. Je kunt dan ook een scheiding maken: deze mensen vinden dit thema leuker dan het andere thema.
Onderzoeker	Je noemt het woord cluster. Kun je iets verder verduidelijken wat dat is?
Leerling	Een cluster is gewoon een groep van punten die je hebt waar mensen heel erg in de buurt van elkaar liggen qua interesses. Ze hebben een bepaald iets gemeen.
Onderzoeker	Kun je een ander voorbeeld verzinnen waar je het K-means algoritme toe kunt passen?
Leerling	Ehm [denkt een tijdje na]. Zoekresultaten op het internet in het algemeen. Dat je een beetje kijkt waar de meeste interesse ligt en wat er veel gezocht wordt. Dat je een idee kunt krijgen wat het meeste wordt gezocht en dat je ook per stad of periode van het jaar wat de meeste mensen in een bepaald gebied zoeken.
Onderzoeker	Wat voor clusters zie je dan voor je?
Leerling	Nouja, je kunt dan bijvoorbeeld groepen maken voor vermaak en zakelijke dingen.

Onderzoeker	Ik heb hier vier punten en twee clustercentrums. Kun je het K-means algoritme uitvoeren op het volgende voorbeeld?
Leerling	Pfoe, ik zal het proberen. Mag ik het boekje erbij gebruiken?
Onderzoeker	Ja, dat kan.
Leerling	Wil je dat ik zoiets ga tekenen? [Wijst naar het voorbeeld uit het boekje]
Onderzoeker	Als je mij daarmee kunt laten zien hoe het algoritme werkt, dan is het goed.
Leerling	Ik heb echt geen idee.
Onderzoeker	Begin bijvoorbeeld eens met de vier punten te tekenen in een assenstelsel.
Leerling	Ah ja, daar zat ik ook aan te denken. [Tekent vier punten (als bolletje) en twee clustercentrums (met een kruisje) in het assenstelsel]
	Wat ik nog weet van de opdracht is dat je eigenlijk kijkt bij welke clusters de punten het dichtst bij in de buurt liggen. Dat is wel een beetje moeilijk te zien in de schets. [Kijkt de onderzoeker vragend aan]. Het eerste punt ligt het dichtst bij het linker clustercentrum. Deze [wijst naar (3,4)] zou één van allebei kunnen zijn.
Onderzoeker	Kun je berekenen bij welke hij dichterbij ligt?
Leerling	Mja, bij deze [wijst naar (1,4)] weet je dat hij een verschil van 1 heeft en hier [wijst naar (2,4) en (3,4)] kun je Pythagoras gebruiken. Deze is dan 2^2 [wijst op de schuine zijde]. Dat is meer dan 1, dus hij ligt dichterbij (2,4). De andere twee punten kun je wel zien dat ze bij het tweede cluster horen.
	Aan de hand daarvan kun je het clustercentrum verplaatsen. Als er maar één punt bij het clustercentrum hoor, dan is dat punt al direct het nieuwe clustercentrum. Het andere cluster heeft drie punten, dus dan pak je eigenlijk het middelste. Het gemiddelde van de afstand van de ... [het blijft stil]
	Je zegt ik neem het gemiddelde? Weet je nog waarvan?
Leerling	Even kijken. Je weet dat ie wat meer richting de rechter twee punten zou komen te liggen. Dan gaat ie ongeveer hierzo [wijst de correcte plek aan] komen te liggen.
Onderzoeker	Ik wil dat graag berekenen. Waarom zou ik dat graag willen?
Leerling	Nou, dan heb je de precieze locatie. Maar ik weet niet meer hoe dat moet.
Onderzoeker	We laten het berekenen dan voor nu even zitten. Stel je hebt dat gedaan, wat zou je daarna doen?
Leerling	Aan de hand daarvan kun je weer kijken bij welk cluster de punten horen. Aangezien de clustercentrums net verplaatst zijn, komt (3,4) nu verder van het tweede clustercentrum te liggen. Allebei de clustercentrums krijgen nu twee punten. Daarna kun je weer opnieuw de clustercentrums verplaatsen.
