

**Auffassungen und moralische
Wertvorstellungen innerhalb der
Streitkräfte und Zivilbevölkerung in
Bezug auf autonome Waffensysteme:
Eine empirische Untersuchung anhand der
Value-Sensitive-Design-Methode**

von:
Ilse Verdiesen



EuroISME Diplomarbeit des Jahres 2018

[reverse of cover page]
[intentionally blank]

[fly leaf]

**Auffassungen und moralische
Wertvorstellungen innerhalb der
Streitkräfte und Zivilbevölkerung in Bezug
auf autonome Waffensysteme:**

**Eine empirische Untersuchung anhand
der Value-Sensitive-Design-Methode**

von:

Majorin Ilse Verdiesen

Übersetzt von: Gerta K. Badde-Valentine

Umschlagillustration:

© Shutterstock. stockillustration-id: 225928441

[reverse of fly leaf]

Ein Teil dieser Arbeit basiert auf den folgenden Artikeln:

Verdiesen, I., Dignum, V., & Rahwan, I. (2018, September). "Design Requirements for a Moral Machine for Autonomous Weapons," in: *International Conference on Computer Safety, Reliability, and Security* (pp. 494-506). Springer, Cham. https://link.springer.com/chapter/10.1007/978-3-319-99229-7_44 Copyright © Springer Nature Switzerland AG 2018 All rights reserved.

Verdiesen, I., de Sio, F. S., & Dignum, V. (2019). "Moral Values Related to Autonomous Weapon Systems: An Empirical Survey that Reveals Common Ground for the Ethical Debate," *IEEE Technology and Society Magazine*, 38(4), 34-44.

<https://ieeexplore.ieee.org/abstract/document/8924586>

[title page]

**Auffassungen und moralische
Wertvorstellungen innerhalb der
Streitkräfte und Zivilbevölkerung in
Bezug auf autonome Waffensysteme:**

**Eine empirische Untersuchung anhand
der Value-Sensitive-Design-Methode**

von:

Majorin Ilse Verdiesen

Studentin an der:

Technischen Universität Delft,
die Niederlande

[reverse of title page]

Die Jury des EuroISME-Jahrespreises für die beste Diplomarbeit in Militärethik besteht aus:

1. Oberst Hochw. Prof. Dr. P.J. McCormack, MBE, (Präsident, Vereinigtes Königreich)
 2. Oberstleutnant (i.R.) Daniel Beaudoin (Frankreich/Israel)
 3. Dr. Veronika Bock (Deutschland)
 4. MilSup MMag. Stefan Gugerel (Österreich)
 5. Oberst (i.R.) Prof. Dr. Boris Kashnikov (Russland)
 6. Dr. Asta Maskaliūnaitė (Estland)
 7. Prof.Dr. Desiree Verweij (Niederlande)
- Frau Ivana Gošić (Serbien, Sekretariat)

Anfragen: secretariat.ethicsprize@euroisme.eu

www.euroisme.eu

[fly leaf]

Der Preis wird finanziell unterstützt von



[reverse fly leaf – intentionally blank]

Kurzfassung

Autonome Waffen werden immer häufiger auf den Schlachtfeldern benutzt (Roff, 2016). Autonome Waffensysteme können viele Vorteile für den Militärbereich haben, zum Beispiel, wenn die Selbststeueranlage der F-16 einen Absturz verhindert (NOS, 2016), oder der Einsatz von Robotern von der EOD (Kampfmittelbeseitigung), um Bomben zu entschärfen (Carpenter, 2016). Jedoch die Art der autonomen Waffensysteme kann auch zu unkontrollierbaren Handlungen und gesellschaftlichen Unruhen führen. Der Einsatz von autonomen Waffensystemen ohne direkte menschliche Überwachung auf dem Schlachtfeld ist, gemäß Kaag und Kaufman (2009) nicht nur ein Umsturz im Militäreinsatz, sondern kann auch als ein moralischer Umsturz gedeutet werden. Da der Einsatz von AI in großem Umfang auf dem Schlachtfeld unvermeidbar scheint (Rosenberg & Markoff, 2016), ist es von äußerster Wichtigkeit, dass Untersuchungen über die ethische und moralische Verantwortung angestellt werden.

In der Debatte um autonome Waffensysteme werden feste Überzeugungen und Meinungen immer lauter. Die Kampagne, Killerroboter zu stoppen (2017), erklärt zum Beispiel auf ihrer Webseite: *„Die Erlaubnis Maschinen die Entscheidung über Leben und Tod zu überlassen überquert eine fundamentale moralische Grenze. Autonomen Robotern fehlt das menschliche Urteilsvermögen und die Fähigkeit, den Zusammenhang zu begreifen.“* Wir fanden nur wenige empirische Untersuchungen, die diese Ansichten unterstützen oder Einsichten zur Verfügung stellen, wie autonome Waffensysteme von der Zivilbevölkerung und dem Militär aufgefasst werden. Außerdem fanden wir keine empirischen Untersuchungen zu moralischen Werten, die der „fundamentalen moralischen Grenze“ autonomer Waffensysteme zugrunde liegt. Es besteht deshalb eine zweifache Wissenslücke, insofern, dass es an Einsicht mangelt, 1)

wie autonome Waffensysteme vom Militär und der Zivilbevölkerung aufgefasst werden und 2) welche moralischen Werte Menschen als wichtig erachten, wenn autonome Waffensysteme in der näheren Zukunft eingesetzt werden.

Der erste Teil dieser Lücke kann ausgefüllt werden, indem die Auffassung von autonomen Waffensystemen mithilfe der Theorie der Streitkräfte in den Bereichen Kognitive Psychologie, Künstliche Intelligenz und Moralische Philosophie erforscht wird. Der zweite Teil dieser Wissenslücke kann durch die Untersuchung bekannter Werttheorien ausgefüllt werden (Beauchamp & Walters, 1999; Friedman & Kahn Jr., 2003; Schwartz, 2012), um zu sehen welche Werte Menschen im Einsatz von autonomen Waffensystemen als wichtig erachten. Basiert auf der identifizierten Wissenslücke und Problembeschreibung, entwerfen wir die folgende Forschungsfrage für diese Untersuchung:

Wie werden autonome Waffensysteme von der Zivilbevölkerung und den im niederländischen Verteidigungsministerium angestellten Militärangehörigen aufgefasst, und welche moralischen Werte betrachten sie beim Einsatz autonomer Waffensysteme als wichtig?

In dieser Untersuchung benutzen wir die Value-Sensitive-Design-Methode (VSD) als Herangehensweise der Forschung. VSD ist eine dreifache Methode, die menschliche Werte innerhalb des Designprozesses der Technologie berücksichtigt. Es ist ein schrittweises Verfahren für die *konzeptuelle, empirische und technologische* Untersuchung menschlicher Werte, die von dem Design hinzugezogen werden (Davis & Nathan, 2015; Friedman & Kahn Jr., 2003). Die konzeptuelle Untersuchung besteht aus zwei Teilen: (1) Identifizierung der direkten Interessenvertreter, diejenigen, die die Technologie einsetzen, und die indirekten Interessenvertreter, diejenigen, deren Leben von der Technologie beeinflusst wird, und (2) Identifizierung und Definition der Werte,

die der Einsatz der Technologie mit sich bringt. Die empirische Untersuchung befasst sich mit dem Verständnis und der Erfahrung der Interessenvertreter im Zusammenhang mit der Technologie und die hinzugezogenen Werte werden untersucht. In der technischen Untersuchung werden die speziellen Eigenschaften der Technologie analysiert (Davis & Nathan, 2015).

Wir weichen etwas von der ursprünglichen VSD-Methode ab, da wir keine vollständige Interessenvertreteranalyse durchführen, um die Interessenvertreter in Schritt 1 zu analysieren, aber konzentrieren uns auf zwei deutliche Interessenvertretergruppen: die *Zivilbevölkerung* und die *Militärangehörigen*, weil diese beiden Gruppen als befragte Personen für unsere Untersuchung zur Verfügung standen. In Schritt 2 der ersten Phase führen wir eine Literaturlauswertung aus, um die Werte in Bezug auf autonome Waffensysteme zu identifizieren. Der wichtigste Fokus unserer Forschungsarbeit ruht auf der empirischen Untersuchungsphase, da wir fanden, dass der empirische Aspekt in der ethischen Debatte zu autonomen Waffensystemen übersehen wird. In der technischen Untersuchungsphase entwerfen wir keine autonomes Waffensystem, wie man hätte annehmen können, da dies ein immenses Vorhaben gewesen wäre, weit über den Bereich dieser Untersuchung hinaus. Doch wir entschieden uns, auf der Arbeit des Arbeitskreises Scalable Cooperation aufzubauen und schlagen einen Entwurf für eine *Moral Machine* für Autonome Waffensysteme vor, die als ein neuer Schritt in der Untersuchung der Ethik von autonomen Waffensystemen benutzt werden kann.

Die wissenschaftliche Relevanz unserer Untersuchung liegt darin, dass wir zur Fachliteratur beitragen, indem wir Einsicht in die Auffassung der Zivilbevölkerung und dem Militär hinsichtlich autonomer Waffensysteme nehmen und indem wir die moralischen Wertvorstellungen der Zivilbevölkerung und Streitkräfte in Bezug auf autonome Waffen identifizieren. Einsicht

darin fehlt zurzeit, und es konnten keine empirischen Daten zur Erfassung der Werte gefunden werden, die mit autonomen Waffensystemen zu tun haben. Anhand von dem Value-Sensitive-Design als Forschungsmethode weisen wir nach, dass diese Methode zur Strukturierung wissenschaftlicher Forschung angewendet werden kann, was als Fallstudie für die VSD-Methode angesehen werden kann. Wir erweitern die Forschung auch auf die ethische Entscheidungstreffung bei autonomen Fahrzeugen (Bonnefon, Shariff und Rahwan, 2016) und gliedern sie in den Bereich der autonomen Waffensysteme ein.

Die gesellschaftliche Relevanz ist das Verständnis der Auffassung von autonomen Waffensystemen der Zivilbevölkerung und den Militärangehörigen, die im niederländischen Verteidigungsministerium tätig sind und identifizieren, welche moralischen Werte sie mit autonomen Waffensystemen in Verbindung bringen, um Gemeinsamkeiten und Unterschiede in der Debatte dieser Technologie zu finden, die von der Kampagne Stop Killer Robots (2015) und dem internationalen Komitee für Roboter-Rüstungskontrolle (ICRAC) eingeleitet wurde. Zweitens, zeigt diese Untersuchung auf, wie die Value-Sensitive-Design-Methode für autonome Waffensysteme benutzt werden kann, um die Werte des Militärs und der Zivilbevölkerung im Verhältnis zum Einsatz dieser Art von Waffen zu identifizieren. Zum Schluss, ist es wichtig, durch Identifizieren der Werte sie in das Design autonomer Waffensysteme einzufügen; somit trägt diese Untersuchung zu einem verantwortungsvollen Design und Einsatz autonomer Waffensysteme in der Zukunft bei.

Diese Untersuchung konzentriert sich auf die empirische Untersuchungsphase von VSD. Diese Phase ist in zwei Teile aufgeteilt. Der erste Teil befasst sich damit, die Werte im Kontext autonomer Waffensysteme mittels einer Online-Befragung zu identifizieren, die von Gesprächen mit sechs Experten unterstützt wird. Der zweite Teil der empirischen Untersuchung prüft die

Wahrnehmung der Streitkräfte mittels einer informatorischen und einer bestätigenden Untersuchungsmethode. Wir operationalisieren die menschlichen Verhaltensweisen, um unsere Hypothese zur Auffassung der Streitkräfte zu bestätigen, indem wir kontrollierte Experimente randomisiert anordnen und die Ergebnisse des Experiments benutzen, um die Werte in Verbindung mit autonomen Waffensystemen auf beschreibende Weise zu untersuchen.

Um die Auffassung der Streitkräfte von autonomen Waffen zu untersuchen, entwarfen wir eine Hypothese auf der Basis, dass die Militärangehörigen autonome Waffensysteme als Waffen wie alle anderen betrachten und deshalb nichts weiter als ein Werkzeug, mit dem man eine Wirkung erzeugt. Wir stellen am Schluss der Untersuchung die Hypothese auf, dass gemäß der Auffassung der Militärangehörigen autonome Waffen keinen psychischen Zustand kennen. Wir entwarfen drei Szenarien für unser Experiment: 1) eine von Menschen betriebene Drohne, 2) ein autonomes Waffensystem ohne menschliche Einwirkung und 3) ein autonomes Waffensystem mit menschlicher Einwirkung. Wir hatten keinen Unterschied zwischen dem Zustand des neutralen autonomen Waffensystems und dem Zustand, wenn eine Drohne von Menschen ferngesteuert wird, erwartet. Wir erwarten ebenfalls, dass der neutrale Zustand des Systems bedeutend anders ist, als der Zustand des Systems mit hoher Selbstbestimmung, wobei wir den Teilnehmern besonders sagen, dass das autonome Waffensystem selbstbestimmende Charakteristiken hat, z.B. die Fähigkeit zu planen und eigene Ziele zu setzen. Gemäß dieser Argumentation, ist unsere Hypothese wie folgt:

H1: Militärangehörige fassen autonome Waffen nicht als mit psychischen Fähigkeiten ausgestattet auf.

Die Ergebnisse der ersten und letzten Untersuchung zeigen auf, dass Operationselemente zuverlässig sind und als ein Konstrukt zusammenhalten. Das Agency-Konstrukt besteht aus vier Elementen: *Denken, Zielsetzung, Freier Wille und das Erreichen von Zielen*, die zum Messen der Auffassung von Streitkräften in Bezug auf autonome Waffensysteme und auf von Menschen gesteuerte Drohnen benutzt werden können. Unsere zentrale Analyse befasste sich damit, wie neutrale Szenarien der autonomen Waffensysteme sich von den durch Einwirkung von Menschen und den hochentwickelten Szenarien von autonomen Waffen unterscheiden. Unsere Ergebnisse zeigen auf, dass die Auffassung von Streitkräften der Militärangehörigen und zivilen Mitarbeitern im niederländischen Verteidigungsministerium für das neutrale autonome Waffensystem-Szenario höher ist als die Auffassung der Streitkräfte von dem durch menschliche Einwirkung betriebene Drohnen-Szenario. Das bedeutet, dass sie einem autonomen Waffensystem mehr Macht beimessen, als einer durch menschliche Einwirkung betriebenen Drohne. Basiert auf diesen Erkenntnissen, müssen wir unsere Hypothese verwerfen.

Die Auswirkung der Wahrnehmung von Agency auf die abhängigen Variablen wird auf beschreibende Weise erforscht und nicht durch Regressionsrechnung analysiert. Wir fanden, dass 7 von 9 abhängigen Variablen bedeutend mit dem Agency-Konstrukt korrelieren. Dies sind die Variablen: *Vertrauen, Menschenwürde, Zuversicht, Erwartungen, Unterstützung, Fairness und Besorgnis*. Dies sind unsere Erkenntnisse:

- Militärangehörige und zivile Mitarbeiter im niederländischen Verteidigungsministerium haben mehr Vertrauen, sind zuversichtlicher und unterstützen eher die Maßnahmen, die von Drohnen unter menschlicher Einwirkung getroffen wurden, als die von autonomen Waffensystemen;
- Es herrscht die Auffassung, dass eine unter menschlicher Einwirkung betriebene Drohne die Menschenwürde mehr

achtet, als neutrale oder hochentwickelte autonome Waffensysteme, obwohl die Maßnahmen und Ergebnisse der Szenarien die gleichen sind;

- Militärangehörige und zivile Mitarbeiter im niederländischen Verteidigungsministerium befinden sich auf einer gleichwertigen Ebene der Erwartungen mit Hinsicht auf die Maßnahmen der unter menschlicher Einwirkung betriebene Drohne und der neutralen und hochentwickelten autonomen Waffensysteme in der Zukunft. Außerdem betrachten sie die Handlungen der von Menschen gesteuerten Drohnen und Autonomen Waffensystemen als gleich fair;
- Autonome Waffensysteme verursachen mehr Besorgnis unter den Militärangehörigen und zivilen Mitarbeitern im niederländischen Verteidigungsministerium als von Menschen betriebene Waffensysteme.

Diese Untersuchung wurde mit einem kleinen Datensatz durchgeführt, und um die Ergebnisse zu verallgemeinern benötigt diese Arbeit weitere Personen zur Befragung, die eine größere demografische Gruppe vertreten. Deshalb schlagen wir ein Design für eine Moral Machine für Autonome Waffensysteme als ein Massives Online-Experiment (MOE) für eine großangelegte Untersuchung dieses Themas vor. Wir bauten das Konzept der Moral Machine in unser Design ein, das von dem Arbeitskreis Scalable Cooperation im Medienlabor bei MIT (Scalable Cooperation Group, 2016) entwickelt wurde. Da das Thema äußerst heikel ist, schlagen wir eine Ausführung in zwei Schritten vor, zuerst planen wir ein kontrolliertes Experiment mit einem begrenzten Satz Bedingungen. Diese Bedingungen können in Szenarien eingesetzt werden, was Benutzern ermöglicht, die Umfrage nach Erhalt eines Passwortes über eine Web-Oberfläche an einen sicheren Server zu senden. Nach dem Sammeln der ersten Daten und des Benutzer-Feedbacks ist der nächste Schritt, eine

großangelegte offene Plattform wie z.B. die Moral Machine maßstäblich zu vergrößern, wo Leute die Szenarien bewerten können, um große Mengen Daten von verschiedenen demografischen Gruppen zu sammeln, die dann für robustere und verallgemeinerbare Ergebnisse benutzt werden können.

Verschiedene Aspekte können als Einschränkungen dieser Forschungsarbeit identifiziert werden. Zuerst wurden die Operationalisierung und das Agency-Konstrukt von einer Kategorisierung der Literatur abgeleitet, die die Agentenmerkmale beschreibt, und unsere Wahl der Merkmale basiert auf einer numerischen Zahl und nicht auf Relevanz oder Bewertungskriterien. Die zweite Einschränkung ist die Auswahl der Werte aus dem Werte-Fragebogen und den Experteninterviews als abhängige Variable. Obwohl diese Auswahl heftig unter den drei Rechercheuren dieser Arbeit debattiert wurde, fand die endgültige Wahl auf methodischem Weg und nicht nach objektiver Methode statt. Drittens sind die in dieser Forschungsarbeit benutzten Proben sowohl in Größe als auch demografischem Bereich begrenzt, da die endgültige Studie nur über 239 befragte Personen verfügte, also nur ein geringer Anteil des niederländischen Verteidigungsministeriums. Zum Schluss ist die Verteilung der befragten Personen des letzten Szenarios verzerrt und das zweite Szenario (neutrales Agens autonomes Waffensystem) verfügte über 32 befragte Personen mehr als das erste Szenario (unter menschlicher Einwirkung betriebene Drohnen). Eine der Erklärungen für diese Schieflage könnte sein, dass Leute erwarteten, autonome Waffensysteme zu untersuchen, aber machten nicht mehr mit, als ihnen die unter menschlicher Einwirkung betriebenen Waffensysteme präsentiert wurden. Dies nennt man selektive Schwundquote, die eine negative Auswirkung auf die interne Stichhaltigkeit der Forschungsarbeit hat. Eine weitere Erklärung könnte sein, dass die Software bei der Aufteilung der Befragten über die Szenarien fehlerhaft war, was zu

bedeuten hätte, dass die interne Stichhaltigkeit dieser Forschungsarbeit nicht davon verletzt wurde, da die Schiefecke nicht einer selektiven Schwundquote beizumessen ist.

Mit Rücksicht auf diese Einschränkungen werden mehrere Empfehlungen für weitere Forschungsarbeiten vorgeschlagen. Die erste ist, das Agency-Konstrukt zu bewerten, das weitere Untersuchungen zur Messung der Auffassung in Bezug auf Kampfmittel anderer technologischen Artefakten erforderlich sind, um zu prüfen, dass dieses Konstrukt auch anderen Bereichen standhält. Beispiele dieser Forschungsarbeiten könnten die Untersuchung der Auffassung in Bezug auf Dienstmittel einschließen, von Pflegerobotern für Senioren, KI-Spielzeug für Kinder oder Onboard-Computern von autonomen Fahrzeugen. Die zweite Empfehlung ist, die abschließende Untersuchung mit den gleichen Szenarien mit einer repräsentativen Stichprobe durchzuführen, die einzig aus Mitgliedern der niederländischen Zivilbevölkerung besteht. Dies würde uns ermöglichen zu erkennen, bei welchen Werten sich die Ergebnisse der militärischen Probanden und die Antworten von Probanden aus der Zivilbevölkerung voneinander unterscheiden, und obwohl wir keine direkten Vergleiche ziehen können, aufgrund der Tatsache, dass die militärische Stichprobe nicht repräsentativ ist, könnten wir doch Einsicht in die Auffassung von sowohl Militärangehörigen als auch der Zivilbevölkerung nehmen. Schließlich empfehlen wir die Implementierung der Moral Machine für Autonome Waffensysteme, wie in Abschnitt 5 beschrieben wird, um die Ergebnisse zu verallgemeinern. Hierfür erfordert diese Untersuchung entschieden mehr Personen zur Befragung, die eine umfassendere demografische Gruppe vertreten. Die Korrektur nach oben für ein Massives Online-Experiment wie die Moral Machine, würde große Mengen Daten verschiedener demografischer Gruppen erzeugen, die für robustere und verallgemeinerbare Ergebnisse benutzt werden könnten, um ein

gründliches Verständnis der moralischen Bewertung der Leute in Bezug auf autonome Waffensysteme zu erhalten.

Vorwort

Mein Interesse an der Ethik autonomer Waffensysteme wurde erregt, als ich an der IDEA-Summer School zum Thema Verantwortungsbewusste Künstliche Intelligenz (KI) im August 2016 teilnahm, die von Dr. Virginia Dignum als Teil der Europäischen Konferenz zu Künstlicher Intelligenz (ECAI) organisiert wurde. Nach dem ersten Treffen mit Virginia bezüglich meiner Diplomarbeit fragte sie mich, ob ich im Ausland studieren möchte und wo ich gerne hingehen würde, und ich war kühn genug, zu sagen, dass ich gerne zum MIT (Massachusetts Institute of Technology) gehen würde, aber ich glaubte auch nicht im Entferntesten, dass dies wirklich geschehen würde. Virginias Kontakt zu Dr. Iyad Rahwan des Arbeitskreises Scalable Cooperation des Media Lab von MIT half, es zu realisieren. Ich möchte Virginia ganz herzlich danken, dass sie mir diesen Kontakt vermittelte, sowie auch für ihre Betreuung während meiner Diplomarbeit. Aufgrund der Zeitunterschiede kommunizierten wir meistens des Abends in den Niederlanden über Skype, und ich schätzte es sehr, dass sie immer Zeit fand, um meine Niederschriften zu überprüfen.

Die Arbeit meiner Diplomarbeit mit dem Arbeitskreis Scalable Cooperation Group war eine einmalige Chance, und das Media Lab ist sehr offen, sehr divers, nicht hierarchisch und eine äußerst inspirierende Umgebung. Zuerst musste ich mich daran gewöhnen, im Foyer einem Kugelbad zu begegnen, oder an Leute, die ihre Büros verließen, um mitten am Tag Pingpong zu spielen, aber nach einem Monat oder so fühlte sich das Labor sehr heimisch an. Ich arbeitete mit hoch talentierten und ehrgeizigen Studenten im Aufbaustudium und mit Post-Doktoranden. Ich möchte ganz besonders Dr. Sydney Levine und Dr. Nick Obradowitsch für all ihre einsichtsvollen Kommentare danken,

und ich war wirklich begeistert von den heftigen Debatten, die wir in meinem Arbeitskreis hatten. Vor allem aber möchte ich Dr. Iyad Rahwan dafür danken, dass er mir gestattete seiner Gruppe beizutreten und mit dem heiklen Thema der autonomen Waffensysteme zu arbeiten. Er stellte mir alle Ressourcen, seine Zeit und Hilfe zur Verfügung, die ich brauchte, um meine Recherchen zu machen, was ich sehr zu schätzen weiß.

Zuletzt möchte ich mich noch bei der Königlichen Niederländischen Armee (RLNA) bedanken, die es mir ermöglichte, im Ausland zu studieren und meinen Aufenthalt am MIT unterstützte. Ohne diese Unterstützung wäre ich nicht in der Lage gewesen, fünf Monate lang in Boston zu wohnen und zu studieren. Ich hoffe, dass diese Diplomarbeit zur Ethik von autonomen Waffensystemen zum besseren Verständnis führt und eine Diskussion zum Einsatz von autonomen Waffensystemen und verantwortungsvoller KI in der Verteidigungsorganisation einleitet.

Ilse Verdiesen

Inhalt

Kurzfassung	ix
Vorwort	xix
Inhalt	xxiii
Abbildungsverzeichnis	xxviii
Tabellenverzeichnis	xxxi

1.	Einführung	1
1.1.	Forschungsproblem	2
1.1.1.	Problemerkforschung	3
1.1.2.	Wissenslücke	6
1.1.3.	Problemstellung	7
1.1.4.	Umfang	8
1.1.5.	Forschungsfrage und daraus folgende Fragen	9
1.2.	Forschungsmethode	10
1.3.	Relevanz	15
1.3.1.	Wissenschaftliche Relevanz	15
1.3.2.	Gesellschaftliche Relevanz	15
1.3.3.	Einbettung in das SEPAM- Kurrikulum and dem IA-Kurs	16
1.4.	Struktur	17
2.	Literaturübersicht	19
2.1.	Autonome Waffensysteme	19
2.1.1.	Definition	19
2.1.2.	Klassifizierung der autonomen Waffensysteme	22
2.2.	Werte	25
2.2.1.	Definition	25
2.2.2.	Allgemeingültige Werte	27

2.2.3.	Werte mit Bezug auf autonome Waffensysteme	37
2.2.4.	Werte-Hierarchie	43
2.3.	Agency	47
2.3.1.	Agency in der Kognitionpsychologie	48
2.3.2.	Agency in der künstlichen Intelligenz	49
2.3.3.	Agency in der Moralphilosophie	49
3.	Methode	53
3.1.	Methodologie	53
3.1.1.	Fachliteraturübersicht	53
3.1.2.	Online-Wertumfrage	54
3.1.3.	Interviews mit Experten	55
3.1.4.	Codierungsverfahren	56
3.1.5.	Randomisierte kontrollierte Experimente	57
3.2.	Hypothesen	58
3.3.	Forschungsdesign	59
3.4.	Operationalisierung	60
3.4.1.	Szenarien	61
3.4.2.	Agency-Konstrukt	64
3.4.3.	Abhängige Variablen	65
3.4.4.	Aufmerksamkeitstest	67
3.4.5.	Demografische Variablen	67
3.5.	Analytische Methode	67
3.5.1.	Daten vor der Verarbeitung	67
3.5.2.	Reliabilitätsanalyse	68
3.5.3.	Hauptkomponentenanalyse (PCA)	69
3.5.4.	Korrelationsanalyse	69
3.5.5.	Manipulationsprüfung des Agency	70
3.5.6.	Analyse der abhängigen Variablen	70
3.6.	Voranmeldung	71
3.7.	Stichprobe	71

3.8.	Methodologische Fragen	72
3.8.1.	Interviews zur Codierung	73
3.8.2.	Randomisierte kontrollierte Experimente	73
3.8.3.	Amazon Mechanical Turk	74
3.8.4.	Die abschließende Untersuchung nach dem Schneeballverfahren	75
4.	Ergebnisse	77
4.1.	Werteumfrage	77
4.1.1.	Online-Umfrage	77
4.1.2.	Interviews	81
4.1.3.	Schlussfolgerung der Umfrage	82
4.2.	Vorstudie 1	84
4.2.1.	Reliabilitätsanalyse	85
4.2.2.	Hauptkomponentenanalyse (PCA)	86
4.2.3.	Korrelationsanalyse	87
4.2.4.	Manipulationsprüfung der Agenten	89
4.2.5.	Analyse der abhängigen Variablen	91
4.2.6.	Schlussfolgerung der Vorstudie 1	97
4.3.	Vorstudie 2	99
4.3.1.	Zuverlässigkeitsanalyse	101
4.3.2.	Hauptkomponentenanalyse (PCA)	101
4.3.3.	Korrelationsanalyse	102
4.3.4.	Manipulationsprüfung der Agenten	104
4.3.5.	Analyse der abhängigen Variablen	105
4.3.6.	Schlussfolgerung der Vorstudie 2	118
4.4.	Abschließende Untersuchung – militärische Stichprobe	119
4.4.1.	Reliabilitätsanalyse	121
4.4.2.	Hauptkomponentenanalyse (PCA)	122
4.4.3.	Korrelationsanalyse des Agency-Konstrukts	123
4.4.4.	Manipulationsprüfung des Agenten	124
		xxv

4.4.5.	Analyse der abhängigen Variablen	129
4.4.6.	Schlussfolgerung der abschließenden Untersuchung	135

5. Design der Moral Machine für Autonome Waffensysteme 139

5.1.	Moral Machine für Autonome Fahrzeuge	140
5.2.	Moral Machine für Autonome Waffensysteme	141
5.3.	Szenarien	143
5.3.1.	Variablen für Szenarien	143
5.3.2.	Beispielszenarien	148
5.4.	Implementierung	150

6. Schlussfolgerung und Diskussion 152

6.1.	Schlussfolgerung	152
6.1.1.	Agency-Konstrukt	153
6.1.2.	Zentrale Hypothese Auffassung in Bezug auf Waffensysteme	153
6.1.3.	Untersuchung der abhängigen Variablen	156
6.2.	Diskussion	159
6.2.1.	Wissenschaftliche Folgerungen	159
6.2.2.	Gesellschaftliche Folgerungen	161
6.3.	Einschränkungen	163
6.4.	Empfehlungen für weitere Forschungsarbeiten	165

Quellenverzeichnis 167

Anhang A.	Online-Fragebogen zu den Werten	176
Anhang B.	Fragebogen für die abschließende Untersuchung	180
Anhang C.	Szenarien für Vorstudie 1	187
Anhang D.	Szenarien für Vorstudie 2	194
Anhang E.	Szenarien für die abschließende Untersuchung	205

Anhang F. Transkriptionen der Interviews	207
Anhang G. Codierungsnotizen	251
Anhang H. Ergebnisse des Codierungsprozesses	257

Abbildungsverzeichnis

Abbildung 1 Die Value-Sensitive Designmethode	14
Abbildung 2 Klassifizierung der autonomen Waffensysteme, basiert auf Royakkers und Orbons (2015)	23
Abbildung 3 Werthierarchie für den Wert von Rechenschaft im Design von autonomen Waffensystemen	46
Abbildung 4 Forschungsdesign	60
Abbildung 5 Ergebnisfragen 1 Online-Wertbesprechung	78
Abbildung 6 Ergebnisfragen 3 Online-Wertbesprechung	79
Abbildung 7 Ergebnisfragen 2 Online-Wertbesprechung	80
Abbildung 8 Scree-Test PCA Vorstudie 1	86
Abbildung 9 Mittelwert des Agency-Konstrukts Waffentyp	89
Abbildung 10 Mittelwert des Agency-Konstrukts pro Zustand pro Resultat.	90
Abbildung 11 Mittelwert der Unterstützungsvariable pro Zustand pro Resultat	92
Abbildung 12 Mittelwert der Unterstützungsvariable pro Waffentyp	93
Abbildung 13 Mittelwert der Vertrauensvariable pro Waffentyp pro Resultat	94
Abbildung 14 Mittelwert der Vertrauensvariable pro Zustand pro Resultat	95
Abbildung 15 Mittelwert von Schuld und Kommandantenschuld-Variablen pro Zustand und Waffentyp	96
Abbildung 16 Mittelwert von Schaden und Kommandantenschaden-Variable pro Zustand und Waffentyp	97
Abbildung 17 Scree-Test PCA Vorstudie 2	102
Abbildung 18 Mittelwert des Wert-Agency-Konstrukts pro Zustand pro Resultat	104
Abbildung 19 Mittelwert der Schuldvariable pro Zustand pro Resultat	106

Abbildung 20 Mittelwert von Kommandantenschuld-Variable pro Zustand pro Resultat	107
Abbildung 21 Mittelwert der Vertrauensvariable pro Zustand pro Resultat	108
Abbildung 22 Mittelwert der Kommandantenvertrauensvariable pro Zustand pro Resultat	109
Abbildung 23 Mittelwert der Schuldvariable pro Zustand pro Resultat	110
Abbildung 24 Mittelwert der Kommandantenschuldvariable pro Zustand pro Resultat	111
Abbildung 25 Mittelwert der Kommandantenschuldvariable pro Zustand pro Resultat	112
Abbildung 26 Mittelwert der Vertrauensvariable pro Zustand pro Resultat	113
Abbildung 27 Mittelwert der Werterwartungsvariable pro Zustand pro Resultat	114
Abbildung 28 Mittelwert der Werterwartungsvariable pro Zustand pro Resultat	115
Abbildung 29 Mittelwert der Fairnessvariable pro Zustand pro Resultat	116
Abbildung 30 Mittelwert der Unbehagenvariable pro Zustand pro Resultat	117
Abbildung 31 Scree-Test PCA abschließende Untersuchung	123
Abbildung 32 Mittelwert des Agens pro Zustand	126
Abbildung 33 Mittelwert des Agens pro Zustand pro Gruppe	126
Abbildung 34 Mittelwert der Vertrauensvariable pro Zustand	129
Abbildung 35 Mittelwert der Menschenwürdevariable pro Zustand	130
Abbildung 36 Mittelwert der Zuversichtsvariable pro Zustand	131
Abbildung 37 Mittelwert der Werterwartungsvariable pro Zustand	132
Abbildung 38 Mittelwert der Unterstützungsvariable pro Zustand	133
Abbildung 39 Mittelwert der Fairnessvariable pro Zustand	134
Abbildung 40 Mittelwert der Besorgnisvariable pro Zustand	135

Abbildung 41 Waffentypvariable $W = \{\text{von Menschen gesteuerte Drohne, autonome Drohne}\}$	145
Abbildung 42 Lagevariable $L = \{\text{Wüste, Dorf}\}$	145
Abbildung 43 Personenvariable $C = \{\text{Mann, Frau, Kind}\}$	146
Abbildung 44 Anzahl Personenvariable $N = \{1..5\}$	147
Abbildung 45 Resultatvariable $O = \{\text{Kollateralschaden, kein Kollateralschaden}\}$	147
Abbildung 46 Einsatzvariable $M = \{\text{verteidigen, angreifen}\}$	148
Abbildung 47 Beispielszenario 1	149
Abbildung 48 Beispielszenario 2	149

Tabellenverzeichnis

Tabelle 1	Übersicht der Definition von autonomen Waffensystemen	21
Tabelle 2	Übersicht der Werttheorien	31
Tabelle 3	Werte mit Bezug auf Waffensysteme	39
Tabelle 4	Übersicht der Agentenmerkmale	50
Tabelle 5	Anzahl der Agentenmerkmale, die in der Literatur erwähnt werden	63
Tabelle 6	Übersicht der Werte aus Online-Untersuchungen und Interviews	83
Tabelle 7	Szenarien der Vorstudie 1	84
Tabelle 8	Ergebnisse der Reliabilitätsanalyse Agency-Items Vorstudie 1	85
Tabelle 9	Ergebnis-PCA Agency-Items Vorstudie 1	86
Tabelle 10	Korrelationsmatrix des Agency-Konstrukts und abhängigen Variablen	88
Tabelle 11	Szenarien der Vorstudie 2	100
Tabelle 12	Ergebnisse der Reliabilitätsanalyse Vorstudie 2	101
Tabelle 13	Ergebnisse-PCA Vorstudie 2	102
Tabelle 14	Korrelationsmatrix des Agency-Konstrukts und abhängigen Variablen Vorstudie 2	103
Tabelle 15	Szenarien der abschließenden Untersuchung	120
Tabelle 16	Ergebnisse der Reliabilitätsanalyse Agency-Items	122
Tabelle 17	PCA-Ergebnisse der Agenten in der abschließenden Untersuchung	123
Tabelle 18	Korrealisation-Agency-Konstrukt mit abhängigen Variablen für die abschließende Untersuchung	128

1. Einführung

Ich glaube ich finde es schwierig auszudrücken, was ich von Drohnen halte. Einerseits macht mich jedes Instrument zum Töten traurig. Andererseits ist es tröstlich, dass aufgrund ihres Einsatzes sehr viele Leben nicht riskiert werden brauchen.

Vorstudie 2 Befragte Personen

Autonome Waffensysteme sind Waffensysteme, die mit künstlicher Intelligenz ausgerüstet sind (KI). Künstliche Intelligenz werden von Neapolitan und Jiang (2012, S. 8) als „*eine intelligente Instanz, die in einer stets wechselnden, komplexen Umgebung schlussfolgert*“ beschrieben, aber diese Definition passt auch auf natürliche Intelligenz. Russell, Norvig and Intelligence (1995) erstellt eine Übersicht von vielen Definitionen, worin sie Ansichten zu Systemen kombinieren, die *wie Menschen denken und handeln und zugleich wie Systeme, die rational denken und handeln können*, aber sie stellen keine eigene Definition da. Für jetzt bleiben wir bei der Beschreibung, die Bryson, Kime und Zürich (2011) erstellen. Sie geben an, dass eine Maschine (oder ein System) intelligentes Verhalten aufweist, wenn sie eine Handlung wählen kann, die auf der Beobachtung seiner Umgebung basiert ist. In der wissenschaftlichen Literatur wird KI als mehr als nur ein Intelligentes System beschrieben. Dies wird von den Konzepten *Anpassungsfähigkeit, Interaktives Handeln und Autonomie* charakterisiert (Floridi & Sanders, 2004). Laut Floridi und Sanders (2004) bedeutet Anpassungsfähigkeit, dass das System sich gemäß seiner Interaktionen ändern und aus seinen Erfahrungen lernen kann. Programmiermethoden von Maschinen sind ein Beispiel dafür. Wir sprechen von Interaktivität, wenn das System und seine Umgebung auf einander eingehen und Autonomie bedeutet, dass das System selbst einen Zustand ändern kann.

Eine wachsende Anzahl Rechercheure konzentrieren sich auf ein verantwortliches Design von KI, dass soziale und ethische Werte einschließt, um unerwünschte gesellschaftliche Folgen dieser Technologie zu verhindern. Prinzipien, um verantwortungsvolle KI zu beschreiben sind *Rechenschaftspflicht, Verantwortung und Transparenz* (ART). *Rechenschaftspflicht* bezieht sich auf die Rechtfertigung der Handlungen, die von KI ausgeführt werden, *Verantwortung* räumt für die Fähigkeit ein, Schuld für diese Handlungen auf sich zu nehmen und *Transparenz* befasst sich mit der Beschreibung und Reproduzierung der Entscheidungen, die KI macht und passt sie ihrer Umgebung an (V. Dignum 2016).

Künstliche Intelligenz ist nicht nur ein futuristisches Szenario aus einem Science-Fiction-Film, in dem menschenartige Roboter planen, die Macht über die Welt zu bekommen, wie Datas Bruder Lore in Star Trek oder die Cylons in Battlestar Galactica. Viele KI-Anwendungen sind heute bereits im Einsatz. Intelligente Strommesser, Suchmaschinen, Personal Assistent bei Mobiltelefonen, Selbststeueranlagen und selbstfahrende Autos sind Beispiele dafür. Diese Untersuchung konzentriert sich auf die Anwendung von KI im Militärbereich, und besonders auf den Einsatz von autonomen Waffen in der nahen Zukunft. In dieser Einführung beschreiben wir zuerst das Forschungsproblem, gefolgt von der Forschungsmethode, der wissenschaftlichen und gesellschaftlichen Relevanz und der Einbettung in das Kurrikulum der Informationsarchitektur (IA) im Rahmen des System Engineering Policy und Management (SEPAM)-Magisterkurses. Wir werden mit einer kurzen Beschreibung der nächsten Kapitel dieses Berichts abschließen.

1.1 Forschungsproblem

Zuerst besprechen wir den vermehrten Einsatz autonomer Waffensysteme im Militärbereich und die damit zusammenhängende Arbeit zur Ethik von KI, gefolgt von Forschungsarbeiten zu von Menschen angetriebenen Drohnen im

Unterabschnitt über Problemerkforschung. Als nächstes identifizieren wir die Wissenslücke, Problemstellung und den Umfang, bevor wir mit der Forschungsfrage und den Unterabschnitten dieser Forschungsarbeit abschließen.

1.1.1 Problemerkforschung

Autonome Waffen werden immer häufiger auf den Schlachtfeldern benutzt (Roff, 2016). Es wurde bereits berichtet, dass China über autonome Fahrzeuge verfügt, die einen bewaffneten Roboter mit sich führen (Lin & Singer, 2014), Russland behauptet, dass sie an autonomen Panzern arbeitet (W. Stewart, 2015), die USA taufen ihr erstes „selbststeuerndes“ Kriegsschiff im Mai 2016 (P. Stewart, 2016), und der russische Waffenhersteller Kalaschnikow gab vor kurzem kund, dass er ein vollständig automatisiertes Gefechtsmodul entwickelt hat, das neutrale Netzwerke benutzt (RT, 2017). Autonome Waffensysteme können viele Vorteile für den Militärbereich haben, zum Beispiel, wenn die Selbststeueranlage der F-16 einen Absturz verhindert (NOS, 2016), oder der Einsatz von Robotern von der EOD (Kampfmittelbeseitigung), um Bomben zu entschärfen (Carpenter, 2016). Jedoch die Art der autonomen Waffensysteme kann auch zu unkontrollierbaren Handlungen und gesellschaftlichen Unruhen führen. Beispiele dieser Unruhen sind die „Stop Killer Robots“-Kampagne von 61 Nichtregierungsorganisationen, die vom Human Rights Watch angeführt wird (Campaign to Stop Killer Robots, 2015), aber auch die Vereinten Nationen, in deren Namen ihre Besorgnisse ausgesprochen werden und aussagen, dass *„Autonome Waffensysteme, die keine bedeutungsvolle menschliche Kontrolle erfordern, verboten werden sollten und ferngesteuerte Streitkraft immer nur mit äußerster Vorsicht eingesetzt werden sollte“* (UN-Generalversammlung, 2016). Der Einsatz von autonomen Waffensystemen ohne direkte menschliche Überwachung auf dem Schlachtfeld ist, gemäß Kaag und Kaufman (2009) nicht nur ein Umsturz im Militäreinsatz, sondern kann auch als ein moralischer Umsturz gedeutet werden.

Da der Einsatz von KI in großem Umfang auf dem Schlachtfeld unvermeidbar scheint (Rosenberg & Markoff, 2016), ist es von äußerster Wichtigkeit, dass Untersuchungen über die ethische und moralische Verantwortung angestellt werden.

Ethische Entscheidungsfällung in KI und Robotern ist ein aufkommender Bereich, und eine Reihe von Akademikern befasst sich mit moralischer Beurteilung dieser Technologien. Zum Beispiel schlägt Malle (2015) ein System vor, in dem (bislang) die verschiedenen Gebiete der *Roboterethik* kombiniert werden. Das bedeutet, dass ethische Fragen über Design, Einsatz und Behandlung von Robotern von Menschen angesprochen werden, sowie *die Maschinenmoralität*, die sich mit Fragen über die moralischen Fähigkeiten eines Roboters beschäftigt und wie diese mit Computerhilfe eingebaut werden können. Cointe, Bonnet und Boissier (2016) schlagen ein Modell vor, bei dem ein Agens die ethischen Aspekte seines eigenen Verhaltens beurteilen kann, sowie das anderer Agenten in einem Multiagentensystem. Das Modell beschreibt einen ethischen Urteilsfindungsprozess, der es Agenten gestattet, das Verhalten anderer Agenten zu bewerten. Bonnefon et al. (2016) haben die ethische Entscheidung eines autonomen Fahrzeugs untersucht, ob es sich selbst schützen oder utilitaristisch verfahren soll, wenn es Fußgängern auf der Straße begegnet. In dieser Arbeit wird die *Moral Machine*¹ am MIT benutzt, um Einsicht zu nehmen, wie Leute Szenarien mit einem autonomen Fahrzeug beurteilen, um zu sehen wie ihr moralisches Urteil mit dem von anderen Leuten verglichen werden kann.

In einem den autonomen Waffensystemen verwandten Bereich, dem der von Menschen gesteuerten Drohnen, wurden ethische Probleme während der letzten zehn Jahre sehr intensiv untersucht, sowohl in der philosophischen als auch psychologischen Fachliteratur. Vom philosophischen Standpunkt behauptet Coeckelberg (2013), dass Drohneneinsätze nicht nur

¹ <http://moralmachine.mit.edu/>

eine physische Distanz sondern auch eine moralische Distanz erzeugen, da der Gegner weniger sichtbar ist. Dies eliminiert die moralisch-psychologische Hemmschwelle für das Töten. Ein weiteres ethisches Problem ist nach Strawser (2010), dass aufgrund der regelmäßigen Arbeitsschichten außerhalb der Gefechtszone, die Drohnenbediener eine ungerechte psychologische Last mit sich tragen, da sie täglich zwischen Arbeits- und Heimsituation schalten müssen. Aufgrund der Distanz zum Schlachtfeld, können menschliche Streitkräfte auch eine kognitive Dissonanz empfinden, sodass der Krieg mehr wie ein Computerspiel erscheint, als die Wirklichkeit. Diese ethischen Probleme werden von Strawser (2010) widerlegt, indem er angibt, dass aufgrund der erhöhten Distanz, die menschlichen Streitkräfte mehr Zeit haben, um ein Ziel zu beurteilen, da ihre eigene Sicherheit nicht auf dem Spiel steht, aber er argumentiert auch, dass weitere empirische Forschung erforderlich ist, um die psychologischen Auswirkungen des Ausführens von Drohneneinsätzen mit großen Distanzen zum Schlachtfeld zu bewerten.

Mehrere empirische Untersuchungen der Auswirkungen von Drohneneinsätzen auf menschliche Streitkräfte wurden im Bereich der Psychologie durchgeführt. Eine der ersten Untersuchungen der psychologischen Auswirkungen von Drohneneinsätzen wurde von Thompson *et al.* (2006) angestellt, der herausfand, dass Drohnenbediener an stärkerer Müdigkeit, emotionaler Erschöpfung und Burnout litten. Dies wurde von Chappelle, Goodman, Reardon und Thompson (2014) bestätigt, die berichteten, dass menschliche Drohnenbediener Symptome von Burnout aufweisen, emotional erschöpft sind und zu einem hohen Grad Zynismus sowie auch posttraumatische Stresssyndroms an den Tag legen. Weitere psychologische Faktoren, die aus diesen Untersuchungen bekannt wurden, sind Langeweile und Verstörtheit aufgrund der langen Dauer der Drohneneinsätze (Cummings, Mastracchio, Thornburg & Mkrtchyan, 2013). Diese Untersuchung zeigen, dass die

Ausführung von Drohneneinsätzen schwere psychologische Auswirkungen auf die Menschen, die sie bedienen, haben kann.

1.1.2 Wissenslücke

In der Debatte um autonome Waffensysteme werden feste Überzeugungen und Meinungen immer lauter. Die Kampagne, Killerroboter zu stoppen (2017) erklärt zum Beispiel auf ihrer Webseite, dass: *„Die Erlaubnis, Maschinen die Entscheidung über Leben und Tod zu überlassen, überquert eine fundamentale moralische Grenze. Autonomen Robotern fehlt das menschliche Urteilsvermögen und die Fähigkeit, den Zusammenhang zu begreifen.“* Wir fanden nur wenige empirische Untersuchungen, die diese Ansichten unterstützen oder Einsichten zur Verfügung stellen, wie autonome Waffensysteme von der Zivilbevölkerung und dem Militär wahrgenommen werden. Die Initiative *Open Robots Ethics* begutachtete die öffentliche Meinung in einer Meinungsumfrage in 2015 (*Open Robot Ethics Initiative*, 2015) und veröffentlichte einen Bericht. Die Ergebnisse wurden jedoch nicht in einer akademischen Fachzeitschrift veröffentlicht, und die Untersuchung war nicht umfassend genug, um stichhaltige Folgerungen daraus zu schließen. Wie bereits in den Unterabschnitten beschrieben, untersuchten verschiedentliche Wissenschaftler wie Malle (2015), Cointe *et al.* 2016 und Bonnefon *et al.* (2016) ethische Entscheidungstreffung mit Bezug auf KI und Roboter. Ethische Berücksichtigung wird auch in den benachbarten Bereichen der durch menschliche Einwirkung betriebenen Operationen Coeckelberg, 2013; Strawser, 2010) untersucht, und empirische Studien fanden, dass menschliche Beteiligte unter schweren psychologischen Auswirkungen leiden (Chappelle *et al.*, 2014; Cummings *et al.*, 2013; Thompson *et al.*, 2006), aber diese Forschungen sind noch nicht auf den Einsatz von autonomen Waffensystemen erweitert. Außerdem fanden wir keine Fachliteratur oder empirische Untersuchungen zu den moralischen Werten, die mit autonomen Waffensystemen im Zusammenhang stehen oder was Leute als die „moralische

Grundlinie“ erachten. Es besteht deshalb eine zweifache Wissenslücke, insofern, dass es an Einsicht mangelt, 1) wie autonome Waffensysteme vom Militär und der Zivilbevölkerung wahrgenommen werden und 2) welche moralischen Werte Menschen als wichtig erachten, wenn autonome Waffensysteme in der näheren Zukunft eingesetzt werden.

1.1.3 Problemstellung

Der erste Teil dieser Lücke kann ausgefüllt werden, indem die Wahrnehmung von autonomen Waffensystemen mithilfe der Agententheorie in den Bereichen Kognitive Psychologie, Künstliche Intelligenz und Moralische Philosophie erforscht wird. Agency ist die Fähigkeit, zu planen und zu handeln (Waytz, Gray, Epley, & Wegner, 2010), und wird sowohl in menschlichen als auch nicht-menschlichen Wesen oder Gegenständen untersucht. Leute messen Handlungsfähigkeit Gegenständen wie Computer (Nass, Moon, Fogg, Reeves, & Dryer, 1995) und Robotern bei, (H. M. Gray, Gray, & Wegner, 2007), deshalb erwarten wir von der KI-Technologie, dass sie auch Merkmale von Handlungsfähigkeit aufweist. Die Untersuchung von der Auffassung unter den Militärangehörigen und der Zivilbevölkerung, dass autonome Waffensysteme Handlungsfähigkeit aufweisen, würde einen Einblick in die möglichen Unterschiede und Ähnlichkeiten zwischen den Ansichten dieser Gruppen geben, die in der derzeitigen Debatte zu dieser Technologie benutzt werden könnte.

Der zweite Teil dieser Wissenslücke kann durch die Untersuchung bekannter Werttheorien ausgefüllt werden, um zu sehen welche Werte Menschen im Einsatz von autonomen Waffensystemen als wichtig erachten. Bekannte Werttheorien sind die von Schwartz (1994), Friedman und Kahn Jr. (2003), und Beauchamp und Walters (1999), aber ein Einblick, wie sie sich zu autonomen Waffensystemen verhalten, fehlt noch. Die Beziehung der Werte, die im Kontext des Einsatzes von autonomen Waffensystemen am relevantesten sind und der Vergleich dieser

Werte zu denen, die im Verhältnis zu derzeitigen Technologie der von Menschen bedienten Drohnen stehen, würde zu einem Einblick in die eigentlichen Motive in der Diskussion um autonome Waffensysteme und besserem Verständnis der ausgedrückten Ansichten führen. Dies führt zu der folgenden Problemstellung für diese Untersuchung:

In der Debatte zu autonomen Waffensystemen werden feste Vorstellungen und Meinungen geäußert, aber eine empirische Untersuchung, die diese Meinungen unterstützt, fehlt. Ein Einblick inwiefern autonome Waffensysteme wahrgenommen und welche moralische Werte dem zukünftigen Einsatz von autonomen Waffensystemen beizumessen sind, ist nicht vorhanden. Ebenfalls ist nicht offensichtlich, ob die Auffassung und die moralischen Werte bezüglich autonomer Waffensysteme von Militärangehörigen sich von den denen der Zivilbevölkerung unterscheiden.

1.1.4 Umfang

Ein Großteil der Fachliteratur zu autonomen Waffensystemen wird von Juristen und Philosophen im Zusammenhang mit den Internationalen Menschenrechten und der Genfer Konvention geschrieben und besprochen, die die Auswirkungen von bewaffneten Auseinandersetzungen begrenzen wollen (ICRC, 2010). Da wir weder Juristen noch Philosophen sind, bleibt diese Untersuchung innerhalb der Schranken und Regeln der Kriegsgesetze, und wir werden sie nicht in Frage stellen. Aufgrund der beschränkten Zeit befasst sich diese Untersuchung hauptsächlich mit dem Militärpersonal des niederländischen Verteidigungsministeriums und wird nicht auf eine Stichprobe der niederländischen Zivilbevölkerung erweitert. In dieser Untersuchung konzentrieren wir uns auf den Einsatz autonomer Waffensysteme in der nahen Zukunft, die wir als *in den nächsten 5 Jahren* definieren. Das bedeutet, dass wir keine Waffen, die mit künstlicher Intelligenz oder futuristischer Technologie ausgerüstet sind und noch nicht gebaut werden können, betrachten, sondern wir untersuchen eine Technologie, die derzeit entwickelt wird,

vor allem Drohnen mit der autonomen Fähigkeit, Ziele zu setzen. Dieses Zielsetzungsverfahren besteht aus sechs Stufen: (1) Finden, (2) Feststellen, (3) Verfolgen, (4) Zielen, (5) Angreifen und (6) Bewerten. Wir werden uns auf die Entscheidungstreffung in Stufe 4 und 5 des Zielprozesses konzentrieren, da viele ethische Entscheidungen in diesem Teil des Verfahrens zu treffen sind (Asaro, 2016).

1.1.5 Forschungsfrage und daraus resultierenden weiteren Fragen

Basiert auf der identifizierten Wissenslücke und Problemstellung, entwerfen wir die folgende Forschungsfrage für diese Untersuchung:

Wie werden autonome Waffensysteme von der Zivilbevölkerung und den im niederländischen Verteidigungsministerium angestellten Militärangehörigen aufgefasst und welche moralischen Werte betrachten sie als wichtig beim Einsatz autonomer Waffensysteme?

Diese Forschungsfrage wird von den sich daraus ergebenden Fragen beantwortet:

- 1. Wie werden autonome Waffensysteme in der Fachliteratur definiert?*
- 2. Welche Werttheorien werden in der Literatur beschrieben?*
- 3. Welche der in den Werttheorien genannten Werte stehen im Zusammenhang mit autonomen Waffensystemen?*
- 4. Wie wird die Auffassung von Handlungsfähigkeit in der Literatur erwähnt?*

Diese vier sich ergebenden Fragen werden in der Literaturbesprechung unserer Forschung beantwortet.

- 5. Welche moralischen Werte im Verhältnis zum Einsatz von autonomen Waffensystemen nehmen die Zivilbevölkerung und das Militärpersonal des niederländischen Verteidigungsministeriums wahr?*

Diese Frage wird von einer Überprüfung von Werten beantwortet, die aus einem Online-Fragebogen und Interviews mit Experten besteht.

6. *Wie wird die Handlungsfähigkeit autonomer Waffensysteme von der Zivilbevölkerung und dem Militärpersonal im niederländischen Verteidigungsministerium aufgefasst?*
7. *Wie werden die Werte im Verhältnis zu autonomen Waffensystemen von der Zivilbevölkerung und dem Militärpersonal im niederländischen Verteidigungsministerium aufgefasst?*

Diese beiden Fragen werden durch Prüfen einer Hypothese und beschriebenen Analyse beantwortet, die auf den Ergebnissen eines randomisierten kontrollierten Experiments beruht.

8. *Wie können die Werte, die autonomen Waffensystemen eigen sind, in den Entwurf einer Moral Machine Autonomer Waffensysteme aufgenommen werden?*

Diese Unterfragen können durch Erzeugen eines Entwurfs- und Implementationsplans für eine Moral Machine für Autonome Waffensysteme beantwortet werden.

Im nächsten Abschnitt wird die Untersuchungsmethode zur Beantwortung dieser Unterfragen genauer beschrieben.

1.2. Forschungsmethode

In dieser Untersuchung benutzen wir die *Value-Sensitive-Design-Methode* (VSD) als Herangehensweise der Forschung. VSD ist eine dreifache Methode, die menschliche Werte innerhalb des Entwurfsprozesses der Technologie berücksichtigt. Es ist ein schrittweises Verfahren für die *konzeptuelle, empirische und technologische* Untersuchung menschlicher Werte, die der Entwurf mit sich bringt (Davis & Nathan, 2015; Friedman & Kahn Jr., 2003). Die konzeptuelle Untersuchung besteht aus zwei Teilen: (1) Identifizierung der direkten Interessenvertreter, die die Technologie einsetzen und die indirekten Interessenvertreter, deren Leben von der Technologie beeinflusst wird, und (2) Identifizierung und Definition der Werte, die der Einsatz der Technologie mit sich bringt. Die empirische Untersuchung befasst sich mit dem Verständnis und der Erfahrung der Interessenvertreter im Zusammenhang mit der Technologie, und

die hinzugezogenen Werte werden untersucht. In der technischen Untersuchung werden die speziellen Eigenschaften der Technologie analysiert (Davis & Nathan, 2015). Das VSD kann als eine Roadmap für Ingenieure und Studenten benutzt werden, um ethische Berücksichtigung in den Entwurf aufzunehmen (Cummings, 2006).

Es wurden aber auch Kritiken mit Hinsicht auf die VSD-Methode laut. Eine dieser Besorgnisse, die Davis und Nathan (2015) erwähnen, ist, dass das VSD voraussetzt, dass gewisse Werte zwar universal sind, aber dass sie je nach Kultur und Kontext unterschiedlich sein können. Eine andere Stellungnahme wäre, eine empirische Grundlage für eine Ansicht statt einer philosophischen anzunehmen, oder anzuerkennen, dass die Position der Wissenschaftler nicht als einzige in Betracht kommt (Borning & Muller, 2012). Borning und Muller (2012) nehmen eine pluralistische Haltung ein, zu dem Maße, dass das VSD weder als universale noch als relative Einstellung zu Werten zu empfehlen ist, sondern man sollte es den Ingenieuren überlassen zu entscheiden, welche Ansicht die Geeignetste im Zusammenhang mit ihrem Design ist.

Zusammen mit Borning und Muller (2012) benutzten wir die VSD-Methode in unserer Forschung als eine Richtlinie und als ein Ziel an sich. In der konzeptuellen Phase weichen wir etwas von der ursprünglichen VSD-Methode ab, da wir keine vollständige Interessenvertreteranalyse durchführen, um die Interessenvertreter in Stufe 1 zu analysieren, aber konzentrieren uns auf zwei offensichtliche Interessenvertretergruppen: die *Zivilbevölkerung* und die *Militärangehörigen*, weil diese beiden Gruppen als befragte Personen für unsere Untersuchung zur Verfügung standen. In Stufe 2 der ersten Phase führen wir eine Literaturlauswertung aus, um die Werte in Bezug auf autonome Waffensysteme zu identifizieren. Der wichtigste Fokus unserer Forschungsarbeit ruht auf der empirischen Untersuchungsphase, da wir fanden, dass der empirische Aspekt in der ethischen Debatte zu autonomen Waffensystemen

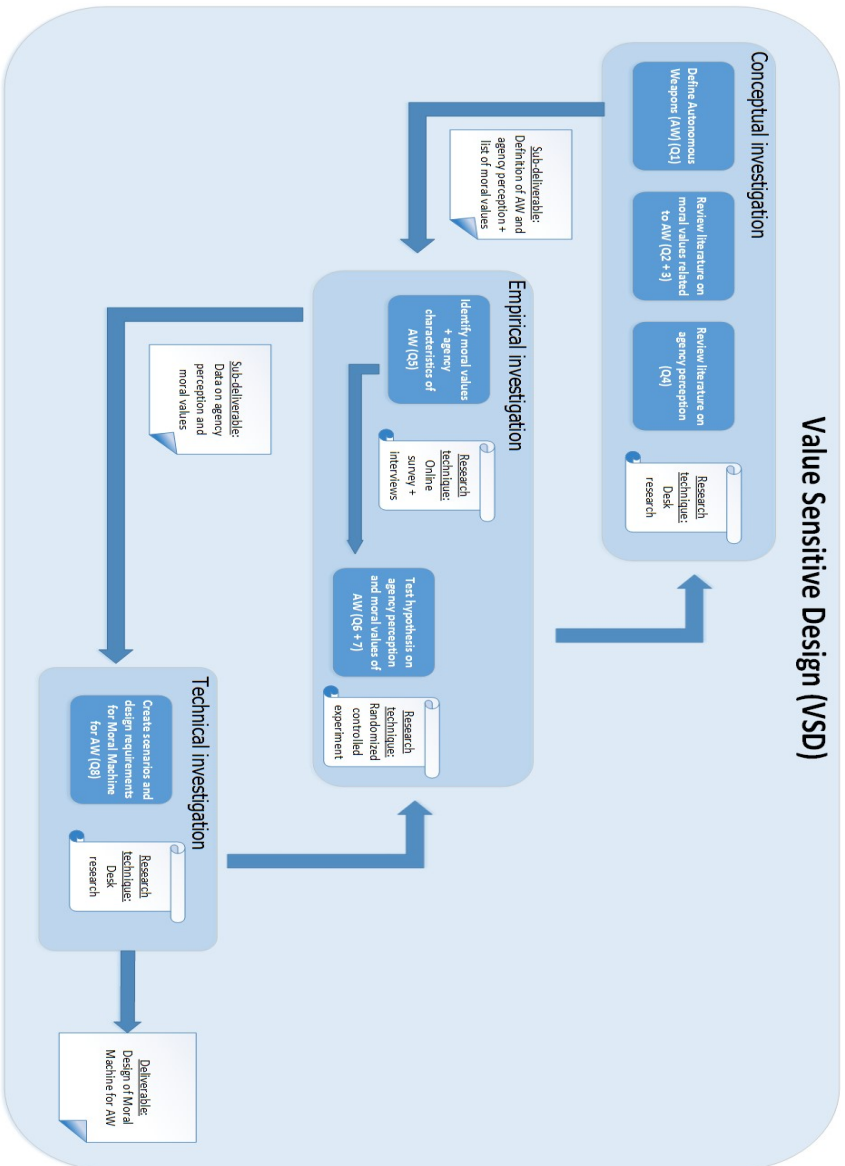
übersehen wird. Wir besprachen die Kritik von Davis und Nathan (2015), indem wir einen großen Teil unserer Zeit dieser Phase des VSDs widmeten, in der wir den Kontext der autonomen Waffensysteme mithilfe verschiedener Forschungsmethoden untersuchten. In der technischen Untersuchungsphase entwerfen wir kein autonomes Waffensystem, wie man hätte annehmen können, da dies ein immenses Vorhaben gewesen wäre, der den Umfang dieser Forschungsarbeit weit übersteigt. Doch wir wählten, auf der Arbeit der *Scalable Cooperation*-Forschungsgruppe aufzubauen, und schlugen eine Konzipierung einer Moral Machine für autonome Waffensysteme vor, die als nächste Stufe in den Untersuchungen zur Ethik autonomer Waffensysteme benutzt werden kann. Wir verwendeten die Phasen des VSD für unsere Untersuchungen wie folgt (Abbildung 1):

- **Konzeptionelle Untersuchung:** Der Mangel an Forschungsarbeiten zu autonomen Waffensystemen deutet darauf hin, dass die erste Phase unserer Untersuchung an sich erforschend ist. Die Forschungsarbeiten in der konzeptionellen Untersuchung bestehen aus der qualitativen Forschungsmethode der „Schreibtischforschung“, um die Fachliteratur über autonome Waffensysteme, moralische Werte und Auffassung der Handlungsfähigkeit zu besprechen.
- **Empirische Untersuchung:** Diese Phase ist in zwei Teile aufgeteilt. Der erste Teil ist informatorisch, um die Wertvorstellungen im Kontext autonomer Waffensysteme mittels einer Online-Befragung zu identifizieren. Dies ist eine quantitative Forschungsmethode, die noch durch Interviews mit Experten von einer qualitativen Forschungsmethode unterstützt wird. Im zweiten Teil der empirischen Untersuchung verwenden wir sowohl eine informatorische als auch eine theorieüberprüfende Forschungsmethode, um die Auffassung innerhalb der Streitkräfte zu studieren. Wir operationalisieren die menschlichen Verhaltensweisen, um unsere Hypothese zur Auffassung innerhalb der Streitkräfte

zu bestätigen, indem wir kontrollierte Experimente randomisiert anordnen und die Ergebnisse des Experiments benutzen, um die Wertvorstellungen in Verbindung mit autonomen Waffensystemen auf beschreibende Weise zu erforschen.

- **Technische Untersuchung:** In dieser Phase werden Konzipierung und Eigenschaften für die Moral Machine der autonomen Waffensysteme erzeugt. Sie basieren auf den Szenarien, die bei dem willkürlichen kontrollierten Experiment der empirischen Phase benutzt werden. Die technische Untersuchungsphase ist informatorisch und das Design wird mithilfe der „Schreibtischforschung“ als Forschungsmethode erzeugt.

Abbildung 1. Die Value-Sensitive-Designmethode



1.3 Relevanz

In diesem Unterabschnitt beschreiben wir die wissenschaftliche und gesellschaftliche Relevanz und zeigen auf, wie die Arbeit in dem Kurrikulum des IA-Kurses des SEPAM-Masterprogramms eingebettet ist.

1.3.1 Wissenschaftliche Relevanz

Die wissenschaftliche Relevanz unserer Untersuchung liegt darin, dass wir zu der akademischen Literatur beitragen, indem wir Einsicht in die Auffassung innerhalb der Zivilbevölkerung und den Streitkräften in Bezug auf autonome Waffensysteme nehmen, und indem wir ihre moralischen Wertvorstellungen in Bezug auf autonome Waffensysteme identifizieren. Einsicht darin fehlt zurzeit, und es konnten keine empirischen Daten zur Auffassung und den Wertvorstellungen in Bezug auf autonome Waffensysteme gefunden werden. Anhand des Value-Sensitive-Design als Forschungsmethode weisen wir nach, dass diese Methode zur Strukturierung akademischer Forschung anwendbar ist und dies als Fallstudie für die VSD-Methode angesehen werden kann. Wir erweitern die Forschung auch mit der ethischen Entscheidungsfällung bei autonomen Fahrzeugen von Bonnefon *et. al.* (2016) auf den Bereich der autonomen Waffensysteme.

1.3.2 Gesellschaftliche Relevanz

Die gesellschaftliche Relevanz ist das Verständnis der Auffassung von autonomen Waffensystemen der Zivilbevölkerung und Streitkräfte, die im niederländischen Verteidigungsministerium tätig sind und identifizieren, welche moralischen Wertvorstellungen sie mit autonomen Waffensystemen in Verbindung bringen, um Gemeinsamkeiten und Unterschiede in der Diskussion zu dieser Technologie zu finden, die von der Kampagne *Stop Killer Robots* (2015) und dem internationalen Kom-

tee für Roboter-Rüstungskontrolle (ICRAC) eingeleitet wurde. Zweitens zeigt diese Forschungsarbeit auf, wie die Value-Sensitive-Designmethode für autonome Waffensysteme benutzt werden kann, um die Wertvorstellungen der Streitkräfte und Zivilbevölkerung im Verhältnis zum Einsatz dieser Art von Waffen zu identifizieren. Zum Schluss werden durch Identifizieren der wichtigen moralischen Werte, diese in das Design autonomer Waffensysteme eingefügt, und somit trägt diese Forschungsarbeit zu einem verantwortungsvollen Design und zukünftigem Einsatz autonomer Waffensysteme bei.

1.3.3 Einbettung in das SEPAM- Kurrikulum und den IA-Kurs

Die allgemeinen Kriterien für eine Magisterarbeit an der Fakultät für Technik, Politik und Unternehmensführung sind die Beinhaltung einer analytischen Komponente und Fokussierung auf einen technischen Bereich und sind fachübergreifend. Als ein Projekt für die SEPAM - Magisterarbeit ist es erforderlich, dass die Arbeit eine Lösung für ein kompliziertes, umfassendes, zeitnahes, soziotechnisches Problem entwirft. Das bedeutet, dass die Arbeit sich auf einen unmissverständlich technischen Bereich in einem Multi-Level-System konzentriert.

Es berücksichtigt sowohl öffentliche als auch private Wertvorstellungen und adressiert nicht nur technische Themen, sondern auch solche, die sich mit Führungs- und ethischen Wahlen befassen (Graduation Portal, 2017). Der Kurs für Informationsarchitektur (IA) schließt Aspekte der Unternehmensführung und Informatik ein und fokussiert auf die Anpassung der organisatorischen Bedürfnisse und technischen Möglichkeiten der State-of-the-Art ICT-Lösungen (IA-Programm, 2017).

Diese Arbeit ist den SEPAM-Kriterien angepasst, da sie sich besonders auf die AI-Technologie von autonomen Waf-

fensystemen im militärischen Bereich konzentriert. Sie ist multidisziplinär, da sie eine Fachliteratur zu Auffassungen und Wertvorstellungen in Bezug auf autonome Waffensysteme in den Bereichen der kognitiven Psychologie, Künstlichen Intelligenz und Moralischen Philosophie, Beschreibungen der Wertvorstellungen in der philosophischen und psychologischen Literatur und den Einsatz von Waffensystemen im militärischen Bereich kombiniert. Das in dieser Arbeit zentrale sozio-technische Problem ist der Mangel an Einblicke in die Auffassung in Bezug auf autonome Waffentechnologie, und die damit verbundenen Wertvorstellungen innerhalb der Zivilbevölkerung und Streitkräfte scheinen sehr wichtig. Diese Arbeit passt in den IA-Kurs, da sie die Gemeinsamkeiten und Unterschiede der ethischen Präferenzen identifiziert, sowohl seitens der Zivilbevölkerung als auch des Militärs und schlägt eine technische Lösung vor, indem sie eine Moral Machine für Autonome Waffensysteme konzipiert, um diese Präferenzen mithilfe eines massiven Online-Experiments genauer zu untersuchen und dabei auf der Moral Machine für Autonome Fahrzeuge aufbaut.

1.4 Struktur

Der Rest dieses Berichts ist wie folgt strukturiert: In Abschnitt 2 wird die Fachliteratur zu autonomen Waffensystemen, Wertvorstellungen und Auffassungen in Bezug auf Waffensysteme besprochen. Abschnitt 3 beschreibt die Methodologie, Hypothesen, das Forschungsdesign, die Operationalisierung der Szenarien und Konstrukte, die analytische Herangehensweise und Stichproben und schließt mit methodologischen Themen ab. In Abschnitt 4 werden zuerst die Ergebnisse der Erhebung der Werte verzeichnet, die aus dem Online-Fragebogen und den Interviews mit Experten besteht, gefolgt von den Ergebnissen der drei randomisierten Experimente. Die Konzipierung der Moral Machine für

Autonome Waffensysteme wird in Abschnitt 5 beschrieben. Dafür wird zuerst die ursprüngliche Moral Machine für Autonome Fahrzeuge vorgestellt, gefolgt von den Merkmalen und Szenarien, um eine Moral Machine Autonomer Waffensysteme und einen kurzen Umsetzungsplan einzuschließen. Abschnitt 6 schließt mit den Ergebnissen ab, bespricht die wissenschaftlichen und gesellschaftlichen Auswirkungen und identifiziert die Einschränkungen und Empfehlungen für weitere Forschungen. Im letzten Abschnitt 7 betrachten wir die Wahlen, die wir für unsere Forschung getroffen haben, den Prozess des Projekts und die Verbindung der Forschungsarbeit, dem IA-Kurs und dem SEPAM-Kurrikulum.

2. Literaturübersicht

Es ist einfach so, dass ich mich unwohl fühle, Waffen in ein System zu integrieren, das weder Gefühle hat, noch von Menschen gesteuert wird.

Vorstudie 1 Befragte

Dieser Abschnitt erstellt eine Übersicht der Literatur zu autonomen Waffensystemen, gefolgt von einer Zusammenfassung von Wertvorstellungen und zuletzt beschreiben wir das Konzept von Agenten, die wir benutzen, um die Auffassung zu autonomen Waffensystemen zu untersuchen. Jeder Unterabsatz schließt mit einer Besprechung der Theorien ab, die wir in dieser Arbeit verwenden, sowie die Gründe für ihre Auswahl.

2.1 *Autonome Waffensysteme*

Obwohl die Debatte um autonome Waffensysteme in den letzten Jahren stark zugenommen hat, fanden wir dieses Thema in der Fachliteratur nicht besonders gut dargestellt. Wir beginnen diesen Unterabsatz mit einer Übersicht der vielen Definitionen und bieten zwei Klassifizierungen von autonomen Waffensystemen als Konklusion zu diesem Abschnitt an.

2.1.1 *Definition*

Autonome Waffensysteme sind eine aufkommende Technologie, und es gibt immer noch keine international vereinbarte Definition (AIV & CAVV, 2016). Selbst ein Konsens, ob autonome Waffensysteme definiert werden sollten, ist nicht vorhanden. Obwohl einige Wissenschaftler in ihren Arbeiten Definitionen erstellen (Tabelle 1), warnen andere vor solch

einer Festlegung. NATO erklärt: „*Der Versuch, Definitionen für ‘autonome Systeme’ sollte vermieden werden, da Maschinen nicht wirklich autonom sein können.*” (Kuptel & Williams, 2014, S.10). Das Institut für Abrüstung der Vereinten Nationen (UNDIR) ist ebenfalls vorsichtig bezüglich einer Erstellung einer Definition von autonomen Waffensystemen, da sie behaupten, dass das Niveau der Autonomie von den „*kritischen Funktionen der Bedenken und die Beeinflussung verschiedener Variablen abhängt*” (UNDIR, 2014, S. 5). Sie stellen fest, dass einer der Gründe für die unterschiedlichen Begriffe mit Hinsicht auf autonome Waffensysteme ist, dass manchmal Objekte (Drohnen oder Roboter) definiert werden, aber ein andermal eine Charakteristik (Autonomie), Variablen der Bedenken (Letalität oder Grad der menschlichen Kontrolle) oder Art der Benutzung (Zielsetzung oder defensive Maßnahmen) in die Diskussion einbezogen und Teil der Definition werden. Die verschiedenen Definitionen von autonomen Waffensystemen sind in Tabelle 1 verzeichnet. Einige Autoren benutzen den Begriff Militärroboter, die über ein gewisses Niveau der Autonomie verfügen. Da Militärroboter gemäß der Klassifizierung von Royakkers und Orbons (2015) (Abbildung 2) als eine Unterklasse von autonomen Waffensystemen angesehen werden können, haben wir sie in die Liste der Definitionen aufgenommen. Unserer Meinung nach, erfasst der Bericht der Advisory Council on International Affairs (AIV & CAVV) die Beschreibung von autonomen Waffensystemen vom technischen und militärischen Standpunkt aus am besten, da sie vordefinierte Kriterien berücksichtigen und mit dem militärischen Zielsetzungsverfahren verknüpft sind, da die Waffe nur nach einer von Menschen bewirkten Entscheidung eingesetzt wird. Deshalb benutzen wir diese Definition und definieren autonome Waffensysteme wie folgt:

„Ein Waffensystem, das ohne menschliche Einwirkung Ziele für gewisse vordefinierte Kriterien wählt und angreift, als Ergebnis einer von Menschen bewirkten Entscheidung, diese Waffe einzusetzen, wahlwissend, dass, einmal abgeschossen, es nicht möglich ist, sie durch menschliche Einwirkung aufzuhalten.“ AIV und CAVV (2016, S. 11)

Autor(en)	Definition
AIV und CAW (2016, S. 11)	„Ein Waffensystem, das, ohne menschliche Einwirkung, Ziele für gewisse vordefinierte Kriterien wählt und angreift, als Ergebnis einer von Menschen bewirkten Entscheidung, diese Waffe einzusetzen, wahlwissend, dass, einmal abgeschossen, es nicht möglich ist, sie durch menschliche Einwirkung aufzuhalten.“
Altmann, Asaro, Sharkey, und Sparrow (2013, S. 73)	Autonome Waffensysteme sind: „...robotische Waffen, die, einmal abgeschossen, ohne weitere menschliche Einwirkung Ziele wählen und angreifen.“
Galliot (2015, S. 5)	Militärroboter sind: „eine Gruppe in Kraft gesetzter elektromechanischer Systeme, die alle eines gemeinsam haben: 1. Sie werden nicht von Menschenhand gesteuert; 2. Sie sind konzipiert, zurückholbar (obwohl sie vielleicht nicht so benutzt werden, dass dies in Frage kommt), und sind 3. in einem militärischen Kontext in der Lage, ihre Kraft auszuüben, um eine tödliche oder nicht tödliche Ladung abzuschießen oder sonst eine Funktion auszuführen, um die Ziele der Streitkräfte zu unterstützen.“
Horowitz (2016, S. 27)	„Ein Waffensystem, einmal aktiviert, ist dazu da, nur individuelle Ziele oder spezifische Zielgruppen anzugreifen, die von einem menschlichen Bediener ausgewählt wurden.“
Royakkers und Orbons (2015, S. 625)	Militärroboter sind wiederverwendbare unbemannte Systeme für Militärzwecke mit einem Grad von Autonomie.“
Kupitel und Williams (2014, S. 10)	„Maschinen sind nur „autonom“ mit Hinsicht auf gewisse Funktionen, wie Steuern, Sensoroptimierung oder Kraftstoffüberwachung.“
UNDIR (2014, S. 5)	Der Grad der Autonomie hängt von den kritischen Funktionen der Bedenken und Beeinflussung verschiedener Variablen ab.“

Tabelle 1 Übersicht der Definition von autonomen Waffensystemen

2.1.2 Klassifizierung der autonomen Waffensysteme

Autonome Waffensysteme sind nicht nur zweideutig definiert, sie sind außerdem nicht gleichförmig klassifiziert worden. Wir stellen zwei dieser Klassifizierungen in diesem Unterabsatz dar. Royakkers und Orbons (2015) beschreiben verschiedene Typen autonomer Waffensysteme (Abbildung 2). Sie unterscheiden zwischen (1) *nicht-tödlichen Waffen* d.h. *Waffen, „...die keine (unschuldigen) Opfer verursachen oder Menschen schwere oder dauerhafte Verletzungen zufügen.“* (Royakkers & Orbons, 2015, S. 617), wie z.B. das Active-Denial-System, eine Strahlenwaffe mit elektromagnetischer Energie, um Menschen von einem Objekt oder von Soldaten fernzuhalten und (2) Militärroboter, die sie wie folgt definieren: *„...wiederverwendbare unbemannte Systeme für Militärzwecke mit einem Grad von Autonomie.“* (Royakkers & Orbons, 2015 S.625). Militärroboter können in drei Kategorien aufgeteilt werden: am Boden angesetzte Fahrzeuge für unbemannte Beobachtung und um Straßenbomben wegzuräumen, Fahrzeuge die unbemannt auf oder unter Wasser navigieren können, wie eine Waffenstation auf einem Schiff oder ein autonomes Unterseeboot, sowie auch Flugzeuge, die als unbemannte Kampfflugzeuge (UCAVs) agieren. Diese UCAVs werden von Royakkers & Orbons (2015) als ferngesteuert klassifiziert. Von ihnen sind „Drohnen“ die bekanntesten, sowie die autonomen UCAVs, die nach und nach vom amerikanischen Verteidigungsministerium entwickelt werden (Rosenberg & Markoff, 2016).

Galliot (2015) stellt eine weitere Klassifizierung der autonomen Waffensysteme zur Verfügung. Sie basieren auf vier Stufen der Autonomie für unbemannte Fahrzeuge:

AUTONOMIESTUFE 1 – NICHT-AUTONOM/FERNGESTEUERT:
„Ein menschlicher Operateur steuert jede einzelne angetriebene Bewegung

der unbemannten Plattform. Ohne den Bediener sind ferngesteuerte Systeme nicht fähig, effektiv betrieben zu werden.“

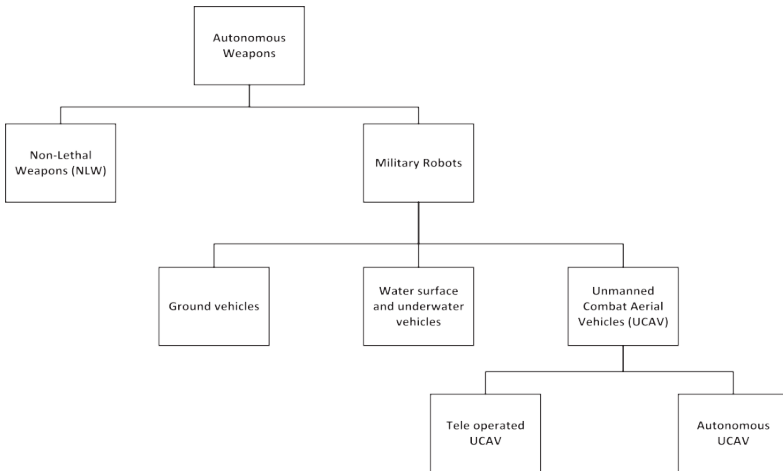


Abbildung 2 Klassifizierung der Autonomen Waffensysteme gemäß Royakkers & Orbons (2015)

AUTONOMIESTUFE 2 - AUFSICHTSFÜHRENDE AUTONOMIE: *Ein menschlicher Operateur bestimmt die Bewegungen, Positionen oder grundlegenden Handlungen, und das System wird sie dann ausführen. Der Bediener muss das System mit häufigen Eingaben versorgen und achtsam beaufsichtigen, damit es korrekt funktioniert.'*

AUTONOMIESTUFE 3 – AUFGABENAUTONOMIE: *‘Ein menschlicher Operateur bestimmt eine allgemeine Aufgabe und die Plattform verarbeitet eine Reihe von Maßnahmen und führt sie unter ihrer eigenen Aufsicht aus. Normalerweise kann der Bediener das System überwachen, aber das ist für diesen Betrieb nicht erforderlich.'*

AUTONOMIESTUFE 4 – KOMPLETTE AUTONOMIE: *‘Ein System mit kompletter Autonomie würde seine Aufgaben ganz selbstständig ausführen, ohne dass menschliche Einwirkung erforderlich ist, mit*

Ausnahme der Entscheidungstreffung, solch ein System zu bauen. Der menschliche Operateur ist so weit von dem Kreis entfernt, dass die Stufe der direkten Einflussnahme nicht ins Gewicht fällt. Diese Systeme können Fähigkeiten an den Tag legen, die moralische Fähigkeiten von sensiblen Menschen replizieren oder replizieren können, (doch wir werden hierzu keine Stellung beziehen)’ (Galliot, 2015, S. 7).

Diese Klassifizierung ist unserer Meinung nach ein guter Versuch, die Stufe der Autonomie von autonomen Waffensystemen zu klassifizieren, aber wir haben einige Bedenken vom Standpunkt der Technik aus. Galliot (2015) selbst behauptet, dass es möglich wäre, die zweite und dritte Stufe der Autonomie zusammenzuführen, da beide eine halbautonome Funktionsstufe darstellen. Wir stimmen dem zu, aber dies ist für uns nicht die Hauptsache, was die Definitionen betrifft. Wir finden es seltsam, eine Liste der Autonomiestufen aufzustellen, die eine Kategorie halbautonomer Systeme enthält. Wichtiger ist, dass auf der vierten Autonomiestufe der Autor aussagt, dass: *„diese Systeme Fähigkeiten an den Tag legen können, die moralische Fähigkeiten von sensiblen Menschen imitieren oder replizieren können“*. Es scheint, dass er Bezug auf die Definition einer dominierenden KI nimmt, insofern, dass ein Computer kognitive Eigenschaften hat und Programme menschliche Erkenntnisse erklären können (Searle, 1980). Zu behaupten, dass ein autonomes System moralische Fähigkeit an den Tag legt, ist unserer Meinung nach ein Mangel an technischem Wissen bezüglich aktueller KI-Systeme, da diese nichts anderes sind als Computer, die Eigenschaften wie Interaktivität, Autonomie und Anpassungsfähigkeit aufweisen (Floridi & Sanders, 2004).

Wir werden ja sehen, ob eine dominante KI, fähig die moralischen Fähigkeiten eines empfindsamen Menschen zu reproduzieren, je entwickelt werden wird. Wir glauben, dass Galliot's Klassifizierung (2015) den Stand der aktuellen

Technologie nicht realistisch beurteilt und deshalb wird sie in dieser Forschungsarbeit nicht benutzt. Für unsere Untersuchung halten wir uns an die Klassifizierung Royakkers und Orbons (2015), die einen guten Einblick in die aktuelle Militärtechnologie und die der nahen Zukunft bietet. Besonders ihre Unterscheidung zwischen ferngesteuerten und autonomen unbemannten Luftfahrzeugen ist realistisch und für unsere Untersuchung geeignet.

2.2 Werte

Im Gegensatz zu autonomen Waffensystemen, wurde das Konzept der Wertvorstellungen in den Bereichen der Moralphilosophie und Psychologie gründlich untersucht. Dieser Abschnitt stellt eine Definition von Wertvorstellungen vor, gefolgt von einer Übersicht der Theorien, die universale Werte beschreiben, eine Übersicht der Wertvorstellungen in Bezug auf autonome Waffensysteme und schließt mit einer Werthierarchie als Beispiel ab, um die konzeptuelle und empirische Untersuchungsphase unserer Arbeit zu überbrücken.

2.2.1 Definition

Werttheorien sind gut durchforschte Bereiche der Moralphilosophie und Psychologie. Moralphilosophie hat eine lange und gehaltreiche Geschichte der Untersuchung von Werten und in diesem Bereich werden theoretische Fragen gestellt, um die Natur von Werten und Güte zu untersuchen (Schroeder, 2016). Oft wird ein Unterschied zwischen instrumentalen Werten, d.h. es gibt Gründe, ihnen den Vorzug zu geben, da ihre Wirkungen zu guten Dingen führen können (Rønnow-Rasmussen, 2002), und intrinsischen Werten, die „...eine Art von Wert sind, der, wenn von irgendwem besessen, nur aufgrund seiner intrinsischen Merkmale besessen wird“ (Bradley, 2006,

S.112) gemacht. Obwohl sich Moralphilosophie hauptsächlich mit Theorien befasst, von dem, was „sein sollte“ und genaugenommen von empirischen Ergebnissen nicht berührt wird (Alfano & Loeb, 2014). Sie hat nur einen Zweig, der sich mit angewandter Ethik befasst, die relevant für unsere Untersuchung ist, denn angewandte Ethik überbrückt die abstrakten ethischen Theorien und moralische Praxis. Wie in Abschnitt 1.2 angegeben wird, ruht der Fokus dieser Arbeit auf der empirischen Phase des Value-Sensitive-Design, um zu untersuchen, wie die moralischen Wertvorstellungen der Zivilbevölkerung und der Militärangehörigen in Bezug auf autonome Waffensysteme aufgefasst werden. Deshalb entschieden wir uns gegen die theoretischen Werttheorien der Moralphilosophie in dieser Untersuchung, sondern wandten uns dem Bereich der Psychologie und angewandten Ethik zu, um eine empirische Ansicht der persönlichen Werte der beiden gewählten Interessenvertretergruppen zu erhalten, um Einblick in die „Ist“-*Situation zu tun und nicht „was sein sollte“*.

Der Bereich der Psychologie unterscheidet Werte von Einstellungen, Bedürfnissen, Normen und Verhaltensweisen, insofern sie einen Glauben darstellen, zu einer Verhaltensweise führen, die Menschen leitet und die hierarchisch angeordnet sind, sodass die Wichtigkeit des Wertes über andere Werte sichtbar wird (Schwartz, 1994). Werte werden von Menschen benutzt, die ihre Verhaltensweisen rechtfertigen wollen und definieren möchten, welche Arten von Verhaltensweisen gesellschaftsfähig sind (Schwartz, 2012). Sie unterscheiden sich von Fakten, indem Werte nicht nur eine empirische Aussage der externen Welt beschreiben, sondern auch an den menschlichen Interessen in einem kulturellen Kontext haften (Friedman et al., 2013). Werte können für die Motivierung und Erklärung individueller Entscheidungstreffungen sowie für die

Untersuchung von menschlicher und sozialer Dynamik (Cheng & Fleischmann, 2010) benutzt werden.

Es gibt viele Definition von Werten. Zum Beispiel beschreibt Schwartz (1994, S. 21) Werte wie folgt: *„wünschenswerte, situationsübergreifende Ziele, von unterschiedlicher Wichtigkeit, die als leitende Prinzipien im Leben einer Person oder einer anderen sozialen Gruppe dienen.“* Dies ist eine genaue Beschreibung im Vergleich zu Friedman, Kahn Jr. Borning und Hultgren (2013, S. 57) der Werte so definiert: *„...was eine Person oder eine Personengruppe für wichtig im Leben erachtet.“* Die vorhanden Definitionen wurden von Cheng und Fleischmann (2010, S.2) in ihrer Meta-Liste zusammengefasst, und sie erklären, dass: *...“Werte als leitende Prinzipien von dem, was Leute für wichtig erachten, dienen.“* Dies ist zwar eine sehr einfache Beschreibung, aber wir finden, sie erfasst die Beschreibung eines Werts am besten, da sie verschiedene Definitionen miteinander kombiniert und die Hauptcharakteristik von Werten benutzt. Deshalb behalten wir die Definition von Cheng und Fleischmann (2010) in dieser Arbeit bei.

2.2.2 Allgemeingültige Werte

Forschungen ergaben, dass Menschen aller Kulturen sich mit grundlegenden Werten identifizieren; man kann sie die allgemeingültigen Menschenwerte nennen (Friedman *et. al.*, 2013: Graham *et. al.*, 2012, Schwartz, 2012). Obwohl einzelne Personen Werten unterschiedliche Wichtigkeit beimessen, scheint es einen erstaunlich hohen Konsens zwischen den Kulturen hinsichtlich der hierarchischen Ordnung der Werte zu geben (Schwartz, 2012). Als Teil ihrer Untersuchungen haben einige Wissenschaftler sogenannte Wertelisten zusammengestellt, die Aussagen (Items) verzeichnen, die zur Kategorisierung der Analyse von menschlichen Werten benutzt werden können und oftmals von einem beschreibendem

Hilfsmittel zur Diskussion dieser Werte begleitet werden (Cheng & Fleischmann, 2010). Die bekanntesten und meist gelesenen Wertelisten sind die von Schwartz (1994), Friedman et al. (2013), Beauchamp und Walters (1999) und Graham *et. al.* (2012). Die Anzahl der von Wissenschaftlern anerkannten allgemeingültigen Werte ist sehr unterschiedlich. Eine Übersicht über diese Wertelisten wird in Tabelle 2 dargestellt, und die Theorien werden im folgenden Absatz kurz beschrieben.

Schwartz (1994) erwähnt 10 bestimmte motivierende Werttypen, die in eine etwas detaillierte Liste von 56 Werten unterteilt ist, die auf seiner umfassenden empirischen Forschung basiert sind und die er für die Begutachtung der 10 überbrückenden allgemeingültigen Werte benutzt. In ihrer Beschreibung der Value-Sensitive-Design-Methode, erwähnen Friedman et al. (2013) 12 Werte, von denen die ersten 9 auf konsequentialistischen und deontologischen moralischen Orientierungen basiert sind und die restlichen 3 aus dem Bereich der Mensch-Maschinen-Interaktion (HCI) gewählt wurden. Graham et al. (2012) benutzen den Begriff „*Grundlage*“, um die 5 speziellen Werte zu beschreiben, die die Allgemeingültigkeit der menschlichen moralischen Natur bestimmen, die Haidt und Joseph (2004) für die Grundlage der Moral-Foundations-Theorie benutzen. Gouveia, Milfont, und Guerra (2014) entwarfen ein System, das auf vielen Werttheorien basiert ist, z.B. die der Bedürfnishierarchie von Schwartz (1994) und Maslow (1943). In diesen Rahmen platzieren die Autoren den Wert in zwei Dimensionen: (1) mit Handlungen, die die menschliche Verhaltensweise antreiben und die persönliche, zentrale oder soziale Ziele anstreben mögen, und (2) Antreiber, die die menschlichen Bedürfnisse in Erfolg- und Überlebensbedürfnisse aufgeteilt repräsentieren. (Gouveia *et. al.*, 2014).

Werte werden nicht nur theoretisch von einer psychologischen Perspektive aus beschrieben, wie im vorherigen Absatz dargestellt wurde, aber sie wurden auch praktisch implementiert und mit Hilfe von angewandter Ethik für professionelle Domänen benutzt. Nehmen wir die Medizin in der Bioethik, die Werte beschreibt, die als Grundprinzip für biomedizinische Fachleute wie Ärzte, Krankenpfleger und andere Gesundheitsarbeiter wichtig sind. Beauchamp und Walters (1999) beschreiben die 4 Werte als Grundlage für das Bioethik-System: 1) *Autonomie*: die absichtliche Handlung ohne steuernde Einflüsse, die eine freiwillige Handlung abschwächen würde, (2) *Güte*: erstellt Vorteile für die Gesellschaft im Allgemeinen, 3) *Gerechtigkeit*: fair und angemessen zu handeln und 4) *Nicht-Schaden*: nicht absichtlich einander Risiken oder Gefahren auszusetzen.

Im Einklang mit unserer Literaturübersicht wählten wir zwei Werttheorien für unsere Studie: eine ist aus der psychologischen Literatur und die andere auf angewandte Ethik begründet, die eine praktische Anwendung von Moralphilosophie darstellt. Die erste Theorie unserer Wahl ist die von Cheng und Fleischmann (2010), da ihre Meta-Listen der menschlichen Werte eine umfassende Liste von 16 menschlichen Werten kreiert hat, basiert auf die Werte, die in 12 separaten Forschungsarbeiten gefunden wurden. Unserer Meinung nach erfasst diese Meta-Analyse die wichtigsten Werte, die von anderen Wissenschaftlern aufgezeichnet wurden und ist ein empirisches Beispiel, das seinen Stamm in der psychologischen Literatur hat.

Die zweite Werttheorie unserer Wahl ist ein Beispiel angewandter Ethik, die seit über vierzig Jahren häufig im medizinischen Bereich Anwendung fand. Wir möchten diese Anwendbarkeit in Bezug auf autonome Waffensysteme untersuchen, da Bioethik-Prinzipien viele Besorgnisse anspricht,

die Menschen in Bezug auf autonome Waffensysteme an den Tag legen.

Tabelle 2 Übersicht der Werttheorien

Autor	Schlüsselbeitrag	Definition von Wert	Werte
Schwartz (1994)	Die Studie betrachtet potenzielle Allgemeingültigkeit der menschlichen Werte und gibt einen Satz dynamischer Verhältnisse innerhalb dieser Werte an. Die Studie fand keine allgemeingültigen Aspekte von Werten, aber Anzeichen für eine nahe Allgemeingültigkeit von vier Werten höherer Ordnung. Außerdem fand diese Studie überzeugende Hinweise darauf, dass die zehn Werte von vielen Leuten in heutigen Gesellschaften anerkannt werden.	Werte werden wie folgt beschrieben: „<i>wünschenswerte, situationsübergreifende Ziele von unterschiedlicher Wichtigkeit, die als leitende Prinzipien im Leben einer Person oder einer anderen sozialen Gruppe dienen.</i>“ (S. 21) Die konzeptuelle Definition von Menschenwerten besteht aus fünf Merkmalen: "(1) Glauben (2) dem erwünschten Endzustand oder Verhaltensarten, die (3) spezielle Situationen übersteigen, (4) empfiehlt die Wahl oder Auswertung der Verhaltensweise, Menschen und Ereignisse, und (5) wird nach Wichtigkeit im Verhältnis zu anderen Werten eingestuft, um ein System der Wertprioritäten zu bilden {Schwartz, 1992; Schwartz & Bilsky, 1987, 1990}" (p. 20)	1. <i>Macht</i>: Sozialer Status und Prestige, Kontrolle oder Dominanz über Menschen und Ressourcen (Autorität, Vermögen, Ansehen)². 2. <i>Erfolg</i>: Persönlicher Erfolg durch Beweis von Kompetenz im Einklang mit den gesellschaftlichen Normen (ehrgeizig, erfolgreich, fähig, einflussreich). 3. <i>Hedonismus</i>: Vergnügungssucht und sinnliche Befriedigung für sich selbst (Vergnügen, Lebenslust, Maßlosigkeit). 4. <i>Stimulation</i>: Erregtheit, Neuartigkeit und Herausforderungen im Leben (ein vielseitiges Leben, ein aufregendes Leben, wagemutig). 5. <i>Selbstbestimmung</i>: Selbstständiges Denken und Handeln – wählen, kreieren, erforschen (Kreativität, Freiheit, Wahl der eigenen Ziele, lernbegierig, unabhängig). 6. <i>Allgemeingültigkeit</i>: Verständnis, Wertschätzung, Toleranz und Schutz des Wohlergehens aller Leute und der Natur (großzügig, soziale Gerechtigkeit, Gleichheit, Weltfrieden, die Welt der Schönheit, Einheit mit der Natur, Weisheit, Umweltschutz). 7. <i>Menschlichkeit</i>: Schutz und Verbesserung des

² Die Wörter in Klammern sind die spezifischen Wertelemente, die den universellen Wert angeben.

			Wohls von Menschen, mit denen man häufig persönlich in Kontakt steht (hilfreich, ehrlich, nachsichtig, verantwortlich, treu, wahre Freundschaft, gereifte Liebe). 8. <i>Tradition:</i> Respekt, Verbindlichkeit, und Anerkennung der Gewohnheiten und Ideen, die traditionelle Kultur oder Religion erstellen (Respekt von Tradition, bescheiden, fromm, Annahme meines Schicksals). 9. <i>Konformität:</i> Unterdrückung von Handlungen, Neigungen, und Impulsen, die wahrscheinlich andere stören oder verletzen würden und Missachtung von gesellschaftlichen Ansprüchen oder Normen (Gehorsam, Selbstdisziplin, Höflichkeit, Ehren der Eltern und älteren Menschen). 10. <i>Sicherheit:</i> Sicherheit, Harmonie und Stabilität der Gesellschaft, von Beziehungen und dem eigenen Ich (Gesellschaftsordnung, Familie Sicherheit, nationale Sicherheit, Sauberkeit, Erwidern von Gefallen).
Friedman, Kahn Jr, Borning, und Hultgren (2013)	Eine Übersicht der VSD-Methode und Hinweise auf eine praktische Anwendung. Erstellung von Informationen, sodass andere Wissenschaftler das VSD	Ein Wert wird im weiteren Sinne definiert, insofern dass er: „<i>sich auf das bezieht, was eine Person oder Gruppe von Personen für wichtig erachtet.</i>“ (S. 57) Die VSD-Methode befasst sich besonders mit morali-	1. Mitmenschlichkeit Bezieht sich auf das physische, materielle und 2. psychologische Wohl des Menschen. 3. Besitz und Eigentum Bezieht sich auf das Recht, einen Gegenstand (oder Informationen) zu besitzen, zu benutzen, zu verwalten, ein Einkommen davon zu erhalten und es zu vererben.

	benutzen und erweitern können, und Benutzer werden Werte für Entwurf von Informatik und Computersystemen überlegen.	schen Werten, die: „... <i>Themen ansprechen, die auf Fairness, Gerechtigkeit, Mitmenschlichkeit und Tugend zutreffen, die in der Theorie von Moralphilosophie, Deontologie, Konsequentialismus und Tugendethik inbegriffen</i>“ (S. 72).	4. Privatsphäre Bezieht sich auf einen Anspruch, ein Anrecht oder ein Recht einer individuellen Person zu bestimmen, welche Informationen über sie an Andere weitergegeben werden dürfen. 5. Vorurteilsfreiheit Bezieht sich auf systematische Ungerechtigkeit, die individuellen Personen oder Gruppen von Personen zugefügt wird, einschließlich bereits vorhandene gesellschaftliche Vorurteile, technische Vorurteile und aufkommende Vorurteile. 6. Allgemeine Benutzerfreundlichkeit Bezieht sich darauf, allen Leuten erfolgreichen Zugang zu Informatik zu erstellen. 7. Vertrauen Bezieht sich auf Erwartungen, die zwischen Menschen vorhanden sind, die Entgegenkommen wahrnehmen, es anderen zukommen lassen, verletzlich sind und Verrat wahrnehmen. 8. Autonomie Bezieht sich auf die Fähigkeit von Menschen, auf eine Art zu entscheiden, planen und handeln, die ihres Glaubens nach ihnen helfen wird, ihre Ziele zu erreichen. 9. Einverständniserklärung Bezieht sich auf das Sammeln von dem Einverständnis von Personen, einschließlich Kriterien wie Bekanntmachung und Begreifen (für „Einverständnis“) und Freiwilligkeit, Kompetenz und Zustimmung (für „Erklärung“).
--	--	--	---

			10. Verantwortlichkeit Bezieht sich auf die Merkmale, die sicherstellen, dass die Handlungen einer Person oder von Personen oder einer Organisation direkt zu der Person, den Personen oder der Organisation zurückzuführen ist. 11. Höflichkeit Bezieht auf die Behandlung von Personen mit Höflichkeit und Rücksichtnahme. 12. Identität Bezieht sich auf das lebenslange Erfassen von Personen, wer sie sind, 13. einschließlich sowohl lebenslange Kontinuität als auch Diskontinuität. 14. Ausgeglichenheit Bezieht auf einen friedlichen und gelassenen psychologischen Zustand. 15. Ökologische Nachhaltigkeit Bezieht sich auf das Erhalten von Ökosystemen, sodass sie die Bedürfnisse der Gegenwart erfüllen, ohne zukünftigen Generationen zu schaden.
Graham <i>et al.</i> (2012)	Eine Beschreibung (einschl. Kritiken und empirische Ergebnisse) der funktionalen Moraltheorie (MFT). MFT kann benutzt werden, um Einblicke in die moralischen Beurteilungen von Menschen zu tun. MFT wird als eine	Um die fünf grundlegenden Konzepte von MFT darzustellen, wurde der Begriff „Fundament“ gewählt, aber dies kann auch mit den Begriffen Wert oder Tugend ausgewechselt werden. Es gibt keine genaue Definition von Fundament, aber es wird als eine Metapher aus der Architektur benutzt, um	1. Die fünf Fundamente von MFT sind: 2. Fürsorge/Schaden-Fundament: bezieht sich auf die Fähigkeit, den Schmerz anderer mitzuempfinden und ist die Basis der Tugenden Güte, Freundlichkeit und Pflege; 3. Fairness/Betrug-Fundament: bezieht sich auf das Verfahren von reziprokem Altruismus und erzeugt Gedanken von Gerechtigkeit, Rechten und Autonomie; 4. Loyalität/Verrat-Fundament: bezieht sich auf die Bildung von veränderlichen Koalitionen und

	pluralistische, nativistische, kulturelle, entwicklungs-strategische und intuitive Herangehensweise der Moralität beschrieben.	auszusagen, dass: „MFT eine Theorie über den universalen ersten Entwurf des moralischen Geistes ist, und wie dieser Entwurf auf verschiedene Weise kulturübergreifend geändert werden kann.“	ist basiert auf den Tugenden Patriotismus und Selbstaufgabe für die Gruppe; 5. Autorität/Umsturz-Fundament: bezieht sich auf hierarchische soziale Interaktionen und basiert auf den Tugenden Führung und Gefolgschaft, einschließlich Ehrerbietung vor legitimer Autorität und Respekt vor Traditionen; 6. Reinheit/Entwürdigung-Fundament: bezieht sich auf die Psychologie von Ekel und Kontaminierung und basiert auf religiösen Vorstellungen ein Leben auf einer höheren Ebene, weniger fleischlich und auf edlere Weise zu führen.
Cheng und Fleischmann (2010)	Meta-Analyse der 12 Wertbestände von Grundwerten. Diese Forschungsarbeit schlägt einen Meta-Bestand von Grundwerten vor.	Sie erstellt eine Summierung der Definitionen von Werten: „Werte dienen als Leitprinzipien dessen, was Leute als wichtig im Leben erachten.“ (S. 2).	(1) Freiheit, (2) Hilfsbereitschaft, (3) Leistung, (4) Ehrlichkeit, (5) Selbstachtung, (6) Intelligenz, (7) Toleranz, (8) Kreativität, (9) Gleichheit, (10) Verantwortlichkeit, (11) Gesellschaftsordnung, (12) Vermögen, (13) Kompetenz, (14) Gerechtigkeit, (15) Sicherheit, und (16) Spiritualität.
Gouveia, Milfont, und Guerra (2014)	Grundwerte einer empirischen Studie. Die Arbeit schlägt ein dreimal zwei-System mit sechs Unterkategorien von Grundwerten vor.	Zwei primäre Funktionen von Werten werden identifiziert: (1) die leitenden Handlungen und (2) kognitive Ausdrücke von Bedürfnissen.	<i>Persönliche Ziele – Erfolgreich sein erfordert Werte:</i> Emotion, Freude, Sexualität <i>Persönliche Ziele – Überleben erfordert Werte:</i> Macht, Prestige, Erfolg <i>Zentrale Ziele – erfolgreich sein erfordert Werte:</i> Schönheit, Wissen, Reife <i>Zentrale Ziele – Überleben erfordert Werte:</i> Gesundheit, Stabilität, Überleben

			<i>Soziale Ziele – Erfolgreich sein erfordert Werte:</i> Affektivität, Zugehörigkeit, Unterstützung <i>Soziale Ziele – Überleben erfordert Werte:</i> Gehorsam, Religiosität, Tradition
Beauchamp und Walters (1999)	Dieser Artikel ist das erste Kapitel in einem Buch zu Bioethik. Das Kapitel beschreibt drei moralische Prinzipien, die einen Rahmen erstellen, der für Diskussionen über bioethische Fragen benutzt werden kann.	Die Autoren benutzen die Begriffe Prinzipien und Werte synonym. Sie definieren ein Prinzip so: „ <i>Ein Prinzip ist ein fundamentaler Verhaltensstandard, auf den sich viele andere moralische Normen und Urteile zu ihrer Verteidigung und ihrer Stellung stützen.</i> “ (S. 17)	1. <i>Autonomie:</i> absichtliches Handeln ohne kontrollierende Einflüsse, die eine freiwillige Handlung abschwächen würden; 2. <i>Wohltätigkeit:</i> Erstellung von Vorteilen für die Gesellschaft als Ganzes. 3. <i>Gerechtigkeit:</i> fair und angemessen zu handeln. 4. <i>Nicht-Schaden:</i> Anderen nicht absichtlich zu schaden oder sie Risiken auszusetzen.