Onderzoeker	Kun je me vertellen wanneer dit stopt?
Leerling	Als de clustercentrums precies in het midden liggen van deze twee punten.
Onderzoeker	Het algoritme geeft niet altijd hetzelfde eindantwoord bij deze punten. Heb je een idee waar dat aan zou kunnen liggen.
Leerling	[Denkt na]. Ik zou het eigenlijk niet weten.
Onderzoeker	Je kunt een aantal dingen kiezen. Zou dat van invloed kunnen zijn?
Leerling	Hoe bedoel je kiezen?
Onderzoeker	Een voorbeeld: wat zou er gebeuren als je de clustercentrums op een andere plek zou laten beginnen? [wijst naar het zojuist gebruikte voorbeeld]

Leerling	Als je bijvoorbeeld de twee clustercentrums gelijk bij twee punten hebt liggen, dan groeien ze gelijk naar die twee punten toe en liggen de clustercentrums bij elkaar. In het voorbeeld liggen ze nu iets meer verspreid en omdat ze continu bewegen komen ze uiteindelijk wel eerlijk verdeeld, maar als je ze allebei in de buurt hebt van dezelfde twee punten dan groeien ze daarnaartoe en komen ze op dezelfde locatie.
Onderzoeker	Soms kan is het algoritme snel klaar, maar soms duurt het ook heel lang voordat je kunt stoppen. Dat ligt natuurlijk aan het aantal punten dat er is; je kunt je vast voorstellen dat het snel gaat als er maar 4 punten zijn. Waar kan de snelheid, naast het aantal punten, aan liggen?
Leerling	Dat ligt ook aan het aantal clustercentrums dat je hebt. Omdat die allemaal aan het verplaatsen zijn, en een nieuw centrum krijgen.
Onderzoeker	Kun je nog een tweede reden verzinnen? Bijvoorbeeld iets met hoe de punten liggen.
Leerling	Nou als de punten heel erg verspreid liggen, en je niet echt van die clusters van punten hebt, dan beweegt het veel meer dan als je allemaal van die clusters hebt. Ik weet niet, kan niet goed uitleggen waarom, maar het lijkt me wel logisch.

Interview leerling 3

Onderzoeker	Hoe vond je het om aan de opdracht te werken?
Leerling	Het was anders dan een normale les. We hebben wel vaker een soort werkblad gehad en daar leer je best veel van. Een soort kleine cursus was het.
Onderzoeker	Had je genoeg tijd om alles te maken?
Leerling	Ja, we waren al eerder klaar.
Onderzoeker	Kun je aangeven welke opgaven lastig waren?
Leerling	Opgave 9! Het was een beetje gezakt hoe je alles in moest voeren. We hebben trouwens die dollartekens niet gebruikt. De rest was eigenlijk vrij eenvoudig.
Onderzoeker	Kun je uitleggen wat een algoritme is?
Leerling	Ja. Een soort van patroon, een cirkelredenering die je elke keer volgt, eigenlijk een stappenplan.
Onderzoeker	Kun je een voorbeeld geven van een algoritme?
Leerling	Ja nou, als het goed is ook de stelling van Pythagoras. Dus $a^2 + b^2 = c^2$, gewoon die stapjes.
Onderzoeker	Zou je elke formule een algoritme noemen?
Leerling	[Twijfelt heel erg]. Ja, nee, nou, bij elke formule heb je een stappenplan, dus in principe wel.
Onderzoeker	We hebben het K-means algoritme bekeken. Kun je uitleggen waarvoor dit algoritme dient?
Leerling	Om zo goed mogelijk clusters, van bepaalde voorkeuren bijvoorbeeld, te maken. Je had verschillende punten, van voorkeur van films bijvoorbeeld. Dan wil je zo nauwkeurig mogelijk een cluster vormen bij mensen die dezelfde interesses hebben.
Onderzoeker	We hebben het nu bij films op Netflix gedaan. Kun je nog een ander voorbeeld verzinnen waarbij je dit algoritme kunt gebruiken?