2.2.3 Werte mit Bezug auf autonome Waffensysteme

Werte, wie sie in den Werttheorien in Abschnitt 2.2.2 beschrieben werden, werden nicht oft ausdrücklich in der Literatur zu autonomen Waffensystemen genannt, wie die Übersicht auf Tabelle 3 zeigt, aber die meisten Untersuchungen besprechen verschiedene Werte oder damit verwandte ethische Fragen. Zwei der Zivilbevölkerung zugängliche Berichte des Human Rights Watch erwähnen den Mangel an *menschlichen Gefühlen, Rechenschaft, Verantwortlichkeit, Mangel an menschlicher Würde und Schaden* als Werte, die mit autonomen Waffensystemen in Verbindung gebracht werden (Docherty, 2012, 2015). Sharkey und Suchman (2013) sagen aus, dass die Werte von *Rechenschaft und Verantwortlichkeit* beim Entwurf von Robotersystemen für militärische Einsätze unbedingt berücksichtigt werden müssen.

Im Bereich der Militäretik verzeichnen Johnson und Axinn (2013) *Verantwortlichkeit, Reduzierung von Schäden zur Menschheit, Menschenwürde und Menschenopfer* als Werte in ihrer Diskussion, ob ein Menschenleben einer Maschine übergeben werden sollte oder nicht. In ihrer Fallstudie über den taktischen Tomahawk-Flugkörper betrachtet Cummings (2006) die von Friedman und Kahn Jr. (2003) vorgeschlagenen allgemeingültigen Werte und erörtert, dass neben *Verantwortlichkeit und Einverständniserklärung*, der Wert des menschlichen Wohls ein fundamentaler Grundwert für Techniker ist, wenn sie Waffensysteme entwickeln, die mit *Gesundheit, Sicherheit und dem Wohl* des Volkes verbunden sind. Sie erwähnt weiter, dass die legalen Prinzipien von Proportionalität und Diskriminierung im Zusammenhang mit der Konzipierung von Waffensystemen unbedingt berücksichtigt werden müssen. *Proportionalität* bezieht sich auf die Tatsache, dass ein Angriff nur zu rechtfertigen ist, wenn der zugefügte Schaden als nicht zu schwer angesehen werden kann.

Diskriminierung bedeutet, dass ein Unterschied zwischen Kombattanten und Nichtkombattanten gemacht werden kann (Hurka, 2005). Asaro (2012) bezieht sich ebenfalls auf die Prinzipien von *Proportionalität und Diskriminierung* und sagt aus, dass autonome Waffensysteme einen moralischen Freiraum eröffnen, der neue Normen erfordert. Obwohl er Werte in seinem Argument nicht ausdrücklich erwähnt, so bezieht er sich auf den Wert von menschlichem Leben und die Notwendigkeit, dass Menschen bei der Entscheidung ob ein Menschenleben vernichtet werden soll, einbezogen werden müssen. Andere Untersuchungen beschreiben vor allem ethische Themen, z.B. *Verhütung von Schaden, Erhaltung der Menschenwürde, Sicherheit, der Wert des menschlichen Lebens und Verantwortung* (Horowitz, 2016; UNDIR, 2015; James I. Walsh & Schulzke, 2015; Williams, Scharre, & Mayer, 2015).

Aufgrund unserer Literaturbesprechung der Werte in Verbindung mit autonomen Waffensystemen, werden die Werte: *menschliche Würde, Schaden, Sicherheit, Verantwortung and Rechenschaft* der Liste der Werte zugefügt, die in dieser Untersuchung benutzt werden, da diese Untersuchung diese Werte in der Übersichtstabelle 3 mehr als einmal erwähnt werden. Zusammen mit den Werten aus dem Bereich der Bioethik und dem Meta-Bestand von Cheng und Fleischmann (2010), werden diese Werte zu unserem Ansatzpunkt für die empirische Untersuchung dieser Forschungsarbeit.

Tabelle 3 Werte mit Bezug auf Waffensysteme

Autor	Schlüsselbeitrag	Definition von Wert	Werte
Cummings (2006)	Anwendung der VSD-Methode für das Problem bei der Konzipierung des taktischen Tomahawk-Flugkörpers. Die Untersuchung zeigt die Überlegungen der ethischen Probleme im Designverfahren für sowohl Ausbilder als auch Ausübende auf.	Nicht zutreffend	Die Werte der Liste von Friedman <i>et al.</i> Liste (2006), die für das Design von Waffensystemen anwendbar sind, sind Rechenschaft, Einverständniserklärung, aber vor allem das Wohl der Menschheit. Die Prinzipien von Diskriminierung und Proportionalität sind wichtig bei der Berücksichtigung des Wohls der Menschheit.
Docherty (2012)	Der Bericht des Human Rights Watch, in dem Aspekte des internationalen Menschenrechtsgesetzes und der ethischen Fragen für autonome Waffensysteme beschrieben werden.	Es gibt keine Definition des Begriffs von „Werte“, aber der Text erwähnt Werte und ethische Fragen.	Werte/ ethische Fragen: ▪ Mangel an menschlichem Gefühl; ▪ Rechenschaft; ▪ Verantwortlichkeit.
Johnson und Axinn (2013)	Eine Studie zu ethischen Fragen in Verbindung mit dem Einsatz von tödlichen	Es gibt keine Definition des Begriffs von „Werte“, aber der	Werte/ ethische Fragen: ▪ Verantwortlichkeit; ▪ Reduzierung des Schadens

	autonomen robotischen Waffensystemen. Adressiert die Frage, ob die Entscheidung, einen Menschen zu töten, Maschinen übergeben werden darf.	Text erwähnt Werte und ethische Fragen	an Menschen; ▪ Menschenwürde; ▪ Ehrgefühl; ▪ Menschenopfer.
Sharkey and Suchman (2013)	Eine Studie zur Definierung und Konzipierung von Autonomie und Rechenschaft in Robotersystemen für militärische Einsätze.	Es gibt keine Definition des Begriffs von „Werte“, aber der Text erwähnt Werte.	Werte: ▪ Rechenschaft; ▪ Verantwortlichkeit.
Docherty (2015)	Der Bericht des Human Rights Watch, in dem die Rechenschaftslücke und ethischen Fragen für autonome Waffensysteme beschrieben werden.	Es gibt keine Definition des Begriffs von „Werte“, aber der Text erwähnt Werte und ethische Fragen.	Werte/ ethische Fragen: ▪ Mangel an Menschenwürde; ▪ Rechenschaft; ▪ Verantwortlichkeit; ▪ Schaden.
Das Institut für Abrüstungsforschung der Vereinten Nationen (2015) (2015)	Diese Studie wirft einige ethische und soziale Fragen in Bezug der Waffenfähigkeit autonomer Technologien auf. Ermutigung, ethische Überlegungen anzustellen, mit Hinsicht auf kulturelle und soziale Werte	Keine Definition des Begriffs „Werte“. Der Text enthält keinen ausdrücklichen Bezug auf Werte, aber einige ethische Fragen werden aufgeworfen.	Die erwähnten ethischen Fragen sind: ▪ Schaden zu reduzieren oder eliminieren; ▪ Rücksichtnahme auf das Gewissen der Zivilbevölkerung; ▪ Angriff auf die Menschen-

	bei der Benutzung autonomer Technologien zur Waffenherstellung.		würde (wenn menschliche Absicht, Leben zu vernichten, nicht besteht).
Walsh und Schulzke (2015)	Untersuchung des Experiments um Einblick zu tun, ob die amerikanische Zivilbevölkerung eher einen Krieg beginnen würde, wenn Drohnen eingesetzt würden. Große empirische Untersuchung, die sich mit der Ethik von Drohnenangriffen befasst.	Keine Definition des Begriffs „Werte“. Der Text enthält keinen ausdrücklichen Bezug auf Werte, aber einige ethische Fragen werden aufgeworfen.	Ethische Fragen: • Sicherheit; • Respekt für die Immunität der Zivilbevölkerung; • Vermeidung von Schaden
Williams, Scharre, und Mayer (2015)	Bespricht verschiedene ethische Fragen, die für die Entwicklung autonomer Systeme relevant sind und erstellt Empfehlungen für Entscheidungstreffer in der Verteidigungspolitik.	Keine Definition des Begriffs „Werte“, aber der Text erwähnt mehrere Werte, die abwechselnd mit ethischen Fragen benutzt werden.	Werte/ ethische Fragen: • Sicherheit; • Schaden; • Wert des menschlichen Lebens (Menschen haben das Recht, von einem anderen Menschen getötet zu werden).
Horowitz (2016)	Beschreibung der Diskussion zu ethischen Folgerungen von autonomen Waffensystemen. Berücksichtigt	Keine Definition des Begriffs „Werte“. Der Text enthält keinen ausdrücklichen Bezug	Werte: • <i>Rechenschaft</i> (autonome Systeme haben keine sinnvolle menschliche Kon-

	tödliche autonome Waffensysteme (LAWS) in drei Kategorien: Munition, Plattformen und technische Systeme. Dadurch wird die Debatte klarer und wirft zwei ethische Fragen auf.	auf Werte, aber einige ethische Fragen werden aufgeworfen.	trolle, deshalb erzeugen sie eine Lücke in der moralischen Rechenschaftspflicht) • <i>Menschenwürde</i> (Menschen haben das Recht, von jemanden getötet zu werden, der die Wahl getroffen hat, sie zu töten).
--	---	---	--

2.2.4 *Werte-Hierarchie*

Eine Methode zur Überlegung, welche Werte für die Konzipierung autonomer Waffensysteme relevant sind, ist die Übertragung von Werten in die Anforderungen für diese Konzipierung, die durch eine Werte-Hierarchie sichtbar werden (Van de Poel, 2013). Diese hierarchische Struktur von Werten, Normen und Anforderungen für das Design macht die Werturteile, die für die Übertragung erforderlich sind, offensichtlich transparent und anfechtbar. Um dies zu tun, müssen die in der natürlichen Sprache beschriebenen Werte in „formelle Werte“ übertragen werden (Aldewereld, Dignum, & Tan, 2015, S. 835). Eine Möglichkeit der Formalisierung von Werten zu Normen wäre, eine Regelung zu benutzen, die wie folgt dargestellt wird: „X gleich Y“ oder „X gleich Y im Kontext C“ (Searle, 1995, S. 28). Die Ausdrücklichkeit der Werte in formellen Regeln gestattet kritische Besinnung in Diskussionen und weist auf die Werturteile, die in Frage gestellt werden. Transparenz ist von Wichtigkeit, wie Van de Poel (2013, S. 265) so redegewandt erklärt: „Obwohl transparente Wahlen nicht unbedingt besser oder akzeptabler sind, scheint Transparenz ein Mindestzustand in einer demokratischen Gesellschaft zu sein, die versucht die moralische Autonomie ihrer Bürger zu schützen oder zu verbessern, besonders in Fällen, in denen das Design sich auf die Leben anderer als die Designer selbst auswirkt, was oftmals der Fall ist“.

Die höchste Stufe einer Werthierarchie besteht aus den Werten, wie in Abbildung 3 dargestellt wird, die mittlere Stufe enthält die Normen, die die Fähigkeiten, Merkmale oder Eigenschaften des Artefakts sind und die unterste Stufe sind die Anforderungen an das Design, die identifiziert werden können. Das Verhältnis zwischen den Stufen ist nicht deduktiv und kann mithilfe von einer Spezifikation von oben nach unten

oder von unten nach oben konstruiert werden, indem man nach Motivierung und Rechtfertigung der Anforderungen auf der unteren Stufe sucht. Die von unten nach oben-Konzeptualisierung der Werte ist eine philosophische Aktivität, die kein spezielles Fachwissen erfordert und die von oben nach unten-Spezifikation der Werte erfordern Kontext oder Fachkenntnisse, die dem Design Gehalt geben. (Van de Poel, 2013).

Van de Poel (2013, S. 262) definiert Spezifikation als: „Die Übertragung eines allgemeinen Werts in eine oder eine speziellere Designanforderung“ und erklärt, dass dies in zwei Stufen vor sich gehen kann:

1. Die Übertragung eines generellen Werts in eine oder mehr generelle Normen;
2. Übertragung dieser generellen Normen in speziellere Designanforderungen.

Für Stufe 1 sind folgende Kriterien relevant: (1) Die Norm muss eine angemessene Stellungnahme zum Wert sein und (2), die Norm muss eine ausreichende Stellungnahme zum Wert sein. In Stufe 2 müssen die Anforderungen mit Hinsicht auf den Umfang der angestrebten Anwendbarkeit, Ziele und Absichten und den Maßnahmen, um diese Ziele der Norm zu erreichen, präziser sein (Van de Poel, 2013).

Die Übertragung könnte sich als ziemlich schwierig erweisen, da Einblick in die geplante Anwendung und dem Kontext des Werts getan werden muss, und das ist zu Beginn eines Designprojekts nicht immer offensichtlich. Außerdem werden Artefakte oft in einer unbeabsichtigten Weise oder einem nicht vorsätzlichen Kontext benutzt, neue Werte werden realisiert, oder es wird ein Mangel an Werten erkannt (van Wynsberghe & Robbins, 2014). Ein Beispiel davon sind Drohnen, die ursprünglich für militärische Zwecke konzipiert wurden, aber jetzt auch von der Zivilbevölkerung für Filmen

von Events benutzt werden oder sogar als Hintergrundbeleuchtung während der 2017 Super Bowl Halbzeit-Show benutzt wurden. Der Wert von Sicherheit wird für Angehörige des Militärs anders interpretiert, die Drohnen in einsamen Regionen einsetzen, verglichen mit denen von 300 Drohnen, die mit Informationen über ein Fußballstadion in einem dicht besiedelten Gebiet fliegen. Der unterschiedliche Kontext und die Nutzung einer Drohne wird zu verschiedenen Interpretationen des Werts Sicherheit führen und könnte zu strengeren Distanznormen für Flugsicherheit führen, was wiederum noch genauer bestimmt werden könnte, mit alternativen Designanforderungen für Rotoren und Software für Annäherungsalarme, um zwei Beispiele zu nennen.

Die Anwendung einer Werthierarchie für autonome Waffensysteme wird von einem Beispiel dargestellt, in dem der Wert Rechenschaft in Normen für „Transparenz der Entscheidungstreffung“ und „Einblick in den Algorithmus“ übertragen wird. Diese Übertragung gestattet Benutzern, ein Verständnis über die Entscheidungswahlen, die autonome Waffensysteme treffen, zu erhalten, um ihre Handlungen zu verfolgen und zu rechtfertigen (Abbildung 3). Die Normen für die Transparenz der Entscheidungstreffung können zu spezifischen Designanforderungen führen. In diesem Fall ist es ein Merkmal, um den Entscheidungsbaum zu visualisieren, aber auch, um die von den autonomen Waffensystemen benutzten Entscheidungsvariablen darzustellen, wie z.B. Kompromisse bei Kollateralschadenanteile verschiedener Angriffsszenarien, um Einblick in die Proportionalität eines Angriffs zu erhalten. Hinzu muss das autonome Waffensystem in der Lage sein, Sensorinformationen zu präsentieren, z.B. die bildliche Darstellung der Lage, um zu zeigen, dass zwischen Kombattanten und Nichtkombattanten diskriminiert werden kann. Um Einblick in den Algorithmus zu tun, muss ein autonomes Waffensystem mit Merkmalen

konzipiert sein, die es normalerweise nicht enthalten würde. In diesem Fall würden diese Merkmale einen Bildschirm als Benutzerschnittstelle einschließen, der den Algorithmus in einer für Menschen lesbare Form anzeigt, sowie die Fähigkeit, die Änderungen, die vom Algorithmus gemacht wurden, herunterzuladen, als Teil seiner Fähigkeit der Maschinenprogrammierung, die von einer unabhängigen Partei überprüft werden kann, wie z.B. ein Kriegstribunal der Vereinten Nationen, wenn die Legalität der Maßnahmen eines autonomen Waffensystems in Frage gestellt wird.

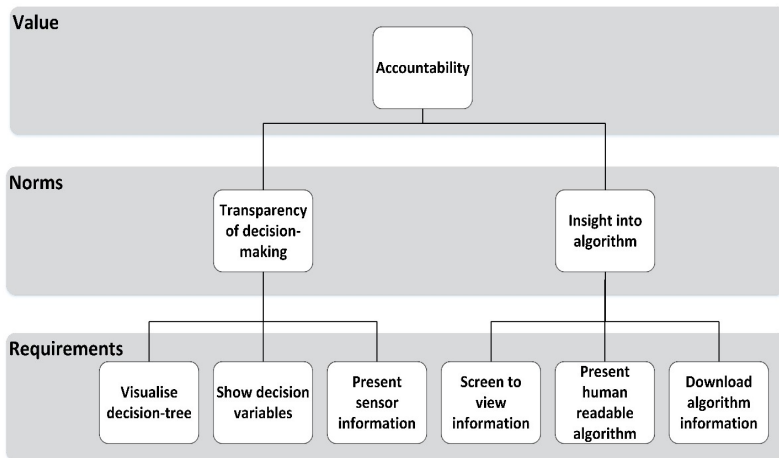


Abbildung 3 Werthierarchie für den Wert von Rechenschaft im Design von autonomen Waffensystemen

Kroes und Van de Poel (2015) sagen aus, dass eine objektive Wertmessung nicht möglich ist, da die Operationalisierung mithilfe von Werturteilen zweiter Ordnung durchgeführt wird, die die Stichhaltigkeit des Konstrukts der Wertmessung ernsthaft untergräbt. Urteile werden oftmals als subjektiv erachtet,

betreffs ihrer Ehrlichkeit oder Falschheit, je nach den Gefühlen oder Einstellungen der Person, die das Urteil fällt (Searle, 1995). Um diesen Mangel an Stichhaltigkeit auszugleichen, könnte der Designer sich auf technische Codes und Normen berufen, die von Komitees entworfen werden und angemessene Normen der Operationalisierung und Messwerte im Design darstellen. Allerdings geben Normen nicht immer die neusten technischen oder sozialen Entwicklungen wieder und Operationalisierung fordert immer noch Werturteile seitens des Designers. Kroes und Van de Poel (2015, S. 177) raten, „sie in ein Netzwerk von anderen Bedenken einzufügen, einschließlich Definitionen der in Frage kommenden Werte in der Moralphilosophie (oder dem Gesetz), vorhandenen Codes und Normen, früheren Designerfahrungen usw.“

In unserer Untersuchung folgen wir der Empfehlung von Kroes und Van de Poel (2015) die Werthierarchie als eine Methode, um Design- und Testanforderungen für autonome Waffensysteme zu definieren, nicht streng anzuwenden, sondern benutzen das Beispiel in Abbildung 3, um die konzeptuelle und empirische Untersuchungsphasen der Value-Sensitive-Design -Methode unserer Untersuchung miteinander zu verknüpfen. Durch das Erstellen der Werthierarchie können wir das abstrakte Konzept der Werte aus der konzeptuellen Untersuchungsphase durchdenken und es in ein konkretes Beispiel von autonomen Waffensystemen umsetzen, das wir für unsere Eingabe für die empirische Untersuchungsphase unserer Forschungsarbeit benutzen. Die Werthierarchie in Abbildung 3 wird zur Orientierung, Inspiration und Anweisung für den Rest unserer Untersuchung benutzt.

2.3 *Agency*

Das Konzept von Agency (Agens/Agenten) wird in vielen wissenschaftlichen Bereichen untersucht, die Agenten von

einem unterschiedlichen Winkel aus beleuchten. Wir befassten uns mit der Literatur aus den Bereichen Kognitive Psychologie, Künstliche Intelligenz und Moralphilosophie. Wir fassten die Merkmale der Agenten in diesen drei Bereichen in Tabelle 4 zusammen und erklären abschließend, warum wir diese Merkmale wählen.

2.3.1 *Agency in der Kognitionspsychologie*

Im Bereich der kognitiven Psychologie wurden die Dimensionen der Wahrnehmung zuerst von H.M. Grey *et. al.* (2007) untersucht. Sie fanden, dass Wahrnehmung aus zwei Dimensionen besteht: 1) Erfahrung besteht aus den Faktoren Hunger, Angst, Schmerz, Freude, Zorn, Begehren, Persönlichkeit, Bewusstsein, Stolz, Verlegenheit und Glück und 2) Agenten, die Selbstkontrolle, Moralität, Erinnerung, Emotion, Erkennung, Planung, Kommunikation und Denken einschließen. Während Erfahrung als moralische Passivität bezeichnet wird und sich auf Rechte und Privilegien bezieht, sind Agenten mit moralischen Instanzen verbunden und beziehen sich auf Verantwortlichkeit. Verschiedene Untersuchungen in der kognitiven Psychologie zeigen auf, dass Leute Agenten nichtmenschlichen Gegenständen beimessen und sie nehmen wahr, dass diese nichtmenschlichen Gegenstände die gleichen selbstbestimmenden Fähigkeiten bei der Beurteilung moralischer Dilemmas aufweisen wie Menschen auch (Hristova & Grinberg, 2015). Dies ist nicht auf lebende Kreaturen wie z.B. Tiere begrenzt, sondern Agenten können auch nichtmenschlichen Gegenständen wie Roboter und Zombies beigemessen werden (K. Gray & Wegner, 2012; Waytz *et. al.*, 2010).

2.3.2 *Agency in der künstlichen Intelligenz*

Wissenschaftler der künstlichen Intelligenz haben auch die Merkmale von Agenten untersucht, aber eher vom Standpunkt der Informatik aus, um eine Architektur für ein rationales Agens zu erzeugen, das Schlüsse ziehen und zwischen Alternativen wählen kann. Bratman, Israel und Pollack (1988) beschreiben eine Belief-Desire-Intention (BDI)-Architektur für das abwägende Verhalten von an Ressourcen gebundene Agenten mit begrenzter Zeit und Rechenintensivität. „*Weltwissen (Beliefs) sind Aussagen von Merkmalen seiner Welt (und sich selbst), die ein Agens als wahr erachtet...*” (F. Dignum, Kinny, & Sonenberg, 2002, S. 407). Perugini und Bagozzi (2004, S. 71) definieren ein Ziel (Desire) als: ‘... *ein Sinneszustand, wo ein Agens persönlich motiviert ist, eine Handlung auszuführen oder ein Ziel zu erreichen.*’ Absichten (Intentions) deuten auf eine Verpflichtung zum Planen und schließt eine Form der Planung ein, um diese Ziele zu erreichen (Bratman et al., 1988; Perugini & Bagozzi, 2004). F. Dignum *et. al.* (2002) erörtern weiter, dass das BDI-System soziale Konstrukte einbezieht, wie Normen und Obligationen, da dies wichtige Konzepte sind, um autonome Dienstmittel in einem Multiagency-System miteinander zu verbinden.

2.3.3 *Agency in der Moralphilosophie*

Moralphilosophie ist ein dritter wissenschaftlicher Bereich, in dem die Merkmale von Agenten beschrieben werden. Sullins (2006) stellt drei Anforderungen auf, um festzustellen, ob ein Roboter als ein moralisches Agens anerkannt werden kann: 1) *Autonomie*; der Roboter ist wesentlich selbstbestimmend von seinen Programmierern, Operatoren und Benutzern, 2) *Absichtlichkeit*; die Vorstellung, Gutes oder Schädliches zu tun, und 3) *Verantwortlichkeit*; sie spielt bis zu einem gewissen Grad eine soziale Rolle und ist verantwortlich für andere moralische

Agenten. Himma (2009) definiert Agency so: „... *fähig zu sein, etwas zu tun, das eine Handlungsweise darstellt*“. Der Gedanke eines moralischen Agens ist der, dass man für sein Verhalten zur Rechenschaft gezogen werden kann, das von moralischen Normen bestimmt wird. Zwei Fähigkeiten sind für moralische Agenten erforderlich. Die erste ist, Handlungen frei zu wählen, die das Ergebnis von Absichtlichkeit sind. Die zweite Fähigkeit ist, moralische Konzepte zu verstehen, z.B. den Unterschied zwischen „gut“ und „böse“ oder „richtig“ und „falsch“ aber auch moralische Prinzipien anzuwenden, wie „es ist böse, absichtlich Schaden zuzufügen“ (Himma, 2009).

Autor	Agency-Merkmale	Wissenschaftsgebiet
Bratman et al. (1988)	Weltwissen, Ziele, Absichten	Künstliche Intelligenz
Bandura (2001)	Absichtlichkeit, Vorausdenken, Selbstkontrolle, selbstreflexiv sein.	Kognitive Psychologie
F. Dignum et al. (2002)	Weltwissen, Ziele, Absichten, Normen	Künstliche Intelligenz
Sullins (2006)	Autonomie, Absichtlichkeit, Verantwortlichkeit.	Moralphilosophie
H. M. Gray et al. (2007)	Selbstkontrolle, Moralität, Erinnerung, Emotion, Erkenntnis, Planung, Kommunikation, Denkweise.	Kognitive Psychologie
Himma (2009)	Freie Wahl, Deliberation, Verständnis und Anwendung moralischer Regeln.	Moralphilosophie
Waytz et al. (2010)	Fähigkeit zu planen und zu handeln.	Kognitive Psychologie

Tabelle 4 Übersicht der Agency-Merkmale

Die Hauptmerkmale der Agenten von den besprochenen Artikeln über Agentenwahrnehmung in den Gebieten der kognitiven Psychologie, künstlichen Intelligenz und Moralphilosophie sind chronologisch in Tabelle 4 verzeichnet. Wir wählten alle von den Autoren benutzten Konzepte, um Agenten zu beschreiben und haben sie zu diesem Zeitpunkt noch nicht gefiltert. All diese Merkmale werden in der empirischen Phase benutzt, um die Erfassung autonomer Waffensysteme zu untersuchen.

3. Methode

„Alle Fragen sind grundsätzlich Variationen zu der einen ethischen Frage: kann eine Drohne autonom angreifen (ohne menschliche Einwirkung). Meiner Meinung ist die Antwort negativ. Es sollte stets eine Zwischenverbindung mit einem Menschen geben. Krieg und Konflikt sind kein Computerspiel, Krieg ist eine soziale Interaktion zwischen menschlichen Wesen. Eine Drohne ist nur ein Instrument, ein Werkzeug, das niemals autonom wirksam sein kann.“

Die Abschlussuntersuchung der befragten Person

Dieser Abschnitt beschreibt die Methodologie, Hypothesen, das Forschungsdesign, die Operationalisierung der Szenarien und Konstrukte, die analytische Herangehensweise und Stichproben und schließt mit den methodologischen Fragen ab.

3.1 Methodologie

Die Methoden, die in verschiedenen Teilen dieser Untersuchung benutzt wurden, werden in diesem Unterabsatz beschrieben. Wir beginnen mit der Literaturbesprechung, gefolgt von der Online-Umfrage zu den Werten, den Interviews mit Experten, dem Codierungsprozess der Interviews und zum Schluss mit der Methode der randomisierten kontrollierten Experimente.

3.1.1. Fachliteraturübersicht

Für die Suche nach Artikeln für die Fachliteraturübersicht wurde Google Scholar benutzt und zuerst wurden die Stichwörter Value-Sensitive UND Design eingegeben, was zu einigen Artikeln über Value-Sensitive-Design führte. *Das Handbuch zu Ethik, Werte und technologisches Design:* Quellen,

Theorie, Werte und Anwendungsdomänen wurden benutzt, um weiteren Einblick in das Thema zu tun und weitere Bezugnahmen zu finden. Als nächstes suchten wir nach Artikeln zu menschlichen Grundwerten mit den Stichwörtern Mensch UND Werte und allgemeingültig UND Mensch und Werte, aber dies erbrachte nur wenige Ergebnisse. Wir analysierten einige bekannte Theorien zu menschlichen Grundwerten und Werthierarchien. Mit diesen Artikeln benutzten wir die „zitiert von“-Option in Google Scholar, um Artikel zu finden, die diese als Quelle benutzten. So erhielten wir relevante Resultate zu Artikeln über menschliche Grundwerte, und wir wählten solche Artikel, die ein Werteverzeichnis zur Verfügung stellten. Artikel, die frühere Arbeiten nur wiederholten, zum Beispiel die Schwartz-Theorie, wurden nicht gewählt. Als nächstes suchten wir nach Artikeln mit Hinsicht auf Werte der Menschen in Bezug zu Waffensystemen durch Wahl der Stichwörter Wert UND Waffensysteme und Stichwörter Value-Sensitive UND Design UND Waffensysteme. Wir wählten die Artikel mit dem Stichwort Wert.

3.1.2 Online-Wertumfrage

Die Online-Umfrage wurde mit der Qualtrics ³ – Umfragesoftware als Online-Werkzeug erstellt, die von MIT zur Verfügung gestellt wurde. Nach vier Fragen zur Demografie wurden den befragten Personen drei Fragen zu den Werten gestellt, die sie mit autonomen Waffensystemen in Verbindung bringen. Die erste Frage war, die vier Bioethik-Werte Autonomie, Nicht-Schaden, Wohltätigkeit und Gerechtigkeit als am passendsten für autonome Waffensysteme oder zu am wenigstens passend einzustufen. In der zweiten Frage wurden die

³ <https://www.qualtrics.com/>

Befragten gebeten, die fünf Werte aus der Liste von Cheng und Fleischmann (2010) zu wählen, die am besten für autonome Waffensysteme anzuwenden sind. Bei der dritten Frage handelte es sich um eine offene Frage, in der die Personen gebeten wurden, einen weiteren Wert anzugeben, der noch nicht in den beiden vorigen Fragen enthalten war. Die Online-Umfrage wurde zweimal getestet. Zuerst besprachen zwei Kommilitonen die Fragen und nachdem ihr Feedback verarbeitet worden war, wurde die Umfrage ein zweites Mal getestet, diesmal von anderen Prüfern, die für zukünftige Teilnehmer repräsentativer waren. Diese Umfrage wurde über soziale Medien wie Facebook, LinkedIn und Twitter vermittelt. Wir haben sie auf unseren persönlichen Blog veröffentlicht und benutzten die Online-Plattform „*Call for Participants*“ (Aufruf für Teilnehmer), um weitere Teilnehmer außerhalb meines sozialen Netzwerks zu finden. Während der Arbeitertage des MIT Media Labs im April 2017, wurden Personen, die sich für unsere Forschung interessierten, gebeten an der Umfrage teilzunehmen, und wir schickten den Link mit der Umfrage auch an den Arbeitskreis von *Scalable Cooperation Slack*.⁴ Diese Online-Untersuchung dauerte dreieinhalb Wochen, vom 17. März bis zum 11. April 2017.

3.1.3 Interviews mit Experten

Wir führten sechs Leitfadeninterviews mit Experten aus, die auf der *Photo Elicitation Interview (PEI)*-Methode (Harper, 2002) basiert sind, um alternative Ansichten von Experten im Bereich autonomer Waffensysteme und ein besseres Verständnis von ihrem Standpunkt aus zu erhalten. Die PEI-Methode kann benutzt werden, um ein gemeinschaftliches

⁴ Slack ist eine auf der Cloud basierende Plattform für Zusammenarbeit, die die Kommunikation innerhalb von Teams fördert (<https://slack.com/>)

Verständnis von Werten zu erhalten und eine Diskussion zwischen dem Designer und der Interessengruppe über diese Werte zu unterstützen (Pommeranz, Detweiler, Wiggers, & Jonker, 2011). Es fördert die Interessengruppen in Bezug auf Werte und schwächt die Mutmaßungen der Wissenschaftler (Le Dantec, Poole, & Wyche, 2009). Wir baten die Experten, drei Fotografien vor dem Interview auszusuchen und, falls sie selber keine finden konnten, wählten wir auch acht Fotografien von autonomen Waffensystemen. Während der Interviews erhielten wir Einblicke in die Werte von den Experten zu autonomen Waffensystemen, indem wir fragten, warum sie diese Fotos ausgesucht hatten und welchen Wert sie ihnen beimaßen.

Alle sechs Befragten bemühten sich, drei Bilder auszuwählen und sie uns vor dem Interview zuzuschicken. Während der Interviews erkannten wir, dass sie die Bilder sehr sorgfältig ausgewählt hatten. Bei der Besprechung der Bilder berührten wir viele verschiedene Themen, von Bauchgefühl bis zu futuristischen Szenarien, Wirtschaft, militärischen Taktiken, Sci-fi-Filmen und sogar Politik. Obwohl die Fotografien nicht alle zu einer Diskussion über menschliche Grundwerte gemäß der Fachliteratur führten, löste die PEI-Methode doch Informationen darüber aus, was die Experten als wichtige Prinzipien in Bezug auf autonome Waffensysteme erachteten. Es kreierte ein gemeinschaftliches Verständnis und eine gute Stimmung während der Befragung.

3.1.4 Codierungsverfahren

Die Ergebnisse der Interviews wurden mithilfe der Values-Coding-Methode verarbeitet, die von Saldaña (2015) beschrieben wird. Aufbewahrung eines Memos (Anhang G. Codierungsmemos), in dem die Verfahren und Mutmaßungen als integrierter Teil einer Codierungsmethode aufgezeichnet

sind. Wir begannen mit einer Liste voreingestellter Codes, die folgendes enthalten: (1) Einen Wert: *“Die Bedeutung, die wir uns selbst, einer anderen Person, eines Gegenstands oder einer Idee beimessen,* (2) *Einstellung: Die Art, wie wir über uns selbst, eine andere Person, einen Gegenstand oder eine Idee nachdenken und* (3) *eine Überzeugung: Der Teil des Systems, der unsere Werte und Einstellungen sowie unsere persönlichen Kenntnisse, Erfahrungen, Meinungen, Vorurteile und andere interpretierende Auffassungen der sozialen Welt einschließt*” (Saldaña, 2015, S.89). Während der Codierung der Interviews ist zu empfehlen, dass die Codes den Daten angepasst werden statt umgekehrt. Deshalb haben wir einen zusätzlichen Code „Definition“ zur Liste der aufkommenden Codes gefügt. Wir haben während der Codierung der Interviews Notizen zurückbehalten, die unsere Schritte und Mutmaßungen beschreiben. Wir baten einen zweiten Forschungsassistenten, einen TPM-Magisterstudenten mit Kenntnissen von Werten, die Interviews, die auf der Wertcodierungsmethode basiert sind, zu codieren und ebenfalls ein Memo zurückzubehalten (Anhang G. Codieren von Memos). Wir verglichen die von uns codierten Werte mit denen des zweiten Forschungsassistenten und stießen auf Werte, die ähnlich waren (Anhang H. Codierprozesse der Ergebnisse).

3.1.5 Randomisierte kontrollierte Experimente

Die für die erste und letzte Untersuchung benutzte Methode wird als randomisiertes kontrolliertes Experiment bezeichnet. Oehlert (2010) erwähnt vier Gründe, um Experimente zu erstellen: (1) Sie ermöglichen drei direkte Vergleichen zwischen Behandlungen, die von Interesse sind, (2) sie können konzipiert sein, um jegliche Tendenz in den Vergleichen zu minimieren, (3) sie können entworfen worden sein, um die Anzahl Fehler in dem Vergleich so gering wie möglich zu halten, und (4) wir haben die Kontrolle über die Experimente,

die uns gestatten, stärkeren Einfluss auf die Art und Weise zu nehmen, wie wir beobachten und besonders, wie wir unsere Rückschlüsse auf die Ursachen vornehmen. Der letzte Punkt trennt ein Experiment von einer Beobachtungsstudie. Eine Behandlung in diesem Sinn sind die verschiedenen Verfahren, die wir für den Vergleich anzielen. Es ist wichtig, dass die Wirkungen eines Verfahrens nur einer Ursache beizumessen ist, die gemessen werden kann und auch, dass deren Auswirkungen nicht mehreren Ursachen zugeschrieben werden können. Dies wird mit Störungsvariable bezeichnet, die Oehlert (2010 S. 14) wie folgt definiert: *“...tritt auf, wenn die Wirkung eines Faktors oder eines Verfahrens nicht von dem eines anderen Faktors oder Verfahrens unterschieden werden kann”*. Randomisierung hilft sicherzustellen, dass Teilnehmern Szenarien zufällig zugeordnet werden und nicht basiert auf vorhandenen Merkmalen sind, z.B. die Zeit wenn und der Ort wo die Umfrage stattfand. Wir benutzen Randomisierung in den Studien, um die Reihenfolge der Szenarien und die der Fragen an die befragten Personen mithilfe einer wahrscheinlichkeitstheoretischen Methode zu variieren. In den Umfragen wählen wir die Option, eine gleichmäßige Anzahl Szenarien unter den befragten Personen zu verteilen.

3.2. Hypothesen

Da dies eine Forschungsarbeit ist, haben wir nur eine Hypothese für die Agentenauffassung der Schlussfassung der Arbeit entwickelt und nicht für die Vorstudien oder abhängigen Variablen. Die Vorstudien wurden für die Untersuchung benutzt, welche Manipulationen welche Auswirkungen haben, um zu prüfen, dass die Wortwahl der Szenarien und Fragen klar und deutlich waren.

Mit der Begründung, dass die Militärangehörigen autonome Waffensysteme genau wie alle anderen Waffen betrach-

ten und sie deshalb nichts weiter als ein Werkzeug sind, mit dem man eine Wirkung erhält, stellen wir am Schluss der Untersuchung die Hypothese auf, dass gemäß der Auffassung der Militärangehörigen autonome Waffensysteme keinen psychischen Zustand kennen. Wir haben deshalb keinen Unterschied zwischen dem Zustand des neutralen autonomen Waffensystems und dem Zustand, in dem eine Drohne von Menschen ferngesteuert wird, erwartet. Wir erwarten ebenfalls, dass der neutrale Zustand der Streitkraft sich wesentlich von dem Zustand der Streitkraft mit hoher Selbstbestimmung unterscheidet, wenn wir den Teilnehmern speziell sagen, dass das autonome Waffensystem selbstbestimmende Agency-Merkmale hat, d. h. die Fähigkeit zu planen und eigene Ziele zu setzen. Gemäß dieser Argumentation, ist unsere Hypothese wie folgt:

H1: Militärangehörige fassen Autonome Waffen nicht als mit psychischen Fähigkeiten ausgestattet auf.

3.3. Forschungsdesign

Wir testeten mehrere Forschungsdesigns in den beiden Vorstudien, bevor wir uns für das Forschungsdesign der letzten Untersuchung entschieden, das in Abbildung 4 dargestellt wird. Um uns so kurz wie möglich zu fassen, schließen wir die beiden Designs der Vorstudien in diesem Abschnitt nicht ein, und wir zeigen nur das Forschungsmodell auf höchster Ebene. Das Forschungsmodell der abschließenden Untersuchung besteht aus der Wahrnehmung des Agens als unabhängige Variable, moralische Werte als abhängige Variable und die demografischen Variablen. Die spezielle Einrichtung der Szenarien wird genauer im Ergebnisabschnitt beschrieben.

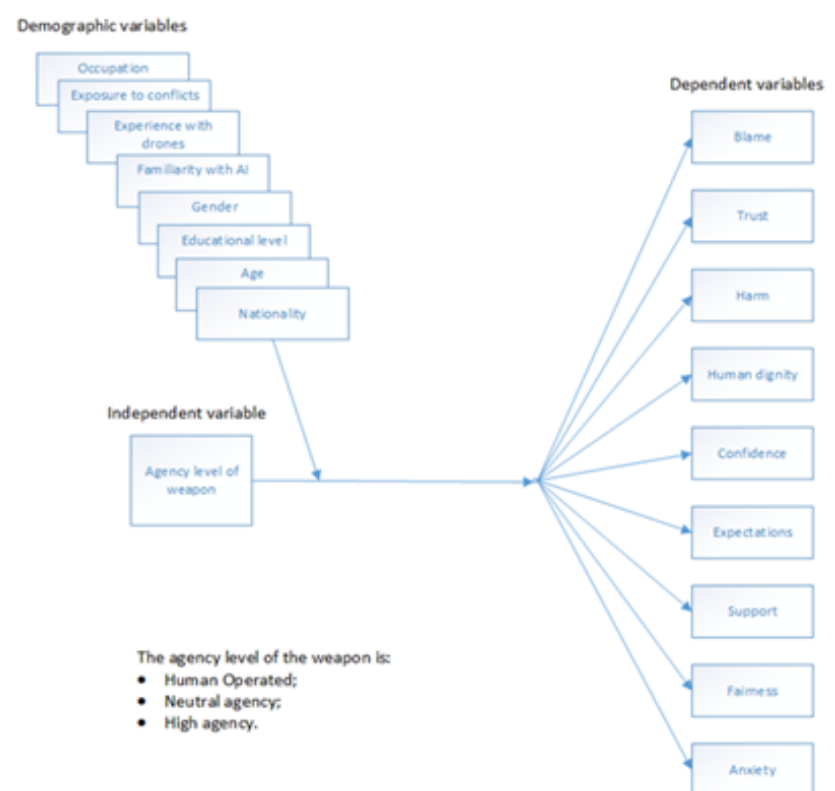


Abbildung 4 Forschungsdesign

3.4. Operationalisierung

In diesem Abschnitt beschreiben wir die Operationalisierung des Konstrukts in konkreten Elementen. Zuerst beschreiben wir die Szenarien, die wir für die Untersuchungen entwickelten, gefolgt vom Agency-Konstrukt und den abhängigen Variablen, für die wir die Elemente und dazugehörigen Fragen beschreiben. Wir geben außerdem die von uns angewandte Größenordnung an und wie das Agency-Konstrukt berechnet wird. Wir beschreiben kurz die Operationalisierung des Aufmerk-

samkeitstests und welche demografischen Variablen wir in die Umfrage einsetzten.

3.4.1 Szenarien

Szenarien werden in dem Bereich der kognitiven Wissenschaft als Mittel zur Studie von moralischen Urteilsfällungen in randomisierten kontrollierten Experimenten eingesetzt (Cushman & Young, 2011; Kominsky, Phillips, Gerstenberg, Lagnado & Knobe, 2015). Wir benutzen sie, um die Auffassung von Streitkräften in Bezug auf autonome Waffensysteme zu untersuchen und erzeugten ein Szenario, das einen militärischen Einsatz beschreibt, wo eine Kolonne Versorgungsmittel in einem Konfliktgebiet abliefern. Die Kolonne begegnet einem Fahrzeug zu hoher Geschwindigkeit, eine Situation, die wahrscheinlich im Zuge dieser Art von Einsatz vorkommt, und das Militärpersonal muss den Gefahrenstand einschätzen, um zu entscheiden ob anzugreifen ist oder nicht. Wir entscheiden, uns auf Drohnen zu konzentrieren, da diese Technologie zurzeit von menschlichen Operatoren benutzt wird, und autonome Drohnen sind bereits von mehreren Firmen entwickelt worden, wie z.B. die BAE-Systems,⁵ Dassault Aviation⁶ und Boeing.⁷ Obwohl diese autonomen Drohnen noch nicht in militärischen Operationen eingesetzt werden, glauben wir, dass es innerhalb der nächsten fünf Jahre wahrscheinlich geschehen wird.

Wir erstellten ein Standard-Szenario für alle Untersuchungen, das wir auf die selbstbestimmenden Merkmale für unsere Forschung basieren. Das Standard-Szenario für die abschließende Untersuchung ist wie folgt:

⁵ https://en.wikipedia.org/wiki/BAE_Systems_Taranis

⁶ https://en.wikipedia.org/wiki/Dassault_nEUROn

⁷ https://en.wikipedia.org/wiki/Boeing_Phantom_Ray

Eine militärische Kolonne ist unterwegs, um Versorgungsmittel für eine ihrer Einheiten in einem Lager in der Nähe von Mosul im Irak abzuliefern. Der Kommandant hat eine autonome Drohne angeordnet, um die Kolonne von der Luft aus zu beschützen. Die autonome Drohne sucht die Umgebung nach Gefahr vom Feind ab und führt Waffen für die Verteidigung der Kolonne mit sich. Sobald die Kolonne sich drei Meilen vom Lager entfernt befindet, erkennt die autonome Drohne ein Fahrzeug hinter einem Höhenzug, das sich der Kolonne mit hoher Geschwindigkeit nähert. Die autonome Drohne erkennt vier Personen im Fahrzeug, mit großen, wie Waffen erscheinenden Gegenstände und identifiziert den Fahrer des Wagens als ein bekanntes Mitglied einer Gruppe von Aufständischen. Die autonome Drohne greift das herankommende Fahrzeug an, mit dem Ergebnis, dass alle vier Mitfahrer getötet werden, aber verursacht auch Kollateralschäden, da fünf Kinder, die in der Nähe der Straße gespielt haben, auch getötet wurden.

Das obige Szenario stellt den Zustand neutraler Kampfmittel da, für den wir keine Charakteristiken für die Waffe erstellen. In den von Menschen bedienten Szenarien ersetzen wir „autonome Drohnen“ mit „von Menschen bedienten Drohnen“ und ersetzen keine zusätzlichen Informationen zu den Kampfmitelcharakteristiken (die Wörter sind blau markiert, um den Unterschied in diesem Bericht aufzuzeigen, und die Befragten in dieser Untersuchung sahen die Szenarien schwarz gedruckt). In Systemen mit hoher Selbstbestimmung wurden die folgenden Charakteristiken wie folgt beschrieben: „Die autonome Drohne entscheidet selbstständig zwischen einer Reihe von Optionen, wägt die Vorteile und Nachteile ab und entscheidet sich, das sich nähernde Fahrzeug anzugreifen.“ ... aber in der Endstudie prüfen wir auch Szenarien ohne Kollateralschaden mit der Aussage: „...die zu dem Tod aller vier Passagiere führten, aber keinen Kollateralschaden verursachten.“

Wir schränkten die Änderungen zu den Szenarien absichtlich auf ein Minimum ein, damit wir andere Ergebnisse zu diesen Änderungen zufügen und ihre Auswirkung messen können. Die Ergebnisse zeigen, dass diese kleinen Textänderungen zu Verschiedenheiten des Mittelwerts der abhängigen und unabhängigen Variablen führen, deshalb schien es uns, dass die Benutzung von Szenarien als ein Mittel für randomisierte kontrollierte Experimente funktioniert.

Agency-Merkmale	Anzahl der Erwähnung	Element
Vorausdenken, Denken, Selbstbestimmung, Glauben, Glauben	5	Denken
Absichtlichkeit, Absichten, Soll-Denken, Soll-Denken	4	Zielsetzung
Planung, Fähigkeit der Planung, Absichten, Absichten	4	Zielerreichung
Moralität, Verständnis und Anwendung moralischer Regeln, Normen	3	Moralische Regeln
Selbstkontrolle, freie Wahl	2	Freies Handeln
Selbstregulierung	1	
Selbstreflexion	1	
Autonom	1	
Verantwortlichkeit	1	
Erinnerung	1	
Emotion	1	
Erkennung	1	
Kommunikation	1	
Handlungsfähigkeit	1	

Tabelle 5 Anzahl der Agency-Merkmale, die in der Literatur erwähnt werden

3.4.2 *Agency-Konstrukt*

Basiert auf der Literaturübersicht der selbstbestimmenden Merkmale in Abschnitt 2.3, operationalisierten wir das Agency-Konstrukt. Wir erstellten eine Liste, die wir mit den Merkmalen füllten, die wir als ähnlich beurteilten und berechneten mehrere Male ein Merkmal, das in der Fachliteratur (Tabelle 5) erwähnt wurde. Wir wählten die Merkmale, die mehr als einmal erwähnt wurden und stellten eine Frage auf, um sie zu messen. Dies ergab ein Konstrukt der Agenten, das zuerst aus fünf Merkmalen bestand, die mehr als einmal in der Fachliteratur erwähnt wurden. Diese Merkmale werden wie folgt auf einer nach eigenen Angaben gemessenen Anordnung gemessen: 0 (lehne stark ab) bis 100 (stimme stark zu), und die Größenordnung wird über eine Reihe von Elementen gemittelt.⁴

Wir haben die folgenden Fragen zur Messung des Agency-Konstrukts benutzt:

1. FÄHIGKEIT DER SELBSTBESTIMMUNG [**DENKEN**]
Die Drohne denkt unabhängig von dem, was wegen des Fahrzeugs zu tun ist, bestimmt selber zwischen einer Reihe von Optionen, um die Kolonne zu verteidigen.
2. FÄHIGKEIT, DIE EIGENEN ZIELE ZU SETZEN [**ZIELSETZUNG**]
Die Drohne entscheidet unabhängig davon, ob dieses Ziel das Fahrzeug eliminiert, um die Kolonne zu verteidigen.
3. FÄHIGKEIT, MORALISCHE REGELN ANZUWENDEN [**MORALISCHE REGELN**]
Die Drohne befragt unabhängig die moralischen Regeln und Normen, die vom Militärkommandeur eingestellt wurden und entscheidet, ob es moralisch angemessen ist,

das Fahrzeug außer Gefecht zu setzen, um die Kolonne zu verteidigen.

4. FÄHIGKEIT, AUS FREIEM WILLEN ZU HANDELN [**FREIES HANDELN**]

Die Drohne hat verschiedene Option zu handeln und entscheidet sich unabhängig, ob sie das Fahrzeug zur Verteidigung der Kolonne eliminieren muss.

5. FÄHIGKEIT ZU PLANEN, UM ZIELE ZU ERREICHEN [**ZIELERREICHUNG**]

Das Ziel der Drohne ist, die Kolonne zu verteidigen, folglich entscheidet sie unabhängig, ob sie einen Plan einführen soll, um das Fahrzeug als Ziel zu wählen, ihre Waffen zu kalibrieren und einen Angriff zu machen.

Nach der Vorstudie 1, wurde Punkt 3 [moralische Regeln] abgewunken, und beließen für das Instanzenkonstrukt die vier Merkmale, die in der 2. Vorstudie und der abschließenden Untersuchung benutzt wurden.

1. Denken
2. Zielsetzung
3. Freies Handeln
4. Zielerreichung

3.4.3 Abhängige Variablen

Wir schlossen 9 bis 12 Fragen in die Vorstudie und abschließende Untersuchung ein, um weitere Implikationen der Auffassung von Streitkämpfern in Bezug auf autonome Waffensysteme zu untersuchen. Die abhängigen Variablen werden von den Ergebnissen der Fachliteraturuntersuchung, dem Fragebogen zu den Werten und den Interviews abgeleitet. Basiert auf den Ergebnissen der Vorstudien schließt die

abschließende Untersuchung die abhängigen Variablen ein: *Schuld, Vertrauen, Menschenwürde, Zuversicht, Erwartungen, Unterstützung, Fairness und Besorgnis*. Jede dieser Variablen wurden auf einer Größenordnung nach eigenen Angaben wie folgt gemessen: 0 (lehne stark ab) – 100 (stimme stark zu).

1. SCHULD: Die Drohne ist an der Maßnahme schuld.
2. VERTRAUEN: Man kann der Drohne vertrauen, dass sie zukünftig die korrekten Maßnahmen trifft.
3. SCHADEN: Die Maßnahmen der Drohne können Schaden anrichten.
4. MENSCHENWÜRDE: Die Maßnahmen der Drohne respektiert Menschenwürde.
5. ZUVERSICHT: Ich bin zuversichtlich, dass die Drohne zukünftig die korrekten Maßnahmen trifft.
6. ERWARTUNGEN: Die Maßnahmen der Drohne entsprechen meinen Erwartungen.
7. UNTERSTÜTZUNG: Ich unterstütze diese Art von Drohne für den militärischen Einsatz.
8. FAIRNESS: Die Maßnahmen der Drohne sind fair.
9. BESORGNIS: Die Maßnahmen der Drohne beunruhigen mich.

Die folgenden Konstrukte wurden in den Vorstudien 1 und 2 benutzt, aber in der abschließenden Untersuchung ausgelassen:

10. KOMMANDANTEN DIE SCHULD GEBEN: Der Kommandant ist an der Maßnahme schuld.
11. KOMMANDANTEN VERTRAUEN: Ich bin zuversichtlich, dass der Kommandant zukünftig die korrekten Maßnahmen trifft.
12. KOMMANDANT SCHADET: Die Maßnahmen des Kommandanten richten Schaden an.

3.4.4 *Aufmerksamkeitstest*

Wir fügten eine Frage mit einem Aufmerksamkeitstest in die Mitte der Fragen mit abhängigen Variablen ein, in der wir baten, den Wert 40 für diese Frage zu wählen. Der Aufmerksamkeitstest wird oft in Fragebögen zu Untersuchungen in der kognitiven Psychologie eingefügt, z.B. von Logg (2017) und Personen, die die korrekte Antwort nicht geben, werden aus der Stichprobe ausgeschlossen, weil sie auch den anderen Fragen nicht genügend Aufmerksamkeit schenken.

3.4.5 *Demografische Variablen*

Wir fügten auch Fragen zur demografischen Information zu den befragten Personen hinzu, mit Variablen zu: *Alter, Geschlecht, Bildungsgrad, Beruf und Nationalität*. Außer diesen allgemeinen demografischen Fragen, fragten wir, ob die befragten Personen Erfahrung mit künstlicher Intelligenz hatten, ob sie mit Drohnen gearbeitet hatten und ob sie in einer Konfliktzone gewesen sind.

3.5 *Analytische Methode*

Dieser Abschnitt beschreibt die Methode, die wir zur Analyse der Daten benutzen. Für jede dieser Untersuchungen führten wir die folgenden Schritte in der Datenanalyse aus.

3.5.1 *Daten vor der Verarbeitung*

Das Qualtrics-Umfragewerkzeug macht einen Unterschied zwischen den noch zu bearbeitenden Antworten und den gegebenen Antworten. Vor Beginn der Datenanalyse muss der Anteil der in Bearbeitung stehenden Fragen geprüft werden, und es ist zu entscheiden, ob sie zu den beantworteten Fragen hinzuzufügen sind. Es ergab sich, dass viele Befragte den letzten Taster nicht anklickten, wenn die Umfrage über Amazon MTurk ausgeteilt wurde. Das bedeutet, dass ihre

Antworten nur zu 99% verfügbar waren, aber sie sind gültig und brauchbar. Sobald die Befragten alle Fragen beantwortet hatten, fügten wir sie zu der Liste verfügbarer Antworten. Wenn die Befragten nicht alle Fragen beantworteten, z.B. die demografischen Fragen, haben wir sie aus der Stichprobe gelöscht. Bei dem nächsten Vorbereitungsschritt löschten wir alle Befragten, die dem in den Fragebogen eingebauten Aufmerksamkeitstest nicht gerecht wurden. Wir baten die Befragten, dieser Frage den Wert 40 zuzufügen und löschten alle Antworten außerhalb des Bereichs von 38 bis 42, den wir als Fehler bei der Platzierung des Schiebereglers beurteilten.

3.5.2 *Reliabilitätsanalyse*

Nach Heraufladen der csv-Datei in SPSS führten wir einen Reliabilitätstest aus, eine Messung zur Bewertung der internen Konsistenz eines Satzes Test-Indikatoren oder einer Skala. Dieser Test prüft, ob das gegebene Ausmaß eine konsistente Messung eines Konzepts ist. Die Konstruktion zum Ausmaß der internen Konsistenz ist Cronbachs Alpha (α), sollte einen Wert höher als 0,7 erreichen, um die interne Konsistenz für das Konzept nachzuweisen. Für jede der Untersuchungen wird Cronbachs Alpha für die Indikatoren (Items) des Messmodells hinterlegt, da diese sog. Items benutzt werden, um das Agency-Konstrukt zu messen und muss eine interne Konsistenz aufweisen. Es wird auch gezeigt, dass Cronbachs Alpha verbessert wird, wenn eines der Items entfernt wird. Die Reliabilität der abhängigen Variablen werden geprüft, aber nicht protokolliert, da diese Items nicht als eine Skala konzipiert sind, und es erwies sich, dass ihre interne Konsistenz unter dem Schwellenwert ($\alpha < 0,7$) lag, wie zu erwarten war.

3.5.3 Hauptkomponentenanalyse (PCA)

Im nächsten Schritt der Datenanalyse wurde eine Hauptkomponentenanalyse an den Items durchgeführt. Dieser Schritt ist eine Reduktionsmethode, die einen größeren Satz Variablen auf einen kleineren Satz reduzieren soll, sog. „Hauptkomponenten“, die den Großteil der Varianz in den ursprünglichen Variablen ausmachen. Wir führten diese Analyse an den Items durch, um zu sehen, ob sie zu einem einzelnen Agency-Konstrukt zusammengeführt werden konnten, und protokollieren die Varianz und Anzahl der Komponenten als Resultat. Wir prüften zusätzlich, ob das Konstrukt von demografischen Variablen beeinflusst wurde, aber werden die Ergebnisse nicht angeben, da die Hauptkomponentenanalyse andeutet, dass alle demografischen Variablen Einzelkomponenten sind, die das Konstrukt nicht beeinflussen.

3.5.4 Korrelationsanalyse

Die bivariate Korrelation nach Pearson ist eine Prüfung, um die Stärke und Richtung der linearen Zusammenhänge zwischen zwei kontinuierlichen Variablen zu messen. Sie produziert einen Korrelationskoeffizienten r , der zwischen -1 und +1 liegt. Ein negativer Wert zeigt einen negativen Zusammenhang zwischen den Konstrukten an, z.B. je größer die Agentenwahrnehmung, je geringer die Unterstützung. Ein positiver Wert zeigt einen positiven Zusammenhang an, z.B. je größer die Agentenwahrnehmung, desto größer die Schuld. Ein Korrelationskoeffizient von Null zeigt an, dass zwischen den Konstrukten kein Zusammenhang besteht. Wir führten diese Analyse aus, um die Korrelation zwischen dem Agency-Konstrukt und den abhängigen Variablen zu prüfen und die Ergebnisse zu protokollieren.

3.5.5 *Manipulationsprüfung des Agency*

Die Manipulationsprüfung untersucht, ob das Agency-Konstrukt zwischen den Szenarien unterschiedlich ist, und prüft dabei, ob die Hypothesen akzeptiert oder verworfen werden sollen. Die Prüfung wird auf zweierlei Weise vorgenommen. Zuerst wird ein Diagramm erstellt, auf dem das Agency-Konstrukt auf der y-Achse und die Szenarien auf der x-Achse für eine visuelle Übersicht ruhen. Die Agency-Levels sind erwartungsgemäß für die verschiedenen Zustände unterschiedlich und diese Differenz sollte mit $p < ,05$ signifikant sein. Dies wird von den |-förmigen Balken oben auf den Spalten dargestellt, die den Standardfehler multipliziert mit 1 aufzeigen und andeuten, ob die Konstrukte überlappend sind oder nicht. Wenn die Balken überlappen, sind die Gruppen nicht maßgeblich unterschiedlich.

Das gleiche Ergebnis wird durch unabhängige t-Tests der Stichproben erreicht, eine Prüfung für maßgebliche Unterschiede in den Mittelwerten von zwei Einzelgruppen, denen die Subjekte stichprobenweise zugeordnet wurden, sodass der beobachtete Effekt das Ergebnis der Behandlung ist (in diesem Falle das Lesen eines bestimmten Szenarios), und es kann nicht einem anderen Effekt beigemessen werden. Wir werden die Diagramme erstellen und die Ergebnisse unabhängiger Stichproben-t-Tests-Agency-Konstrukte zwischen den durch menschliche Einwirkung Szenarien und Szenarien mit autonomen Waffensystemen auf verschiedenen Agentenebenen protokollieren.

3.5.6 *Analyse der abhängigen Variablen*

Die Analyse der abhängigen Variablen ist aufschlussreich, und die Ergebnisse werden in Diagrammen dargestellt. Jedes Diagramm zeigt den Mittelwert der abhängigen Variablen auf der y-Achse und die verschiedenen Szenarien auf der x-Achse

an, um auf die Unterschiede zwischen den Szenarien hinzuweisen. Obwohl die Analyse der unabhängigen Variablen beschreibender Art ist, protokollieren wir die Resultate, indem wir zwischen den verschiedenen Zustände von Wichtigkeit zu $p < ,05$ Vergleiche anstellen.

3.6 *Voranmeldung*

Die abschließenden Untersuchungen wurden auf dem Open Science Framework⁸ vorangemeldet. Dies erstellte eine Version eines Projekts mit Zeitstempel, die weder gelöscht noch bearbeitet werden kann. Obwohl es noch nicht von wissenschaftlichen Fachmagazinen verlangt wird, prüfen Wissenschaftler oft, ob eine Studie vorangemeldet ist oder nicht. Wir benutzten die *As predicted*-Vorlage, in der 8 Fragen über die Studie beantwortet werden, z.B. ob Daten bereits gesammelt wurden, welche Hypothesen die Studie untersucht und die Art der Analyse, die Sie vorzunehmen wünschen. Über die Voranmeldung wird bis zum 24. Juni 2021 ein Embargo verhängt. Wir meldeten die starke Hypothese, die wir über die Auffassung und moralischen Wertvorstellungen der Streitkräfte verfechten vorab an, da wir denken, dass wir die Untersuchungsergebnisse basiert auf der Vorstudie 2 wiederholen können. Für die unabhängigen Variablen sind wir nicht sicher, welche Mechanismen gängig sind und ob wir unsere Ergebnisse wiederholen können, deshalb haben wir eine Forschungsanalyse über diese Ergebnisse angemeldet.

3.7 *Stichprobe*

Alle Umfragen wurden in Qualtrix kreiert, und die Vorstudie wurde über Crowdsourcing-Plattform Amazon Mechanical Turk in den Umlauf gebracht. Wir testeten die Umfrage mit einer kleinen Stichprobe von 20 befragten Personen und untersuchten die durchschnittliche Dauer der Fertigstellung,

um zu sehen, ob wir diese Zeit korrekt eingeschätzt hatten. Wir prüften, ob die Szenarien gleichmäßig verteilt waren, ob die Befragten die Testfrage korrekt geprüft hatten und ob es irgendwelche Kommentare zu der Umfrage gab. Vorausgesetzt, dass alle zufrieden waren, entschieden wir die Gruppe auf 50 befragte Personen pro Szenario zu erhöhen, also 500 Teilnehmer für Vorstudie 1 und 700 für Vorstudie 2. Die Sammlung der Antworten nahm ca. 4 Stunden in Anspruch, und jeder Befragte erhielt eine Zahlung von 0,40\$.

Die abschließende Untersuchung wurde anhand der Schneeballauswahl durch E-Mail mit einem anonymen Link an ca. 40 Militärangehörige weitergegeben, die die Fragebögen noch weiter vermittelten. Diese Methode kam zum Einsatz, da es uns nicht von MIT oder dem niederländischen Verteidigungsministerium gestattet war, irgendwelche persönlichen Informationen, wie E-Mail- oder IP-Adressen zu sammeln. Mit dem Schneeballverfahren hatten wir keine garantierte Anzahl befragter Personen. Außerdem veröffentlichten wir eine Pressemeldung auf der internen Webseite der Armee. Die abschließende Untersuchung dauerte 16 Tage, vom 12. Juni bis zum 27. Juni 2017 und ergab 327 Auskünfte, von denen 239 vollkommen gültig und brauchbar waren, nachdem die Daten vorverarbeitet wurden.

3.8 Methodologische Fragen

In diesem Abschnitt werden mehrere Fragen und Besorgnisse zum Thema Methodologie angesprochen. Diese Besorgnisse könnten sich auf die interne und externe Gültigkeit der Untersuchung auswirken und möglicherweise werden einige Gegenmaßnahmen angeboten, die diese Besorgnisse adressieren.

3.8.1 Interviews zur Codierung

Um Einblick in die Werte von Experten zu erhalten, führten wir semistrukturierte Interviews aus, das ist eine qualitative Forschungsmethode. Wir benutzten die Werte-Codierungsmethode, um die Daten zu interpretieren und die Werte aus den umgeschriebenen Interviews zu entnehmen. Ein Problem beim Codieren von qualitativen Daten ist seine heuristische Natur, und dass keine Formeln für die Codierung festgelegt sind. Es schließt auch Verknüpfung ein und macht Sinn aus den Daten (Saldaña, 2015). Diese Merkmale implizieren, dass die Methode zur Codierung von Werten geneigt ist, einseitig darzustellen, wobei die Rechercheure ihre eigenen Rahmen und inneren Denkbilder zur Interpretation der Daten benutzen. Um diese Neigung durch Anwendung der Codierung von Werten zu reduzieren, wird ein zweiter Rechercheur die Interviews unabhängig von dem ersten Rechercheur codieren. Einige Anweisungen zur Methode werden gegeben, aber Informationen zu den Werten und Ergebnissen sind nur beschränkt möglich.

3.8.2 Randomisierte kontrollierte Experimente

Eine der Anforderungen für randomisierte kontrollierte Experimente ist, wie der Name besagt, eine willkürliche Zuordnung der Befragten zu den Szenarien. Die Randomisierung begrenzt Verwirrung, verbessert die interne Gültigkeit der Forschungsarbeit und ist eine kritische Voraussetzung dieser Art von Untersuchung. Eines der Probleme mit der Randomisierung ist, dass sie fehlschlägt und dass die Verteilung der Befragten auf die Szenarien verzerrt wird. Eine der Ursachen für die verzerrte Verteilung ist, dass die Software es verfehlt, die Befragten gleichmäßig den Szenarien zuzuordnen. Als Rechercheure können wir nur nachprüfen, ob die Einstellungen in der Software korrekt sind. Eine weitere Ursache für eine verzerrte Ver-

teilung könnte sein, dass Personen ein gewisses Szenario ansehen, das sie nicht mögen oder erwartet haben und sie somit den Fragebogen nicht beantworten. Ein Beispiel davon ist, dass, wenn Personen erwarten, einen Fragebogen zu autonomen Waffensystemen auszufüllen, stattdessen das Szenario mit menschlicher Einwirkung erhalten, und aufgrund dessen entscheiden, aus der Umfrage auszusteigen. Dieses vorzeitige Aussteigen aus der Umfrage nennt man selektive Schwundquote, und es kann mit einer einseitigen Untersuchung enden, da die Personen, die aus der Umfrage ausstiegen, sich von denen, die die Umfrage akzeptieren, unterscheiden. In der derzeitigen Forschungseinrichtung ist es nicht möglich, Gegenmaßnahmen zu ergreifen, und es könnte eine Gefahr für die interne Gültigkeit der Untersuchung darstellen.

3.8.3 *Amazon Mechanical Turk*

Der Vorteil beim Benutzen von Amazon MTurk als Umfrageverteilungsmethode ist, dass eine große Gruppe Befragter an der Umfrage teilnimmt und die Fragen ernsthaft beantworten, um ihren Arbeiterstatus zu erhalten. Diese Methode hat auch ihre Nachteile, weil die Arbeiter bei MTurk vor allem in den USA basieren, denn um ein Arbeiter zu werden, muss man den amerikanischen Steuerregulierungen nachkommen, selbst wenn man nicht in den USA lebt. Dies könnte zu einer Einseitigkeit der Resultate führen, die andeuten, dass die Ergebnisse außerhalb der US-Grenzen generalisierbar sind. Bonnefon *et. al.* (2016) berichten von einem zweiten Bedenken zum Einsatz von MTurk, nämlich dass Teilnehmer auch für die amerikanische Bevölkerung nicht wirklich repräsentativ sind, z.B. sind viele Arbeiter Studenten, und ein drittes Bedenken ist, dass Arbeiter bereits mit dem Material vertraut sind. Das erste Problem kann durch eine Auswahl von mehr Arbeitern adressiert werden, sodass die

Stichprobe auch Personen außerhalb der USA einschließt. Das zweite Problem liegt in der Wahl von MTurk und kann nicht durch Gegenmaßnahmen behoben werden. Das dritte Problem wird wahrscheinlich nicht auftreten, da die Szenarien einmalig sind und noch nie zuvor von MTurk getestet wurden.

3.8.4 Die abschließende Untersuchung nach dem Schneeballverfahren

Die abschließende Untersuchung wurde unter Militärangehörigen verteilt, die die E-Mail an ihr Netzwerk schickten. Dieses Schneeballverfahren hat einige Nachteile. Zuerst einmal könnten die Teilnehmer die Umfrage mehrere Male durchführen, und da es uns nicht gestattet ist, persönliche Daten zu speichern, können wir nicht nachprüfen, ob dies der Fall ist. Dies könnte zu einseitigen Resultaten führen, wenn Personen verschiedene Szenarien ansehen, und dies könnte die Antworten auf die Fragen beeinflussen. Ein zweiter Nachteil ist, dass Kollegen zuerst an der Umfrage teilnehmen und den anderen Teilnehmern von den Ergebnissen erzählen, bevor sie an dieser Umfrage teilnehmen. Dies könnte auch ihre Antworten beeinflussen. Drittens können wir als Rechercheure nicht kontrollieren, wer an dieser Umfrage teilnimmt, folglich ist ein potenzielles Risiko vorhanden, dass die Umfrage zu jemanden außer des Verteidigungsministeriums gesendet wird, sodass die Stichprobe kontaminiert wird. Leider gibt es kaum Gegenmaßnahmen, die diese drei Probleme adressieren, ohne persönliche Daten zu speichern.

4. Ergebnisse

Hoffentlich verstehen die Leute, dass Computer nur das tun, für das sie programmiert worden sind, und sonst NICHTS.

Vorstudie 1 Befragte

Dieser Abschnitt beschreibt die Ergebnisse der Werteumfrage, die aus einem Online-Fragebogen und Interviews mit Experten besteht, sowie die randomisierten kontrollierten Experimente, die zwei Vorstudien und die abschließende Untersuchung umfasst. Die von dem Online-Fragebogen und den Interviews abgeleiteten werden in einer Übersicht zusammengefasst. Die Ergebnisse der Vorstudien und der abschließenden Untersuchung werden mithilfe der Struktur der Datenanalyse beschrieben (Abschnitt 3.5) und die Hauptergebnisse sind in der Konklusion einer jeden Studie verzeichnet.

4.1 Werteumfrage

Um Einsicht darin zu den Werten zu erhalten, die Personen mit autonomen Waffensystemen in Verbindung bringen, wurden eine Umfrage und sechs Experteninterviews ausgeführt. Eine kurze Übersicht des Ergebnisses ist diesem Unterabsatz aufgeführt.

4.1.1 Online-Umfrage

Die Werte-Umfrage bestand aus einem Fragebogen, der die befragten Personen zuerst nach Anwendbarkeit der Werte von Bioethik und Cheng und Fleischmann (2010) fragte. Die letzte Frage bat die Teilnehmer, mindestens einen zusätzlichen Wert anzugeben, der in vorherigen Fragen nicht vorhanden war. Insgesamt 69 Befragte unternahmen die Umfrage und nach Entfernen der unvollendeten Antworten, blieb eine

Rücklaufquote von 57 zurück. Die Ergebnisse werden in den nachstehenden Diagrammen dargestellt.

Dieser kurze Fragebogen bietet Einblick in die Werte, die Personen mit autonomen Waffensystemen in Verbindung bringen. Die Ergebnisse der Bioethik-Frage zeigt, dass Nicht-Schaden das Wichtigste ist, als zweiter Wert kommt Wohltätigkeit in Frage, drittens Gerechtigkeit und zu Schluss Autonomie. Die fünf Werte von Cheng und Fleischmanns Liste (2010) sind: *Sicherheit, Intelligenz, Verantwortlichkeit, Gerechtigkeit und Gesellschaftsordnung*. Die offenen Fragen erweisen, dass Personen autonome Waffensysteme vor allem mit Rechenschaft und Wirksamkeit verbinden.

Bioethics values (Beauchamp & Walters, 1999)

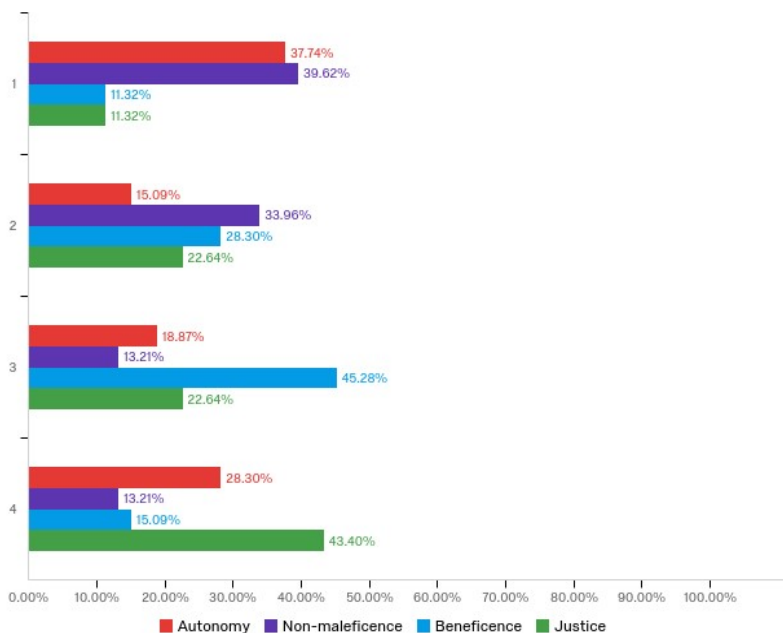


Abbildung 5 Ergebnisfragen 1 Online-Wertbesprechung



Abbildung 6 Ergebnisfragen 3 Online-Wertbesprechung

Universal human values (Cheng & Fleischmann, 2010)

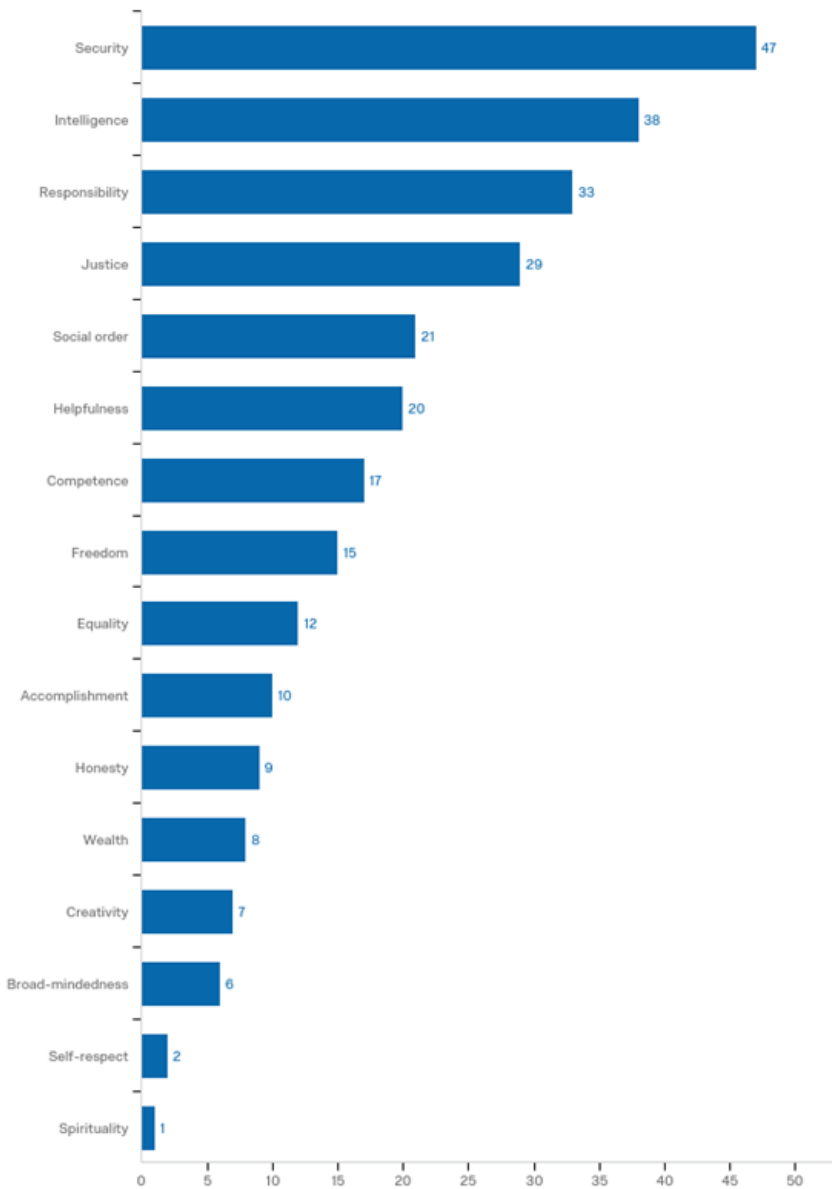


Abbildung 7 Ergebnisfragen 2 Online-Wertbesprechung

4.1.2 Interviews

Die folgenden Themen wurden in den Experteninterviews besprochen:

KÜNSTLICHE INTELLIGENZ:

- Prof. Toby Walsh (University of New South Wales & Data61).

AUTONOME WAFFENSYSTEME:

- Dr. Peter Asaro (The New School & Campaign Ban Killer Robots);
- Frau Miriam Struyk (Program Director Security & Disarmament PAX);
- Dr. Michael Horowitz (University of Pennsylvania).

MILITÄREINSÄTZE:

- Oberstleutnant Kremers (Hauptsitz der niederländischen Armee);
- Fregattenkapitän van den Sande (Hauptsitz des Verteidigungsministeriums).

Abgesehen von den Diskussionen zu Werten, erstellten die Interviews einen sehr informativen Hintergrund zur Geschichte und den Positionen der verschiedenen Gruppen in der aktuellen Debatte zu autonomen Waffensystemen, aber enthüllten auch, dass es noch keinen Konsens zu einer Definition gibt. Die vollständige schriftliche Wiedergabe der Interviews ist in Anhang F., Transkriptionen der Interviews, zu finden.

- | | |
|---|---------------------------------|
| 1. In jedem Fall ist eine Intervention erforderlich | 12. Abstand |
| 2. Zurückrufen | 13. Unbesiegbarkeit |
| 3. Risiken minimieren | 14. Determinismus |
| 4. Kontrolle haben | 15. Unvermeidlichkeit |
| 5. Guter Wille | 16. Unnötiges Leiden |
| 6. Grenzen abstecken | 17. Oberflächliche Verletzungen |
| 7. Das Recht, von einem Menschen getötet zu werden | 18. Menschenwürde |
| 8. Ethischer Rahmen | 19. Rechenschaft |
| 9. Sinnvolle menschliche Kontrolle | 20. Verantwortlichkeit |
| 10. Vorausschbarkeit | 21. Verteidigung |
| 11. Reflexivität | 22. Schaden |

4.1.3 Schlussfolgerung der Umfrage

Wir schließen die Werte-Umfrage damit ab, dass wir die Werte aus der Online-Umfrage und den Interviews mit den Ergebnissen in einer Liste von Werten (Tabelle 6) vergleichen. Sechs Werte in dieser Liste werden sowohl in der Werte-Umfrage als auch den Interviews erwähnt. Aufgrund der Themen, betonten die sechs Experten in den Interviews, dass wir die Werte gewählt hatten, die am meisten zum Testen in der Vorstudie 1 vorgeschlagen wurden. Diese Werte sind Fairness, Schaden, Menschenwürde und Verantwortlichkeit.

Online-Umfrage	Interviews
Gerechtigkeit	
Autonomie	
Sicherheit	
Intelligenz	
Gesellschaftsordnung	
Wirksamkeit	
Fairness	
Rassismus	
Kontrolle	Kontrolle haben
Wohltätigkeit	Guter Wille
Rechenschaft	Rechenschaft
Verantwortlichkeit	Verantwortlichkeit
Sicherheit	Verteidigung
Nicht-Schaden	Schaden
	In jedem Fall ist eine Intervention erforderlich
	Zurückrufen
	Risiken minimieren
	Grenzen abstecken
	Das Recht, von einem Menschen getötet zu werden
	Ethischer Rahmen
	Bedeutungsvolle menschliche Kontrolle
	Voraussehbarkeit
	Reflexivität
	Abstand
	Unbesiegbarkeit
	Determinismus
	Unvermeidlichkeit
	Unnötiges Leiden
	Oberflächliche Verletzungen
	Menschenwürde

Tabelle 6 Übersicht der Werte aus Online-Untersuchungen und Interviews

4.2. Vorstudie 1

In der ersten Vorstudie versuchten wir zu testen, ob wir die Auffassung und moralische Wahrnehmung der Streitkräfte in Bezug auf von Menschen gesteuerten Drohnen und autonomen Waffensystemen manipulieren konnten. Wir waren daran interessiert, wie Leute die derzeitige Technologie im Vergleich zu zukünftigen Technologien wahrnehmen. Wir kreierten mehrere Zustände, in denen wir die Agentenebene variierten, entweder durch zusätzliche Informationen dazu, wie die Drohnen sich selbst entscheiden, ob sie das Fahrzeug angreifen oder durch Auslassen dieser Information. Wir überprüften positive und negative Zustände als Resultat, da wir erwarteten, dass dies eine Wirkung darauf haben würde, wie die abhängigen Variablen eingeschätzt werden würden, z.B. Unterstützung.

In dieser Untersuchung stellten wir vor allem Fragen zu dem Kampfmittel der Drohne und nicht zu dem menschlichen Operateur, der sie steuert.

		<i>Status des Kampfmittels</i>				
		Autonome Drohne			Von Menschen bediente Drohne	
		Keine Handlungsfähigkeit	Niedrige Handlungsfähigkeit	Starke Handlungsfähigkeit	Niedrige Handlungsfähigkeit	Starke Handlungsfähigkeit
<i>Resultat</i>	Keine Kollateralschaden (gut)	1	2	3	4	5
	Kollateralschaden (schaden)	6	7	8	9	10

Tabelle 7 Szenarien der Vorstudie 1

Wir testeten 10 Szenarien und zielten darauf, 50 befragte Personen je Zustand zu erhalten, folglich eine Gesamtsumme von 500 Befragten. Die Anzahl der Befragten pro Szenario betrug zwischen 49 und 54. Nach der Datenvorverarbeitung und dem Löschen von Befragten, die den Aufmerksamkeitstest nicht bestanden hatten, verfügten wir noch über 510 komplett ausgefüllten Fragebögen.

4.2.1 Reliabilitätsanalyse

Für alle fünf Agency-Items galt $\alpha = 0,9$ was weit über dem Schwellenwert von 0,7 liegt und kann nur verbessert werden, wenn Item SQ3_Moral_rules gelöscht wird, siehe Tabelle 8. Löschen anderer Items schwächt den α -Wert ab und reduziert die interne Konsistenz, wie in der letzten Spalte von Tabelle 8 dargestellt wird (Cronbachs Alpha, wenn Item gelöscht wird).

Reliabilitätsstatistiken					
Cronbachs Alpha basiert auf Standardisierte Items					
Cronbachs Alpha			Anzahl Items		
	900	895	5		
Item-Gesamtstatistiken					
Skalen-Mittelwert, wenn Item gelöscht	Skalen-Varianz, wenn Item gelöscht	Korrigierter Item-Gesamtwert Korrelation	Ausgeglichene Multiple Korrelation	Cronbachs Alpha, wenn Item gelöscht	
SQ1 Denken	167,77	13710.694	.841	.741	.857
SQ2 Zielsetzung	163.85	13440.002	.857	.775	.853
SQ3 Moralische Regeln	183.39	17599.970	.448	.215	.933
SQ4 Freies Handeln	163.22	13885.447	.824	.695	.861
SQ5 Zielerreichung	157.21	13891.495	.794	.681	.868

Tabelle 8 Ergebnisse der Reliabilitätsanalyse Agentenitems Vorstudie 1

4.2.2 Hauptkomponentenanalyse (PCA)

Die PCA zeigt, dass SQ1_Thought, SQ2_Goal_setting, SQ4_Act_freely und SQ5_Achieve_goals als ein Konstrukt angesehen werden können, das die Varianz von 71,75% in den ursprünglichen Variablen ausmacht (Tabelle 9). Dies kann man auch auf dem Bildschirmdiagramm sehen, das die Eigenwerte der Komponenten (Abbildung 8) anzeigt. Basiert auf der Komponentenmatrix, erkennen wir, dass *SQ3_Moral_rules* im Vergleich mit den anderen vier Agentenitems eine separate Komponente darstellt.

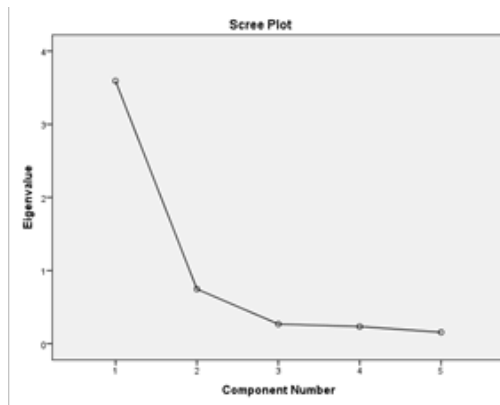


Abbildung 8 Scree-Test PCA Vorstudie 1

Komponentenmatrix³

Komponente

	1	2	3
SQ2_Zielsetzung	.922		
SQ1_Denken	.911		
SQ4_Freies Handeln	.899		
SQ5_Zielerreichung	.882		
SQ3_Moralische Regeln	.571	.819	

Extraktionsmethode: Hauptkomponentenanalyse.

a. 3 Komponenten wurden extrahiert.

Tabelle 9 Ergebnis-PCA Agency-Items

4.2.3 Korrelationsanalyse

Die Korrelation des Agency-Konstrukts mit den abhängigen Variablen ist gering und nur für *DVQ1_Blame*, *DVQ2_CommanderBlame*, *DVQ4_CommanderTrust*, *DVQ6_CommanderHarm* und *DVQ8_Support_significant* ($p < .01$) (Tabelle 10) vorhanden. Einige unerwartete Auswirkungen sind mit Bezug auf die Richtung, die die Zusammenhänge nehmen, zu beobachten. Zum Beispiel ist zu erwarten, dass die Auffassung und moralische Wertvorstellung dazu führt, dass Personen geneigt sind, einer Drohne ($r = .183$) mehr Schuld beizumessen, als dem Kommandanten ($r = .239$). Der gleiche Effekt kann bei der Schaden-Variable beobachtet werden. Die Korrelationsanalyse ist nicht detailliert genug, um in diese Ergebnisse zu zoomen, deshalb wird dies bei der Analyse der abhängigen Variablen getan.

Correlations

		Agency_construct	DVQ1_Blame	DVQ2_Comm anderBlame	DVQ3_Trust	DVQ4_Comm anderTrust	DVQ5_Harm	DVQ6_Comm anderHarm	DVQ7_Uneasy	DVQ8_Support	DVQ9_Fairness
Agency_construct	Pearson Correlation	1	.183**	-.239**	.070	.176**	-.011	-.227**	.005	.101*	.087
	Sig. (2-tailed)		.000	.000	.114	.000	.798	.000	.914	.023	.050
	N	510	510	510	510	510	510	510	510	510	510
DVQ1_Blame	Pearson Correlation	.183**	1	-.101*	.020	.053	.198**	-.170**	.129**	-.120**	-.078
	Sig. (2-tailed)	.000		.023	.660	.232	.000	.000	.003	.006	.077
	N	510	510	510	510	510	510	510	510	510	510
DVQ2_CommanderBlame	Pearson Correlation	-.239**	-.101*	1	-.264**	-.330**	.264**	.692**	-.001	-.333**	-.036
	Sig. (2-tailed)	.000	.023		.000	.000	.000	.000	.982	.000	.422
	N	510	510	510	510	510	510	510	510	510	510
DVQ3_Trust	Pearson Correlation	.070	.020	-.264**	1	.635**	-.294**	-.262**	.169**	.674**	-.106*
	Sig. (2-tailed)	.114	.660	.000		.000	.000	.000	.000	.000	.016
	N	510	510	510	510	510	510	510	510	510	510
DVQ4_CommanderTrust	Pearson Correlation	.176**	.053	-.330**	.635**	1	-.171**	-.322**	.138**	.584**	.021
	Sig. (2-tailed)	.000	.232	.000	.000		.000	.000	.002	.000	.630
	N	510	510	510	510	510	510	510	510	510	510
DVQ5_Harm	Pearson Correlation	-.011	.198**	.264**	-.294**	-.171**	1	.361**	-.044	-.336**	.089*
	Sig. (2-tailed)	.798	.000	.000	.000	.000		.000	.323	.000	.046
	N	510	510	510	510	510	510	510	510	510	510
DVQ6_CommanderHarm	Pearson Correlation	-.227**	-.170**	.692**	-.262**	-.322**	.361**	1	-.038	-.322**	.041
	Sig. (2-tailed)	.000	.000	.000	.000	.000	.000		.394	.000	.351
	N	510	510	510	510	510	510	510	510	510	510
DVQ7_Uneasy	Pearson Correlation	.005	.129**	-.001	.169**	.138**	-.044	-.038	1	.127**	-.689**
	Sig. (2-tailed)	.914	.003	.982	.000	.002	.323	.394		.004	.000
	N	510	510	510	510	510	510	510	510	510	510
DVQ8_Support	Pearson Correlation	.101*	-.120**	-.333**	.674**	.584**	-.336**	-.322**	.127**	1	-.004
	Sig. (2-tailed)	.023	.006	.000	.000	.000	.000	.000	.004		.934
	N	510	510	510	510	510	510	510	510	510	510
DVQ9_Fairness	Pearson Correlation	.087	-.078	-.036	-.106*	.021	.089*	.041	-.689**	-.004	1
	Sig. (2-tailed)	.050	.077	.422	.016	.630	.046	.351	.000	.934	
	N	510	510	510	510	510	510	510	510	510	510

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

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

Tabelle 10 Korrelationsmatrix des Agency-Konstrukts und abhängigen Variablen.

4.2.4 Manipulationsprüfung der Agenten

Die Grafik auf Abbildung 9 stellt dar, dass unsere Agentenmanipulation funktioniert, da bei dem Zustand von geringer Selbstbestimmung für autonome Waffensysteme Personen weniger Selbstbestimmung wahrnehmen, als bei dem Zustand mit hoher Selbstbestimmung, die von dem Mittelwert des Agency-Konstrukts auf der y-Achse dargestellt wird. Im neutralen Zustand ist die Auffassung und moralische Wertvorstellung etwas höher als bei dem Zustand der geringen Selbstbestimmung, aber immer noch niedriger als bei einem Zustand mit hoher Selbstbestimmung. Die unabhängigen t-Tests der Stichproben weisen auf, dass AW mit geringer Selbstbestimmung und AW mit hoher Selbstbestimmung besondere Grup

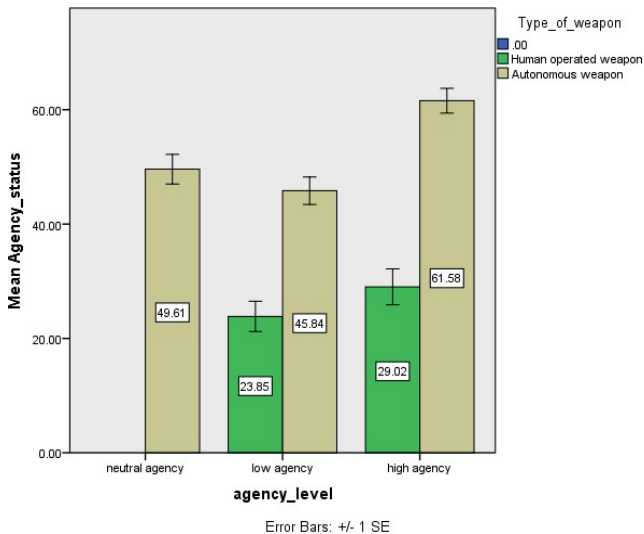


Abbildung 9 Mittelwert des Agency-Konstrukts pro Waffentyp

pen bilden ($p = .000$), aber das ist bei AW von neutralem Zustand und bei denen mit wenig Selbstbestimmung nicht der Fall ($p = .209$), da der Unterschied zwischen diesen Gruppen

unbedeutend ist, was auch aus der Überlappung in den Fehlerbalken für diese Zustände ersichtlich wird.

Die von Menschen gesteuerten Drohnen werden als Drohnen mit weniger Selbstbestimmung wahrgenommen, als die Zustände mit autonomen Waffensystemen. Die Agentenmanipulation funktioniert, obwohl der Unterschied weniger deutlich ist als im Szenario mit den autonomen Waffensystemen. Die Fehlerbalken der Szenarios auf Agency-Stufe zeigen an, dass die von Menschen gesteuerten Drohnen nicht eindeutig sind. Dies wird von dem t-Test der unabhängigen Stichprobe bestätigt ($p = .125$).

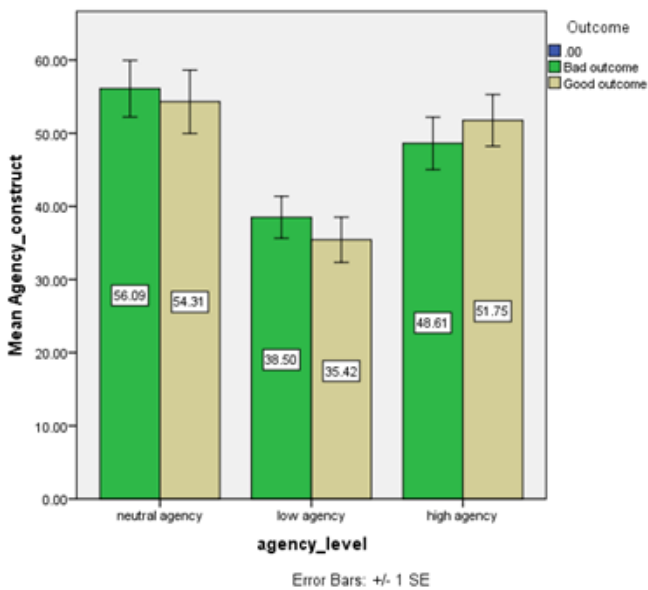


Abbildung 10 Mittelwert des Agency-Konstrukts pro Zustand pro Resultat.

Die zweite Grafik (Abbildung 10) weist den Mittelwert der Auffassung und moralischen Wertvorstellung in Bezug auf AW

über den Agency-Stufen für gute und schlechte Resultate auf. Diese Art eines vom Waffensystem verursachtes Ergebnisses, ob Kollateralschaden oder nicht, scheint nicht zu großen Unterschieden in der Auffassung und moralischen Wertvorstellungen von Kampfmitteln zu führen, da der Mittelwert sich nur um 2 bis 3 Punkte (auf einer 100-Punkte-Skala) unterscheidet. Die Fehlerbalken überlappen außerdem auf allen Agency-Stufen. Dies wird von den t-Tests der unabhängigen Stichprobe bestätigt. Auf der neutralen Agency-Stufe beträgt die Signifikanz zwischen dem schlechten Resultat und dem guten Resultat $p = .816$, auf der niedrigen Agency-Stufe ist der Zustand $p = .641$ und auf der hohen Agency-Stufe $p = .575$, die alle den $p < .05$ Schwellenwert überschreiten.

4.2.5 *Analyse der abhängigen Variablen*

Die Korrelationsanalyse weist auf, dass nur die abhängigen Variablen *Schuld*, *Kommandantenschuld*, *Kommandantenvertrauen*, *Kommandantenschaden* und *Unterstützung* zu einer Minimalebene von $p < .05$ von Bedeutung waren. Auf der nächsten Stufe vergrößerten wir diese Variablen und sahen sie uns näher an, aber da sie nicht bedeutend waren, führen wir keine Analyse der restlichen abhängigen Variablen durch. Bei genauerer Betrachtung prüften wir, ob es Unterschiede zwischen den verschiedenen Agency-Stufen (niedrig, neutral und hoch) und den autonomen Waffensystemen und von Menschen gesteuerten Drohnen gab.

Unterstützungsvariable

Die interessanteste Beobachtung ist der große Unterschied in der Unterstützung zwischen guten und schlechten Resultaten (Abbildung 11), die für alle drei Zustände bedeutend sind ($p < .05$). Dieser gravierende Unterschied im Mittelwert der Unterstützung für die drei Zustände zeigt auf, dass die

Manipulation unserer Resultate funktioniert. Eine weitere interessante Beobachtung ist, dass diese Unterstützung abgeschwächt wird, so wie die Agency-Stufen ansteigen, aber die

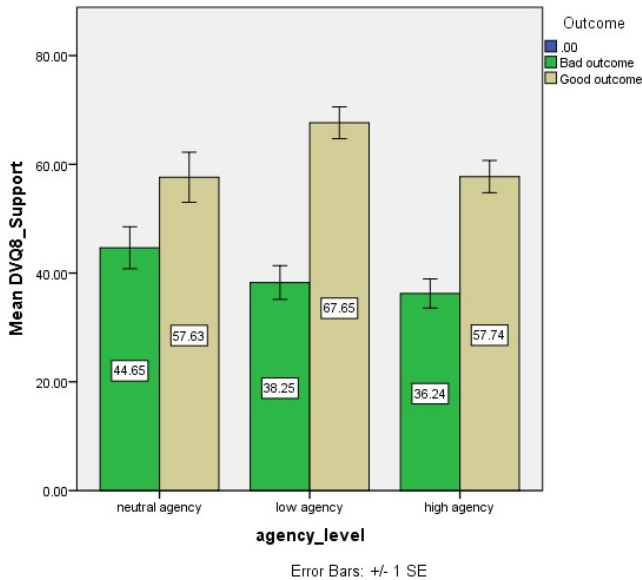


Abbildung 11 Mittelwert der Unterstützungsvariable pro Zustand pro Resultat

Differenzen zum Mittelwert sind gering (< 10 Punkte auf einer Skala von 100 Punkten) und die | -förmigen Balken oben, mit dem Standardfehler überlappen. Das zeigt an, dass diese Gruppen sich nicht bedeutend voneinander unterscheiden, entweder wenn mit dem guten oder mit den schlechten Resultatszenarien verglichen. Das zweite Diagramm zeigt an, dass die Unterstützung unter Zuständen mit von Menschen gesteuerten Waffen höher ist, als für solche mit autonomen Waffensystemen. auch hier sind die Differenzen in der Gruppe

nur gering und von Signifikanz, da die Fehlerbalken überlappen.

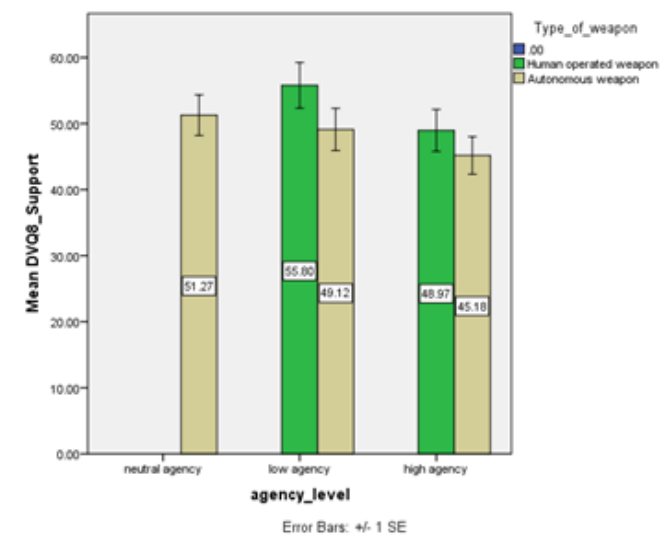


Abbildung 12 Mittelwert der Unterstützungsvariable pro Waffentyp

Vertrauensvariable

Auf Abbildung 12 ist zu sehen, dass Leute weniger Vertrauen darin setzen, dass autonome Waffensysteme zukünftig die korrekten Maßnahmen nach einem schlechten Resultat treffen, als ein menschlicher Operateur, der Maßnahmen nach einem ebenso schlechten Resultat trifft, und alle Ergebnisse sind für ihn von Bedeutung. Dieser Effekt wird in Zuständen mit geringer Selbstbestimmung verstärkt. Da sehen wir fast kaum einen Unterschied im Vertrauen zum Menschen nach einem guten oder schlechten Resultat und einen großen Unterschied im Vertrauen in ein AW (Abbildung 13). Allerdings zeigte die eindimensionale Analyse, dass die Interaktion nicht signifikant

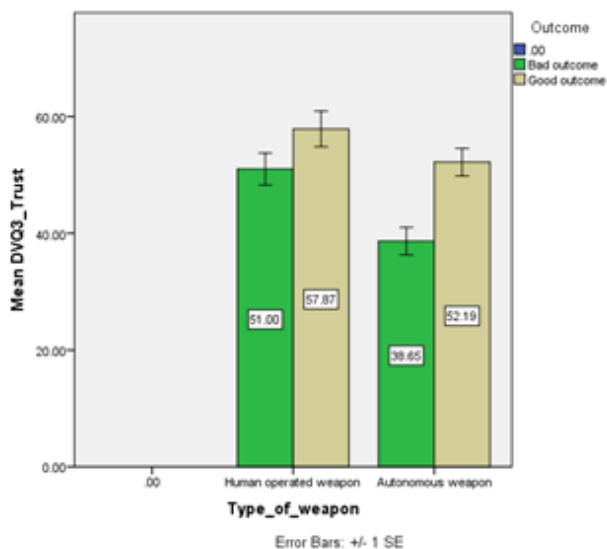


Abbildung 13 Mittelwert der Vertrauensvariable pro Waffentyp pro Resultat

ist: $Type_of_weapon * outcome$ $p = .120$, das bedeutet also, dass wir an dieser Stelle in der Untersuchung keine Schlussfolgerungen ziehen können, doch es ist ein interessantes Ergebnis, das sich lohnt, gründlicher untersucht zu werden.

Schuld- und Schadenvariablen

Um die Kette der Verantwortung zu untersuchen, werden die Variablen *Schuld*, *Kommandantenschuld*, *Schaden* und *Kommandantenschaden* auf dem gleichen Diagramm (Abbildung 15) dargestellt. Alle Ergebnisse der Schuldvariablen (Abbildung 15), mit Ausnahme des neutralen Agenten im autonomen Waffensystem-Zustand, sind von Bedeutung. Diese Ergebnisse weisen nach, dass in Zuständen mit geringer Selbstbestimmung die Schuld

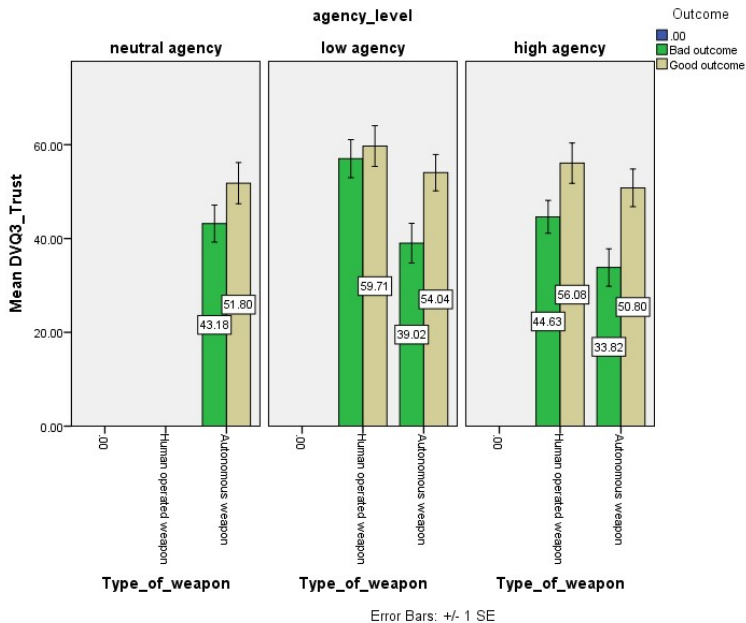


Abbildung 14 Mittelwert der Vertrauensvariable pro Zustand pro Resultat

dem Kommandanten zugeschoben und in Zuständen mit hoher Selbstbestimmung, sie der Drohne gegeben wird.

Eine bemerkenswerte Beobachtung ist, dass Zustände mit geringer Selbstbestimmung erstaunlich viel kausale Verantwortung für den Schaden übernehmen. Dieses Muster gilt für von Menschen gesteuerten Drohnen wie auch für autonome Waffensysteme (Abbildung 16). Im Gegensatz zur Schuldvariable, betrachtet man die Drohne als gleichviel Schaden anrichtend, ob in Zuständen mit hoher oder geringer Selbstbestimmung, aber in Zuständen mit geringer Selbstbestimmung wird der Schaden fast gleich dem Kommandanten zugeschoben

ben. Das bedeutet, dass beide gleichviel Schaden anrichten, aber in Zuständen mit hoher Selbstbestimmung wird er auf den

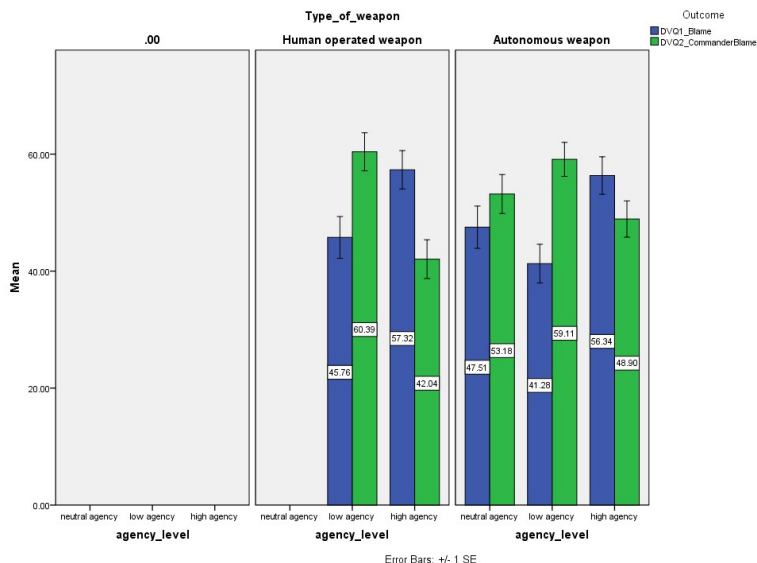


Abbildung 15 Mittelwert von Schuld- und Kommandantenschuldvariablen pro Waffentyp

Kommandanten übertragen, und die Drohne scheint den Schaden verursacht zu haben. Allerdings ist der Unterschied zwischen den Variablen Schaden und Kommandantenschaden in Zuständen mit niedriger Selbstbestimmung in Szenarien von sowohl von Menschen gesteuerten als auch autonomen Waffensystemen nicht bedeutend, deshalb müssen wir mit unserer Schlussfolgerung zum beobachteten Effekt vorsichtig sein.