Leerling	Misschien wat vergezocht, maar bijvoorbeeld in de supermarkt. Stel iedereen zou een eigen identificatiecode zou hebben, bijvoorbeeld een bonuskaart. De supermarkt kan dan precies zien wat je koopt en daarop kunnen ze dan het één en ander baseren. Ze maken dan bijvoorbeeld groepen van producten.
Onderzoeker	Waarom willen we eigenlijk clusteren met een algoritme? En niet gewoon met de hand?
Leerling	Soms heb je ook zoiets als de smiley, dan is het sowieso heel lastig. Bij de grafiek uit hoofdstuk 2, heb je eigenlijk ook maar één soort van cluster, alleen dat is niet de bedoeling. Er zit namelijk nog best wat verspreiding in wat mensen leuk vinden. En dan zou dus iedereen hetzelfde aanbevolen krijgen. En dat wil je niet, dus dat is zo een algoritme wel handig.
Onderzoeker	Je zou ook nog kunnen kijken naar hogere dimensies? Zie je dan nog meer nut van dit algoritme?
Leerling	Bij drie-dimensionaal?
Onderzoeker	Ja, of bij vier-, vijf of zesdimensionaal.
Leerling	Ik weet niet hoe vier en vijf zich zouden uiten in beeld. Maar bij drie is het sowieso handiger, want dan heb je ook nog diepte erbij en dat maakt het lastig om te zien.
Onderzoeker	Dan gaan we nu nog even kijken naar een voorbeeld. Ik heb hier vier punten en twee clustercentrum. Kun je mij laten zien hoe je twee clusters zou vormen met het k-means algoritme?
Leerling	Waarmee moet ik rekenen, gewoon uit mijn hoofd?

Onderzoeker	Ja, dat kan wel. Het zijn niet te ingewikkelde berekeningen.
Leerling	[Begint een schets]. Even kijken. Nou, ik kan eigenlijk wel gelijk gaan rekenen. Ik maak een kleine tabel met punten en clustercentrums. Ik moet nu gaan kijken hoever de punten van de clustercentrums afliggen. Dan moet ik hemelsbreed gaan kijken.
Onderzoeker	Weet je nog hoe dat heette?
Leerling	Iets met een E.
Onderzoeker	Misschien weet je de andere afstandmaat nog?
Leerling	Ja, dat was Manhattan. Zoals de straten.
Leerling	Even denken, ik weet niet waarom ik dit nu opeens zo lastig vindt.
Onderzoeker	Hoe bereken je de afstanden?
Leerling	Ah, nou de x-coördinaten min elkaar in het kwadraat en de y-coördinaten min elkaar in het kwadraat. Maar dat is nul hier. Oh, dus de afstand is gewoon 1. Dan is de afstand hier dus wortel 2. Ja, misschien had ik toch die schets moeten maken.
Leerling	Dan ligt die dus het dichtste bij... Ik noem het eerste clustercentrum A en het andere clustercentrum B. 1 is minder dan wortel 2, dus hij ligt dicht bij A.
	[Leerling rekent de andere afstanden correct uit] Ah, ik had het eigenlijk ook wel direct kunnen aflezen, nouja.
Onderzoeker	Goed, en dan?
Leerling	Dan bereken ik de nieuwe gemiddeldes. Dus voor de x van B $(4+5+2)/3$ en hetzelfde doe ik met de y. Cluster A heeft er maar 1 dus die blijft staan.
Onderzoeker	Waar blijft ie staan?
Leerling	Nouja, waar die was. Onee, hij gaat naar (1,4), want er is maar één punt.
Onderzoeker	En dan? Ben je dan klaar?
Leerling	Voor de zekerheid doe je het dan nog een keer hetzelfde, omdat je niet weet of het nog gaat veranderen.
Onderzoeker	Wanneer stop je dan?
Leerling	Als het niet meer veranderd.