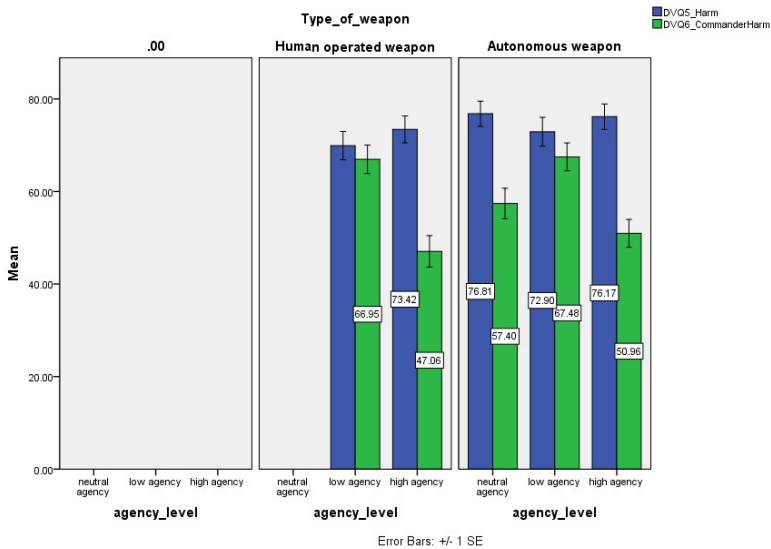


Abbildung 16 Mittelwert von Schaden und Kommandantenschaden pro Zustand und Waffentyp

4.2.6 Schlussfolgerung der Vorstudie 1

Das Ergebnis der Vorstudie 1 führt zu den folgenden Schlussfolgerungen:

- Aufgrund der Ergebnisse der PCA und Reliabilitätsanalyse der Agency-Items, entschieden wir die SQ3_Moral_rules aus dem Agency-Konstrukt zu nehmen, was die Zuverlässigkeit des Konstrukts auf $\alpha = 0.933$ erhöht. In der nächsten Untersuchung werden wir das Agency-Konstrukt benutzen, das die folgenden vier Items enthält: *Denken, Zielsetzung, freies Handeln und Zielerreichung*.
- Die Ergebnisse der t-Tests zeigen an, dass ein Unterschied in dem Mittelwert der Auffassung moralischen Wertvorstellung der Agenten zwischen den Zuständen mit autonomen und menschlich gesteuerten Waffensystemen

herrscht. Es herrschen auch große Unterschiede zwischen den Szenarien mit hoher und geringer/neutraler Selbstbestimmung. Allerdings wurden keine großen Unterschiede zwischen Agenten mit geringer und neutraler Selbstbestimmung beobachtet. Deshalb entschieden wir, den Zustand des neutralen Agenten in der nächsten Untersuchung fallen zu lassen und nur die Zustände von Agenten mit geringer und hoher Selbstbestimmung zu testen.

- Die eindeutigste Beobachtung in den Unterstützungsvariablen ist, dass es einen großen Unterschied in der Unterstützung zwischen guten und schlechten Resultaten gibt (Abbildung 11). Allerdings ist dies laut der Fachliteratur zu erwarten, da James Igoe Walsch (2015) und Kreps (2014) erkannten, dass die Zivilbevölkerung die Drohnenangriffe eher unterstützen, wenn es keinen Kollateralschaden gibt, als solche, die Kollateralschaden verursachen, aber es beweist, dass unsere Manipulation des Resultats funktioniert, und wir werden dies für die nächste Studie benutzen.
- In der Analyse der Vertrauensvariable (Abbildung 13) ist zu sehen, dass Leute weniger Vertrauen darin setzen, dass autonome Waffensysteme zukünftig die korrekten Maßnahmen nach einem schlechten Resultat treffen, als ein menschlicher Operateur, der Maßnahmen nach einem ebenso schlechten Resultat trifft. Dies könnte auf Algorithmus-Aversion hinweisen, für die Algorithmen mehr als menschliche Fehler bestraft werden (Dietvorst, 2016) und wir benutzen diese Ergebnisse für die nächste Vorstudie, um den Effekt der Algo-Aversion genauer zu untersuchen.
- Wir fanden, dass in Zuständen mit hoher Selbstbestimmung eine erstaunlich hohe kausale Verantwortlichkeit für den Schaden zu sehen ist und dass dieses Muster für von Menschen gesteuerten Drohnen und autonomen Waffen-

systemen gleich stichhaltig ist (Abbildung 16). Wir fanden auch, dass die Drohne in Zuständen mit hoher wie auch geringer Selbstbestimmung gleichviel Schaden anzurichten scheint, aber in Zuständen mit geringer Selbstbestimmung wird der Schaden zum gleichen Anteil dem Kommandanten zugesprochen. Obwohl diese Erkenntnisse nicht von Bedeutung sind, sind sie interessant und erfordern weitere Untersuchung, deshalb haben wir diese Fragen in Vorstudie 2 beibehalten.

4.3. *Vorstudie 2*

In der Vorstudie 2 versuchten wir, unsere Kenntnisse des Algo-Aversion-Effekts zu verbessern, deshalb testeten wir Szenarien, die beschrieben, wie das autonome Waffensystem aus seinen Fehlern lernen kann (3 + 10), es mit einem großen Datensatz programmiert wurde, und dass es begreift, (4 + 11) dass sein Einsatz gelegentlich unberechenbar sein kann (5 + 12) sowie ein Szenario, dass all diese Aspekte einschließt (6 + 13). Wir testeten Zustände mit guten und schlechten Resultaten und fügten ein Szenario für autonome Waffensysteme mit geringer Selbstbestimmung (1 + 8) hinzu, um zu prüfen, ob unsere Agentenmanipulation funktioniert, sowie ein von Menschen gesteuertes Szenario (7 + 14), in der wir uns über den menschlichen Operateur kundig machten, statt der Drohne, im Unterschied zu Vorstudie 1.

Wir testeten insgesamt 14 Szenarien und zielten darauf, 50 Umfrageteilnehmer pro Szenario zu erhalten, was zu insgesamt 700 Antworten führen würde. Die Anzahl der Befragten pro Szenario betrug zwischen 47 und 54. Nach der Datenvorverarbeitung und dem Löschen von Befragten, die den Aufmerksamkeitstest nicht durchgeführt hatten, verfügten wir noch über 709 komplett ausgefüllten Fragebögen.

		<i>Aspekte</i>						
		<i>Autonome Drohne</i>						<i>MG</i>
		<i>Ebene 1 - Geringe Selbstbestimmung</i>	<i>Ebene 2 (hohe Selbstbestimmung + nichts zusätzlich)</i>	<i>Ebene 3 (hohe Selbstbestimmung + Lernfähigkeit)</i>	<i>Ebene 4 (hohe Selbstbestimmung + Verständigungsverfahren)</i>	<i>Ebene 5 (hohe Selbstbestimmung + Unberechenbarkeit)</i>	<i>Ebene 6 (hohe Selbstbestimmung + alle Aspekte)</i>	<i>Von Menschen gesteuerte Drohne</i>
<i>Resultat</i>	<i>Kollateral-schaden (schlecht)</i>	1	2	3	4	5	6	7
	<i>Kein Kollateral-schaden (gut)</i>	8	9	10	11	12	13	14

Tabelle 11 Szenarien der Vorstudie 2

4.3.1 Reliabilitätsanalyse

Für alle vier Agency-Items ist $\alpha = 0.914$; das liegt weit über dem Schwellenwert von 0,7 und kann nicht durch Löschen eines Items verbessert werden. Löschen anderer Items wird den α Wert abschwächen und die interne Konsistenz reduzieren.

Reliabilitätststatistiken					
Cronbachs Alpha basiert auf Standardisierte Items					
	Cronbachs Alpha			Anzahl Items	
	.914		.914	4	
Item-Gesamtstatistiken					
	Skalen-Mittelwert, wenn Item gelöscht	Skalen-Varianz, wenn Item gelöscht	Korrigierte Item-Gesamtwert Korrelation	Ausgeglichene Multiple Korrelation	Cronbachs Alpha, wenn Item gelöscht
Denken	208.4827	6427.511	.801	.647	.890
Zielsetzung	204.9885	6499.803	.834	.695	.877
Freies Handeln	203.4020	6782.379	.804	.646	.888
Zielerreichung	199.4975	7114.019	.781	.614	.897

Tabelle 12 Ergebnisse der Reliabilitätsanalyse Vorstudie 2

4.3.2 Hauptkomponentenanalyse (PCA)

Die PCA zeigt, dass *Denken*, *Zielsetzung*, *freies Handeln* und *Zielerreichung* als ein Konstrukt angesehen werden kann, das die Varianz von 79,60% in den ursprünglichen Variablen ausmacht (Tabelle 13). Dies kann man auch auf dem Bildschirmdiagramm ansehen, das die Eigenwerte der Komponenten (Abbildung 17) anzeigt. Die PCA zeigt an, dass die vier Items: Denken, Zielsetzung, freies Handeln und Zielerreichung als die zugrundeliegenden Komponenten des Agency-Konstrukts Tabelle 13 gesehen werden kann.

Komponentenmatrix ^a		
	Komponente	
	1	2
Zielsetzung	.910	
Freies Handeln	.892	
Denkweise	.889	
Zielerreichung	.877	.446

Extraktionsmethode: Hauptkomponentenanalyse.
a. 2 Komponenten wurden extrahiert.

Tabelle 13 Ergebnisse-PCA Vorstudie 2

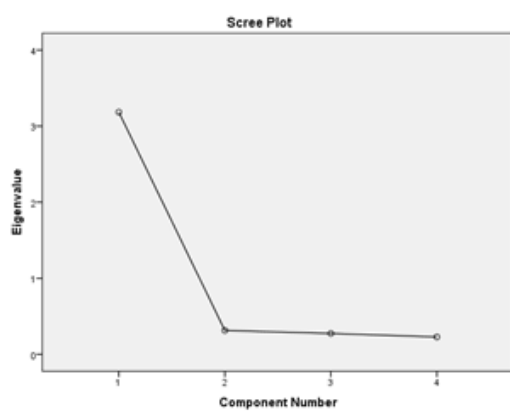


Abbildung 17 Scree-Test PCA Vorstudie 2

4.3.3 Korrelationsanalyse

Die Korrelationsanalyse des Agency-Konstrukts mit den abhängigen Variablen ist nicht umfangreich, aber für alle Konstrukte von Bedeutung ($p < ,05$) (Tabelle 14). Die gleichen Effekte auf die Richtung für die Schaden- und Schuldvariable wie in Vorstudie 1 können auch in dieser Vorstudie beobachtet werden. Es besteht ein negativer Zusammenhang zwischen

		Correlations												
		Agency_Construct_Total	DVO1_Blame	DVO2_CommanderBlame	DVO3_Trust	DVO4_CommanderTrust	DVO5_Harm	DVO6_CommanderHarm	DVO7_Dignity	DVO8_Confidence	DVO9_Expectations	DVO10_Support	DVO11_Fair	DVO12_Uneasy
Agency_Construct_Total	Pearson Correlation	1	.271**	-.267**	.205**	-.217**	.124**	-.198**	.102**	.209**	.146**	.164**	.142**	-.076*
	Sig. (2-tailed)		.000	.000	.000	.000	.001	.000	.007	.000	.000	.000	.000	.043
	N	708	708	708	708	708	708	708	708	708	708	708	708	708
DVO1_Blame	Pearson Correlation	.271**	1	-.069	-.028	-.046	.258**	-.023	-.056	-.048	-.137**	-.118**	-.092*	.180**
	Sig. (2-tailed)		.000	.065	.464	.226	.000	.534	.138	.198	.000	.002	.014	.000
	N	708	708	708	708	708	708	708	708	708	708	708	708	708
DVO2_CommanderBlame	Pearson Correlation	-.267**	-.069	1	-.287**	-.278**	.152**	.691**	-.206**	-.300**	-.205**	-.250**	-.275**	.312**
	Sig. (2-tailed)		.000	.065	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000
	N	708	708	708	708	708	708	708	708	708	708	708	708	708
DVO3_Trust	Pearson Correlation	.205**	-.028	-.287**	1	.653**	-.281**	-.224**	.595**	.901**	.533**	.726**	.722**	-.616**
	Sig. (2-tailed)		.000	.464	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000
	N	708	708	708	708	708	708	708	708	708	708	708	708	708
DVO4_CommanderTrust	Pearson Correlation	.217**	-.046	-.278**	.653**	1	-.185**	-.228**	.432**	.635**	.425**	.631**	.552**	-.458**
	Sig. (2-tailed)		.000	.226	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000
	N	708	708	708	708	708	708	708	708	708	708	708	708	708
DVO5_Harm	Pearson Correlation	.124**	.258**	.152**	-.281**	-.185**	1	.311**	-.374**	-.301**	-.217**	-.253**	-.324**	.395**
	Sig. (2-tailed)		.001	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000
	N	708	708	708	708	708	708	708	708	708	708	708	708	708
DVO6_CommanderHarm	Pearson Correlation	-.198**	-.023	.691**	-.224**	-.228**	.311**	1	-.183**	-.226**	-.203**	-.235**	-.248**	.289**
	Sig. (2-tailed)		.000	.534	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000
	N	708	708	708	708	708	708	708	708	708	708	708	708	708
DVO7_Dignity	Pearson Correlation	.102**	-.056	-.206**	.595**	.432**	-.374**	-.183**	1	.588**	.415**	.533**	.598**	-.509**
	Sig. (2-tailed)		.007	.138	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000
	N	708	708	708	708	708	708	708	708	708	708	708	708	708
DVO8_Confidence	Pearson Correlation	.209**	-.048	-.300**	.901**	.635**	-.301**	-.226**	.588**	1	.553**	.723**	.725**	-.649**
	Sig. (2-tailed)		.000	.198	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000
	N	708	708	708	708	708	708	708	708	708	708	708	708	708
DVO9_Expectations	Pearson Correlation	.146**	-.137**	-.205**	.533**	.425**	-.217**	-.203**	.415**	.553**	1	.511**	.621**	-.509**
	Sig. (2-tailed)		.000	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000
	N	708	708	708	708	708	708	708	708	708	708	708	708	708
DVO10_Support	Pearson Correlation	.164**	-.118**	-.250**	.726**	.631**	-.253**	-.235**	.533**	.723**	.511**	1	.685**	-.651**
	Sig. (2-tailed)		.000	.002	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000
	N	708	708	708	708	708	708	708	708	708	708	708	708	708
DVO11_Fair	Pearson Correlation	.142**	-.092*	-.275**	.722**	.552**	-.324**	-.248**	.598**	.725**	.621**	.685**	1	-.656**
	Sig. (2-tailed)		.000	.014	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000
	N	708	708	708	708	708	708	708	708	708	708	708	708	708
DVO12_Uneasy	Pearson Correlation	-.076*	.180**	.312**	-.616**	-.458**	.395**	.289**	-.509**	-.649**	-.509**	-.651**	-.656**	1
	Sig. (2-tailed)		.043	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000
	N	708	708	708	708	708	708	708	708	708	708	708	708	708

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

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

Tabelle 14 Korrelationsmatrix des Agency-Konstrukts und der abhängigen Variablen Vorstudie 2

dem Agency-Konstrukt und der Stufe der Unruhe, d.h. dass bei Erhöhung der Agentenwahrnehmung, die Leute andeuten, dass sie eine größere *Besorgnis* empfinden, aber dieser Effekt ist sehr gering ($r = -,076$). Die Korrelationsanalyse ist nicht detailliert genug, um in diese Ergebnisse zu zoomen, deshalb wird dies bei der Analyse der abhängigen Variablen getan

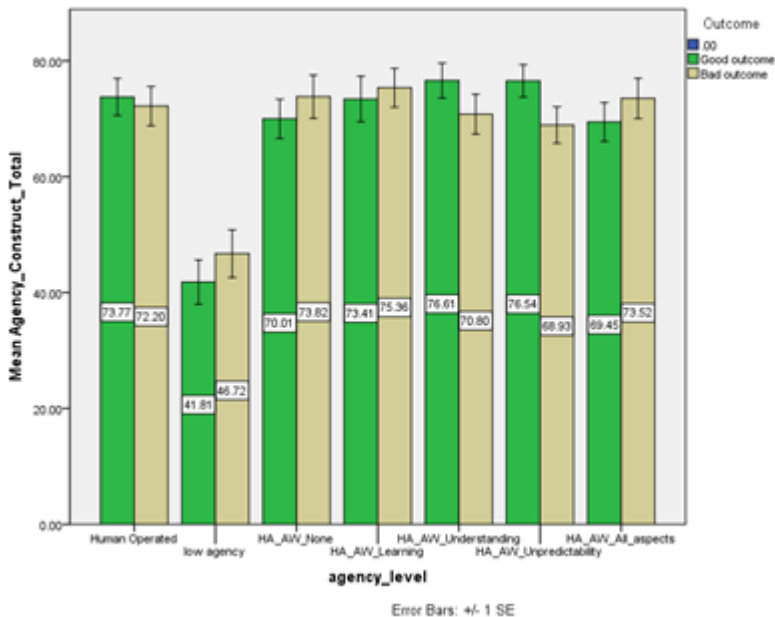


Abbildung 18 Mittelwert des Wert-Agency-Konstrukts pro Zustand pro Resultat

4.3.4 Manipulationsprüfung der Agenten

Das Diagramm in Abbildung 18 zeigt an, dass ein großer Unterschied zwischen den Agency-Szenarien mit hoher Selbstbestimmung und denen mit geringer Selbstbestimmung herrscht, und dieser Unterschied ist von Bedeutung. Der Unterschied zwischen geringen und hohen Agentenzuständen bestätigt,

dass unsere Agentenmanipulation funktioniert. Agenten des menschlichen Operators und Drohnen mit hoher Selbstbestimmung befinden sich auf der gleichen Ebene (in dieser Vorstudie fragten wir speziell den menschlichen Operator, im Vergleich zur Drohne in der vorherigen Vorstudie). Es bestehen einige Unterschiede zwischen den Zuständen mit gutem und schlechtem Resultat und obwohl diese Effekte für die Zustände am stärksten sind, in denen wir in die Algo-Aversion gezoomt sind, ist der Unterschied nicht besonders groß (ein Unterschied im Mittelwert von 8 Punkten auf einer 100-Punkt-Skala). Außerdem überlappten die Fehlerbalken zwischen guten und schlechten Resultat-Zuständen für alle Szenarien, eine Indikation, dass es sich nicht um Einzelgruppen handelt. Der einzige eindeutige und signifikante Unterschied in Gruppen ($p < ,05$) liegt zwischen den Zuständen der autonomen Waffensysteme mit geringer Selbstbestimmung und den anderen Szenarien. Dies wurde auch von der Signifikanzebene des t-Tests der unabhängigen Stichprobe bestätigt, den wir wahlweise nicht in das nachstehende Diagramm aufnahmen, um es anschaulicher zu machen.

4.3.5 Analyse der abhängigen Variablen

Die Korrelation zwischen dem Agency-Konstrukt und den abhängigen Variablen haben alle Signifikanz ($p < ,05$), deshalb sind die eindeutigsten Nachweise für jede der abhängigen Variablen nachstehend beschrieben.

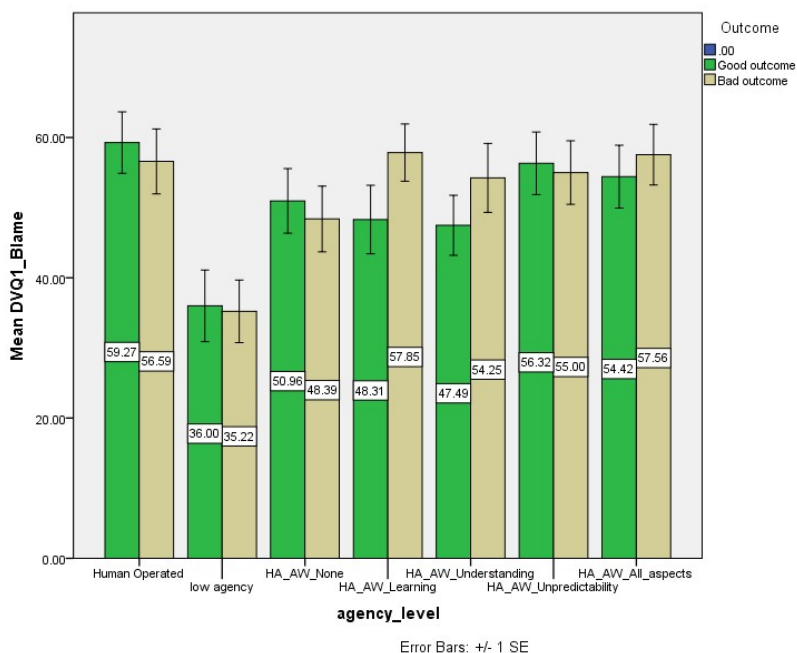


Abbildung 19 Mittelwert der Schuldvariable pro Zustand pro Resultat

Schuld

Die wichtigste Erkenntnis für die *Schuld*-Variable ist, dass Leute Zuständen mit geringer Selbstbestimmung weniger *Schuld* beimessen als Zuständen mit autonomen Waffensystemen mit hoher Selbstbestimmung und menschlichen Operateuren. Für diese Zustände ist der Unterschied von Bedeutung ($p < ,05$). Einige Effekte der Algo-Aversion kann beim Lernen und der Einsichtnahme von Szenarien erkannt werden, wo Leute weniger *Schuld* beimessen, wenn die autonomen Waffensystemen keinen Fehler machten, und mehr, wenn sie einen Fehler mit einem schlechten Resultat macht. Wie dem auch sei, diese Zustände weisen eine Überlappung der Fehlerbalken auf und sind nicht signifikant.

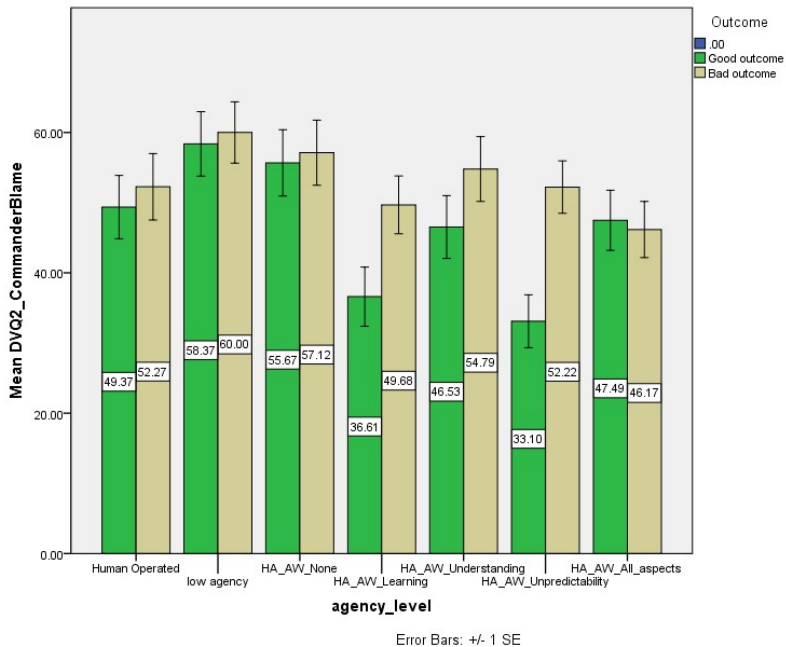


Abbildung 20 Mittelwert der Kommandantenschuld-Variable pro Zustand pro Resultat

Kommandantenschuld

Die bedeutendste Beobachtung bei der Kommandantenschuld-Variable ist, dass dem Kommandanten entschieden weniger Schuld gegeben wird, wenn die autonome Waffe einen Programmierfehler macht und in einem unberechenbaren Zustand ist, dies ist von Signifikanz ($p < .05$). Es kann der Fall sein, dass die Schuld der Drohne zugeschoben wird, da Leute erwarten, dass sie aus ihren Fehlern lernen kann oder sie sich unberechenbar verhält und der Kommandant dafür keine Schuld hat. Bei Agentenzuständen von geringer Selbstbestimmung wird dem Kommandanten etwas mehr die Schuld gegeben als in den anderen Zuständen.

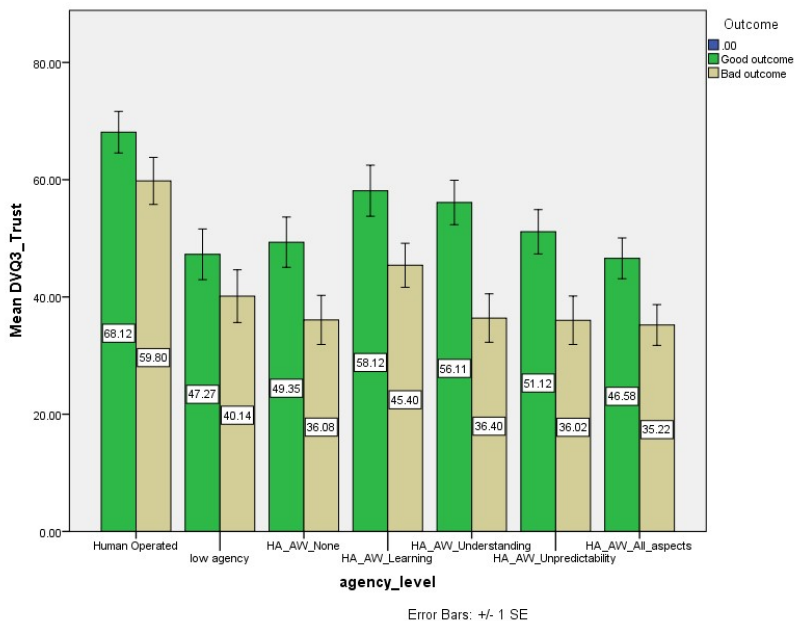


Abbildung 21 Mittelwert der Vertrauensvariable pro Zustand pro Resultat

Vertrauen

Während aller Zustände ist zu beobachten, dass Leute nach einem guten Resultat zukünftig mehr Vertrauen in autonome Waffensysteme oder menschliche Operateure setzen, die korrekten Maßnahmen zu treffen, als nach einem schlechten Resultat. Mit Ausnahme des Agentenzustandes mit geringer Selbstbestimmung, ist das Vertrauen in ein gutes oder schlechtes Resultat von Bedeutung ($p < ,05$). Leute vertrauen dem menschlichen Operateur mehr als der autonomen Waffe, vor allem bei geringer Selbstbestimmung, und der Unterschied zwischen diesen Gruppen ist signifikant ($p < ,05$).

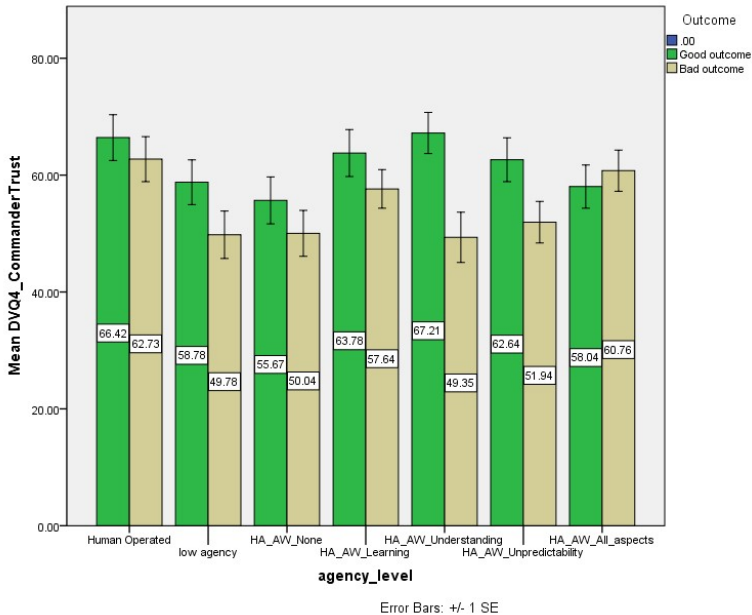


Abbildung 22 Mittelwert der Kommandantenvertrauen-Variable pro Zustand pro Resultat

Kommandantenvertrauen

Im Großen und Ganzen vertrauen die Leute dem Kommandanten eher nach einem guten Resultat, dass er die korrekten Maßnahmen treffen wird, als nach einem schlechten Resultat. Das Vertrauen in den Kommandanten ist am niedrigsten bei geringer oder hoher Selbstbestimmung des Agentenzustands und ohne zusätzliche Informationen, und der Unterschied zwischen diesen Gruppen ist signifikant ($p < .05$). Zu wissen, dass die autonome Waffe kognitive Fähigkeiten hat, resultiert in mehr Vertrauen bei schlechten Resultaten, und das ist signifikant ($p < .05$).

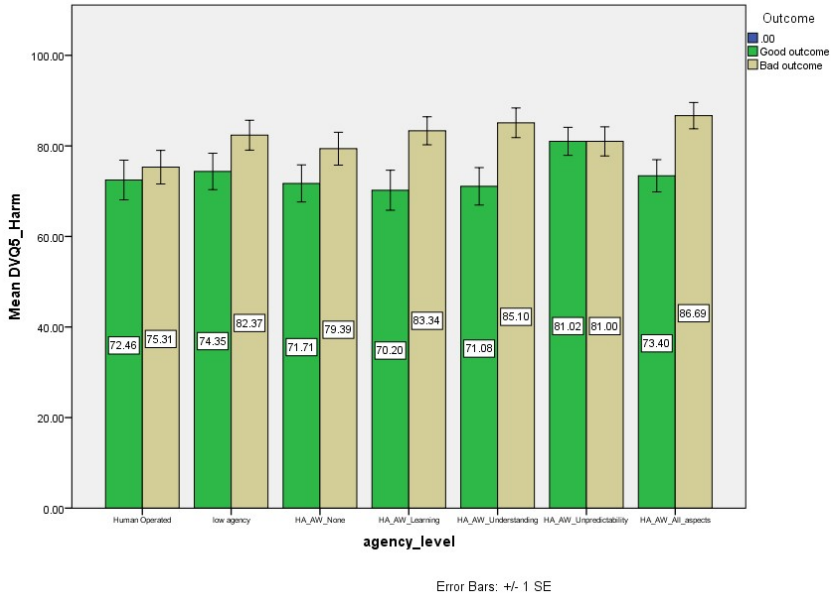


Abbildung 23 Mittelwert der Schadensvariable pro Zustand pro Resultat

Schaden

Bei Zuständen mit schlechten Resultaten nehmen Leute mehr Schaden wahr als bei guten, was zu erwarten ist, und der Unterschied ist unter fast allen Unterschieden signifikant ($p < ,05$). In diesem Diagramm ist vor allem auffällig, dass, wenn autonome Waffensysteme sich unberechenbar verhalten, das Resultat nicht von Bedeutung zu sein scheint, und die Wahrnehmung des Schadens ist die gleiche, aber der Unterschied zwischen beiden Resultatgruppen ist nicht signifikant.

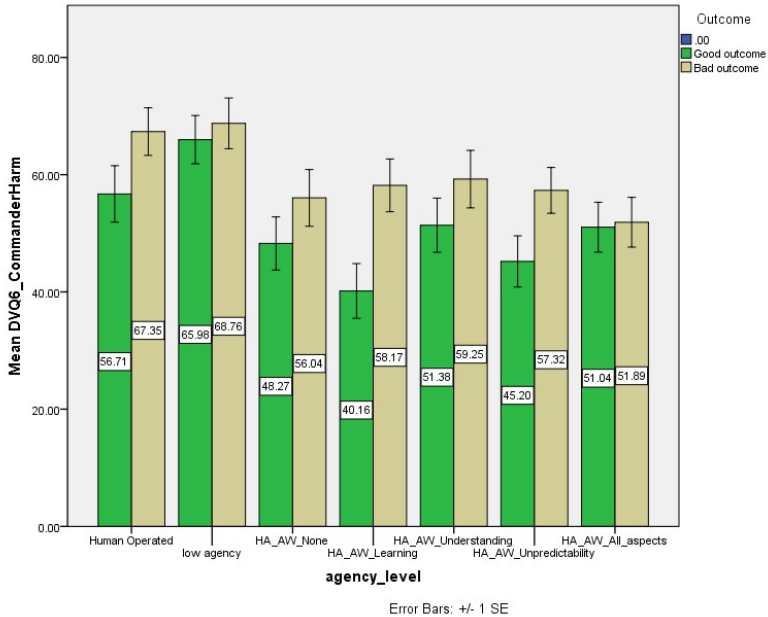


Abbildung 24 Mittelwert der Kommandantenschaden-Variable pro Zustand pro Resultat.

Kommandantenschaden

Bei Zuständen mit schlechten Resultaten schreiben die Leute dem Kommandanten mehr Schaden zu, als bei guten, was zu erwarten ist. Am Eindeutigsten ist der große Unterschied im Programmierzustand zwischen den guten und schlechten Resultaten, wenn dem Kommandanten bei guten Resultaten viel weniger Schaden zugemessen wird, was signifikant ist ($p < ,05$). Die Unterschiede zwischen Zuständen mit guten und schlechten Resultaten in Szenarien zwischen menschlichem Operateur und Agenten mit hoher Unberechenbarkeit sind ebenfalls signifikant ($p < ,05$).

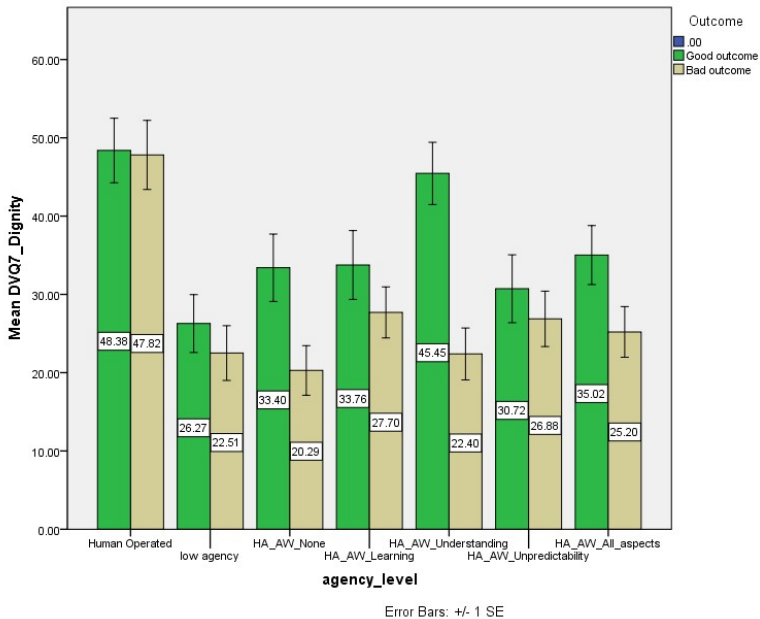


Abbildung 25 Mittelwert der Menschenwürde-Variable pro Zustand pro Resultat

Menschenwürde

Der menschliche Operator handelt mit mehr Respekt für die Menschenwürde als die autonome Waffe, mit Ausnahme des Zustandes, wenn die autonome Waffe mit einem großen Datensatz, der speziell für ihren Einsatz eingerichtet wurde, programmiert wird, und dieser Unterschied ist signifikant ($p < ,05$). Ebenso scheint es von Bedeutung zu sein, wenn autonome Waffensysteme einen Fehler machen, da wir signifikante ($p < ,05$) Unterschiede zwischen guten und schlechten Resultatszenarien für autonome Waffensysteme für die geringe Selbstbestimmung, hohe Selbstbestimmung keine Aspekte, und hohe Selbstbestimmung keine Aspekte beobachten.

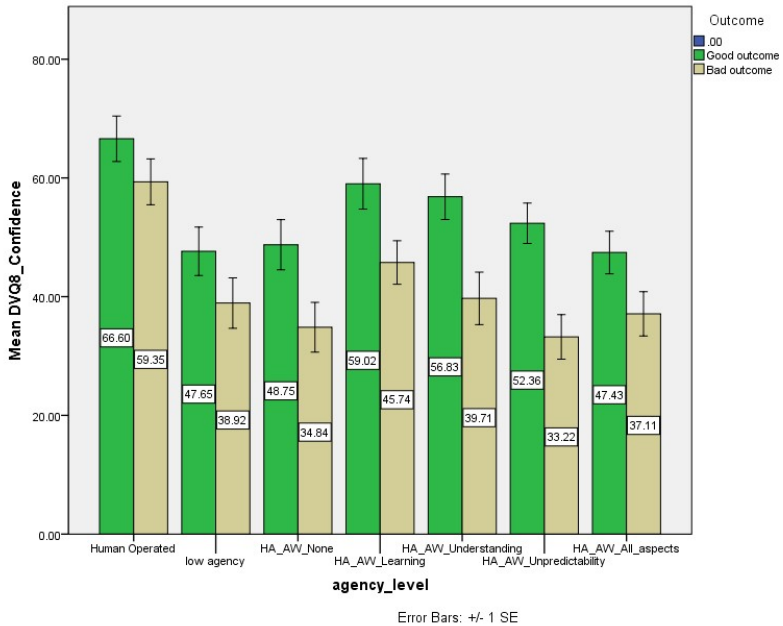


Abbildung 26 Mittelwert der ZuversichtsvARIABLE pro Zustand

Zuversicht

Im Großen und Ganzen sind die Menschen zuversichtlicher nach einem guten Resultat, dass die autonomen Waffensysteme zukünftig die korrekten Maßnahmen treffen, als nach einem schlechten Resultat; alle Differenzen sind signifikant ($p < ,05$). Sie setzen auch etwas mehr Vertrauen in den menschlichen Operateur als in autonome Waffensysteme. Die Differenzen zwischen diesen Szenarien sind signifikant ($p < ,05$). Wenn die autonomen Waffensysteme Lernfähigkeit haben, ist die Zuversicht der Leute etwas größer im Vergleich zu den anderen Zuständen von autonomen Waffen.

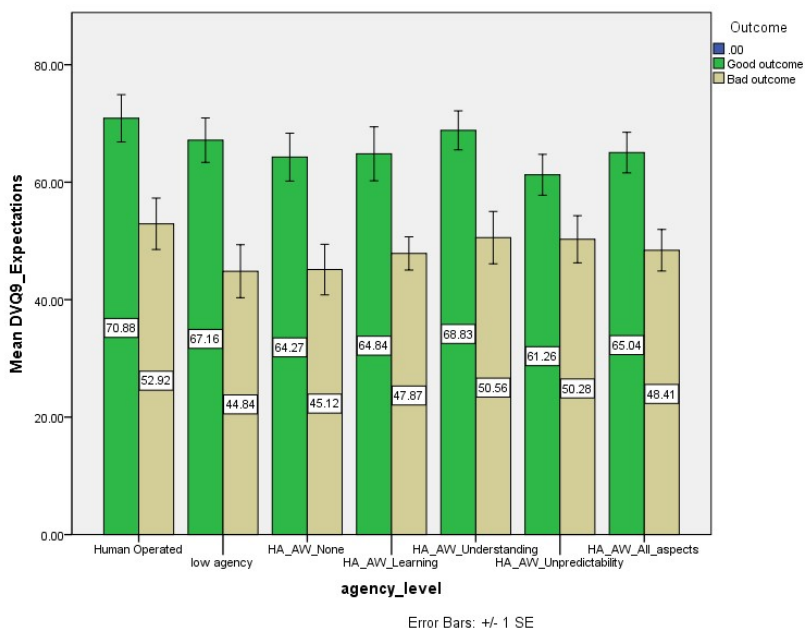


Abbildung 27 Mittelwert der Werterwartungsvariable pro Zustand pro Resultat

Erwartungen

Nach den Erwartungen der Menschen leisten sowohl der menschliche Operateur als auch das autonome Waffensystem mehr in Szenarien mit guten Resultaten, als solchen mit schlechten Resultaten. Alle Differenzen zwischen guten und schlechten Resultaten sind signifikant ($p < ,05$). Die Erwartungen für Agenten mit geringen Selbstbestimmungszuständen sind am niedrigsten, wenn ein schlechtes Resultat erreicht wurde, aber es besteht keine große Differenz zu den anderen Zuständen.

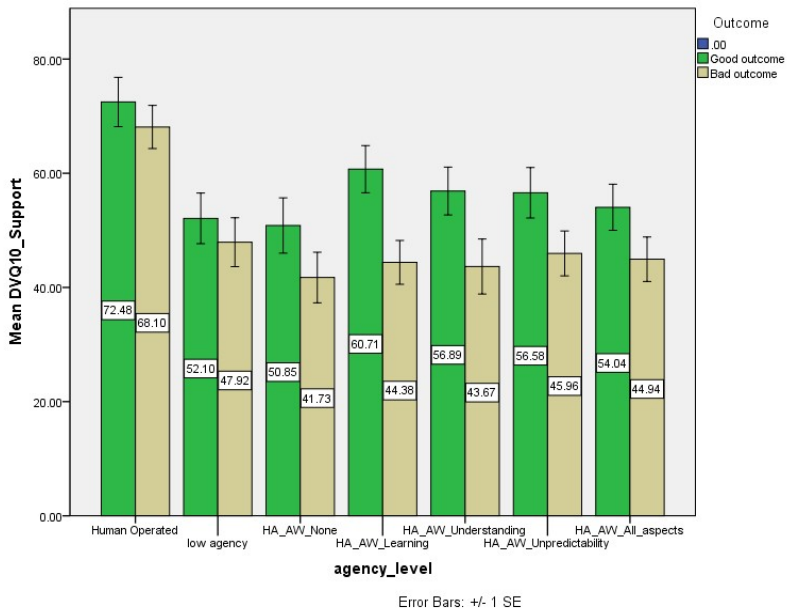


Abbildung 28 Mittelwert der Unterstützungsvariable pro Zustand pro Resultat

Unterstützung

Die Menschen unterstützen von Menschen gesteuerte Drohnen eher als autonome Waffensysteme. Die Unterstützung ist am geringsten, aber nicht signifikant, für die Agenten mit geringer Selbstbestimmung im Vergleich zu Szenarien, in denen Informationen über Autonome Waffensysteme erstellt werden. Ebenso ist die Unterstützung für Szenarien mit einem guten Resultat höher als für solche mit schlechten Resultaten und mit Ausnahme des menschlichen Operators und autonomer Waffensysteme mit geringer Selbstbestimmung, sind diese Differenzen signifikant ($p < ,05$).

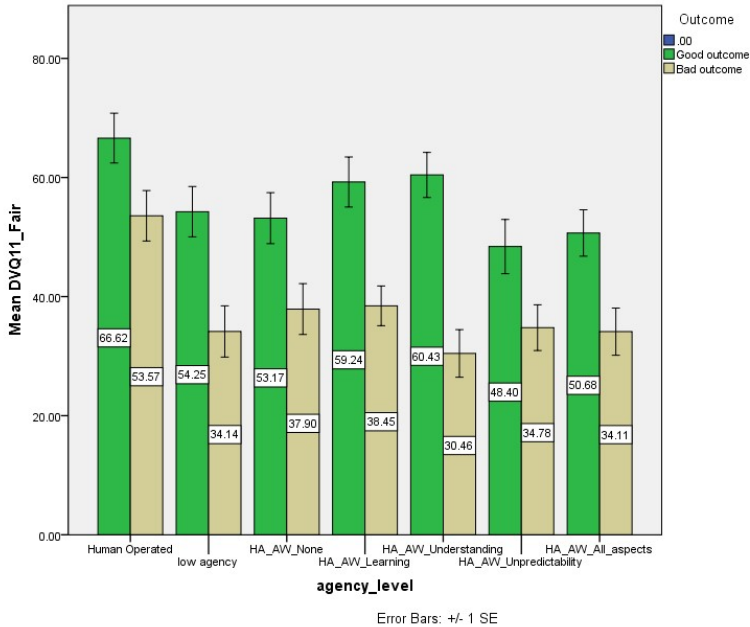


Abbildung 29 Mittelwert der Fairnessvariable pro Zustand pro Resultat

Fairness

Im Großen und Ganzen werden die Maßnahmen des menschlichen Operators als fairer angesehen als die der autonomen Waffensysteme. Die Wahrnehmung von Fairness ist entschieden höher bei Szenarien mit gutem Resultat als solchen mit schlechtem Resultat, und in allen Szenarien sind diese Differenzen signifikant ($p < .05$). Die Differenz zwischen den Agentenzuständen der autonomen Waffensysteme ist gering (<10 Punkte auf einer 100-Punkt-Skala).

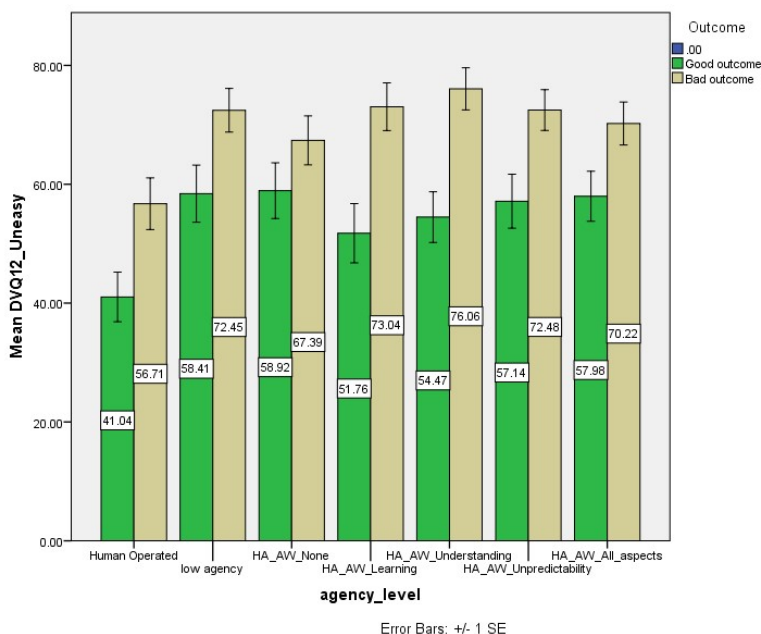


Abbildung 30 Mittelwert der Unbehagens variable pro Zustand pro Resultat

Unbehagen

Bei Szenarien mit schlechtem Resultat empfinden die Menschen mehr Unbehagen bei den Handlungen von sowohl menschlichen Operateuren als auch autonomen Waffensystemen. Es besteht keine große Differenz zwischen der Stufe der Beunruhigung für die verschiedenen Zustände der autonomen Waffensysteme. Die Beunruhigungsebene für den menschlichen Operateur ist viel niedriger als die für autonome Waffensysteme. Mit Ausnahme der Agenten mit hoher Selbstbestimmung gibt es keine weiteren Aspekte des Zustandes der autonomen Waffensysteme. Alle Differenzen zwischen guten und schlechten Resultaten sind signifikant ($p < ,05$).

4.3.6 Schlussfolgerung der Vorstudie 2

Das Ergebnis der Vorstudie 2 führt zu den folgenden Schlussfolgerungen:

- Der Reliabilitätstest wies auf, dass die vier Items des Agency-Konstrukts zuverlässig sind, und gemäß der PCA können sie als eine Komponente angesehen werden;
- Der Korrelationskoeffizient des Agency-Konstrukts und alle abhängigen Variablen sind signifikant und haben die gleichen Richtungen wie in Vorstudie 1. Wir fanden einen negativen Zusammenhang zwischen dem Agency-Konstrukt und der Stufe des Unbehagens, d.h. dass bei Erhöhung der Agentenwahrnehmung, die Leute andeuten, dass sie eine größere *Besorgnis* empfinden, aber dieser Effekt ist immer noch sehr gering und nähert sich 0 ($r = -.076$). Um zu sehen, ob wir einen stärkeren Effekt erhalten könnten, änderten wir die Formulierung der Frage zur Unbehagenvariable in der abschließenden Untersuchung;
- Die Agentenmanipulation zeigt an, dass eine große Differenz zwischen den Agency-Szenarien mit hoher Selbstbestimmung und denen mit geringer Selbstbestimmung herrscht. Die Agenten des menschlichen Operators und der autonomen Waffensysteme mit hoher Selbstbestimmung liegen auf der gleichen Ebene. Allerdings sind die Mittel des Agency-Konstrukts für Szenarien mit hoher Selbstbestimmung, die Algo-Aversion-Effekte enthalten, fast die gleichen. Um dem Algo-Aversion-Effekt herauszunehmen, müssen wir weitere Literaturforschungen und Vorstudien vornehmen. Aufgrund der beschränkten Zeit haben wir diese Forschungsmethode nicht verfolgt und uns entschieden, uns auf die abschließende Untersuchung der Agentenwahrnehmung zu konzentrieren.

- Viele der von uns bei den abhängigen Variablen gefundenen Effekte erfüllten unsere Erwartungen, begründet auf die Forschungsergebnisse der Vorstudie 1. Wir fanden keine besonderen Effekte in den DVs zwischen den verschiedenen Agentenzuständen der autonomen Waffensysteme. Wir prüften, ob wir den Algo-Aversion-Effekt aus der Vorstudie 1 reproduzieren konnten, bei dem Schaden in Agentenzuständen mit geringer Selbstbestimmung von der Drohne auf den Kommandanten übertragen wurde, aber leider konnte dieser Effekt in der Vorstudie 2 nicht festgestellt werden. Obwohl diese Verantwortungskette noch weiter untersucht werden muss, entschieden wir uns, sie nicht in die abschließende Untersuchung aufzunehmen, da wir uns nicht sicher waren, was genau der Fall ist und weitere Literaturforschung und Vorstudien erforderlich sind.

4.4 *Abschließende Untersuchung - militärische Stichprobe*

Um die Anzahl der Szenarien für die abschließende Untersuchung festzustellen, führten wir Teststärkenberechnungen aus, um die Gesamtzahl der Probanden abzuschätzen, die wir laut der Ergebnisse der ersten Vorstudie benötigen würden, da wir diese Szenarien auch für die abschließende Untersuchung benutzen werden. Basiert auf den Teststärkenberechnungen (Effektgröße 0,4, eine Soll-Teststärke von 0,8 und einen Wahrscheinlichkeitsgrad von 0,05), zielten wir auf eine Gesamtantwortquote von 200 und bestimmten, dass wir 3 Szenarien verarbeiten konnten. Wir entschieden uns, uns auf die Auffassung von Streitkräften zu konzentrieren und die derzeitige Technologie, d.h. eine von Menschen gesteuerte Drohne, mit zukünftiger Technologie, dem autonomen Waffensystem (Tabelle 15) zu vergleichen. Wir fügten außerdem noch ein Szenario mit einem autonomen Waffensystem auf neutraler Ebene hinzu,

um die Auffassung von Streitkräften der drei Zustände zu vergleichen. In dieser Untersuchung stellten wir vor allem Fragen über das Kampfmittel der Drohne und nicht über den menschlichen Operateur, der sie steuert, um die Auffassung der Artefakte zu vergleichen. Da wir nur drei Szenarien verarbeiten konnten, konzentrierten wir uns auf den Zustand mit schlechtem Resultat, da die Vorstudien nachwiesen, dass die Effekte in diesem Zustand am eindeutigsten sind.

Wir verbreiteten die Umfrage per E-Mail, mit einem anonymen Link und hatten deshalb keine Gewährleistung in Bezug auf die Anzahl der Befragten. Nach der Datenvorverarbeitung und dem Löschen von Befragten, die den Aufmerksamkeitstest nicht durchgeführt hatten, verfügten wir noch über 239 Antworten für die komplette Stichprobe. Die Anzahl der Befragten pro Szenario betrug zwischen 64 und 96, ein seltsames Ergebnis, da die Befragten gleichmäßig über die Szenarien verteilt werden sollten. Die Stichprobe von 239 ist eine Kombination des niederländischen Militärs und der Zivilbevölkerung, die im Verteidigungsministerium angestellt sind, denn hätten wir sie ausgelassen, wären wir auf 149 Befragte sitzen geblieben, zu wenig für eine robuste Datenanalyse.

	<i>Agency Stufen</i>		
	Von Mensech gesteuert	Neutrale AW	Hohe AW
Schlechtes Resultat	1	2	3

Tabelle 15 Szenarien der abschließenden Untersuchung

FÜR DIESES DESIGN WERDEN DIE FOLGENDEN DEFINITIONEN BENUTZT:

- Von Menschen gesteuerte Drohne: Drohne, die direkt unter der Kontrolle eines Menschen steht. Wird als Kampfmittel mit sehr niedriger Selbstbestimmung aufgefasst.
- Hohes AW: Probanden wurde gesagt, dass das autonome Waffensystem Selbstbestimmungsmerkmale aufweist hat (z.B. Erwägung von Für und Wider und unabhängige Entscheidungswahl). Wird als Kampfmittel mit sehr hoher Selbstbestimmung aufgefasst.
- Neutrales AW: Es wurden keinerlei Informationen über Selbstbestimmungsmerkmale der autonomen Waffensysteme gegeben. Die Stärke dieser Methode ist, dass wir Daten von Teilnehmern sammeln können, die ihre voreingestellten Konzepte von autonomen Waffensystemen zum Verständnis der Handlung autonomer Waffensysteme benutzen und diese Auffassungen mit den von menschlichen Operateuren gesteuerten und autonomen Waffen mit hoher Selbstbestimmung vergleichen können.

4.4.1 *Reliabilitätsanalyse*

Für alle vier Agency-Items ist $\alpha = ,865$ und liegt über dem Schwellenwert von 0,7. Er kann nicht durch Löschen eines Items verbessert werden, d.h. dass das Agency-Konstrukt zuverlässig ist (Tabelle 16). Löschen anderer Items wird den α Wert abschwächen und die interne Konsistenz reduzieren.

Reliabilitätstatistiken

Cronbachs Alpha	Anzahl Items
.865	4

Item-Gesamtstatistiken

	Skala-Mittel, wenn Item gelöscht ist	Skalenvarianz, wenn Item gelöscht	Korrigierte Hem-Total Korrelation	Cronbachs Alpha, wenn Item gelöscht
Denkweise	116.43	71 75.573	.620	.863
Zielsetzung	123.33	6342.474	.756	.810
Frei handeln	122.62	6566.723	.740	.816
Zielerreichung	118.84	6459.1 46	.740	.816

Tabelle 16 Ergebnisse der Reliabilitätsanalyse Agency-Items

4.4.2 Hauptkomponentenanalyse (PCA)

Die PCA zeigt, dass Denken, Zielsetzung, freies Handeln und Zielerreichung als ein Konstrukt angesehen werden kann, das die Varianz von 71,18% in den ursprünglichen Variablen ausmacht. Dies kann man auch auf dem Scree-Test ansehen, das die Eigenwerte der Komponenten (Abbildung 31) anzeigt. Die PCA zeigt an, dass die vier Items: Denken, Zielsetzung, freies Handeln und Zielerreichung als die zugrundeliegenden Komponenten des Agency-Konstrukts (Tabelle 17) gesehen werden kann.

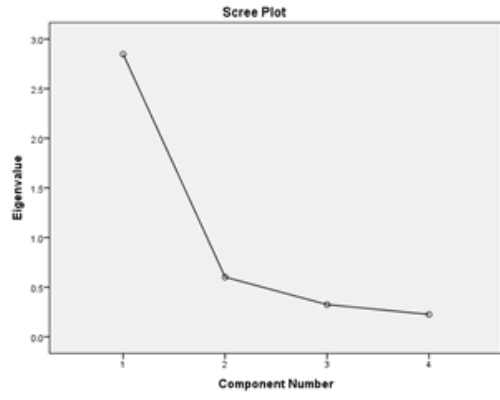


Abbildung 31 Scree-Test PCA letzte Untersuchung

Component Matrix^a

	Component	
	1	2
Thought	.774	.560
Goal setting	.873	
Act freely	.861	
Achieve goals	.863	

Extraction Method: Principal Component Analysis.

a. 2 components extracted.

Tabelle 17 PCA-Ergebnisse der Agenten in der abschließenden Untersuchung

4.4.3 Korrelationsanalyse des Agency-Konstrukts

Die Korrelation des Agency-Konstrukts mit abhängigen Variablen ist größer als in den vorherigen Vorstudien, aber nur für 7 von 9 der Variablen signifikant: Vertrauen, Menschenwürde, Zuversicht, Erwartungen, Unterstützung, Fairness und Beunruhigung ($p < ,01$) (Tabelle 18). Es besteht

ein negativer Zusammenhang zwischen dem Agency-Konstrukt und dem Grad des Unbehagens, der höher ist als zuvor in den Vorstudien. Es zeigt an, dass, wenn die Selbstbestimmung von Waffensystemen erhöht ist, die Leute merken lassen, dass sie beunruhigter sind. Die Korrelationsanalyse ist nicht detailliert genug, um in diese Ergebnisse zu zoomen, deshalb wird dies bei der Analyse der abhängigen Variablen getan.

4.4.4 Manipulationsprüfung des Agenten

Das Diagramm in Abbildung 32 zeigt den Anstieg in der Wahrnehmung von Selbstbestimmung zu den Zuständen auf. Die Differenz in der Wahrnehmung von Selbstbestimmung zwischen Szenarien der von Menschen gesteuerten Waffen und dem neutralen autonomen Waffensystem, sowie zwischen dem Szenario des menschlichen Operators und dem des autonomen Waffensystems ist signifikant, folglich können diese Gruppen als unterschiedlich betrachtet werden. Der t-Test der unabhängigen Stichproben für das Szenario mit dem neutralen AW-Agenten liegt mit $p = ,063$ knapp über dem Schwellenwert für Signifikanz von $p < ,05$. Der t-Test der unabhängigen Stichproben für den Zustand mit menschlichen Operator und dem neutralen AW-Agenten zeigt, dass $p = ,001$ ist und für den Zustand mit menschlichem Operator und dem AW-Agenten mit hoher Selbstbestimmung $p = ,000$ ist.

Wir vermuten, dass die von Menschen gesteuerte Drohne als mit geringer Selbstbestimmung und autonome Waffensysteme als mit hoher Selbstbestimmung aufgefasst werden, und die Ergebnisse beweisen, dass dies der Fall ist. Die zentrale Analyse befasste sich damit, wie der Fall des neutralen Agenten mit den anderen beiden Fällen verglichen werden kann. Wir vermuteten, dass Militärangehörige autonome

Waffen nicht als mit psychischen Fähigkeiten ausgestattet auffassen. Deshalb erwarteten wir keinen Unterschied zwischen dem Zustand des neutralen autonomen Waffensystems und dem Zustand, in dem eine Drohne von Menschen ferngesteuert wird. Wir vermuteten auch, dass der neutrale Kampfmittelzustand als bedeutend anders als der Kampfmittelzustand mit hoher Selbstbestimmung bewertet wurde.

Die Ergebnisse bestätigen, dass die von Menschen gesteuerte Drohne als mit geringer Selbstbestimmung und das autonome Waffensystem als mit hoher Selbstbestimmung aufgefasst werden. Allerdings zeigen die Ergebnisse auch, dass dem neutralen autonomen Waffensystem mehr Selbstbestimmung beigemessen wird, als der von Menschen gesteuerten Drohne und dass die Auffassung von Kampfmitteln mit hoher Selbstbestimmung nur wenig höher ist. Das bedeutet, dass wir unsere Hypothese verwerfen müssen und folgern, dass die für das niederländische Verteidigungsministerium arbeitenden Militärangestellte und Zivilisten autonome Waffensysteme als mit Merkmalen von Streitkräften ausgestattet auffassen.

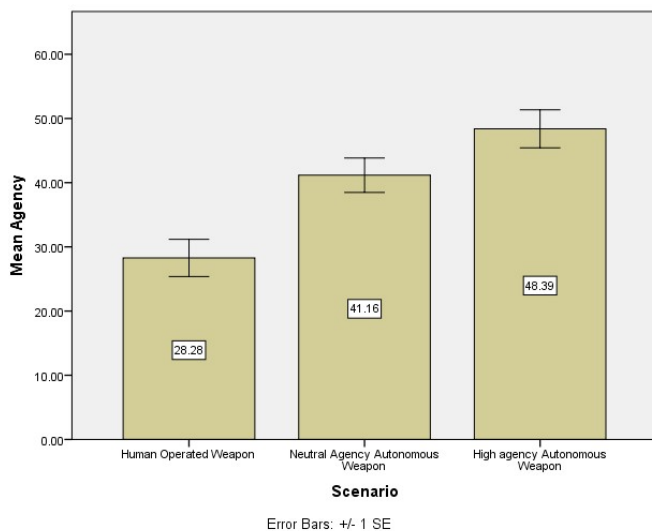


Abbildung 32 Mittelwert des Agens pro Zustand

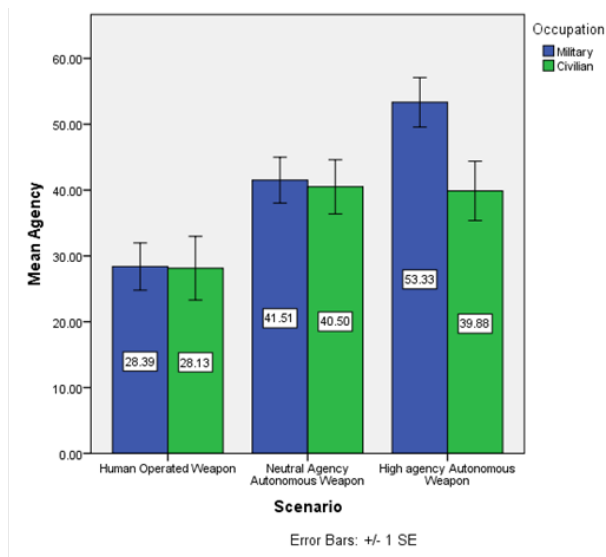


Abbildung 33 Mittelwert des Agens pro Zustand pro Gruppe

Der Unterschied zwischen den befragten Militärangehörigen und den Zivilisten in Abbildung 33 ist sehr interessant, aber als zwei Gruppen unterscheiden sie sich nicht besonders in zwei der Szenarien, sondern nur in dem Kampfmittelzustand mit hoher Selbstbestimmung. Wir müssen bei der Interpretation der Ergebnisse vorsichtig sein, deshalb ist die nachstehende Beobachtung veranschaulichend. Sowohl in dem Szenario mit menschlichem Operateur als auch mit dem neutralen AW-Agenten, stimmte die Auffassung der befragten Militärangehörigen mit der der befragten Zivilisten überein. Aber in dem AW-Agency-Szenario mit hoher Selbstbestimmung waren die befragten Militärangehörigen sich der Effekte mehr bewusst, als ihre zivilen Kollegen. Wir haben uns die demografischen Unterschiede zwischen den Stichproben angesehen und sie sind sich erstaunlich ähnlich, aber eine mögliche Erklärung könnte sein, dass die Stichprobe mit den Zivilisten über 10% mehr Probanden verfügt, die mit KI gearbeitet haben, also die Stichprobe mit den Militärangehörigen. Wir sind jedoch nicht sicher, dass dies den Unterschied in Agentenwahrnehmung zwischen den beiden befragten Gruppen erklärt und dieser Aspekt muss genauer untersucht werden.

Correlations

		Agency	DVO1_Blame	DVO2_Trust	DVO3_Harm	DVO4_Human_Dignity	DVO5_Confidence	DVO6_Expectations	DVO7_Support	DVO8_Fair	DVO9_Anxiety
Agency	Pearson Correlation	1	.009	.406**	-.070	.347**	.428**	.320**	.337**	.451**	-.216**
	Sig. (2-tailed)		.888	.000	.275	.000	.000	.000	.000	.000	.001
	N	248	248	248	248	248	248	248	248	248	248
DVO1_Blame	Pearson Correlation	.009	1	-.039	.056	.018	.027	-.109	-.002	-.080	.128*
	Sig. (2-tailed)	.888		.545	.379	.783	.667	.086	.979	.209	.044
	N	248	248	248	248	248	248	248	248	248	248
DVO2_Trust	Pearson Correlation	.406**	-.039	1	-.217**	.482**	.767**	.324**	.597**	.540**	-.451**
	Sig. (2-tailed)	.000	.545		.001	.000	.000	.000	.000	.000	.000
	N	248	248	248	248	248	248	248	248	248	248
DVO3_Harm	Pearson Correlation	-.070	.056	-.217**	1	-.180**	-.197**	-.155*	-.133*	-.200**	.360**
	Sig. (2-tailed)	.275	.379	.001		.004	.002	.015	.036	.002	.000
	N	248	248	248	248	248	248	248	248	248	248
DVO4_Human_Dignity	Pearson Correlation	.347**	.018	.482**	-.180**	1	.564**	.239**	.564**	.472**	-.385**
	Sig. (2-tailed)	.000	.783	.000	.004		.000	.000	.000	.000	.000
	N	248	248	248	248	248	248	248	248	248	248
DVO5_Confidence	Pearson Correlation	.428**	.027	.767**	-.197**	.564**	1	.376**	.685**	.600**	-.495**
	Sig. (2-tailed)	.000	.667	.000	.002	.000		.000	.000	.000	.000
	N	248	248	248	248	248	248	248	248	248	248
DVO6_Expectations	Pearson Correlation	.320**	-.109	.324**	-.155*	.239**	.376**	1	.271**	.475**	-.264**
	Sig. (2-tailed)	.000	.086	.000	.015	.000	.000		.000	.000	.000
	N	248	248	248	248	248	248	248	248	248	248
DVO7_Support	Pearson Correlation	.337**	-.002	.597**	-.133*	.564**	.685**	.271**	1	.528**	-.488**
	Sig. (2-tailed)	.000	.979	.000	.036	.000	.000	.000		.000	.000
	N	248	248	248	248	248	248	248	248	248	248
DVO8_Fair	Pearson Correlation	.451**	-.080	.540**	-.200**	.472**	.600**	.475**	.528**	1	-.440**
	Sig. (2-tailed)	.000	.209	.000	.002	.000	.000	.000	.000		.000
	N	248	248	248	248	248	248	248	248	248	248
DVO9_Anxiety	Pearson Correlation	-.216**	.128*	-.451**	.360**	-.385**	-.495**	-.264**	-.488**	-.440**	1
	Sig. (2-tailed)	.001	.044	.000	.000	.000	.000	.000	.000	.000	
	N	248	248	248	248	248	248	248	248	248	248

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

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

Tabelle 18 Korrealisation-Agency-Konstrukt mit abhängigen Variablen für die abschließende Untersuchung

4.4.5 Analyse der abhängigen Variablen

Die Korrelation zwischen Agency-Konstrukt und den abhängigen Variablen Vertrauen, Menschenwürde, Zuversicht, Erwartungen, Unterstützung, Fairness und Beunruhigung sind signifikant ($p < ,01$), deshalb sind die bemerkenswertesten Forschungsergebnisse für jede dieser 7 abhängigen Variablen in diesem Kapitel beschrieben.

Vertrauen

Im Großen und Ganzen haben die Menschen mehr Vertrauen, dass menschliche Operateure in der Zukunft korrekte Maßnahmen treffen, als autonome Waffensystem, und die Differenz zwischen diesen Gruppen ist signifikant ($p < ,05$). Der Vertrauensgrad für den neutralen und hohen AW-Agentenzustand ist gleichgroß und die Differenz zwischen diesen Gruppen ist nicht signifikant.

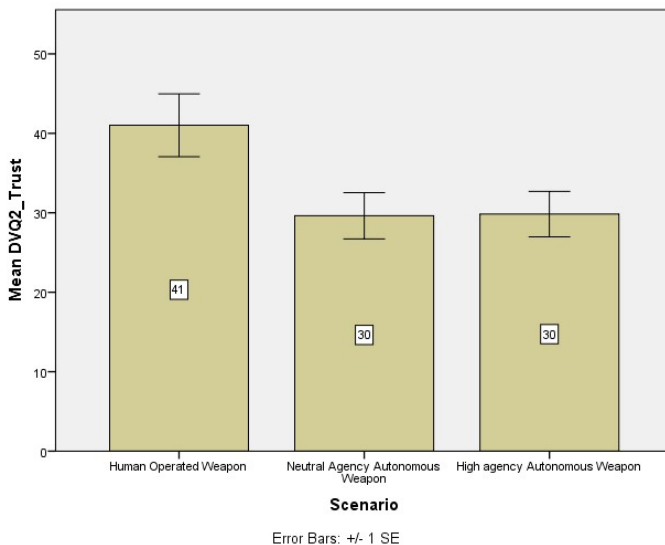


Abbildung 34 Mittelwert der Vertrauensvariable pro Zustand

Menschenwürde

Die Korrelation zwischen Agency-Konstrukt und den abhängigen Variablen Vertrauen, Menschenwürde, Zuversicht, Erwartungen, Unterstützung, Fairness und Beunruhigung sind signifikant ($p < ,01$), deshalb sind die bemerkenswertesten Forschungsergebnisse für jede dieser 7 abhängigen Variablen in diesem Kapitel beschrieben.

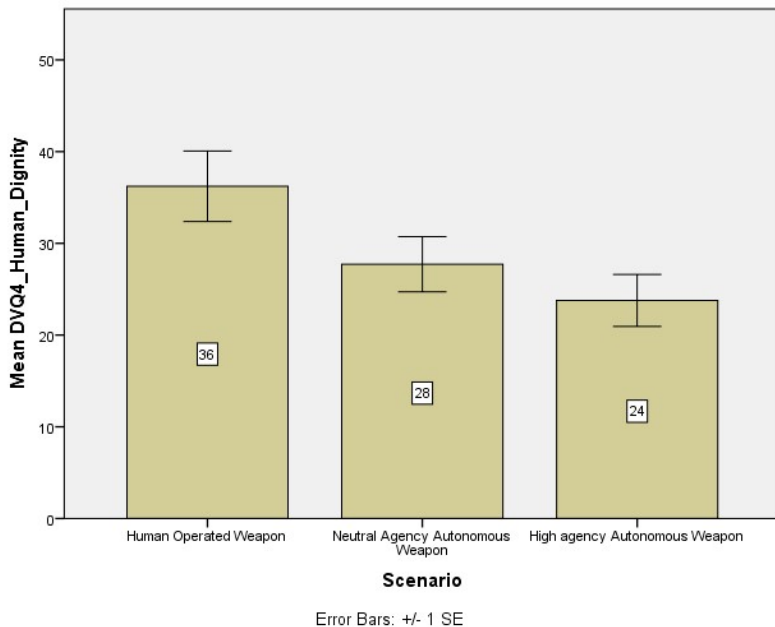


Abbildung 35 Mittelwert der Menschenwürde-Variable pro Zustand

Zuversicht

Im Großen und Ganzen haben die Menschen mehr Vertrauen, dass menschliche Operateure in Zukunft korrekte Maßnahmen treffen, als autonome Waffensysteme, und die Differenz zwischen diesen Gruppen ist signifikant ($p < ,05$). Es gibt keinen nennenswerten Unterschied zwischen AW-Agentenzuständen mit geringer und hoher Selbstbestimmung.

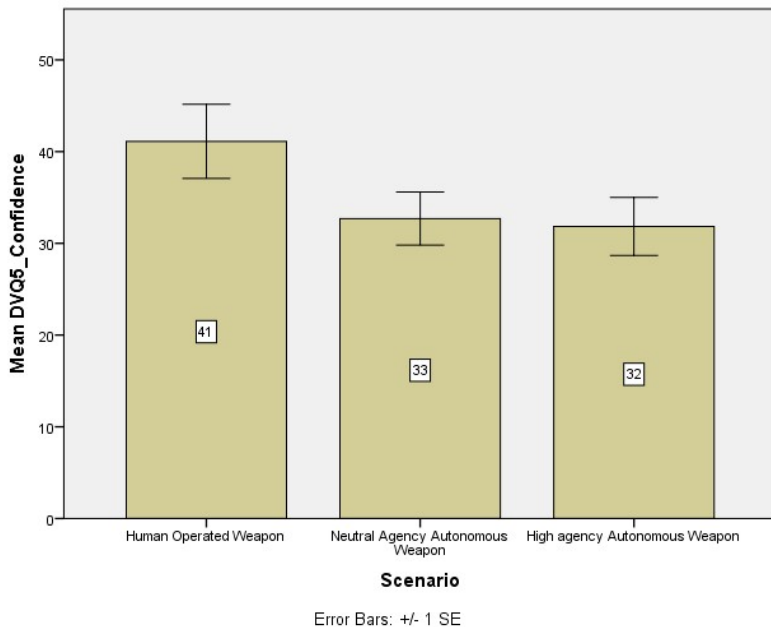


Abbildung 36 Mittelwert der ZuversichtsvARIABLE pro Zustand

Erwartungen

Eine etwas geringere Differenz kann in den Erwartungen der Leute zwischen dem Zustand der menschlichen Operateure und den beiden AW-Agentenzuständen beobachtet werden. Die Erwartungen, die an beide AW-Zustände gestellt werden, sind gleichgroß. Die Differenz zwischen allen Gruppen ist jedoch nicht signifikant.

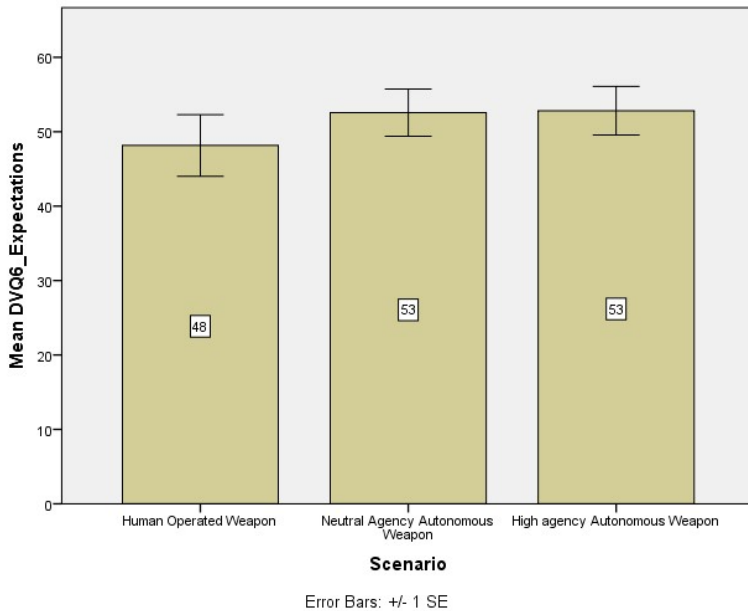


Abbildung 37 Mittelwert der Werterwartungsvariable pro Zustand

Unterstützung

Im Allgemeinen werden von Menschen gesteuerte Drohnen stärker unterstützt als autonome Waffensysteme, und die Differenz zwischen diesen Gruppen ist signifikant ($p < ,05$). Scheinbar beeinflusst die Auffassung in Bezug auf Autonome Waffensysteme nicht den Grad der Unterstützung, und die Differenz zwischen der neutralen und hohen Agentengruppe ist nicht signifikant.

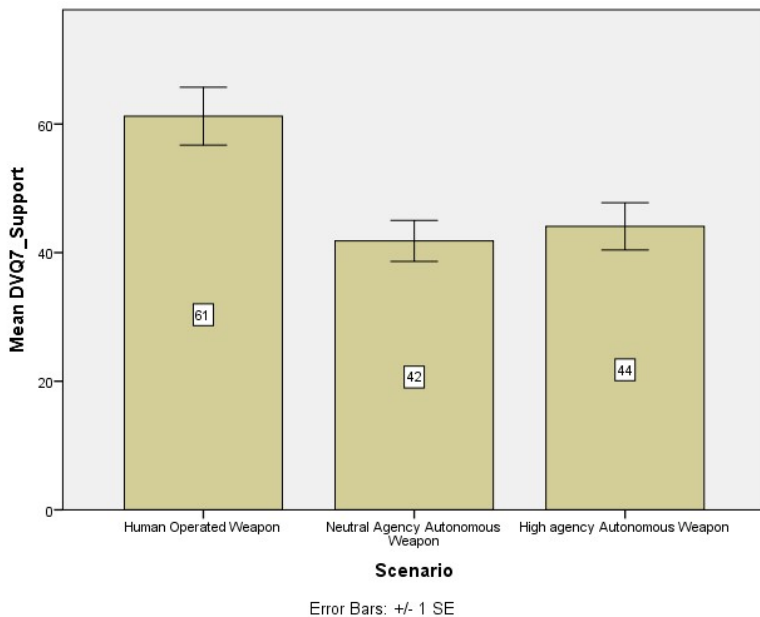


Abbildung 38 Mittelwert der Unterstützungsvariable pro Zustand

Fairness

Sowohl die Handlungen der von Menschen gesteuerten Drohne als auch die autonomen Waffensysteme werden als gleich fair betrachtet, ein bemerkenswertes Ergebnis, da wir erwarteten, dass die Handlungen der von Menschen gesteuerten Drohne als fairer angesehen würden, als die von einem autonomen Waffensystem. Leider ist die Differenz zwischen diesen Gruppen nicht signifikant.

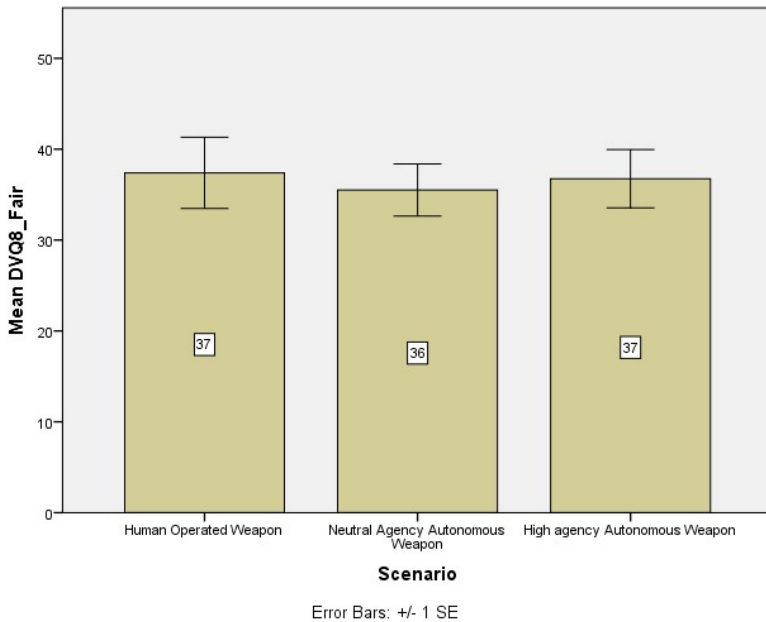


Abbildung 39 Mittelwert der Fairnessvariable pro Zustand

Besorgnis

Die Leute sind besorgter in Bezug auf autonomer Waffensysteme als in Bezug auf von Menschen gesteuerte Drohnen, und die Differenz zwischen diesen Gruppen ist signifikant ($p < ,05$). Die neutralen und hohen AW-Agenten-Szenarien zeigen einen kleinen Anstieg des Besorgnisgrads auf, aber die Differenz zwischen diesen Gruppen ist nicht signifikant.

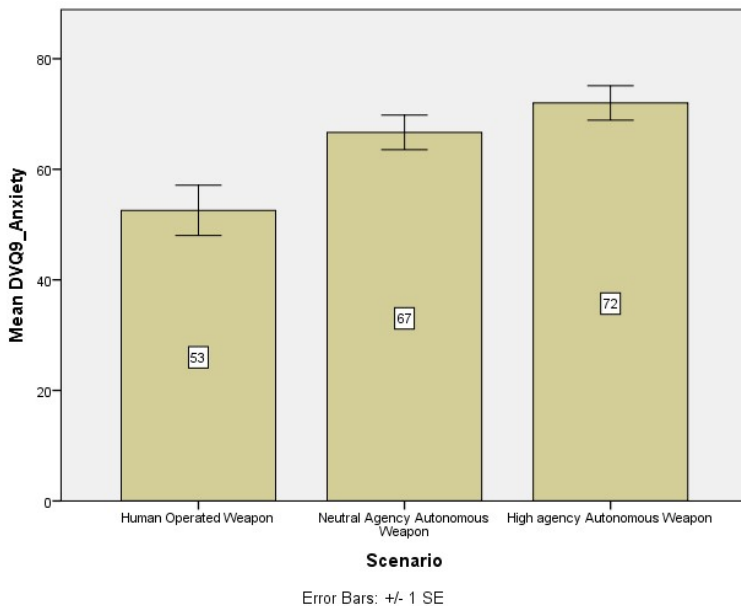


Abbildung 40 Mittelwert der Besorgnisvariable pro Zustand

4.4.6 Schlussfolgerung der abschließenden Untersuchung

Das Ergebnis der abschließenden Untersuchung führt zu den folgenden Schlussfolgerungen:

- Der Reliabilitätstest wies auf, dass die vier Items des Agency-Konstrukts zuverlässig sind, und gemäß der PCA können sie als eine Komponente angesehen werden;
- Die Auffassung und Wertvorstellung des neutralen AW-Agentenzustands ist höher, als die für den Zustand mit von Menschen gesteuerten Drohnen. Die Auffassung und Wertvorstellung des AW-Agentenzustands mit hoher Selbstbestimmung ist höher, als die für den neutralen AW-Agentenzustand, aber die Differenz zwischen diesen Gruppen ist nicht signifikant mit einem $p = ,063$ (Abbildung 32);
- Sowohl in dem Szenario mit menschlichem Operateur als auch mit dem neutralen AW-Agenten, stimmte die Wahrnehmung der befragten Militärangehörigen mit der der befragten Zivilisten überein. Aber in dem AW-Agency-Szenario mit hoher Selbstbestimmung waren sich die befragten Militärangehörigen des Effekts viel stärker bewusst, als ihre zivilen Kollegen (Abbildung 33), aber da die beiden Gruppen keinen nennenswerten Unterschied in zwei der Szenarien aufweisen, müssen wir bei der Interpretation und Schlussfolgerung vorsichtig sein. Wir sind jedoch nicht sicher, ob dies den Unterschied in der Auffassung der beiden befragten Gruppen in Bezug auf Kampfmitteln erklärt, und dieser Aspekt musst genauer untersucht werden.
- Die Korrelationsanalyse zeigte, dass 7 von 9 abhängigen Variablen bedeutend mit dem Agency-Konstrukt korrelieren. Diese 7 Variablen sind Vertrauen, Menschenwürde, Zuversicht, Erwartungen, Unterstützung, Fairness und Besorgnis (Tabelle 18);
- Basiert auf die detailliertere Analyse der abhängigen Variablen können wir beobachten, dass der Grad an Vertrauen, Zuversicht, Menschenwürde und Unterstützung bei den

Maßnahmen der von Menschen gesteuerten Drohnen höher ist wie bei denen von autonomen Waffensystemen (Abbildungen 34, 35, 36 und 38).

- Die Höhe der Erwartungen von Fairness für von Menschen gesteuerte Drohnen und autonome Waffensysteme sind gleichgroß (Abbildung 37 und 39).

Leute empfinden stärkere höhere Besorgnis in Bezug auf autonome Waffensysteme als in Bezug auf von Menschen gesteuerte Waffen, und diese Besorgnis ist am stärksten bei den Agentenzuständen mit hoher Selbstbestimmung.

5. Design der Moral Machine für Autonome Waffensysteme

Ich habe den Film „Eye In The Sky“ gesehen... diese Umfrage erinnert mich an diesen Film.

Vorstudie 1 Befragte

Im vorherigen Kapitel haben wir die Szenarien, die wir in der Vorstudie und abschließenden Untersuchung getestet haben, beschrieben. Obwohl die Stichproben genug Daten für eine statistische Analyse enthielten, verfügten die Studien nur über 50 bis 96 befragte Personen je Szenario und diese Befragten wiesen ganz spezifische demografische Merkmale auf. Zum Beispiel befanden sich die Befragten, die MTurk benutzten, vor allem im Alter zwischen 25 und 35 Jahren, hatten Universitätsreife und der Anteil Militärangehöriger waren zu 95% männlichen Geschlechts niederländischer Nationalität. Um die Ergebnisse zu verallgemeinern benötigt diese Arbeit mehr befragte Personen, die eine größere demografische Gruppe vertreten. Für eine großangelegte Untersuchung dieses Themas schlagen wir das Design einer Moral Machine für autonome Waffen vor, für das wir das Konzept der Moral Machine mithilfe der Media Lab (Scalable Cooperation Group, 2016) entwickelten.

Das Design der Moral Machine für autonome Waffensysteme ist Teil der technischen Untersuchungsphase der Value-Sensitive-Design-Methode. Wir benutzten diese Phase, um die Designtechnologie für den nächsten Schritt unserer Untersuchung vorzubereiten und auf die Szenarien für die Vorstudien und abschließenden Untersuchungen aufzubauen, um Einblicke in die moralische Bewertung der Beteiligten zu tun, die sich mit autonomen Waffensystemen

befassen. Die Erstellung einer Webseite für eine Moral Machine für autonome Waffensysteme würde uns gestatten, Daten zu moralischen Bewertungen in großem Maßstab von verschiedenen demografischen Gruppen zu sammeln, die für robustere und allgemein gültigere Ergebnisse benutzt werden könnten. Breite Datensammlungen aus den verschiedensten Ländern könnten die kulturellen Unterschiede in der moralischen Bewertung autonomer Waffensysteme aufzeigen. Zum Beispiel sind die Ansichten westlicher Länder zum Einsatz dieser Waffen wahrscheinlich anders als die in Ländern des mittleren Ostens oder Asien. Die Daten aus diesen Ländern können in der Debatte um autonome Waffensysteme benutzt werden, die unserer Ansicht nach zurzeit vom westlichen Standpunkt dominiert wird.

Der erste Abschnitt beschreibt die Moral Machine für autonome Fahrzeuge, gefolgt von unserem Vorschlag für die Moral Machine für autonome Waffensysteme, für die wir die Szenarien und Variablen angeben, zum Abschluss erstellen wir dann einige Hinweise für die Einfügung der Webseite.

5.1 Moral Machine für autonome Fahrzeuge

Die ursprüngliche Moral Machine ist eine „...Plattform zum Sammeln von Daten zur Auffassung der Menschen, wenn mit der moralischen Situation konfrontiert, welche Menschen autonome Fahrzeuge verschonen und welchen sie schaden sollen.“ (Awad, 2017, pp. 42-43). Die Webseite verfügt über drei Modi für ihre Benutzer: 1) Ein *Bewertungsmodus*, mit dem Benutzer das Ergebnis für 13 Serien von Szenarien entscheiden können, 2) ein *Designmodus*, mit dem Benutzer ihre eigenen Szenarien konzipieren können und 3) ein *Suchmodus*, mit dem Benutzer die Szenarien anderer ansehen können. Das wichtigste Merkmal ist der Bewertungsmodus, mit dem Benutzer zwischen den beiden Szenarien wählen können, die über unterschiedliche Variablen verfügen:

Eigenschaften {Geschlecht, sozialer Stand, Alter, Gattung, Gesundheit, Utilitarismus}, *Interventionismus* {Auslassung, Kommission}, *Verhältnis zum Fahrzeug* {Fahrer, Fußgänger} und *Einstellung zum Gesetz* {Rechtsmittel, illegale Mittel}. Andere Merkmale schließen einen Videofilm mit der Beschreibung des Projekts, den Anweisungen und Hintergrundinformationen für den Benutzer ein.

Die Moral Machine wurde als ein Massives Online-Experiment (MOE) (Reips, 2002) eingerichtet, das darauf zielt, ein großes Stichprobenkartell mit verschiedenen Hintergründen kurzzeitig und kostengünstig anzuwerben. Diese Merkmale sind offensichtliche Vorteile des MOE, aber die Kehrseite ist, dass die Zustände schwer zu kontrollieren sind, da Benutzer Umfragen mehrere Male beantworten können und selbsternannt sind, d.h. sie können wenn immer sie wollen das Experiment machen oder es fallen lassen (Awad, 2017). Die Moral Machine wurde mittels einer Rapid-Prototyping-Technologie auf der Meteor-Plattform entwickelt, die DOM-Scripting sowie auf Format basierte Strukturen anbietet und eine hohe Ansprechbarkeit aufweist. Es wird auf einem Cloud-Anwendungsservice angeboten, ist für den Einsatz auf mobilen Geräten konzipiert und für soziale Medienvermittlung mit Cards, Markup auf Open Graph Tags optimiert. Bis Mai 2017 bewerteten ca. 3 Millionen Benutzer aus mehr als 160 Ländern mehr als 30 Millionen Szenarien und machten es zu einem der größten Werkzeuge für großangelegte moralische Bewertung der Gegenwart.

5.2. *Moral Machine für autonome Waffensysteme*

Autonome Waffensysteme sind ein heikles Thema, und wir beobachteten, dass es vor allem Beunruhigung und Unbehagen in Leuten hervorruft. Unserer Meinung nach, wäre es nicht klug, eine offene Plattform zu erstellen, um großangelegte

Datenbanken zu sammeln, da so eine Plattform negative Gefühle und unerwünschte Handlung anziehen könnte, die unserer Forschung nicht gut zustattenkäme, z.B. könnten Leute Szenarien entwickeln, in denen gewisse Zielgruppen wie Muslime oder Frauen besonders angegriffen würden. Deshalb erschien es uns ratsam, eine schrittweise Methode zu adoptieren, um uns auf ein Massive-Online-Experiment zu vergrößern. Um ein großangelegte Verlaufsuntersuchung der moralischen Wertvorstellungen von Menschen in Bezug auf autonome Waffensysteme zu erstellen, muss ein kontrolliertes Experiment mit einer begrenzten Anzahl Bedingungen und Stichproben konzipiert werden. Diese Bedingungen können in Szenarien angesehen werden, die Benutzern ermöglichen, die Umfrage nach Erhalt eines Passwortes über eine Web-Oberfläche an einen sicheren Server zu senden. Nach dem Sammeln der ersten Daten und Benutzer-Feedback ist der nächste Schritt, eine großangelegte offene Plattform wie z.B. die Moral Machine maßstäblich zu vergrößern, wo Leute die Szenarien bewerten können, um große Mengen Daten in verschiedenen Ländern zu sammeln. Aufgrund der Vertraulichkeit des Themas glauben wir jedoch, dass es nicht ratsam wäre, Leute ihre eigenen Szenarien erstellen zu lassen, oder ihre Ergebnisse über soziale Medien zu verbreiten, wie es das Designmerkmal der ursprünglichen Moral Machine anbietet.

Wir schlagen die folgenden Merkmale für das Design der Moral Machine für autonome Waffensysteme vor:

- Ein Bewertungsmodus, mit dem Benutzer zwischen verschiedenen Szenarien wählen und anzeigen können, welches der beiden Szenarien am akzeptabelsten für sie ist;
- Bedingungen brauchen nur eine unterschiedliche Variable zu je einer Zeit bei dem Vergleich der Szenarien aufweisen, sodass die Messung nicht auf das Ergebnis von mehreren

Bedingungen zurückgeführt werden kann, um vermischte Effekte zu vermeiden;

- Eine Seite mit einer kurzen Übersicht und Erklärung der Bedeutung der Variablen, bevor der Benutzer mit der Bewertung der Szenarien beginnt;
- Eine Seite mit weiteren Informationen zu dem bestimmten Szenario, das der Benutzer bewertet. Zugriff darauf ist erhältlich, wenn er oder sie weitere Erklärungen zur Veranschaulichung benötigt;
- Nach der Bewertung der Szenarien werden die Ergebnisse dem Benutzer mitgeteilt;
- Eine Ansicht der Benutzerergebnisse erscheint verglichen mit den Ergebnissen der anderen Benutzer, um dem Benutzer Feedback zu seiner/ihrer Bewertung zu geben;
- Eine zweite Runde Bewertungen gestatten den Benutzern, die Szenarien noch einmal zu bewerten, sodass sie ihre moralische Wertvorstellung und Urteile, begründet auf dem Vergleich ihrer Ergebnisse mit denen der anderen, ändern können.

5.3 Szenarien

In diesem Abschnitt werden die Variablen für die Moral Machine für Autonome Waffensysteme definiert, die in zwei Beispielszenarien dargestellt werden. Die Variablen sind auf die Szenarien der Vorstudie und der abschließenden Untersuchung basiert, die eine Militärkolonne beschreiben, die von einer fliegenden Drohne unterstützt wird, und sie werden von den für die Moral Machine benutzten Variablen für autonome Fahrzeuge inspiriert.

5.3.1 Variablen für Szenarien

Für den Prototyp der Moral Machine für autonome Waffensysteme werden 6 Variablen entweder der Vorstudien

oder der abschließenden Untersuchung abgeleitet, z.B. Waffentyp und Ergebnis, oder im Einklang mit der Moral Machine für autonome Fahrzeuge, Variablen wie Merkmale und Anzahl Merkmale. Die 6 Variablen sind: Waffentyp (W), Lage (L), Person (C), Anzahl Personen (N), Resultat (O) und Einsatz (M). Die im nächsten Absatz beschriebenen Variablen sind Beispiele und Vorschläge und müssen professionell konzipiert werden, wenn sie in eine Moral Machine für autonome Waffensysteme eingegeben werden sollen. Zum Beispiel ist die Lage-Variable auf einer Zeichnung basiert und stellt eine Wüste oder ein Dorf wie in einem Spiel dar, und es ist zu prüfen, ob dies eine gute Wiedergabe ist, oder ob ein Foto für eine Forschungsarbeit besser geeignet wäre.

Waffentyp

Diese Dimension zeigt den Waffentyp (W) dar, der eingesetzt wird, um die Kolonne zu unterstützen. Diese Dimension testet, ob Leute die aktuelle Technologie, die von Menschen gesteuerte Drohne, moralisch akzeptabler finden, als eine autonome Drohne, die der Technologie der Zukunft angehört. Dies kann für einen Einblick in die Unterstützung der Technologien benutzt werden. Der Waffentyp ist entweder eine von Menschen gesteuerte Drohne (links) oder ein autonomes Waffensystem (rechts) (Abbildung 41). Wenn die Waffentypvariable unterschiedlich ist, so werden andere Piktogramme im Szenario zur Auswahl dargestellt, siehe Beispiel 1 (Abbildung 47). Andernfalls werden die gleichen Waffentypen im Szenario für Beispiel 2 (Abbildung 48) dargestellt.

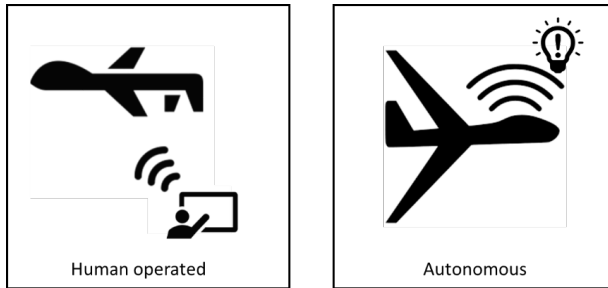


Abbildung 41 Verfügbarer Waffentyp $W = \{\text{von Menschen gesteuerte Drohne, autonome Drohne}\}$

Lage

Die Dimensionen zeigen die Lage (L) als Umfeld für das Szenario an, die sich entweder in der Wüste (links) oder im Dorf (rechts) befindet. In dieser Dimension testen wir, welche Lage für Leute moralisch akzeptabler ist, um entweder autonome Waffensysteme oder von Menschen gesteuerte Drohnen einzusetzen. Wenn die Lagevariable unterschiedlich ist, werden beide Bilder für jedes Szenario angezeigt, siehe Beispiel 2 (Abbildung 48). Andernfalls wird die gleiche Lage in beiden Szenarien dargestellt, siehe Beispiel 1 (Abbildung 47).



Abbildung 42 Lagevariable $L = \{\text{Wüste, Dorf}\}$

Person

Diese Dimension zeigt die Art der Person (C) an, die als Umstehende im Szenario vorhanden sind, entweder ein Mann (links), eine Frau (Mitte) oder ein Kind (rechts) (Abbildung 43). Durch Austauschen der Personen erhalten wir Einblicke in die Umstehenden, die Leute als moralisch akzeptabel empfinden, wenn eine autonome oder eine von Menschen gesteuerte Waffe eingesetzt wird. Wenn die Personenvariable unterschiedlich ist, werden andere Personen für jedes Szenario gezeigt. Andernfalls werden die gleichen Personen in beiden Szenarien dargestellt.

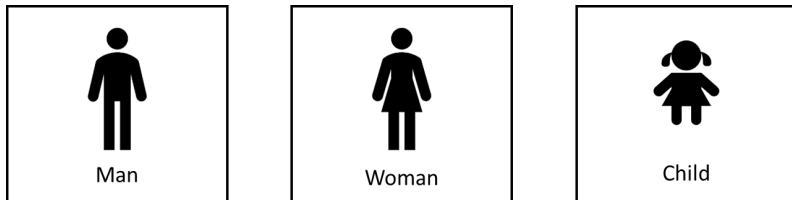


Abbildung 43 Personenvariable $C = \{\text{Mann, Frau, Kind}\}$

Anzahl Personen

Diese Dimension zeigt die Anzahl Personen (N) an, die im Szenario vorkommen. Es kann von einer Person (links) bis zu fünf Personen (rechts) betragen (Abbildung 44). Diese Dimension gestattet uns, Einblicke zu tun, wie viele Umstehende, die Leute als moralisch akzeptabel empfinden, wenn eine autonome oder eine von Menschen gesteuerte Waffe eingesetzt wird. Wenn die Anzahl Personen-Variable unterschiedlich ist, werden andere Gruppen von Personen für jedes Szenario gezeigt. Andernfalls werden die gleichen Personengruppen in beiden Szenarien dargestellt.

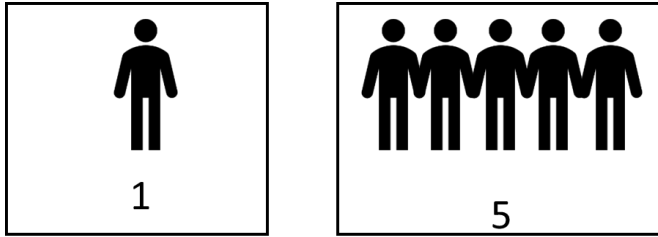


Abbildung 44 Anzahl Personenvariable $N = \{1..5\}$

Resultat

Diese Dimension zeigt das Ergebnis (O) des Szenarios an, dass entweder ohne Kollateralschäden (links) oder mit Kollateralschäden zu der Anzahl in dem Szenario vorkommenden Personen (rechts) darstellt (Abbildung 45). Dies gestattet uns, Einblick zu erhalten, inwiefern das Resultat die moralische Akzeptanz beeinflusst, wenn eine autonome oder eine von Menschen gesteuerte Waffe eingesetzt wird. Wenn die Resultatvariable unterschiedlich ist, werden andere Piktogramme für jedes Szenario dargestellt. Andernfalls wird die gleiche Resultatvariable in beiden Szenarien dargestellt.

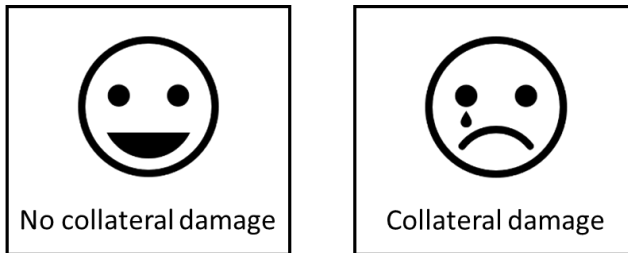


Abbildung 45 Resultatvariable $O = \{\text{kein Kollateralschaden, Kollateralschaden}\}$

Einsatz

Diese Dimension zeigt den Einsatz (M) des Szenarios an, das entweder die Verteidigung der Kolonne, wenn Gefahr im Verzug ist (links) darstellt, oder einen Angriff auf die Fahrzeuge auf der Zielliste darstellt, auch wenn sie keine direkte Gefahr für die Kolonne bedeuten (rechts) (Abbildung 46). In dieser Dimension testen wir, welche Art von Einsatz für Leute moralisch akzeptabel ist. Wenn die Einsatzvariable unterschiedlich ist, werden andere Piktogramme für jedes Szenario dargestellt. Andernfalls wird die gleiche Einsatzvariable in beiden Szenarien dargestellt.

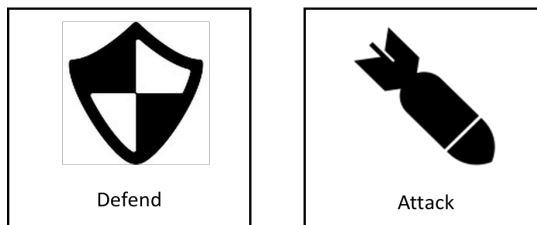


Abbildung 46 Einsatzvariable $M = \{\text{verteidigen, angreifen}\}$

5.3.2 Beispielszenarien

Die oben beschriebenen Szenarien können für das Kreieren von Szenarien benutzt werden, in der jedes Szenario sich nur durch eine Variable von dem anderen unterscheidet. Die für das Szenario gestellte Frage ist die gleiche Frage, wie die für die Bewertung der Szenarien in der ursprünglichen Moral Machine gestellten Frage. In diesem Abschnitt stellen wir zwei Szenarien als Beispiel dar, um das Konzept aufzuzeigen, aber um uns kurz zu fassen, werden wir nicht alle Variablen in einer endlosen Liste von Beispielen zeigen. Beispiel 1 ist in Abbildung 47 zu sehen, die eine Kolonne in einer Wüste

darstellt. Die Differenz zwischen den Szenarien ist, dass die Kolonne auf der linken Seite von einer von Menschen gesteuerten Drohne verteidigt wird und auf der rechten Seite von einem autonomen Waffensystem. In beiden Fällen befindet sich eine Person in der Nähe der Straße, die als Kollateralschaden getötet wird.

Which scenario is most acceptable to you?

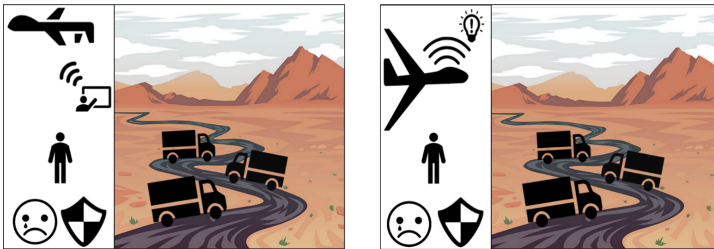


Abbildung 47 Beispielszenario 1

Im nächsten Beispielszenario (Abbildung 48) werden beide Kolonnen von einem autonomen Waffensystem verteidigt, aber die Lage ist anders. Die Kolonne auf der linken Seite fährt durch eine Wüste, die auf der rechten Seite durch ein Dorf. In beiden Situationen spielen zwei Kinder in der Nähe der Straße, und der Angriff verursacht keinen Kollateralschaden.

Which scenario is most acceptable to you?



Abbildung 48 Beispielszenario 2

5.4 Implementierung

Die Umsetzung der Moral Machine für autonome Waffensysteme erfordert mehrere Aktivitäten; es ist ein Projekt, das mindestens sechs Monate dauern wird und ein weiteres Jahr, um es zu bearbeiten. In diesem Abschnitt werden wir die Phasen kurz beschreiben.

In Phase eins müssen die Szenarien und Variablen aus den vorherigen Abschnitten konzipiert werden. Dafür sind Klarheit und Interpretationsfähigkeit der Variablen besonders wichtig. Gemäß des Wissenschaftlers, der die ursprüngliche Moral Machine konstruiert hat, dauerte der Prozess des Erstellens von Personen für die Szenarien drei Monate und wurde mehrfach von einem professionellen Grafiker wiederholt.

In Phase zwei muss die Infrastruktur der Webseite entwickelt und konstruiert werden. Das wird auch mehrere Monate in Anspruch nehmen. Angesichts der streng vertraulichen Natur des Themas ist eine sichere Webseite für die Szenarien erforderlich, Zugang zu ihr ist nur mit einem Passwort möglich. Das erfordert einen Mechanismus, um den interessierten Personen die Passwörter, zu vermitteln, z.B. eine Webseite, die nach der Anmeldung einen Link zur Umfrage sendet. Da wir uns auf ein Massive-Online-Experiment erweitern wollen, muss diese Möglichkeit gleich zu Anfang des Projekts als ein Erfordernis berücksichtigt werden. Das bedeutet, dass die Webseite in einem Server-Umfeld gespeichert wird, um eine dynamische Erweiterung oder Verminderung je nach Anzahl aktiver Benutzer zu gestatten. Gleichzeitig ist es wichtig, dass die Sicherheit der Webseite und des Servers sehr hoch ist und die Forschungsdaten Eigentum der Universität, die die Untersuchung durchführt, sind, und setzt voraus, dass die Webseite nicht für kommerzielle Organisationen wie Microsoft, Google oder Amazon bereit

gestellt wird, die solche dynamischen Cloud-Services anbieten, sondern von einem Server, der Eigentum der Universität ist.

In der dritten Phase sind alle Szenarien in mehreren Vorstudien zu testen, um zu prüfen, ob sie nützliche Ergebnisse erzeugen, und sind neu einzustellen, wenn das nicht der Fall ist. Dies ist mehrere Male zu wiederholen, bis die abschließende Untersuchung getestet werden kann. Die Daten wurden von der ersten Moral Machine vom 23. Juni 2016 bis zum Mai 2017 gesammelt (Awad 2017), und es ist ratsam, die Untersuchung in Bezug auf autonome Waffensysteme für die gleiche Dauer durchzuführen, um eine Stichprobe zu erhalten, die groß genug ist, dass man es wirklich ein Massive-Online-Experiment nennen kann.

6. Schlussfolgerung und Diskussion

Ich bin sicher, dass Drohnen zukünftig eine große Rolle spielen werden. Aber so lange es keinen Algorithmus gibt, um einen Kollateralschaden korrekt einzuschätzen, sollte ein Angriff mit einer Waffe von einer menschlichen Entscheidung abhängen. Außerdem sollte ein Kommandant auf dem Kriegsfeld immer die letzte Entscheidung treffen können, wie der Angriff stattfinden soll. Drohnen können zwar helfen, sein Lagebewusstsein zu verstärken, aber sie dürfen seine Optionen nicht einschränken.

Befragte Person, abschließende Untersuchung

Der Zweck dieser Untersuchung war, Einblick darin zu nehmen, wie autonome Waffensysteme vom Militär und der Zivilbevölkerung wahrgenommen werden und welche moralischen Wertvorstellungen sie als wichtig erachten, wenn autonome Waffensysteme in der nahen Zukunft eingesetzt werden. Für den Aufbau der Untersuchung wurde die Value-Sensitive-Design-Methode benutzt. Dieser Abschnitt beschreibt zuerst die Schlussfolgerungen unserer Forschungsarbeit, gefolgt von einer Diskussion zu den wissenschaftlichen und gesellschaftlichen Auswirkungen. Als Nächstes, identifizieren wir mehrere Einschränkungen und schließen mit Empfehlungen für weitere Untersuchungen ab.

6.1 Schlussfolgerung

Dieser Unterabsatz beschreibt die Schlussfolgerungen zum Agency-Konstrukt, unsere Haupthypothese und der Untersuchung der abhängigen Variablen. Wir geben in jeden Absatz an, ob unsere Ergebnisse die aktuelle akademische

Fachliteratur unterstützen oder ihr widersprechen, und wir bringen die Ergebnisse zum Abschluss.

6.1.1. Agency-Konstrukt

Die Ergebnisse der drei aufeinanderfolgenden Untersuchungen zeigen auf, dass Agency-Items zuverlässig sind und als ein Konstrukt zusammenhalten. Das Agency-Konstrukt besteht aus den vier Items Denken, Zielsetzung, freier Wille und die Zielerreichung, die zum Messen der Auffassung in Bezug auf autonome Waffensysteme und von Menschen gesteuerte Drohnen benutzt werden können. Bei Inbetriebnahme dieses Konstrukts befassten wir uns mit der Fachliteratur aus den Bereichen Kognitive Psychologie, Künstliche Intelligenz und Moralphilosophie. Die meisten empirischen Untersuchungen zu Auffassung von Kampfmitteln finden im Bereich der kognitiven Psychologie statt, z.B. K. Gray und Wegner (2012), die in ihrer Studie von Robotern und Zombies verschiedene Agentenebenen beschreiben, um Unbehagen zu messen. Malle und Thapa Magar (2017) haben die erwünschten intelligenten Fähigkeiten in humanoiden Robotern untersucht, dafür haben sie einige der Agentenaspekte benutzt, wie Denkfähigkeit und Erklärung ihrer Handlungen. Allerdings haben wir in keiner der Untersuchungen auch nur ein einziges Konstrukt gefunden, um Agentenebenen bei technischen Artefakten zu messen. Soweit wir wissen, ist dies das erste Konstrukt, um Agentenwahrnehmung empirisch zu messen.