Onderzoeker	Je komt, bij dezelfde dataset, niet altijd tot dezelfde eindoplossing. Waar kan dat aan liggen?
Leerling	Hoe bedoel je?
Onderzoeker	Nou, op de website die je bezocht was er na lang klikken niet altijd dezelfde eindoplossing. Waar lag dat aan?
Leerling	Het hangt er vanaf waar je de clustercentrums plaatst!
Onderzoeker	Is dat niet een beetje raar?
Leerling	Nee nou [krabbelt een voorbeeld op een kladblaadje]. Het is logisch dat je het afhangt van waar je begint. Als de clusters alle drie bij elkaar starten, dan is het logisch dat ze niet allemaal verspreiden.
Onderzoeker	Stel je zou een automatische methode moeten bedenken om bij de initialisatiestap de clustercentrums neer te zetten. Wat zou je dan bedenken?
Leerling	Je moet de clustercentrums neerzetten waar de meeste punten zijn.
Onderzoeker	Het algoritme komt soms snel tot een oplossing en soms kan het heel lang duren. Dat ligt natuurlijk liggen aan de hoeveelheid punten. Maar waaraan nog meer?
Leerling	Als de punten veel meer verspreid zijn zal het ook langer duren. Dan gaan de clustercentrums veel meer bewegen.

Interview leerling 4

Onderzoeker	Wat vond je leuk aan de opdracht?
Leerling	Nou ik vond het best wel lastig eigenlijk. Ik kwam er op een gegeven moment niet echt meer uit. Meestal als het lukt vind ik het allemaal leuk om te doen, maar als ik er niet uitkom raak ik gefrustreerd.
Onderzoeker	Hoe was de opgave in vergelijking met het boek?
Leerling	Ik vind het boek leuker, want dat snap ik vaak wel in één keer. Je kunt daar het voorbeeld nadoen. Nu moest ik nadenken en dan is het irritant als ik het niet snap.
Onderzoeker	Kun je aangeven welke opgaven je goed te doen vond?
Leerling	De eerste vier opgaven waren best goed te doen. Daarna werd het lastiger.
Onderzoeker	Welke opgaven waren lastig?
Leerling	Bij opgave 7 had ik wat hulp nodig, want ik snapte het algoritme niet. Nadat u het had uitgelegd snapte ik het wel. Bij opgave 9 snapte ik die formule niet zo goed.
Onderzoeker	Had je genoeg tijd?
Leerling	Ik heb het niet afgekregen. Dat lag niet alleen aan de lengte hoor, ik was ook niet zo gemotiveerd.
Onderzoeker	Kun je mij uitleggen wat een algoritme is?
Leerling	Geen idee eigenlijk. Algoritme? Is dat gewoon met die bolletjes?
	Iets met elke keer dezelfde stappen dat uiteindelijk een uitkomst geeft. Dat je elke keer doorberekent todat er een oplossing is. Elke keer dezelfde berekening met nieuwe resultaten.
Onderzoeker	Kun je misschien een voorbeeld geven van een algoritme?
Leerling	De normaalverdeling bijvoorbeeld. Die gaat elke keer omhoog in het midden en stopt dan ergens.
Onderzoeker	En iets meer praktisch?
Leerling	Opstaan bijvoorbeeld op een school. Wakker worden, aankleden, haar doen, tanden poetsen, onbijten, spullen pakken.
Onderzoeker	Wat is het doel van het k-means algoritme?
Leerling	Het was dat die bolletje steeds naar het clustercentrum gingen. Het doel is dan om te bepalen waar het centrum van de clusters is. En welke punten er dan bij horen.
Onderzoeker	Kun je dan uitleggen hoe dat precies gaat?
Leerling	Je had die clustercentrums. Dan worden de punten opgedeeld in bij welke ze hoorden. En dan werden ze opnieuw verdeeld.
Onderzoeker	Wat werd er dan opnieuw verdeeld?
Leerling	Nou, uhm, de punten. Onee, dan wordt er een gemiddelde genomen waar die clustercentrums kwamen. Daarna werden die punten weer opnieuw opgedeeld. En dan weer het gemiddelde. Uiteindelijk kwam het stil te staan.