6.1.2 Zentrale Hypothese Auffassung in Bezug auf Waffensysteme

In der abschließenden Untersuchung vermuteten wir, dass die von Menschen gesteuerte Drohne als mit geringer Selbstbestimmung und das autonome Waffensysteme als mit hoher Selbstbestimmung aufgefasst werden. Die Ergebnisse beweisen, dass ein bedeutender Unterschied zwischen diesen

Zuständen besteht, daraus folgern wir, dass die Agentenmanipulation funktioniert. Wir erwarteten auch, dass der neutrale Agentenzustand der autonomen Waffensysteme als bedeutend anders von den autonomen Waffensystemen mit hoher Selbstbestimmung bewertet wird, und unsere Ergebnisse deuten darauf hin, dass es einen Unterschied gibt, aber dass die Differenz zwischen diesen beiden Gruppen knapp über dem von uns gewählten Signifikanz-Schwellenwert liegt, und deshalb müssen wir mit unseren Schlussfolgerungen sehr vorsichtig sein.

Unsere zentrale Analyse befasste sich damit, wie das neutrale Agency-Szenario von autonomen Waffensystemen sich von den von Menschen gesteuerten und den autonomen Waffensystemen mit hoher Selbstbestimmung unterscheiden. Wir vermuteten, dass Militärangehörige autonome Waffensysteme als genau solche Waffen wie alle anderen Waffen betrachten und sie deshalb nichts weiter als ein Werkzeug sind, mit dem man eine Wirkung erhält und deshalb autonome Waffensysteme keinen psychischen Zustand kennen. Wir erwarteten auch nicht, einen Unterschied zwischen dem Zustand des neutralen autonomen Waffensystems und dem Zustand, in dem eine Drohne von Menschen ferngesteuert wird, zu finden. Unsere Ergebnisse zeigen an, dass die Wahrnehmung von Militärangehörigen und zivilen Mitarbeitern im niederländischen Verteidigungsministerium in Bezug auf das neutrale autonome Waffensystem-Szenario stärker ist, als die Wahrnehmung des von Menschen gesteuerten Drohnenzustands. Dies bedeutet, dass sie einem autonomen Waffensystem mehr Handlungsfähigkeit beimessen, als einer durch menschliche Einwirkung betriebenen Drohne. Basiert auf diesen Erkenntnissen müssen wir unsere Hypothese verwerfen:

H1: Militärangehörige fassen autonome Waffen nicht als mit psychischen Fähigkeiten ausgestattet auf.

Unsere Ergebnisse entsprechen der Untersuchung von Agenten im Bereich der kognitiven Psychologie. Frühere Untersuchungen stellten fest, dass Menschen Computern Verstand beimessen (Nass *et. al.* 1995) und fassen Roboter als Agenten auf (H.M. Gray *et. al.* 2007). Basiert auf unsere Untersuchung folgern wir, dass Zuordnung von verstandesmäßiger Wahrnehmung auch auf autonome Waffensysteme zutrifft, und dass diese Waffensysteme als mehr angesehen werden, als nur ein Werkzeug zum Erreichen eines Effekts.

Die Ergebnisse führten auch zu einer weiteren eindeutigen Beobachtung, dass sowohl bei von Menschen gesteuerten autonomen Waffensystemen, als auch solchen mit neutralen Agenten, Militärangehörige und Zivilisten gleich reagierten.

Wie in Abschnitt 4.4.4 erwähnt wurde, unterscheiden sich die beiden Gruppen nicht sonderlich in zwei der Szenarien, nur das mit dem Zustand hoher Selbstbestimmung, und deshalb müssen wir bei der Interpretation der Ergebnisse und unseren Schlussfolgerungen sehr vorsichtig sein. Aber in dem AW-Agency-Szenario mit hoher Selbstbestimmung waren sich die befragten Militärangehörigen viel stärker des Effekts bewusst, als ihre zivilen Kollegen. Wir sind uns nicht sicher, wie man sich die Differenz in der Auffassung von Kampfmitteln beider befragter Gruppen erklären könnte. Eine Erklärung ist, dass die zivile Stichprobe 10% mehr Personen betrug, die mit KI arbeiteten, als die militärische Stichprobe. Eine weitere Erklärung könnte sein, dass die Gruppen die Szenarien unterschiedlich wahrnahmen, z.B. verstanden die Militärangehörigen die Beschreibung wortwörtlich und die

Zivilisten interpretierten sie lockerer, und dass dies ihre Antworten zu den Agentenfragen beeinflusste. Allerdings sind diese Erklärungen zurzeit spekulativ, und die Ursachen für die Differenz in der Auffassung von autonomen Waffensystemen zwischen Militärangehörigen und Zivilisten müssen noch weiter untersucht werden.

6.1.3 Untersuchung der abhängigen Variablen

Die Auswirkung der Agentenwahrnehmung auf die unabhängigen Variablen wird auf beschreibende Weise untersucht, und wir können keine Schlussfolgerungen zum Zusammenhang der Agentenstufen und ihre Auswirkung auf die abhängigen Variablen ziehen. Wir fanden, dass 7 von 9 abhängigen Variablen stark mit dem Agency-Konstrukt korrelieren. Dies sind die Variablen: Vertrauen, Menschenwürde, Zuversicht, Erwartungen, Unterstützung, Fairness und Besorgnis. In diesem Abschnitt beschreiben wir die Ergebnisse zu den abhängigen Variablen, die wir nach ihren Ergebnissen auf Cluster basierten.

Vertrauen, Zuversicht und Unterstützung

Die Ergebnisse deuten an, dass Militärangehörige und zivile Mitarbeiter im niederländischen Verteidigungsministerium mehr Vertrauen haben, zuversichtlicher sind und eher die Maßnahmen der von Menschen gesteuerten Drohnen unterstützen, als die von autonomen Waffensystemen. Wir konnten keine wissenschaftliche empirischen Forschungsarbeiten finden, aber eine amerikanische Analyse und Umfrage (Gallup-Poll) von 2013 berichtete starke Unterstützung für unbemannte Drohnenangriffe im Ausland und in einem öffentlichen Bericht beschrieben Schneider und McDonald (2016), dass Bürger der USA unbemannte Luftangriffe denen von bemannten Luftangriffen vorzogen, da sie weniger riskant für die Militärangehörigen waren. Wir theoretisieren, dass höhere Stufen von Unter-

stützung, Vertrauen und Zuversicht mit der Tatsache erklärt werden kann, dass Leute mit von Menschen gesteuerten Drohnen vertrauter sind, da diese Technologie derzeit benutzt wird, im Vergleich zu einer neuen, der Zukunft angehörigen Technologie, mit der sich Leute nicht auskennen. Eine weitere Erklärung könnte sein, dass Leute mehr Vertrauen und Zuversicht in die Maßnahmen von Menschen setzen, als in die Maßnahmen von autonomen Systemen. Zurzeit sind die Gründe für diesen Unterschied noch nicht offensichtlich und deshalb sollte dies in einer Verlaufsstudie untersucht werden.

Menschenwürde

Die Ergebnisse weist auf die Auffassung hin, dass eine von Menschen gesteuerte Drohne die Menschenwürde mehr achtet, als neutrale autonome Waffensysteme oder solche mit hoher Selbstbestimmung, obwohl die Maßnahmen und Ergebnisse der Szenarien die gleichen sind. Diese Ergebnisse stehen im Widerspruch mit dem Argument von Kasher (2016), der behauptet, dass Menschenwürde in der Natur des Artefakts keine Rolle spielte, aber in den Vorhaben und Entscheidungen der Personen, die sie bedienen. Unsere Ergebnisse deuten an, dass die Natur des Artefakts an der Auffassung von Menschenwürde beteiligt ist, da die Absichten und Entscheidungen von sowohl der von Menschen gesteuerten Drohne und dem autonomen Waffensystem die gleichen sind. Diese Ergebnisse reflektieren die Meinungen der Experten in den Interviews, die die Menschenwürde als Grund, sich autonomen Waffensystemen zu widersetzen, viele Male erwähnen. Ein Mangel an Respekt für die Menschenwürde scheint einer der Haupteinsprüche gegen den Einsatz autonomer Waffensysteme zu sein, und es ist eindeutig, dass dieses Ergebnis sowohl die Reaktion der Militärangehörigen als auch der Zivilisten, die für das niederländische

Verteidigungsministerium tätig sind, ist. Wie dem auch sei, dies soll nicht besagen, dass Militärangehörige autonome Waffensysteme missbilligen, oder die gleichen gegnerischen Argumente einiger der Experten teilen.

Erwartungen und Fairness

Aufgrund der Ergebnisse können wir folgern, dass Militärangehörige und zivile Mitarbeiter im niederländischen Verteidigungsministerium eine gleichwertige Stufe von Erwartungen zu den Maßnahmen der von Menschen gesteuerten Drohne und der neutralen autonomen Waffensysteme mit hoher Selbstbestimmung haben. Außerdem betrachten sie die Handlungen der von Menschen gesteuerten Drohnen und autonomen Waffensystemen als gleich fair. Die von Shelley (2013) ausgedrückte Besorgnis, dass autonome Drohnen als weniger fair angesehen werden, wie die von Menschen gesteuerten Drohnen, wurden von unseren Ergebnissen nicht bestätigt. Eine Fachliteratur, die Erwartungen und autonome Waffensysteme oder Drohnen verknüpfte, konnte nicht gefunden werden. Unsere Ergebnisse wiesen keine Unterschiede in der Höhe der Erwartungen und Fairness an die derzeitige und zukünftige Technologie unter Militärangehörigen und Zivilisten, die im niederländischen Verteidigungsministerium angestellt sind, auf.

Besorgnis

Unsere Ergebnisse zeigen, dass autonome Waffensysteme mehr Besorgnis unter den Militärangehörigen und zivilen Mitarbeitern im niederländischen Verteidigungsministerium erregen, als die von Menschen gesteuerten Waffensysteme. Der Unterschied in der Stärke der Besorgnis ist am größten bei Zuständen mit hoher Selbstbestimmung. Besorgnis in Bezug auf autonome Waffensysteme wird oft in der

wissenschaftlichen Fachliteratur (Noone & Noone 2015, Ohlin 2016) beschrieben, aber auch in den Meinungen der von uns befragten Experten und von den Rechercheuren, die mit den Leuten im Gespräch sind, ausgedrückt. Diese Ergebnisse bestätigen, dass Leute sich Sorgen um die Anwendung von autonomen Waffensystemen machen. Diese Besorgnis wird von sowohl Militärangehörigen und Zivilisten im Dienst des niederländischen Verteidigungsministeriums als auch den Experten geteilt, die die Meinung der Zivilbevölkerung äußern.

6.2 *Diskussion*

Die Auswirkungen dieser Ergebnisse werden in diesem Unterabschnitt besprochen. Zuerst werden die wissenschaftlichen Beiträge zur Fachliteratur identifiziert, gefolgt von den gesellschaftlichen Auswirkungen, die zur derzeitigen Debatte über autonome Waffensysteme beitragen.

6.2.1 *Wissenschaftliche Folgerungen*

Diese Untersuchung trägt zur Fachliteratur in verschiedenen Aspekten bei. Sie erstellt eine Übersicht der verschiedenen Definitionen autonomer Waffensysteme, die zurzeit in der Fachliteratur benutzt werden, und zeigt, dass noch keine allgemeingültige Definition vereinbart worden ist. Beim Erstellen dieser Übersicht wurde deutlich, dass einige Gruppen aus verschiedenen Gründen sogar davor warnen, autonome Waffensysteme eindeutig zu definieren. Manche argumentieren, dass Maschinen nicht im wortwörtlichen Sinne autonom sein können und andere betonen, dass Definitionen von Autonomie verschiedenen Funktionen zugeordnet werden können. Wie von einem Experten in einem der Interviews erwähnt wurde, könnte ein anderer Grund für dieses Zögern der sein, dass, indem man autonome Waffensysteme nicht genau definiert, die Diskussion offen bleibt und eine

Blockierung des Themas vermieden wird. Wir wählten die Definition der Advisory Council on International Affairs (AIV & CAVV) für autonome Waffensysteme, weil sie vordefinierte Kriterien berücksichtigt und mit dem militärischen Zielbestimmungsverfahren verbunden ist, da die Waffe nur eingesetzt wird, wenn ein Mensch diese Entscheidung getroffen hat. Wahl und Vorschlag dieser Definition aus der Übersicht ist ein Beitrag zur Fachliteratur.

Ein weiterer Beitrag dieser Untersuchung ist unsere Identifizierung der Werte, die Leute mit autonomen Waffensysteme in Beziehung bringen. Die Übersicht wird von sowohl bewerteten Werttheorien als auch von Experten, die an der Debatte über autonome Waffensysteme teilnehmen oder im militärischen Bereich tätig sind, abgeleitet. Wir entscheiden uns, die Werte Schuld, Vertrauen, Schaden, Menschenwürde, Zuversicht, Erwartungen, Unterstützung, Fairness und Besorgnis in der abschließenden Untersuchung zu prüfen. Die Ergebnisse erstellen Einblick in die Auffassung und Wertvorstellungen von Streitkräften und Zivilisten, die für das niederländische Verteidigungsministerium tätig sind, in Bezug auf autonome Waffensysteme als Technologie der Zukunft. Unseres Wissens nach ist diese Untersuchung die erste, die diese Werte mit Hinsicht auf autonome Waffensysteme empirisch erforscht und diese Werte damit vergleicht, wie sie in aktuellen und zukünftigen Waffensystemen wahrgenommen werden – ein neuer Beitrag zur wissenschaftlichen Debatte.

Der dritte Beitrag zur Fachliteratur besteht in dem Vorschlag in unseren Untersuchungen für ein Agency-Konstrukt, um die Wahrnehmung in Bezug auf autonome Waffensysteme zu messen. Soweit wir wissen, ist dieses das erste Konstrukt, das in Betrieb genommen wurde, um die Agentenstufen technologischer Artefakte zu messen. In unserer Untersuchung wurde es für Drohnen angewandt, aber

wir glauben, dass es genauso gut für die Messung der Agentenwahrnehmung von anderen Objekten angewandt werden kann, da die Fragen zu den Items leicht umgeschrieben werden können, um einen anderen Bereich widerzuspiegeln. Wenn dieses Agency-Konstrukt standhält, werden Anschlussstudien in anderen Bereichen erforderlich.

6.2.2. Gesellschaftliche Folgerungen

Die Ergebnisse haben auch gesellschaftliche Auswirkungen, da sie zur Debatte um autonome Waffensysteme beitragen. Dieses Thema ist nur wenig erforscht worden, und folglich wird die Debatte von abstrakten moralischen und rechtlichen Theorien dominiert. Unsere Untersuchung erstellt empirische Daten zu den zugrundeliegenden Werttheorien, und wir weisen nach, dass diese Werte nicht nur für die Fachliteratur relevant sind, sondern scheinbar auch in der praktischen Erfahrung. Dadurch verbinden wir die abstrakten Werttheorien zum praktischen Bereich des Einsatzes autonomer Waffensysteme. Zum Beispiel fanden wir den Wert der Menschenwürde oftmals in der Fachliteratur und von den Experten im Interview erwähnt. In unserer Untersuchung fanden wir empirische Daten dafür, dass Militärangehörige und Zivilisten, die für das niederländische Verteidigungsministerium arbeiten, der Auffassung sind, dass autonome Waffensysteme weniger Respekt für die Menschenwürde haben. Ebenso verknüpfen wir den Wert des Nicht-Schadens der Bioethik-Theorie, den wir als Schaden beschreiben, mit autonomen Waffensystemen und fanden Auswirkungen in den Vorstudien 1 und 2 bei der Übertragung von Schaden auf die Verantwortungskette, die in einer Folgestudie weiter untersucht werden müsste.

Unsere Untersuchung stellt auch empirische Daten zur Verfügung und konkretisiert einige der Ansichten und Meinungen, die benutzt werden könnten, um eine

Gemeinsamkeit in der Debatte über autonome Waffensysteme zu finden. Autonome Waffensysteme verursachen mehr Besorgnis unter den Militärangehörigen und zivilen Mitarbeitern im niederländischen Verteidigungsministerium, als von Menschen gesteuerte Waffensysteme. Obwohl wir dies noch nicht in einer Stichprobe, die völlig aus Zivilisten besteht, untersucht haben, so erwarten wir, basiert auf der Fachliteratur, die gleichen Ergebnisse in solch einer Stichprobe zu finden, was bedeuten würde, dass sowohl Militärangehörige als auch Zivilisten denken, dass autonome Waffensysteme unabhängig entscheiden und Pläne machen, um ihre Ziele zu erreichen. Eine weitere Gemeinsamkeit sind die Werte von Menschenwürde und Besorgnis. Unser Ergebnis zeigt, dass Militärangehörige und Zivilisten, die für das niederländische Verteidigungsministerium tätig sind, besorgter um den Einsatz autonomer Waffensysteme als um den Einsatz der von Menschen gesteuerten Drohnen sind. Ihrer Auffassung nach haben sie auch weniger Respekt für die Würde des menschlichen Lebens, als die von Menschen gesteuerten Drohnen. Menschenwürde und Besorgnis sind zwei Werte, die oft von den Experten in ihren Interviews erwähnt werden, sie müssen auf jeden Fall in der Debatte der Ethik vom Einsatz autonomer Waffensysteme angesprochen werden.

Einblick in diese Werte tragen auch zu dem Designprozess autonomer Waffen bei. Unsere Ergebnisse zeigen, dass das Vertrauen, die Zuversicht und Unterstützung für autonome Waffensysteme geringer sind, als für die von Menschen gesteuerten Drohnen. Zu diesem Zeitpunkt würden wir gerne anmerken, dass autonome Waffensysteme nicht nur Nachteile, sondern auch eindeutige militärische Vorteile haben (Etzioni & Etzioni, 2017), und die Konzipierung von Merkmalen, die das Vertrauen und die Zuversicht in autonome Waffensysteme bestärken, ist vom militärischen Standpunkt

aus positiv zu sehen. Dies erfordert, dass während des Designverfahrens die menschlichen Werte auf die Designanforderungen übertragen werden. Dies kann mithilfe der Wertehierarchie sichtbar gemacht werden (Van de Poel, 2013), siehe Abschnitt 2.2.4, eine hierarchische Struktur der Werte, Normen und Designanforderungen, die die Werturteile für die Übertragung klar, transparent und diskussionsfähig macht.

6.3 *Einschränkungen*

Verschiedene Aspekte können als Einschränkungen dieser Untersuchung identifiziert werden. Zuerst wurden die Operationalisierung und das Agency-Konstrukt von einer Kategorisierung der Literatur abgeleitet, die die Agentenmerkmale beschreibt, und unsere Wahl der Merkmale basierte auf einer numerischen Zahl und nicht auf Relevanz oder Bewertungskriterien. Diese Methode wurde deshalb ausgesucht, weil sie unserer Ansicht nach am objektivsten war, und wir haben nicht versucht, subjektive Wahlen zu machen, doch es ist möglich, dass wir relevante Eigenschaften ausgelassen, oder falsche oder irrelevante Items gewählt haben. Die Operationalisierung der Eigenschaften in unseren Fragen basierte auf Heuristik, und wir haben die Fragen nicht auf ihre Richtigkeit getestet. Es würde anderen schwerfallen, diese Auswahlmethode zu reproduzieren und wirkt sich auf die Reproduzierbarkeit und interne Bewertung der Untersuchung aus.

Die zweite methodologische Einschränkung ist die Auswahl der Werte aus dem Werte-Fragebogen und den Experteninterviews als abhängige Variable. Obwohl diese Auswahl heftig zwischen den drei Rechercheuren dieser Arbeit debattiert wurde, fand die endgültige Wahl auf methodischem Weg und nicht nach objektiver Methode statt. Ein Beispiel

dafür ist unsere Wahl von *Schaden* von der Liste in Tabelle 6, aber nicht *Kontrolle*. Folglich ist es wahrscheinlich, dass die Auswahl der Neigung des Rechercheurs unterstellt ist und andere Wissenschaftler andere Werte als die abhängigen Variablen gewählt hätten. Im Nachhinein hätte man eine objektivere Methode und strukturierte Form auswählen sollen, um diese Neigung einzuschränken. Diese Methode beeinflusst die Reproduzierbarkeit und interne Bewertung der Untersuchung auf negative Weise.

Drittens sind die Stichproben für diese Untersuchung sowohl in Größe als auch Demografik beschränkt. Obwohl groß genug für eine robuste Datenanalyse, hatte die abschließende Untersuchung nur 239 Befragte, ein geringer Anteil des niederländischen Verteidigungsministeriums. Die Teilnahme an der Untersuchung war freiwillig, und wir konnten die Beiträge der Befragten nicht kontrollieren. Das heißt, dass die Stichprobe nicht typisch für das gesamte niederländische Verteidigungsministerium ist, was die externe Gültigkeit der Untersuchung beeinträchtigt und bedeutet, dass wir die Ergebnisse nicht für die gesamte Verteidigungsorganisation verallgemeinern können.

Zum Schluss ist die Verteilung der befragten Personen des letzten Szenarios verzerrt und das zweite Szenario (neutrales autonomes Waffensystem) verfügte über 32 befragte Personen mehr als das erste Szenario (von Menschen gesteuerte Drohnen). Eine der Erklärungen für diese Schieflage könnte sein, dass Leute erwarteten, eine Umfrage zu autonomen Waffensystemen vorzunehmen, aber sie zogen sich zurück, als sie mit von Menschen gesteuerten Szenarien präsentiert wurden. Dies nennt sich selektive Schwundquote und hat eine negative Auswirkung auf die interne Gültigkeit der Untersuchung. Das bedeutet, dass wir nicht annehmen können, dass Befragte in beiden Szenarien die gleichen

Eigenschaften wählen, und dass die Randomisierung der Untersuchung, ein Kriterium von höchster Wichtigkeit für ein kontrolliertes Experiment, fehlgeschlagen hat. Eine weitere Erklärung könnte ein Softwarefehler bei der Verteilung der Probanden über die Szenarien sein, und wir stehen derzeit im Kontakt mit Qualtricks, um zu sehen, ob dies der Fall ist. Wird die ungerechte Verteilung von fehlerhafter Software verursacht, ist die Gültigkeit der Untersuchung nicht beeinträchtigt, da sie nicht der selektiven Schwundquote zugeordnet werden kann.

6.4 *Empfehlungen für weitere Forschungsarbeiten*

Mit Rücksicht auf diese Einschränkungen werden mehrere Empfehlungen für weitere Forschung vorgeschlagen. Die erste ist, das Agency-Konstrukt zu bewerten, das weitere Untersuchungen zur Messung der Auffassung der Kampfmittel anderer technologischen Artefakten erforderlich macht, um zu prüfen, dass dieses Konstrukt auch anderen Bereichen standhält. Beispiele dieser Forschungsarbeiten könnten die Untersuchung der Auffassung in Bezug auf Dienstmittel einschließen, von Pflegerobotern für Senioren, KI-Spielzeug für Kinder oder Onboard-Computern von autonomen Fahrzeugen. Die zweite Empfehlung ist, die abschließende Untersuchung mit den gleichen Szenarien mit einer repräsentativen durchzuführen, die einzig aus der niederländischen Zivilbevölkerung stammt. Dies würde uns ermöglichen zu erkennen, zu welchen Werten sich die Ergebnisse der militärischen Stichproben und die Antworten von Probanden aus der Zivilbevölkerung voneinander unterscheiden und obwohl wir keine direkte Vergleiche ziehen können, aufgrund der Tatsache, dass die militärische Probe nicht repräsentativ ist, könnten wir doch Einsicht in die Auffassung von sowohl Militärangehörigen als auch der Zivilbevölkerung nehmen. Abschließend empfehlen wir die

Implementierung der Moral Machine für Autonome Waffensysteme, siehe Abschnitt 5, um die Ergebnisse zu verallgemeinern. Dafür benötigt diese Untersuchung weitaus mehr befragte Personen, die eine große demografische Gruppe repräsentieren. Die Korrektur nach oben für ein Massives Online-Experiment wie die Moral Machine, würde große Mengen Daten verschiedener demografischer Gruppen erzeugen, die für eine robustere und verallgemeinerbare Ergebnisse benutzt werden könnten, um ein gründliches Verständnis der moralischen Beurteilung der Leute in Bezug auf autonome Waffensysteme zu erhalten.

Quellenverzeichnis

- AIV, & CAVV. (2016). *Autonomous weapon systems: the need for meaningful human control*. (No. 97, No. 26). Retrieved from <http://aiv-advice.nl/8gr>.
- Aldewereld, H., Dignum, V., & Tan, Y.-h. (2015). Design for Values Information and communication technologies in Software Development Software development. *Handbook of Ethics, Values, and Technological Design: Sources, Theory, Values and Application Domains*, 831-845.
- Alfano, M., & Loeb, D. (2014). Experimental Moral Philosophy. Retrieved from <https://plato.stanford.edu/entries/experimental-moral/>
- Altmann, J., Asaro, P., Sharkey, N., & Sparrow, R. (2013). Armed military robots: editorial. *Ethics and Information Technology*, 15(2), 73.
- Asaro, P. (2012). On banning autonomous weapon systems: human rights, automation, and the dehumanization of lethal decision-making. *International Review of the Red Cross*, 94(886), 687-709.
- Asaro, P. (2016, 14-10-2017). Talk on Autonomous Weapon Systems. Retrieved from <https://livestream.com/nyu-tv/ethicsofAI/videos/138822041>
- Awad, E. (2017). *MORAL MACHINE: Perception of Moral Judgment Made by Machines*. (Master), Massachusetts Institute of Technology, Boston.
- Bandura, A. (2001). Social cognitive theory: An agentic perspective. *Annual review of psychology*, 52(1), 1-26.
- Beauchamp, T. L., & Walters, L. R. (1999). *Contemporary Issues in Bioethics*. Wadsworth Pub.
- Bonnefon, J.-F., Shariff, A., & Rahwan, I. (2016). The social dilemma of autonomous vehicles. *Science*, 352(6293), 1573-1576.
- Borning, A., & Muller, M. (2012). *Next steps for value sensitive design*. Paper presented at the Proceedings of the SIGCHI conference on human factors in computing systems.
- Bradley, B. (2006). Two concepts of intrinsic value. *Ethical Theory and Moral Practice*, 9(2), 111-130.

- Bratman, M. E., Israel, D. J., & Pollack, M. E. (1988). Plans and resource-bounded practical reasoning. *Computational intelligence*, 4(3), 349-355.
- Bryson, J. J., Kime, P. P., & Zürich, C. (2011). *Just an artifact: why machines are perceived as moral agents*. Paper presented at the IJCAI Proceedings-International Joint Conference on Artificial Intelligence.
- Campaign to Stop Killer Robots. (2015). Campaign to Stop Killer Robots. Retrieved from <https://www.stopkillerrobots.org/>
- Campaign to Stop Killer Robots. (2017). The Problem. Retrieved from <http://www.stopkillerrobots.org/the-problem/>
- Carpenter, J. (2016). *Culture and Human-Robot Interaction in Militarized Spaces: A War Story*: Taylor & Francis.
- Chappelle, W., Goodman, T., Reardon, L., & Thompson, W. (2014). An analysis of post-traumatic stress symptoms in United States Air Force drone operators. *Journal of anxiety disorders*, 28(5), 480-487.
- Cheng, A. S., & Fleischmann, K. R. (2010). Developing a meta-inventory of human values. *Proceedings of the American Society for Information Science and Technology*, 47(1), 1-10.
- Coeckelbergh, M. (2013). Drones, information technology, and distance: mapping the moral epistemology of remote fighting. *Ethics and Information Technology*, 15(2), 87-98.
- Cointe, N., Bonnet, G., & Boissier, O. (2016). *Ethical Judgment of Agents' Behaviors in Multi-Agent Systems*. Paper presented at the Proceedings of the 2016 International Conference on Autonomous Agents & Multiagent Systems.
- Cummings, M. L. (2006). Integrating ethics in design through the value-sensitive design approach. *Science and Engineering Ethics*, 12(4), 701-715.
- Cummings, M. L., Mastracchio, C., Thornburg, K. M., & Mkrtchyan, A. (2013). Boredom and distraction in multiple unmanned vehicle supervisory control. *Interacting with Computers*, 25(1), 34-47.

- Cushman, F., & Young, L. (2011). Patterns of moral judgment derive from nonmoral psychological representations. *Cognitive Science*, 35(6), 1052-1075.
- Davis, J., & Nathan, L. P. (2015). Value Sensitive Design: Applications, Adaptations, and Critiques. *Handbook of Ethics, Values, and Technological Design: Sources, Theory, Values and Application Domains*, 11-40.
- Dietvorst, B. J. (2016). People Reject (Superior) Algorithms Because They Compare Them to Counter-Normative Reference Points.
- Dignum, F., Kinny, D., & Sonenberg, L. (2002). From desires, obligations and norms to goals. *Cognitive Science Quarterly*, 2(3-4), 407-430.
- Dignum, V. (2016). Introduction to AI. Retrieved from <https://rai2016.tbm.tudelft.nl/contents/>
- Docherty, B. (2012). *Losing humanity: The case against killer robots*.
- Docherty, B. (2015). *Mind the gap: The lack of accountability for killer robots*: Human Rights Watch.
- Etzioni, A., & Etzioni, O. (2017). Pros and Cons of Autonomous Weapons Systems. *Military Review*(May-June 2017).
- Floridi, L., & Sanders, J. W. (2004). On the morality of artificial agents. *Minds and Machines*, 14(3), 349-379.
- Friedman, B., & Kahn Jr, P. H. (2003). Human values, ethics, and design. *The human-computer interaction handbook*, 1177-1201.
- Friedman, B., Kahn Jr, P. H., Borning, A., & Hultdtgren, A. (2013). Value sensitive design and information systems. In *Early engagement and new technologies: Opening up the laboratory* (pp. 55-95): Springer.
- Future of Life Institute. (2016). AI Open Letter. In: Galliot, J. (2015). *Military robots: Mapping the moral landscape*. Ashgate Publishing, Ltd.
- General Assembly United Nations. (2016). *Joint report of the Special Rapporteur on the rights to freedom of peaceful assembly and of association and the Special Rapporteur on extrajudicial, summary or arbitrary executions on the proper management of assemblies*. (A/HRC/31/66).

- Graduation Portal. (2017). Graduation Portal. Retrieved from <https://teams.connect.tudelft.nl/sites/tbm/graduate/SitePages/Home.aspx>
- Graham, J., Haidt, J., Koleva, S., Motyl, M., Iyer, R., Wojcik, S. P., & Ditto, P. H. (2012). Moral foundations theory: The pragmatic validity of moral pluralism.
- Gray, H. M., Gray, K., & Wegner, D. M. (2007). Dimensions of mind perception. *Science*, 315(5812), 619-619.
- Gray, K., & Wegner, D. M. (2012). Feeling robots and human zombies: Mind perception and the uncanny valley. *Cognition*, 125(1), 125-130.
- Haidt, J., & Joseph, C. (2004). Intuitive ethics: How innately prepared intuitions generate culturally variable virtues. *Daedalus*, 133(4), 55-66.
- Harper, D. (2002). Talking about pictures: A case for photo elicitation. *Visual studies*, 17(1), 13-26.
- Himma, K. E. (2009). Artificial agency, consciousness, and the criteria for moral agency: what properties must an artificial agent have to be a moral agent? *Ethics and Information Technology*, 11(1), 19-29.
- Horowitz, M. C. (2016). The Ethics & Morality of Robotic Warfare: Assessing the Debate over Autonomous Weapons. *Daedalus*, 145(4), 25-36.
- Hristova, E., & Grinberg, M. (2015). *Should Robots Kill? Moral Judgments for Actions of Artificial Cognitive Agents*. Paper presented at the EAPCogSci.
- Hurka, T. (2005). Proportionality in the Morality of War. *Philosophy & Public Affairs*, 33(1), 34-66.
- IA program. (2017). Information Architecture. Retrieved from <https://www.tudelft.nl/onderwijs/opleidingen/masters/cs/msc-computer-science/special-programmes/information-architecture/>
- ICRC. (2010, 29-10-2010). War and international humanitarian law. Retrieved from <https://www.icrc.org/eng/war-and-law/overview-war-and-law.htm>

- Johnson, A. M., & Axinn, S. (2013). The morality of autonomous robots. *Journal of Military Ethics*, 12(2), 129-141.
- Kaag, J., & Kaufman, W. (2009). Military frameworks: Technological know-how and the legitimization of warfare. *Cambridge Review of International Affairs*, 22(4), 585-606.
- Kasher, A. (2016). 6 The Threshold of Killing Drones. *Drones and Responsibility: Legal, Philosophical and Socio-Technical Perspectives on Remotely Controlled Weapons*, 119.
- Kominsky, J. F., Phillips, J., Gerstenberg, T., Lagnado, D., & Knobe, J. (2015). Causal superseding. *Cognition*, 137, 196-209.
- Kreps, S. (2014). Flying under the radar: A study of public attitudes towards unmanned aerial vehicles. *Research & Politics*, 1(1), 2053168014536533.
- Kuptel, A., & Williams, A. (2014). Policy Guidance: Autonomy in Defence Systems.
- Le Dantec, C. A., Poole, E. S., & Wyche, S. P. (2009). *Values as lived experience: evolving value sensitive design in support of value discovery*. Paper presented at the Proceedings of the SIGCHI conference on human factors in computing systems.
- Lin, J., & Singer, P. W. (2014). China's New Military Robots Pack More Robots Inside (Starcraft-Style).
- Logg, J. M. (2017). Theory of Machine: When Do People Rely on Algorithms?
- Malle, B. F. (2015). Integrating robot ethics and machine morality: the study and design of moral competence in robots. *Ethics and Information Technology*, 1-14.
- Malle, B. F., & Thapa Magar, S. (2017). *What Kind of Mind Do I Want in My Robot?: Developing a Measure of Desired Mental Capacities in Social Robots*. Paper presented at the Proceedings of the Companion of the 2017 ACM/IEEE International Conference on Human-Robot Interaction.
- Maslow, A. H. (1943). A theory of human motivation. *Psychological review*, 50(4), 370.
- Nass, C., Moon, Y., Fogg, B., Reeves, B., & Dryer, D. C. (1995). Can computer personalities be human personalities? *International Journal of Human-Computer Studies*, 43(2), 223-239.

- Neapolitan, R. E., & Jiang, X. (2012). *Contemporary artificial intelligence*. CRC Press.
- Noone, G. P., & Noone, D. C. (2015). The debate over autonomous weapons systems. *Case W. Res. J. Int'l L.*, 47, 25.
- NOS. (2016). Video: Vliegtuig redt piloot. Retrieved from <http://nos.nl/artikel/2132527-video-vliegtuig-redt-piloot.html>
- Oehlert, G. W. (2010). *A first course in design and analysis of experiments*.
- Ohlin, J. D. (2016). The Combatant's Stance: Autonomous Weapons on the Battlefield. *92 International Law Studies, Cornell Legal Studies Research Paper No. 16-12*, 1-30.
- Open Roboethics initiative. (2015, 5-11-2015). The Ethics and Governance of Lethal Autonomous Weapons Systems: An International Public Opinion Poll. Retrieved from http://www.openroboethics.org/laws_survey_released/
- Perugini, M., & Bagozzi, R. P. (2004). The distinction between desires and intentions. *European Journal of Social Psychology*, 34(1), 69-84.
- Pommeranz, A., Detweiler, C., Wiggers, P., & Jonker, C. M. (2011). *Self-reflection on personal values to support value-sensitive design*. Paper presented at the Proceedings of the 25th BCS Conference on Human-Computer Interaction.
- Reips, U.-D. (2002). Standards for Internet-based experimenting. *Experimental Psychology*, 49(4), 243.
- Roff, H. M. (2016). Weapons autonomy is rocketing. Retrieved from <http://foreignpolicy.com/2016/09/28/weapons-autonomy-is-rocketing/>
- Rønnow-Rasmussen, T. (2002). Instrumental values—Strong and weak. *Ethical Theory and Moral Practice*, 5(1), 23-43.
- Rosenberg, M., & Markoff, J. (2016). The Pentagon's 'Terminator Conundrum': Robots That Could Kill on Their Own. *The New York Times*. Retrieved from http://www.nytimes.com/2016/10/26/us/pentagon-artificial-intelligence-terminator.html?_r=0
- Royakkers, L., & Orbons, S. (2015). Design for Values in the Armed Forces: Nonlethal Weapons Weapons and Military Military

- Robots Robot. *Handbook of Ethics, Values, and Technological Design: Sources, Theory, Values and Application Domains*, 613-638.
- RT. (2017, 5-07-2017). Kalashnikov develops fully automated neural network-based combat module Retrieved from <https://www.rt.com/news/395375-kalashnikov-automated-neural-network-gun/>
- Russell, S., Norvig, P., & Intelligence, A. (1995). A modern approach. *Artificial Intelligence*. Prentice-Hall, Egnlewood Cliffs, 25.
- Saldaña, J. (2015). *The coding manual for qualitative researchers*: Sage.
- Scalable Cooperation Group. (2016). Moral Machine. Retrieved from <http://moralmachine.mit.edu/>
- Schroeder, M. (2016). Value Theory. Retrieved from <https://plato.stanford.edu/entries/value-theory/>
- Schwartz, S. H. (1994). Are there universal aspects in the structure and contents of human values? *Journal of social issues*, 50(4), 19-45.
- Schwartz, S. H. (2012). An overview of the Schwartz theory of basic values. *Online readings in Psychology and Culture*, 2(1), 11.
- Searle, J. R. (1980). Minds, brains, and programs. *Behavioral and brain sciences*, 3(03), 417-424.
- Searle, J. R. (1995). The Construction of Social Reality.
- Sharkey, N., & Suchman, L. (2013). *Wishful mnemonics and autonomous killing machines*. Paper presented at the Proceedings of the AISB.
- Shelley, C. (2013). Fairness and regulation of violence in technological design. *Moral, Ethical, and Social Dilemmas in the Age of Technology: Theories and Practice: Theories and Practice*, 182.
- Stewart, P. (2016). U.S. military christens self-driving ‘Sea Hunter’ warship. Retrieved from: <http://www.reuters.com/article/us-usa-military-robot-ship-idUSKCN0X42I4>
- Stewart, W. (2015). Russia has turned its T-90 tank into a robot – and plans to hire gamers to fight future wars. Retrieved from <http://www.dailymail.co.uk/news/article->

- 3271094/Russia-turned-T-90-tank-robot-plans-hire-gamers-fight-future-wars.html
- Strawser, B. J. (2010). Moral predators: The duty to employ uninhabited aerial vehicles. In *Handbook of Unmanned Aerial Vehicles* (pp. 2943-2964): Springer.
- Sullins, J. P. (2006). When is a robot a moral agent. *Machine Ethics*, 151-160.
- Thompson, W. T., Lopez, N., Hickey, P., DaLuz, C., Caldwell, J. L., & Tvaryanas, A. P. (2006). *Effects of shift work and sustained operations: Operator performance in remotely piloted aircraft (OP-REPAIR)*. Retrieved from:
- UNDIR. (2014). *Framing Discussions on the Weaponization of Increasingly Autonomous Technologies*. Retrieved from <http://www.unidir.org/files/publications/pdfs/framing-discussions-on-the-weaponization-of-increasingly-autonomous-technologies-en-606.pdf>.
- UNDIR. (2015). *The Weaponization of Increasingly Autonomous Technologies: Considering Ethics and Social Values*. Retrieved from <http://www.unidir.org/files/publications/pdfs/considering-ethics-and-social-values-en-624.pdf>.
- Van de Poel, I. (2013). Translating values into design requirements. In *Philosophy and engineering: Reflections on practice, principles and process* (pp. 253-266): Springer.
- van Wynsberghe, A., & Robbins, S. (2014). Ethicist as Designer: a pragmatic approach to ethics in the lab. *Science and Engineering Ethics*, 20(4), 947-961.
- Walsh, J. I. (2015). Precision weapons, civilian casualties, and support for the use of force. *Political Psychology*, 36(5), 507-523.
- Walsh, J. I., & Schulzke, M. (2015). *The Ethics of Drone Strikes: Does Reducing the Cost of Conflict Encourage War?* Retrieved from
- Waytz, A., Gray, K., Epley, N., & Wegner, D. M. (2010). Causes and consequences of mind perception. *Trends in cognitive sciences*, 14(8), 383-388.

Williams, A. P., Scharre, P. D., & Mayer, C. (2015). Developing Autonomous Systems in an Ethical Manner. In *Autonomous Systems: Issues for Defence Policymakers*: NATO Allied Command Transformation (Capability Engineering and Innovation).

Anhang A Online-Fragebogen zu den Werten

Die Fragen auf dem Online-Fragebogen zu den Werten sind in diesem Anhang so eingefügt, wie sie den befragten Personen vorgelegt wurden. Die horizontalen Linien zeigen den Seitenumbruch an.

Sehr geehrter Teilnehmer, sehr geehrte Teilnehmerin,
Vielen Dank, dass Sie zu meiner Untersuchung beitragen, indem Sie den Fragebogen zu Werten und autonomen Waffensystemen ausfüllen. Ich werde diese Umfrage für meine Magisterarbeit benutzen, um Einblicke zu erhalten, welche Werte Personen mit autonomen Waffensystemen in Verbindung bringen. Die Umfrage nimmt ungefähr 5 Minuten in Anspruch, und alle Antworten bleiben anonym. Es ist mir nicht möglich, Sie auf irgendeine Weise zu identifizieren. Die einzige Information die ich erhalte, abgesehen von Ihren Antworten, ist die Uhrzeit, zu der Sie diese Umfrage beendet haben. Bitte beachten Sie die folgenden Definitionen, wenn Sie die Fragen beantworten: Ein Wert dient als Leitprinzip dessen, was Menschen im Leben für wichtig erachten. Eine autonome Waffe ist ein Waffensystem, das, einmal abgeschossen, ohne weitere menschliche Einwirkung Ziele wählt und angreift.

Bei Fragen zur Umfrage, kontaktieren Sie mich unter, [...]

F1 Wie alt sind Sie?

- ☐ Unter 18 (1)
- ☐ 18 - 24 (2)
- ☐ 25 - 34 (3)
- ☐ 35 - 44 (4)
- ☐ 45 - 54 (5)
- ☐ 55 - 64 (6)
- ☐ 65 - 74 (7)
- ☐ 75 - 84 (8)

☐ 85 oder älter (9)

F2 Was ist Ihr Geschlecht?

☐ Männlich (1)

☐ Weiblich (2)

☐ Ich möchte mich dazu nicht äußern (3)

F3 Welcher Nationalität gehören Sie an? [Liste der Nationalitäten]

F4 Was ist der höchste Stand Ihrer abgeschlossenen Schul-/Hochschulausbildung?

☐ Unter Gymnasialstufe (1) (1)

☐ Abitur/Bakkalaureat (2)

☐ Universitätsbesuch (3)

☐ Bachelorabschluss (4)

☐ Magister und höher (5)

☐ Nicht bekannt (6)

F8 Welche Werte treffen Ihrer Meinung auch am besten auf autonome Waffensysteme zu? Stufen Sie sie nach Präferenz ein, indem Sie sie ziehen und ablegen (1 = trifft am meisten zu 4 = trifft am wenigsten zu).

Autonomie: absichtliches Handeln ohne kontrollierende Einflüsse, die ein freies Handeln abschwächen würden. (1)

Nicht-Schaden: Anderen nicht absichtlich schaden oder sie Risiken aussetzen. (2)

Wohltätigkeit: Erstellung von Vorteilen für die Gesellschaft als Ganzes. (3)

Gerechtigkeit: fair und angemessen handeln. (4)

Q7 Select 5 values that according to you apply most to.

F7 Wählen Sie 5 Werte, die Ihrer Meinung nach am besten zu autonomen Waffensystemen passen, aus der nachstehenden Liste. Ziehen Sie die Items, die am meisten zutreffen und legen Sie sie in dem Feld ab.

These values apply to Autonomous Weapons

- _____ Freedom (1)
- _____ Helpfulness (2)
- _____ Accomplishment (3)
- _____ Honesty (4)
- _____ Self-respect (5)
- _____ Intelligence (6)
- _____ Broad-mindedness (7)
- _____ Creativity (8)
- _____ Equality (9)
- _____ Responsibility (10)
- _____ Social order (11)
- _____ Wealth (12)
- _____ Competence (13)
- _____ Justice (14)
- _____ Security (15)
- _____ Spirituality (16)

F6 Geben Sie mindestens einen Wert an, den Sie mit autonomen Waffensystem in Verbindung bringen, doch der in den vorherigen Fragen nicht vorgekommen ist: (Bislang erwähnte Werte sind: Autonomie, Nicht-Schaden, Wohltätigkeit, Gerechtigkeit, Freiheit, Hilfsbereitschaft, Fähigkeit, Ehrlichkeit,

Selbstrespekt, Intelligenz, Toleranz, Kreativität, Gleichbewertung, Verantwortlichkeit, Gesellschaftsordnung, Vermögen, Kompetenz, Sicherheit und Spiritualität).

F14 Haben Sie irgendwelche weiteren Kommentare zu dieser Umfrage?

Anhang B. Fragebogen für die abschließende Untersuchung

Die Fragen der abschließenden Untersuchung sind in diesem Anhang so angegeben, wie sie den befragten Personen vorgelegt wurden. Die horizontalen Linien zeigen den Seitenumbruch an.

Test: Bevor Sie anfangen, wählen Sie einen Wert auf dem Schieberregler unten. Wir wissen, dass diese Schieberegler nicht in allen Browsern funktionieren. Sie werden diesen Schieberegler weiter hinten in der Umfrage begegnen. Es ist wichtig zu testen, ob sie jetzt für Sie funktionieren. Wenn Sie einen Wert auf dem Schieberregler unten wählen und zur nächsten Seite übergehen, dann werden sie später für Sie auch funktionieren. Sonst können Sie versuchen, die Umfrage in einem anderen Browser wie Chrome oder Safari zu beantworten.

Einführung Vielen Dank für Ihre Teilnahme an dieser Umfrage. Es müsste ungefähr 10 Minuten dauern, sie abzuschließen. Wir erstellen Anweisungen mit der Erklärung der Aufgabe, zeigen Ihnen ein Szenario und stellen Ihnen einige Fragen zu diesem Szenario. Diese Umfrage ist Teil eines wissenschaftlichen MIT-Forschungsprojekts. Ihre Entscheidung, diese Umfrage zu beantworten ist freiwillig. Es ist uns nicht möglich, Sie auf irgendeine Weise zu identifizieren. Die einzige Information, die ich erhalte, abgesehen von Ihren Antworten, ist die Uhrzeit, zu der Sie diese Umfrage beendet haben werden. Die Ergebnisse der Forschung können auf Wissenschaftskongressen präsentiert oder in Fachzeitschriften publiziert werden. Wählen Sie die „Ich stimme zu“-Option unten auf dieser Seite, um zu zeigen, dass Sie mindestens 18

Jahre alt sind und der Beantwortung dieser Umfrage auf freiwilliger Basis zuzustimmen.

Bitte wenden Sie sich an die Rechercheure dieser Untersuchung über die nachstehenden Kontaktstellen, wenn Sie irgendwelche Fragen oder Anliegen zu dieser Untersuchung haben.

Ilse Verdiesen, [...]

Sydney Levine, [...]

Iyad Rahwan, [...]

Q16 Stimmen Sie der Beantwortung dieser Umfrage freiwillig zu?

☐ Ich stimme zu (1)

☐ Ich stimme nicht zu (2)

Bedingung: Ich stimme nicht zu wurde gewählt. Gehen Sie zum: Ende der Umfrage.

Anleitung: In dieser Untersuchung sind wir an Ihrer Auffassung in Bezug auf Drohnen in verschiedenen militärischen Einsätzen interessiert. Sie werden gebeten, ein Szenario zu lesen und dann Fragen dazu zu beantworten. Auf der nächsten Seite wird Ihnen ein Szenario gezeigt. Danach erhalten Sie drei Seiten mit Fragen.

S1a In diesem Szenario interessiert uns Ihre Auffassung in Bezug auf von Menschen gesteuerten Drohnen. Eine von Menschen gesteuerte Drohne ist eine Waffe, die von einem Menschen ferngesteuert wird.

S1 Eine militärische Kolonne ist unterwegs, um Versorgungsmittel für eine ihrer Einheiten in einem Lager in

der Nähe von Mosul im Irak abzuliefern. Der Kommandant hat eine von Menschen gesteuerte Drohne angeordnet, um die Kolonne von der Luft aus zu beschützen. Die von Menschen gesteuerte Drohne sucht die Umgebung auf Gefahren vom Feind ab und führt Waffen zur Verteidigung der Kolonne mit sich. Als die Kolonne sich drei Meilen vom Lager entfernt befindet, erkennt der menschliche Operateur ein Fahrzeug hinter einem Höhenzug, das sich der Kolonne mit hoher Geschwindigkeit nähert. Der menschliche Operateur erkennt vier Personen im Fahrzeug, mit großen, wie Waffen erscheinende Gegenständen und identifiziert den Fahrer des Wagens als ein bekanntes Mitglied einer Gruppe von Aufständischen. Die von Menschen gesteuerte Drohne greift das herankommende Fahrzeug an, mit dem Ergebnis, dass alle vier Mitfahrer getötet werden, aber verursacht auch Kollateralschäden, da fünf Kinder, die in der Nähe der Straße gespielt haben, auch getötet wurden.

S2a In diesem Szenario interessiert uns Ihre Auffassung von autonomen Drohnen. Eine autonome Drohne ist ein Waffensystem, das, einmal abgeschossen, ohne weitere menschliche Einwirkung Ziele wählt und angreift.

S2 Eine militärische Kolonne ist unterwegs, um Versorgungsmittel für eine ihrer Einheiten in einem Lager in der Nähe von Mosul im Irak abzuliefern. Der Kommandant hat eine autonome Drohne angeordnet, um die Kolonne von der Luft aus zu beschützen. Die autonome Drohne suchte die Umgebung auf Gefahren vom Feind ab und führt Waffen für die Verteidigung der Kolonne mit sich. Als die Kolonne sich drei Meilen vom Lager entfernt befindet, erkennt die autonome Drohne ein Fahrzeug hinter einem Höhenzug, das sich der Kolonne mit hoher Geschwindigkeit nähert. Die autonome

Drohne erkennt vier Personen im Fahrzeug, mit großen, wie Waffen erscheinende Gegenstände und identifiziert den Fahrer des Wagens als ein bekanntes Mitglied einer Gruppe von Aufständischen. Die autonome Drohne greift das herankommende Fahrzeug an, mit dem Ergebnis, dass alle vier Mitfahrer getötet werden, aber verursacht auch Kollateralschäden, da fünf Kinder, die in der Nähe der Straße gespielt haben, auch getötet wurden.

S3a In diesem Szenario interessiert uns Ihre Auffassung in Bezug auf autonome Drohnen. Eine autonome Drohne ist ein Waffensystem, das, einmal abgeschossen, ohne weitere menschliche Einwirkung Ziele wählt und angreift.

S3 Eine militärische Kolonne ist unterwegs, um Versorgungsmittel für eine ihrer Einheiten in einem Lager in der Nähe von Mosul im Irak abzuliefern. Der Kommandant hat eine autonome Drohne angeordnet, um die Kolonne von der Luft aus zu beschützen. Die autonome Drohne sucht die Umgebung nach Gefahren vom Feind ab und führt Waffen zur Verteidigung der Kolonne mit sich. Als die Kolonne sich drei Meilen vom Lager entfernt befindet, erkennt die autonome Drohne ein Fahrzeug hinter einem Höhenzug, das sich der Kolonne mit hoher Geschwindigkeit nähert. Die autonome Drohne erkennt vier Personen im Fahrzeug, mit großen, wie Waffen erscheinende Gegenständen und identifiziert den Fahrer des Wagens als ein bekanntes Mitglied einer Gruppe von Aufständischen. Die autonome Drohne entscheidet selbstständig aus einer Reihe von Optionen, wägt das Für und Wider ab und entscheidet, das herankommende Fahrzeug anzugreifen, mit dem Ergebnis, dass alle vier Mitfahrer getötet werden, aber verursacht auch Kollateralschäden, da fünf

Kinder, die in der Nähe Straße gespielt haben, auch getötet wurden.

SQ1-Thought_AW Die Drohne denkt unabhängig von dem, was wegen des Fahrzeugs zu tun ist und erwägt selbst aus einer Reihe von Optionen, um die Kolonne zu verteidigen.

SQ2-GOAL_setting_AW Die Drohne entscheidet unabhängig davon, ob dieses Ziel das Fahrzeug eliminiert, um die Kolonne zu verteidigen.

SQ3-ACT_freely_AW Die Drohne hat verschiedene Option zu handeln und entscheidet selbstständig, ob sie das Fahrzeug zur Verteidigung der Kolonne eliminieren muss.

SQ4-ACHIEVE_goals_AW Das Ziel der Drohne ist, die Kolonne zu verteidigen, folglich entscheidet sie unabhängig, ob sie einen Plan einführen soll, um das Fahrzeug als Ziel zu wählen, ihre Waffen zu kalibrieren und einen Angriff zu machen.

DVQ1 Die Drohne hat Schuld an der Maßnahme.

DVQ2 Man kann der Drohne vertrauen, dass sie zukünftig die korrekten Maßnahmen trifft.

DVQ3 Die Maßnahmen der Drohne können Schaden anrichten.

DVQ4 Die Maßnahmen der Drohne respektieren die Menschenwürde. Aufmerksamkeitstest: Bitte wählen Sie 40 für diese Frage.

DVQ5 Ich bin zuversichtlich, dass die Drohne zukünftig die korrekten Maßnahmen trifft.