Onderzoeker	En dan stop je?
Leerling	Ja dan stop je. Dan heb je de verdeling.
Onderzoeker	Stel je hebt deze vier punten en deze twee clustercentrum [wijst naar het voorbeeld]. Kun je het k-means algoritme uitvoeren bij dit voorbeeld?
Leerling	[Tekent een assenstelsel met daarin punten en clustercentrums]. Dan hoort het meest linker punt bij het eerste clustercentrum en die andere drie bij het tweede clustercentrum.
Onderzoeker	Hoe weet je dat het eerste punt bij het eerste clustercentrum hoort?

Leerling	Nou, omdat de afstand naar eerste clustercentrum 1 is.
Onderzoeker	En waar vergelijk je dat mee?
Leerling	Nou, met deze afstand [Wijst naar de afstand tussen het eerste punt en het tweede clustercentrum] Dat is een schuin blokje. Is dat kleiner?
Onderzoeker	Je hebt geleerd hoe je de afstand kunt berekenen.
Leerling	Met eh, Pythagoras. $a^2+b^2=c^2$. Dus 1^2+1^2 , dat is eh 2^2 . Dus dat is 4. Uhm, nee, dat kan niet. Er klopt iets niet.
Onderzoeker	Ik zal je helpen. [Tekent een driehoek]. Welke zijde gaan we berekenen?
Leerling	De schuine zijde.
Onderzoeker	Dat is dus c^2 .
Leerling	Dus de wortel van 2. Dat is toch 1? Ik weet niet wat de wortel van 2 is.
Onderzoeker	Dat is ongeveer 1,41.
Leerling	Oke. Dan vergelijk ik dus 1 met 1,41. En dat is 1 kleiner dan 1,41.
Onderzoeker	En nu de volgende stap.
Leerling	Het midden van de punten berekenen. Dat is met een berekening: $2+4+5$ is 11. Dan deel ik 11 door 3, want het zijn drie punten. Dat geeft $3\frac{2}{3}$. [wil het clustercentrum bij $(3\frac{2}{3}, 3\frac{2}{3})$ tekenen, maar bedenkt zich dan]. Onee, je moet ook de andere coördinaat doen. Dat geeft $2\frac{1}{3}$. Dus het clustercentrum gaat naar $(3\frac{2}{3}, 2\frac{1}{3})$. Het andere clustercentrum komt op het ene punt te liggen.
Onderzoeker	Kun je me vertellen hoe het daarna verder gaat?
Leerling	Nu gaan weer de punten verspreid worden over de nieuwe clustercentrums waar ze het dichtste bij liggen, dat moet je weer uitrekenen. Dan stopt het wanneer het niet meer zal veranderen.
Onderzoeker	Zullen er in de volgende stap nog punten verplaatsen?
Leerling	Ja, er komen twee punten in het eerste cluster.
Onderzoeker	Soms is het algoritme snel klaar en soms duurt het heel lang. Waaraan kan dat liggen?
Leerling	Ehm, ik denk dat als alle punten heel dicht bij elkaar liggen, het de hele tijd heen en weer gaat. Dan hoort een punt eerst bij het ene cluster en dan weer bij het andere cluster. Als de punten ver uit elkaar liggen en de clustercentrums zijn eenmaal verdeeld, dan liggen ze zo ver uit elkaar dat de punten ook niet meer naar een ander cluster gaan.
Onderzoek	Je kunt kijken naar de snelheid, maar ook naar de uitkomst. Waar kan het aan liggen dat je, bij dezelfde dataset, niet altijd dezelfde oplossing krijgt?
Leerling	Als alles op dezelfde plek begint?
Onderzoek	Nou, dat hoeft niet.
Leerling	Oh, dan als het clustercentrum op een andere plek begint. Dan krijg je een andere uitkomst.
Onderzoek	Krijg je dan elke keer een andere oplossing?
Leerling	Ja, ik denk altijd.