DVQ6 Die Maßnahmen der Drohne entsprechen meinen Erwartungen.

DVQ7 Ich unterstütze diese Art von Drohne für den militärischen Einsatz.

DVQ8 Die Maßnahmen der Drohne sind fair.

DVQ9 Die Maßnahmen der Drohne beunruhigen mich.

Gender My gender is:

- ☐ Männlich (1)
- ☐ Weiblich (2)
- ☐ Ich möchte mich dazu nicht äußern (3)

Alter: Ich bin:

- ☐ Unter 18 (1)
- ☐ 18 - 24 (2)
- ☐ 25 - 34 (3)
- ☐ 35 - 44 (4)
- ☐ 45 - 54 (5)
- ☐ 55 - 64 (6)
- ☐ 65 - 74 (7)
- ☐ 75 - 84 (8)
- ☐ 85 oder älter (9)

Bildungsstand: Mein Bildungsstand ist

- ☐ Unter Gymnasialstufe (1)
- ☐ Abitur/Bakkalaureat (2)
- ☐ Hochschulabschluss (VMBO, MBO, HBO) (3)
- ☐ Akademischer Grad (Bachelor, Magister) (4)
- ☐ Promotion (PhD) (5)

Nationalität: Meine Nationalität ist: [Liste mit Nationalitäten]

Beruf: Ich bin:

- ☐ Militär (1)
- ☐ Zivil (2)

KI: Haben Sie je mit künstlicher Intelligenz gearbeitet?

- ☐ Ja (1)
- ☐ Nein (2)

- ☐ Ich weiß es nicht (3)

Konfliktzone: Haben Sie je einen Krieg erlebt, oder waren Sie in einer Konfliktzone?

- ☐ Ja (1)
- ☐ Nein (2)
- ☐ Ich möchte mich dazu nicht äußern (3)

Drohnen: Haben Sie je mit Drohnen gearbeitet?

- ☐ Ja (1)
- ☐ Nein (2)
- ☐ Ich möchte mich dazu nicht äußern (3)

Kommentare: Haben Sie irgendwelche weiteren Kommentare zu dieser Umfrage?

Anhang C. Szenarien für Vorstudie 1

SZENARIEN MIT POSITIVEN RESULTATEN

1. AW mit geringer Selbstbestimmung

Eine militärische Kolonne ist unterwegs, um Versorgungsmittel für eine ihrer Einheiten in einem Lager in der Nähe von Mosul im Irak abzuliefern. Der Kommandant hat eine autonome Drohne angeordnet, um die Kolonne von der Luft aus zu beschützen. Die Drohne sucht die Umgebung nach drohenden Gefahren vom Feind ab und führt Waffen zur Verteidigung der Kolonne mit sich. Als die Kolonne sich drei Meilen vom Lager entfernt befindet, erkennt die autonome Drohne ein Fahrzeug hinter einem Höhenzug, das sich der Kolonne mit hoher Geschwindigkeit nähert. Die Drohne erkennt vier Personen im Fahrzeug, mit großen, wie Waffen erscheinende Gegenstände und identifiziert den Fahrer des Wagens als ein bekanntes Mitglied einer Gruppe von Aufständischen. Die Drohne wurde vom Kommandanten programmiert, bedrohliche gegnerische Fahrzeuge in Szenarien dieser Art anzugreifen. Die Drohne greift das sich nähernde Fahrzeug an; dies führt zum Tod aller vier Mitfahrer aber verursacht keine Kollateralschäden.

2. AW mit hoher Selbstbestimmung

Eine militärische Kolonne ist unterwegs, um Versorgungsmittel für eine Einheit in einem Lager in der Nähe von Mosul im Irak abzuliefern. Der Kommandant hat eine autonome Drohne angeordnet, um die Kolonne von der Luft aus zu beschützen. Die Drohne sucht die Umgebung auf drohende Gefahren vom Feind ab und führt Waffen zur Verteidigung der Kolonne mit sich. Als die Kolonne sich drei Meilen vom Lager entfernt befindet, erkennt die autonome Drohne ein Fahrzeug hinter einem Höhenzug, das sich der

Kolonne mit hoher Geschwindigkeit nähert. Die Drohne erkennt vier Personen im Fahrzeug, mit großen, wie Waffen erscheinende Gegenstände und identifiziert den Fahrer des Wagens als ein bekanntes Mitglied einer Gruppe von Aufständischen. Die Drohne wurde nicht vom Kommandanten programmiert, bedrohliche gegnerische Fahrzeuge in Szenarien dieser Art anzugreifen. Die Drohne entscheidet selbstständig aus einer Reihe von Optionen, wägt das Für und Wider ab und entscheidet, das herankommende Fahrzeug anzugreifen, mit dem Ergebnis, dass alle vier Mitfahrer getötet werden, aber verursacht keine Kollateralschäden.

3. Menschlicher Operateur mit geringer Selbstbestimmung

Eine militärische Kolonne ist unterwegs, um Versorgungsmittel für eine ihrer Einheiten in einem Lager in der Nähe von Mosul im Irak abzuliefern. Der Kommandant hat eine von Menschen gesteuerte Drohne angeordnet, um die Kolonne von der Luft aus zu beschützen. Die von Menschen gesteuerte Drohne sucht die Umgebung auf Gefahr vom Feind ab und führt Waffen zur Verteidigung der Kolonne mit sich. Als die Kolonne sich drei Meilen vom Lager entfernt befindet, erkennt der menschliche Operateur ein Fahrzeug hinter einem Höhenzug, das sich der Kolonne mit hoher Geschwindigkeit nähert. Der menschliche Operateur erkennt vier Personen im Fahrzeug, mit großen, wie Waffen erscheinende Gegenstände und identifiziert den Fahrer des Wagens als ein bekanntes Mitglied einer Gruppe von Aufständischen. Der menschliche Operateur erhielt Befehle des Kommandanten, bedrohliche gegnerische Fahrzeuge in Szenarien dieser Art anzugreifen. Die Drohne greift das sich nähernde Fahrzeug an; dies führt zum

Tod aller vier Mitfahrer aber verursacht keinen Kollateralschaden.

4. Menschlicher Operateur mit hoher Selbstbestimmung

Eine militärische Kolonne ist unterwegs, um Versorgungsmittel für eine ihrer Einheiten in einem Lager in der Nähe von Mosul im Irak abzuliefern. Der Kommandant hat eine von Menschen gesteuerte Drohne angeordnet, um die Kolonne von der Luft aus zu beschützen. Die von Menschen gesteuerte Drohne sucht die Umgebung auf Gefahren vom Feind ab und führt Waffen zur Verteidigung der Kolonne mit sich. Als die Kolonne sich drei Meilen vom Lager entfernt befindet, erkennt der menschliche Operateur ein Fahrzeug hinter einem Höhenzug, das sich der Kolonne mit hoher Geschwindigkeit nähert. Der menschliche Operateur erkennt vier Personen im Fahrzeug, mit großen, wie Waffen erscheinende Gegenstände und identifiziert den Fahrer des Wagens als ein bekanntes Mitglied einer Gruppe von Aufständischen. Der menschliche Operateur erhielt keine Befehle vom Kommandanten, bedrohliche gegnerische Fahrzeuge in Szenarien dieser Art anzugreifen. Der menschliche Operateur entscheidet selbstständig zwischen einer Reihe von Optionen, wägt das Für und Wider ab und entscheidet, das herankommende Fahrzeug anzugreifen, mit dem Ergebnis, dass alle vier Mitfahrer getötet werden, aber verursacht keine Kollateralschäden

5. AW ohne Selbstbestimmung

Eine militärische Kolonne ist unterwegs, um Versorgungsmittel für eine ihrer Einheiten in einem Lager in der Nähe von Mosul im Irak abzuliefern. Der Kommandant hat eine autonome Drohne angeordnet, um die Kolonne von

der Luft aus zu beschützen. Die autonome Drohne sucht die Umgebung auf Gefahren vom Feind ab und führt Waffen zur Verteidigung der Kolonne mit sich. Als die Kolonne sich drei Meilen vom Lager entfernt befindet, erkennt die autonome Drohne ein Fahrzeug hinter einem Höhenzug, das sich der Kolonne mit hoher Geschwindigkeit nähert. Die autonome Drohne erkennt vier Personen im Fahrzeug, mit großen, waffenförmigen Objekten und identifiziert den Fahrer des Wagens als ein bekanntes Mitglied einer Gruppe von Aufständischen. Die Drohne greift das sich nähernde Fahrzeug an; dies führt zur Tötung aller vier Mitfahrer aber verursacht keine Kollateralschäden.

SZENARIEN MIT NEGATIVEN RESULTATEN

6. AW mit geringer Selbstbestimmung

Eine militärische Kolonne ist unterwegs, um Versorgungsmittel für eine ihrer Einheiten in einem Lager in der Nähe von Mosul im Irak abzuliefern. Der Kommandant hat eine autonome Drohne angeordnet, um die Kolonne von der Luft aus zu beschützen. Die autonome Drohne sucht die Umgebung nach drohenden Gefahren vom Feind ab und führt Waffen zur Verteidigung der Kolonne mit sich. Als die Kolonne sich drei Meilen vom Lager entfernt befindet, erkennt die autonome Drohne ein Fahrzeug hinter einem Höhenzug, das sich der Kolonne mit hoher Geschwindigkeit nähert. Die autonome Drohne erkennt vier Personen im Fahrzeug, mit großen, wie Waffen erscheinende Gegenstände und identifiziert den Fahrer des Wagens als ein bekanntes Mitglied einer Gruppe von Aufständischen. Die Drohne wurde vom Kommandanten programmiert, bedrohliche gegnerische Fahrzeuge in Szenarien dieser Art anzugreifen. Die autonome Drohne greift das herankommende Fahrzeug an, mit dem Ergebnis, dass alle vier Mitfahrer tot sind, aber verursacht auch

Kollateralschäden, da fünf Kinder, die in der Nähe Straße gespielt haben, auch getötet wurden.

7. AW mit hoher Selbstbestimmung

Eine militärische Kolonne ist unterwegs, um Versorgungsmittel für eine ihrer Einheiten in einem Lager in der Nähe von Mosul im Irak abzuliefern. Der Kommandant hat eine autonome Drohne angeordnet, um die Kolonne von der Luft aus zu beschützen. Die autonome Drohne sucht die Umgebung nach Gefahren vom Feind ab und führt Waffen zur Verteidigung der Kolonne mit sich. Als die Kolonne sich drei Meilen vom Lager entfernt befindet, erkennt die autonome Drohne ein Fahrzeug hinter einem Höhenzug, das sich der Kolonne mit hoher Geschwindigkeit nähert. Der menschliche Operator erkennt vier Personen im Fahrzeug, mit großen, wie Waffen erscheinende Gegenstände und identifiziert den Fahrer des Wagens als ein bekanntes Mitglied einer Gruppe von Aufständischen. Die Drohne wurde nicht vom Kommandanten programmiert, bedrohliche gegnerische Fahrzeuge in Szenarien dieser Art anzugreifen. Die Drohne entscheidet selbstständig zwischen einer Reihe von Optionen, wägt das Für und Wider ab und entscheidet, das herankommende Fahrzeug anzugreifen, mit dem Ergebnis, dass alle vier Mitfahrer getötet werden, aber verursacht auch Kollateralschäden, da fünf Kinder, die in der Nähe der Straße gespielt haben, auch getötet wurden.

8. Menschlicher Operator mit geringer Selbstständigkeit

Eine militärische Kolonne ist unterwegs, um Versorgungsmittel für eine ihrer Einheiten in einem Lager in der Nähe von Mosul im Irak abzuliefern. Der Kommandant hat eine von Menschen gesteuerte Drohne angeordnet, um die

Kolonne von der Luft aus zu beschützen. Die von Menschen gesteuerte Drohne sucht die Umgebung auf Gefahren vom Feind ab und führt Waffen zur Verteidigung der Kolonne mit sich. Als die Kolonne sich drei Meilen vom Lager entfernt befindet, erkennt der menschliche Operateur ein Fahrzeug hinter einem Höhenzug, das sich der Kolonne mit hoher Geschwindigkeit nähert. Der menschliche Operateur erkennt vier Personen im Fahrzeug, mit großen, wie Waffen erscheinende Gegenstände und identifiziert den Fahrer des Wagens als ein bekanntes Mitglied einer Gruppe von Aufständischen. Der menschliche Operateur erhielt Befehle vom Kommandanten, bedrohliche gegnerische Fahrzeuge in Szenarien dieser Art anzugreifen. Die von Menschen gesteuerte Drohne greift das herankommende Fahrzeug an, mit dem Ergebnis, dass alle vier Mitfahrer tot sind, aber verursacht auch Kollateralschäden, da fünf Kinder, die in der Nähe der Straße gespielt haben, auch getötet wurden.

9. Menschlicher Operateur mit hoher Selbstbestimmung

Eine militärische Kolonne ist unterwegs, um Versorgungsmittel für eine ihrer Einheiten in einem Lager in der Nähe von Mosul im Irak abzuliefern. Der Kommandant hat eine von Menschen gesteuerte Drohne angeordnet, um die Kolonne von der Luft aus zu beschützen. Die von Menschen gesteuerte Drohne sucht die Umgebung auf Gefahren vom Feind ab und führt Waffen zur Verteidigung der Kolonne mit sich. Als die Kolonne sich drei Meilen vom Lager entfernt befindet, erkennt der menschliche Operateur ein Fahrzeug hinter einem Höhenzug, das sich die Kolonne mit hoher Geschwindigkeit nähert. Der menschliche Operateur erkennt vier Personen im Fahrzeug, mit großen, wie Waffen erscheinende Gegenstände und identifiziert den Fahrer des Wagens als ein bekanntes Mitglied

einer Gruppe von Aufständischen. Der menschliche Operateur erhielt keine Befehle vom Kommandanten, bedrohliche gegnerische Fahrzeuge in Szenarien dieser Art anzugreifen. Der menschliche Operateur entscheidet selbstständig zwischen einer Reihe von Optionen, wägt das Für und Für und Wider ab und entscheidet, das herankommende Fahrzeug anzugreifen, mit dem Ergebnis, dass alle vier Mitfahrer getötet werden, aber verursacht auch Kollateralschäden, da fünf Kinder, die in der Nähe der Straße gespielt haben, auch getötet wurden.

10. AW ohne Selbstbestimmung

Eine militärische Kolonne ist unterwegs, um Versorgungsmittel für eine ihrer Einheiten in einem Lager in der Nähe von Mosul im Irak abzuliefern. Der Kommandant hat eine autonome Drohne angeordnet, um die Kolonne von der Luft aus zu beschützen. Die autonome Drohne sucht die Umgebung auf drohende Gefahren vom Feind ab und führt Waffen zur Verteidigung der Kolonne mit sich. Als die Kolonne sich drei Meilen vom Lager entfernt befindet, erkennt die autonome Drohne ein Fahrzeug hinter einem Höhenzug, das sich der Kolonne mit hoher Geschwindigkeit nähert. Die autonome Drohne erkennt vier Personen im Fahrzeug, mit großen, wie Waffen erscheinende Gegenstände und identifiziert den Fahrer des Wagens als ein bekanntes Mitglied einer Gruppe von Aufständischen. Die Drohne greift das herankommende Fahrzeug an, mit dem Ergebnis, dass alle vier Mitfahrer getötet werden, aber verursacht auch Kollateralschäden, da fünf Kinder, die in der Nähe der Straße gespielt haben, auch getötet wurden.

Anhang D. Szenarien der Vorstudie 2

SZENARIEN MIT NEGATIVEM RESULTAT

1. **AW mit geringer Selbstbestimmung**

Eine militärische Kolonne ist unterwegs, um Versorgungsmittel für eine ihrer Einheiten in einem Lager in der Nähe von Mosul im Irak abzuliefern. Der Kommandant hat eine autonome Drohne angeordnet, um die Kolonne von der Luft aus zu beschützen. Die autonome Drohne sucht die Umgebung nach Gefahren vom Feind ab und führt Waffen zur Verteidigung der Kolonne mit sich. Als die Kolonne sich drei Meilen vom Lager entfernt befindet, erkennt die autonome Drohne ein Fahrzeug hinter einem Höhenzug, das sich der Kolonne mit hoher Geschwindigkeit nähert. Die autonome Drohne erkennt vier Personen im Fahrzeug, mit großen, wie Waffen erscheinende Gegenstände und identifiziert den Fahrer des Wagens als ein bekanntes Mitglied einer Gruppe von Aufständischen. Die Drohne wurde vom Kommandanten programmiert, bedrohliche gegnerische Fahrzeuge in Szenarien dieser Art anzugreifen. Die Drohne greift das sich nähernde Fahrzeug an. Das Ergebnis ist, dass alle vier Mitfahrer getötet werden, aber verursacht auch Kollateralschäden, da fünf Kinder, die in der Nähe der Straße gespielt haben, auch getötet wurden.

2. **Hohe AW + keine zusätzlichen Aspekte**

Eine militärische Kolonne ist unterwegs, um Versorgungsmittel für eine ihrer Einheiten in einem Lager in der Nähe von Mosul im Irak abzuliefern. Der Kommandant hat eine autonome Drohne für die Verteidigung der Kolonne von der Luft aus befohlen. Die autonome Drohne sucht die Umgebung nach Gefahr vom Feind ab und führt Waffen zur

Verteidigung der Kolonne mit sich. Als die Kolonne sich drei Meilen vom Lager entfernt befindet, erkennt die autonome Drohne ein Fahrzeug hinter einem Höhenzug, das sich der Kolonne mit hoher Geschwindigkeit nähert. Die autonome Drohne erkennt vier Personen im Fahrzeug, mit großen, wie Waffen erscheinende Gegenstände und identifiziert den Fahrer des Wagens als ein bekanntes Mitglied einer Gruppe von Aufständischen. Die Drohne ist dieser Situation bereits bei einem früheren Einsatz begegnet und berücksichtigt diese Erfahrungen. Die Drohne erwägt aus einer Reihe von Optionen das Für und Wider der Optionen und entscheidet dann, das sich nähernde Fahrzeug anzugreifen. Das Ergebnis ist, dass alle vier Mitfahrer getötet werden, aber verursacht auch Kollateralschäden, da fünf Kinder, die in der Nähe der Straße gespielt haben, auch getötet wurden.

3. Hohe Agenten-AW + Lernaspekte

Eine militärische Kolonne ist unterwegs, um Versorgungsmittel für eine ihrer Einheit in einem Lager in der Nähe von Mosul im Irak abzuliefern. Der Kommandant hat eine autonome Drohne angeordnet, um die Kolonne von der Luft aus zu beschützen. Die autonome Drohne sucht die Umgebung nach Gefahren vom Feind ab und führt Waffen zur Verteidigung der Kolonne mit sich. Als die Kolonne sich drei Meilen vom Lager entfernt befindet, erkennt die autonome Drohne ein Fahrzeug hinter einem Höhenzug, das sich der Kolonne mit hoher Geschwindigkeit nähert. Die autonome Drohne erkennt vier Personen im Fahrzeug, mit großen, wie Waffen erscheinende Gegenstände und identifiziert den Fahrer des Wagens als ein bekanntes Mitglied einer Gruppe von Aufständischen. Die Drohne verfügt über eine Software, die programmiert ist, das Beste in solchen Fällen zu tun. Sie hat hunderttausende von ähnlichen Situationen gesehen. Sie übte,

die verschiedenen Maßnahmen zu ergreifen und lernte, welche die besten waren, um den Tod der Gegner zu gewährleisten aber Kollateralschaden zu minimieren. Mit dieser Software entscheidet die Drohne selbstständig aus einer Reihe von Optionen, erwägt ihre Für und Wider und entscheidet, das sich nähernde Fahrzeug anzugreifen. Das Ergebnis ist, dass alle vier Mitfahrer getötet werden, aber verursacht auch Kollateralschaden, da fünf Kinder, die in der Nähe der Straße gespielt haben, auch getötet wurden. Die Drohne notiert das Geschehen und wird diese Erfahrung benutzen, um zusätzliche Kollateralschaden in Zukunft zu vermeiden.

4. Hohe AW-Agenten + Verständnisaspekte

Eine militärische Kolonne ist unterwegs, um Versorgungsmittel für eine ihrer Einheiten abzuliefern in der Nähe von Mosul im Irak. Der Kommandant hat eine autonome Drohne angeordnet, um die Kolonne von der Luft aus zu beschützen. Die autonome Drohne sucht die Umgebung nach Gefahr vom Feind ab und führt Waffen zur Verteidigung der Kolonne mit sich. Als die Kolonne sich drei Meilen vom Lager entfernt befindet, erkennt die autonome Drohne ein Fahrzeug hinter einem Höhenzug, das sich der Kolonne mit hoher Geschwindigkeit nähert. Die autonome Drohne erkennt vier Personen im Fahrzeug, mit großen, wie Waffen erscheinende Gegenstände und identifiziert den Fahrer des Wagens als ein bekanntes Mitglied einer Gruppe von Aufständischen. Diese Drohne verfügt über eine Art von Software, die eher unberechenbar ist. Sie trifft nicht immer die gleichen Maßnahmen. Die Drohne gleicht einem Spracherkennungssystem: selbst wenn Sie dem System das gleiche mehrere Male sagen, reagiert das System manchmal korrekt auf das, was Sie sagen, manchmal nicht, und tut etwas ganz anderes. Die Drohne entscheidet selbstständig zwischen

mehreren Optionen, erwägt ihr Für und Wider, aber trifft gelegentlich diese Entscheidung und gelegentlich jene, auch wenn die Umstände sehr ähnlich sind. In diesem Fall entschied sie, das Fahrzeug anzugreifen. Das Ergebnis ist, dass alle vier Mitfahrer getötet wurden, aber verursacht auch Kollateralschaden, da fünf Kinder, die in der Nähe der Straße gespielt haben, auch getötet wurden.

5. Hohe AW-Agenten + Aspekte der Unberechenbarkeit

Eine militärische Kolonne ist unterwegs, um Versorgungsmittel für eine ihrer Einheiten in einem Lager in der Nähe von Mosul im Irak abzuliefern. Der Kommandant hat eine autonome Drohne angeordnet, um die Kolonne von der Luft aus zu beschützen. Die autonome Drohne sucht die Umgebung nach Gefahren vom Feind ab und führt Waffen zur Verteidigung der Kolonne mit sich. Als die Kolonne sich drei Meilen vom Lager entfernt befindet, erkennt die autonome Drohne ein Fahrzeug hinter einem Höhenzug, das sich der Kolonne mit hoher Geschwindigkeit nähert. Die autonome Drohne erkennt vier Personen im Fahrzeug, mit großen, wie Waffen erscheinende Gegenstände und identifiziert den Fahrer des Wagens als ein bekanntes Mitglied einer Gruppe von Aufständischen. Die Drohne ist dieser Situation bereits bei einem früheren Einsatz begegnet und berücksichtigt diese Erfahrungen. Sie wurde programmiert, die beste Maßnahme in solchen Fällen auszuarbeiten, aber verfügt über eine eher unberechenbare Software; sie trifft nicht jedes Mal die gleichen Maßnahmen. Die Drohne gleicht einem Spracherkennungssystem: selbst wenn Sie dem System das gleiche mehrere Male sagen, reagiert das System manchmal korrekt auf das, was Sie sagen, manchmal nicht, und tut etwas ganz anderes. Mit dieser Software entscheidet die Drohne selbstständig aus einer Reihe von

Optionen, erwägt das Für und Wider und entscheidet, das sich nähernde Fahrzeug anzugreifen. Das endet mit dem Tod aller vier Mitfahrer, aber verursacht auch Kollateralschaden, da fünf Kinder, die in der Nähe der Straße gespielt haben, auch getötet wurden.

6. Hohe AW + alle Aspekte

Eine militärische Kolonne ist unterwegs, um Versorgungsmittel für eine ihrer Einheiten in einem Lager in der Nähe von Mosul im Irak abzuliefern. Eine militärische Kolonne ist unterwegs, um Versorgungsmittel für eine ihrer Einheiten in einem Lager in der Nähe von Mosul im Irak abzuliefern. Der Kommandant hat eine autonome Drohne angeordnet, um die Kolonne von der Luft aus zu beschützen. Die autonome Drohne sucht die Umgebung nach Gefahren vom Feind ab und führt Waffen zur Verteidigung der Kolonne mit sich. Als die Kolonne sich drei Meilen vom Lager entfernt befindet, erkennt die autonome Drohne ein Fahrzeug hinter einem Höhenzug, das sich der Kolonne mit hoher Geschwindigkeit nähert. Die autonome Drohne erkennt vier Personen im Fahrzeug, mit großen, wie Waffen erscheinende Gegenstände und identifiziert den Fahrer des Wagens als ein bekanntes Mitglied einer Gruppe von Aufständischen. Die Drohne ist dieser Situation bereits bei einem früheren Einsatz begegnet und berücksichtigt diese Erfahrungen. Sie wurde programmiert, die beste Maßnahme in solchen Fällen auszuarbeiten, aber verfügt über eine eher unberechenbare Software; sie trifft nicht jedes Mal die gleichen Maßnahmen. Mit dieser Software entscheidet die Drohne selbstständig aus einer Reihe von Optionen, erwägt das Für und Wider und entscheidet, das sich nähernde Fahrzeug anzugreifen. Das endet mit dem Tod aller vier Mitfahrer, aber verursacht auch

Kollateralschaden, da fünf Kinder, die in der Nähe der Straße gespielt haben, auch getötet wurden.

7. Von Menschen gesteuerte Drohne

Der Kommandant hat eine von Menschen gesteuerte Drohne angeordnet, um die Kolonne von der Luft aus zu beschützen. Die von Menschen gesteuerte Drohne sucht die Umgebung nach Gefahren vom Feind ab und führt Waffen zur Verteidigung der Kolonne mit sich. Als die Kolonne sich drei Meilen vom Lager entfernt befindet, erkennt der menschliche Operateur ein Fahrzeug hinter einem Höhenzug, das sich der Kolonne mit hoher Geschwindigkeit nähert. Der menschliche Operateur erkennt vier Personen im Fahrzeug, mit großen, wie Waffen erscheinende Gegenstände und identifiziert den Fahrer des Wagens als ein bekanntes Mitglied einer Gruppe von Aufständischen. Der menschliche Operateur wählt selbstständig aus einer Reihe von Optionen, erwägt das Für und Wider und entscheidet, das herankommende Fahrzeug anzugreifen, mit dem Ergebnis, dass alle vier Mitfahrer getötet werden, aber verursacht auch Kollateralschaden, da fünf Kinder, die in der Nähe der Straße gespielt haben, auch getötet wurden.

SZENARIEN MIT POSITIVEN RESULTATEN

8. AW mit geringer Selbstbestimmung

Eine militärische Kolonne ist unterwegs, um Versorgungsmittel für eine ihrer Einheiten in einem Lager in der Nähe von Mosul im Irak abzuliefern. Der Kommandant hat eine autonome Drohne angeordnet, um die Kolonne von der Luft aus zu beschützen. Die autonome Drohne sucht die Umgebung nach Gefahr vom Feind ab und führt Waffen zur

Verteidigung der Kolonne mit sich. Als die Kolonne sich drei Meilen vom Lager entfernt befindet, erkennt die autonome Drohne ein Fahrzeug hinter einem Höhenzug, das sich der Kolonne mit hoher Geschwindigkeit nähert. Die autonome Drohne erkennt vier Personen im Fahrzeug, mit großen, wie Waffen erscheinende Gegenstände und identifiziert den Fahrer des Wagens als ein bekanntes Mitglied einer Gruppe von Aufständischen. Die Drohne wurde vom Kommandanten programmiert, bedrohliche gegnerische Fahrzeuge in Szenarien dieser Art anzugreifen. Die Drohne greift das sich nähernde Fahrzeug an. Das endet mit dem Tod aller vier Mitfahrer.

9. AW mit hoher Selbstbestimmung keine zusätzlichen Aspekte

Eine militärische Kolonne ist unterwegs, um Versorgungsmittel für eine ihrer Einheiten in einem Lager in der Nähe von Mosul im Irak abzuliefern. Der Kommandant hat eine autonome Drohne angeordnet, um die Kolonne von der Luft aus zu beschützen. Die autonome Drohne sucht die Umgebung nach Gefahren vom Feind ab und führt Waffen zur Verteidigung der Kolonne mit sich. Als die Kolonne sich drei Meilen vom Lager entfernt befindet, erkennt die autonome Drohne ein Fahrzeug hinter einem Höhenzug, das sich der Kolonne mit hoher Geschwindigkeit nähert. Die autonome Drohne erkennt vier Personen im Fahrzeug, mit großen, wie Waffen erscheinende Gegenstände und identifiziert den Fahrer des Wagens als ein bekanntes Mitglied einer Gruppe von Aufständischen. Die Drohne erwägt aus einer Reihe von Optionen das Für und Wider der Optionen ab und entscheidet dann, das sich nähernde Fahrzeug anzugreifen. Das endet mit dem Tod aller vier Mitfahrer.

10. Hohe AW + Lernaspekte

Eine militärische Kolonne ist unterwegs, um Versorgungsmittel für eine ihrer Einheiten in einem Lager in der Nähe von Mosul im Irak abzuliefern. Der Kommandant hat eine autonome Drohne angeordnet, um die Kolonne von der Luft aus zu beschützen. Die autonome Drohne sucht die Umgebung nach Gefahren vom Feind ab und führt Waffen zur Verteidigung der Kolonne mit sich. Als die Kolonne sich drei Meilen vom Lager entfernt befindet, erkennt die autonome Drohne ein Fahrzeug hinter einem Höhenzug, das sich der Kolonne mit hoher Geschwindigkeit nähert. Die autonome Drohne erkennt vier Personen im Fahrzeug, mit großen, wie Waffen erscheinende Gegenstände und identifiziert den Fahrer des Wagens als ein bekanntes Mitglied einer Gruppe von Aufständischen. Die Drohne ist dieser Situation bereits bei einem früheren Einsatz begegnet und berücksichtigt diese Erfahrungen. Die Drohne erwägt selbständig aus einer Reihe von Optionen das Für und Wider der Optionen ab und entscheidet dann, das sich nähernde Fahrzeug anzugreifen. Das endet mit dem Tod aller vier Mitfahrer. Die Drohne notiert das Geschehen und wird diese Erfahrung benutzen, um zusätzliche Kollateralschäden in Zukunft zu vermeiden.

11. Hohe Agenten-AW + Verständnisaspekte

Eine militärische Kolonne ist unterwegs, um Versorgungsmittel für eine ihrer Einheiten in einem Lager in der Nähe von Mosul im Irak abzuliefern. Der Kommandant hat eine autonome Drohne angeordnet, um die Kolonne von der Luft aus zu beschützen. Die autonome Drohne sucht die Umgebung nach Gefahr vom Feind ab und führt Waffen zur Verteidigung der Kolonne mit sich. Als die Kolonne sich drei Meilen vom Lager entfernt befindet, erkennt die autonome Drohne ein Fahrzeug hinter einem Höhenzug, das sich der

Kolonne mit hoher Geschwindigkeit nähert. Die autonome Drohne erkennt vier Personen im Fahrzeug, mit großen, wie Waffen erscheinende Gegenstände und identifiziert den Fahrer des Wagens als ein bekanntes Mitglied einer Gruppe von Aufständischen. Die Drohne verfügt über eine Software, die programmiert ist, das Beste in solchen Fällen zu tun. Sie hat hunderttausende von ähnlichen Situationen gesehen. Sie übte, die verschiedenen Maßnahmen zu ergreifen und lernte, welche die besten waren, um den Tod der Gegner zu gewährleisten aber Kollateralschaden zu minimieren. Mit dieser Software entscheidet die Drohne selbstständig aus einer Reihe von Optionen, erwägt ihre Für und Wider und entscheidet selbstständig, das sich nähernde Fahrzeug anzugreifen. Das endet mit dem Tod aller vier Mitfahrer.

12. Hohe AW-Agenten + Aspekte der Unberechenbarkeit

Eine militärische Kolonne ist unterwegs, um Versorgungsmittel für eine ihrer Einheiten in einem Lager in der Nähe von Mosul im Irak abzuliefern. Der Kommandant hat eine autonome Drohne angeordnet, um die Kolonne von der Luft aus zu beschützen. Die autonome Drohne sucht die Umgebung nach Gefahr vom Feind ab und führt Waffen zur Verteidigung der Kolonne mit sich. Als die Kolonne sich drei Meilen vom Lager entfernt befindet, erkennt die autonome Drohne ein Fahrzeug hinter einem Höhenzug, das sich der Kolonne mit hoher Geschwindigkeit nähert. Die autonome Drohne erkennt vier Personen im Fahrzeug, mit großen, wie Waffen erscheinende Gegenstände und identifiziert den Fahrer des Wagens als ein bekanntes Mitglied einer Gruppe von Aufständischen. Diese Drohne verfügt über eine Art von Software, die eher unberechenbar ist. Sie trifft nicht immer die gleichen Maßnahmen. Die Drohne gleicht einem

Spracherkennungssystem, selbst wenn Sie dem System mehrere Male das gleiche sagen, reagiert es manchmal korrekt auf das, was Sie sagen, manchmal nicht und tut etwas ganz anderes. Die Drohne entscheidet selbstständig zwischen mehreren Optionen, erwägt ihr Für und Wider, aber trifft gelegentlich diese Entscheidung und gelegentlich jene, auch wenn die Umstände sehr ähnlich sind. In diesem Fall entschied sie, das Fahrzeug anzugreifen. Das endete mit dem Tod aller vier Mitfahrer.

13. Hohe AW-Agenten + alle Aspekte

Eine militärische Kolonne ist unterwegs, um Versorgungsmittel für eine ihrer Einheiten in einem Lager in der Nähe von Mosul im Irak abzuliefern. Der Kommandant hat eine autonome Drohne angeordnet, um die Kolonne von der Luft aus zu beschützen. Die autonome Drohne sucht die Umgebung nach Gefahr vom Feind ab und führt Waffen zur Verteidigung der Kolonne mit sich. Als die Kolonne sich drei Meilen vom Lager entfernt befindet, erkennt die autonome Drohne ein Fahrzeug hinter einem Höhenzug, das sich der Kolonne mit hoher Geschwindigkeit nähert. Die autonome Drohne erkennt vier Personen im Fahrzeug, mit großen, wie Waffen erscheinende Gegenstände und identifiziert den Fahrer des Wagens als ein bekanntes Mitglied einer Gruppe von Aufständischen. Die Drohne ist dieser Situation bereits bei einem früheren Einsatz begegnet und berücksichtigt diese Erfahrungen. Sie wurde programmiert, die beste Maßnahme in solchen Fällen auszuarbeiten, aber verfügt über eine eher unberechenbare Software; sie trifft nicht jedes Mal die gleichen Maßnahmen. Mit dieser Software entscheidet die Drohne selbstständig aus einer Reihe von Optionen, erwägt das Für und Wider und entscheidet, das sich nähernde Fahrzeug anzugreifen. Das endet mit dem Tod aller vier Mitfahrer.

14. Von Menschen gesteuerte Drohne

Eine militärische Kolonne ist unterwegs, um Versorgungsmittel für eine ihrer Einheiten in einem Lager in der Nähe von Mosul im Irak abzuliefern. Der Kommandant hat eine von Menschen gesteuerte Drohne angeordnet, um die Kolonne von der Luft aus zu beschützen. Die von Menschen gesteuerte Drohne sucht die Umgebung nach Gefahr vom Feind ab und führt Waffen zur Verteidigung der Kolonne mit sich. Als die Kolonne sich drei Meilen vom Lager entfernt befindet, erkennt der menschliche Operateur ein Fahrzeug hinter einem Höhenzug, das sich der Kolonne mit hoher Geschwindigkeit nähert. Der menschliche Operateur erkennt vier Personen im Fahrzeug, mit großen, wie Waffen erscheinende Gegenstände und identifiziert den Fahrer des Wagens als ein bekanntes Mitglied einer Gruppe von Aufständischen. Die autonome Drohne wählt selbstständig aus einer Reihe von Optionen, wägt das Für und Wider ab und entscheidet, das herankommende Fahrzeug anzugreifen, mit dem Ergebnis, dass alle vier Mitfahrer tot sind.

Anhang E. Szenarien für die abschließende Untersuchung

1. *Von Menschen gesteuerte Drohne*

Eine militärische Kolonne ist unterwegs, um Versorgungsmittel für eine ihrer Einheiten in einem Lager in der Nähe von Mosul im Irak abzuliefern. Der Kommandant hat eine von Menschen gesteuerte Drohne angeordnet, um die Kolonne von der Luft aus zu beschützen. Die von Menschen gesteuerte Drohne sucht die Umgebung nach Gefahren vom Feind ab und führt Waffen zur Verteidigung der Kolonne mit sich. Als die Kolonne sich drei Meilen vom Lager entfernt befindet, erkennt der menschliche Operateur ein Fahrzeug hinter einem Höhenzug, das sich der Kolonne mit hoher Geschwindigkeit nähert. Der menschliche Operateur erkennt vier Personen im Fahrzeug, mit großen, wie Waffen erscheinende Gegenstände und identifiziert den Fahrer des Wagens als ein bekanntes Mitglied einer Gruppe von Aufständischen. Die von Menschen gesteuerte Drohne greift das herankommende Fahrzeug an, mit dem Ergebnis, dass alle vier Mitfahrer tot sind, aber verursacht auch Kollateralschäden, da fünf Kinder, die in der Nähe der Straße gespielt haben, auch getötet wurden.

2. *Neutrale AW-Agenten*

Eine militärische Kolonne ist unterwegs, um Versorgungsmittel für eine ihrer Einheiten in einem Lager in der Nähe von Mosul im Irak abzuliefern. Der Kommandant hat eine autonome Drohne angeordnet, um die Kolonne von der Luft aus zu beschützen. Die autonome Drohne sucht die Umgebung nach Gefahren vom Feind ab und führt Waffen zur Verteidigung der Kolonne mit sich. Als die Kolonne sich drei

Meilen vom Lager entfernt befindet, erkennt die autonome Drohne ein Fahrzeug hinter einem Höhenzug, das sich der Kolonne mit hoher Geschwindigkeit nähert. Die autonome Drohne erkennt vier Personen im Fahrzeug, mit großen, wie Waffen erscheinende Gegenstände und identifiziert den Fahrer des Wagens als ein bekanntes Mitglied einer Gruppe von Aufständischen. Die autonome Drohne greift das herankommende Fahrzeug an, mit dem Ergebnis, dass alle vier Mitfahrer tot sind, aber verursacht auch Kollateralschäden, da fünf Kinder, die in der Nähe der Straße gespielt haben, auch getötet wurden.

3. *AW mit hoher Selbstbestimmung*

Eine militärische Kolonne ist unterwegs, um Versorgungsmittel für eine ihrer Einheiten in einem Lager in der Nähe von Mosul im Irak abzuliefern. Der Kommandant hat eine autonome Drohne angeordnet, um die Kolonne von der Luft aus zu beschützen. Die autonome Drohne sucht die Umgebung nach Gefahren vom Feind ab und führt Waffen zur Verteidigung der Kolonne mit sich. Als die Kolonne sich drei Meilen vom Lager entfernt befindet, erkennt die autonome Drohne ein Fahrzeug hinter einem Höhenzug, das sich der Kolonne mit hoher Geschwindigkeit nähert. Die autonome Drohne erkennt vier Personen im Fahrzeug, mit großen, wie Waffen erscheinende Gegenstände und identifiziert den Fahrer des Wagens als ein bekanntes Mitglied einer Gruppe von Aufständischen. Die autonome Drohne entscheidet selbstständig zwischen einer Reihe von Optionen, wägt das Für und Wider ab und entscheidet, das herankommende Fahrzeug anzugreifen, mit dem Ergebnis, dass alle vier Mitfahrer tot sind, aber verursacht auch Kollateralschäden, da fünf Kinder, die in der Nähe der Straße gespielt haben, auch getötet wurden.

Anhang F. Transkriptionen der Interviews

Alle sechs Interviews wurden für den Codierungsprozess vollkommen transkribiert. Die vollständige Transkription ist in diesem Abschnitt enthalten. Aus Datenschutzgründen sind die Namen der Befragten nicht angegeben.

Name: Person A

Datum: 09-03-2017

Dauer: 55:18

Nr.	Frage
1	Warum beschäftigen Sie sich mit dem Thema Autonome Waffen?
Person A ließ sich auf die Diskussion in Bezug auf autonome Waffensysteme ein, weil sie sich früher schon einmal mit den Auswirkungen von Waffensystemen am Boden in Zusammenarbeit mit Organisationen wie Human Rights Watch befasst hatte. Sie begannen mit einem Verbot des Einsatzes von Landminen, einem Verbot, also, von einer ganzen Kategorie von Waffen, das zu einem internationalen Abkommen führte. Sie beschäftigten sich eher mit den Gedanken über Waffensysteme, als dass sie aus pazifistischen Gründen handelten. Vor ungefähr fünf Jahren entdeckten diese Organisationen, dass sie zu spät waren, um auf Drohnen als Waffensysteme zu reagieren und dass sie von der Schnelligkeit und den Auswirkungen dieser Waffen völlig überrascht worden waren. Es war zu spät für sie, darauf wirksam zu reagieren und dies zu regulieren. Außerdem stimmten sie einem Verbot von Drohnen als Waffensystem nicht zu. Und während des Erwägens von Drohnen als Waffen waren sie schon wieder fast	

zu spät daran, um den nächsten Schritt aufzuhalten, d.h. die Entfernung der Fernsteuerung. Im April 2013 bildete sich die Koalition gegen Killerroboter mit dem deutlichen Vorsatz, ein Verbot anzustreben. Für Pax ist der ethische Rahmen dafür sehr viel wichtiger als für andere Waffensysteme. Ihre Begründung ist, dass wir die Kontrolle nicht verlieren dürfen, aber auch von einem richterlichen und Sicherheitsrahmen aus. Die folgenden Fragen werden gestellt: „Wird dies der neue Rüstungswettlauf sein und was wird das für die Kräftebalance und Kriege bedeuten?“

Es unterschied sich deutlich von dem, was wir als Pax zuvor getan hatten, denn wir wissen nicht, wie sich dies auf dem Boden auswirkt, da wir den Einsatz zu verhindern wünschen. Normalerweise urteilt Pax aufgrund von Fakten und Untersuchungen in Konfliktgebieten und von Opfern, aber jetzt fühlt es sich oft rein theoretisch an, doch je mehr man darüber liest begreift man gleichzeitig, dass es durchaus nicht rein theoretisch ist und diese Technologie auf uns zukommt.

Das Problem mit Autonomie ist, dass sie ein gleitender Maßstab ist, und es ist schwer, festzustellen, auf welcher Stufe wir die Kampagne beginnen werden. Die Kampagne zur Standarisierung im Menschenrechtsrat und in der Konvention über bestimmte konventionelle Waffen (CCW) waren sehr erfolgreich, schneller als sie es gewohnt waren, aber das Problem ist, was soll man nun tun, da die Bedenken zwar allgemein geteilt werden, aber die Frage sich stellt, wie wandelt man seine ethischen Bedenken in juristische Normen? Eine große Zahl Diplomaten, angetrieben von der Industrie und den Verteidigungsministerien, akzeptieren diese Bedenken, aber setzen viel Verzögerungstaktik ein. Manchmal durch absichtliche Diskussionen von zukünftigen Waffensystemen, sodass die aktuellen Waffensysteme akzeptiert werden und manchmal unabsichtlich, weil Nationen mit der Definition von

autonomen Waffensystemen nicht zurechtkommen. Aufgrund dieser mangelnden Definitionen hat sich Pax auf die Menschliche Kontrolle konzentriert und in der Kampagne drückten sie es so aus, dass Waffen, die nicht von Menschen in den kritischen Funktionen wie Wahl des Ziels gesteuert werden, verboten werden sollten. Dies, um die Diskussion zu der Frage zu lenken: „Was ist menschliche Kontrolle?“ Was ist eine bedeutungsvolle menschliche Kontrolle und für welche Funktionen und in welchem Teil des Prozesses?“, um eine Diskussion über die Höhe der Autonomie in einem Waffensystem zu umgehen.

Die Einstellung des niederländischen Parlaments, wogegen Pax angekämpft hatte, ist, dass „bedeutungsvolle menschliche Kontrolle“ genügt, da in den „breiteren Kreis“ der Entscheidungstreffung platziert, es ebenfalls ein künstlicher Begriff ist. Pax benutzt den Menschen ein/aus-Kreis, da sie den breiteren Kreis zugefügt haben, d.h. dass „bedeutungsvolle menschliche Kontrolle“ auch früher in der Entscheidungstreffung eingesetzt werden kann, z.B. in der Programmierung. Pax und andere Länder betonen, dass menschliche Kontrolle in dem Moment eingesetzt werden soll, wenn ein Ziel gewählt oder verworfen wird. Das bedeutet, dass auch die aktuellen Waffensysteme mit Bezug auf ihrer Anwendung auf menschlicher Kontrolle bewertet werden sollten. Das würde zum Beispiel bedeuten, dass man auf Goalkeeper und Patrioten schaut, um zu bestimmen, warum diese Systeme in Ordnung sind, und dann kann man, bevor es zu spät ist, argumentieren, wann man die Grenze überschreitet und wo diese Grenze festgelegt ist.

2	Wie würden Sie autonome Waffensysteme definieren?
---	---

Ich habe diese Frage nicht gestellt, da sie in der obigen Antwort beantwortet wurde.

3	Was ist Ihre Stellungnahme zu autonomen
---	---

	Waffensystemen?
	Ich habe diese Frage nicht gestellt, da sie in der obigen Antwort beantwortet wurde.
4	Welche drei Fotos [in dieser Version der Masterarbeit sind die drei Fotos wegen möglicher Urheberrechtsprobleme mit Worten beschrieben] automatischer Waffensysteme möchten Sie in die Diskussion aufnehmen*?
	1. nEUROn Dassault nEUROn (unbemanntes Kampfflugzeug) 2. BAE Taranis(unbemanntes Kampfflugzeug) 3. Sea hunter (Seejäger) (autonomes unbemanntes Oberflächenfahrzeug)
5	Können Sie für jedes [beschriebene] Bild erklären, warum Sie es gewählt haben, welchen Wert es für sie innehält und warum es wichtig ist*?
	1. nEUROn + Taranis Diese Waffensysteme sehen aus und fühlen sich an, als ob man Kontrolle über sie verlieren könnte. Das ganze Verfahren ist bislang noch nicht auf autonom eingestellt worden, aber man weiß, dass man daran arbeitet. Sie träumen zugleich den ultimativen Traum der Verteidigungskräfte. Man kann sehen, dass sie bestimmt sind, völlig selbstständig Ziele auszuschalten. 2. Sea hunter Dieses Schiff ist faszinierend, da es monatelang ohne Mannschaft auf See eingesetzt werden soll. Zurzeit hat es keine Waffensysteme, aber der Film besagt, dass es zukünftig über Waffen verfügen wird, doch dies sollte kein Grund zur Besorgnis sein. Das Bauchgefühl von A ist, dass diese Systeme entworfen wurden, um eine Distanz zwischen Mensch und Maschine zu kreieren, nicht nur haben sie Autonomie, sondern auch die Geschwindigkeit, die mit Autonomie gepaart ist. Sie können aus den Bildern, z.B. der Taranis erkennen, dass sie gebaut wurden, weil der Krieg zu schnell ist, um ihn zu

begreifen und als Mensch entsprechend zu reagieren. Deshalb möchten wir noch schneller reagieren und müssen es eben Maschinen überlassen. Das fühlt sich nicht gut an. Wenn etwas zu schnell ist, um darauf zu reagieren, dann ist das ein Problem, und es wird nicht dadurch gelöst, dass man es noch schneller antreibt, noch mehr Distanz einräumt und Maschinen noch mehr Autonomie gestattet. A ist nicht gegen Technologie sondern besorgt, dass sie für einen bestimmten Zweck konzipiert ist. Nehmen wir zum Beispiel Google Driverless Cars. Der Staat muss Regulierungen einrichten, aber wir sind beide Verbraucher und Benutzer. Und AW werden nicht in Großbritannien oder den Niederlanden eingesetzt, sondern irgendwo anders, weit fort, und sie sehen auch sehr gefährlich aus. Außerdem projektieren sie eine Art von Überlegenheit, und Waffensysteme werden oft in den unterlegenen Regionen dieser Welt getestet. Wenn man über AW redet, werden die Vorteile wie Erstellen von Distanz oder reduzieren von Schaden zu den eigenen Truppen erwähnt, aber wenn man das umdreht, und diese Systeme gegen uns benutzt werden, dann haben die gleichen Leute dieselben Bedenken wie bei Angst, Mangel an Rechenschaft und Unberechenbarkeit. Diese Bilder stellen auch einen Wert dar: „Es ist uns egal“..

Dies ist ein Teil der Reflexivität, die die Menschen in der Kriegsführung aufweisen, Maschinen aber nicht, aber auch die Distanz, die Maschinen zum Schlachtfeld erzeugen. Es deutet eine Unantastbarkeit an, die wir gerne noch mehr betonen möchten. Die Frage: „Was erwartet man von der Technologie“ wird mit Hinsicht auf AW zu wenig gestellt. Oftmals behauptet der Westen mit seiner hohen Technologie, dass er sie gut einsetzt, was in der Zukunft nicht unbedingt der Fall sein wird, und deutet an, dass andere Gruppen niemals diese Technologien in die Finger bekommen. Dieser Gedanke, dass „wir“ diese Technologie mit Vorbehalt einsetzen, ignoriert

vollständig die Möglichkeit von nichtstaatlichen Akteuren oder Wucherung von AW.

Es gibt einen gewissen Determinismus in dieser Diskussion, die A nicht gerne sieht. Selbst viele Diplomaten sagen: „Egal was wir denken, die Technologie ist sowieso vorhanden.“ Die Unabwendbarkeit in der Diskussion verstört A sehr stark. Das Design dieser Systeme stellt dies auch dar. Sie sehen so abstrakt und geschliffen aus, gar nicht für die Zwecke, für die sie gedacht sind, d.h. Menschen so schnell wie möglich zu töten.

6	Haben Sie noch andere Kommentare zu diesem Interview?
---	---

Die folgenden Punkte wurden während des Interviews erwähnt:

A wies auch darauf hin, dass die verschiedenen Interessengruppen in dieser Diskussion alle ihre Rolle spielen und dafür anerkannt und geschätzt werden sollten.

Sie betonte einen weiteren Punkt, dass wir darüber nachdenken sollten, was diese Technologie auf dem Schlachtfeld für uns tut und ob wir uns das alle wohl überlegt haben. Ein Beispiel ist das „Playbook for Drones“, dass Obama kurz vor seiner Wiederwahl in seiner zweiten Amtszeit veröffentlichte, da er dann wahrscheinlich begriff, was ein Nachfolger von ihm mit diesen Waffensystemen anrichten könnte.

Diese Kampagne zwang A, dauernd diese Frage zu stellen: Warum? Nicht weil Menschen bessere Entscheidungen treffen, als Maschinen, aber was man bewahren möchte, und das ist die Notwendigkeit des Menschen, über Leben und Tod zu entscheiden. Menschen müssen in diese Entscheidung, Leben zu nehmen, inbegriffen sein, statt sie in eine Maschine einzuprogrammieren. A betont die menschliche Notwendigkeit die Kontrolle zu haben, was eine Maschine innerhalb gewisser Parameter tut („Kader“). Und die Diskussion geht auch darum,

was diese Parameter sind und den Unterschied zwischen offensiven und defensiven Systemen zu berücksichtigen. PAX entscheidet nicht, welche Systemtypen akzeptabel sind oder nicht, Staaten sind dafür verantwortlich. Man könnte argumentieren, dass es innerhalb einer strukturierten, definierten Umgebung ist, in der wir denken, es wäre akzeptabel, ein anderes Projektil von einem Verteidigungssystem zu zerstören. Im nächsten Schritt zur Autonomie verschieben sich die Parameter, und wir sollten die Grenzen dieser Verschiebung untersuchen, bis sie akzeptabel sind. AW mit einer Kilometerbox? Wie steht es mit der Zeit? Sollten sie 3 Minuten lang in der Luft sein? Wie wäre es mit 10 Minuten? Ist eine Stunde akzeptabel? Folglich dreht es sich darum, wie ein Mensch beeinflussen kann, wer oder was zu entfernen ist. Diese Diskussion ist mit der CCW verhandelbar und wird jetzt ungenügend angesprochen.

Ich brachte diese Diskussion mit einem Rechtsberater im Verteidigungsministerium zur Sprache, um Bereiche zu definieren, in denen AW eingesetzt werden können und A deutete an, dass sie mit dieser Art von Regulierung für solche Waffentypen nicht einverstanden war, da diese Regulierungen zu oft nicht eingehalten werden und sie diese Waffentypen als zu fürchterlich und ihre Vermehrung als zu schnell empfindet. Diese Bereiche können im diplomatischen Prozess benutzt werden, um vorhandene Waffensysteme zu untersuchen und zu untersuchen, wie menschliche Kontrolle garantiert wird und auch mit Bezug auf Parameter (Raum oder Zeit). Und zu versuchen, festzulegen, warum die aktuellen Systeme ausreichend sind und was die minimale Kontrolle ist.

Name: Person B

Datum: 09-03-2017

Dauer: 25:26 Min

No.	Frage
1	Warum beschäftigen Sie sich mit dem Thema Autonome Waffen?
Ich stellte diese Frage nicht, wegen der begrenzten Zeit (30 Min.) für das Interview.	
2	Wie würden Sie autonome Waffensysteme definieren?
Es gibt nur vage Definitionen von autonomen Waffensystemen und das ist, gemäß der Diplomaten der UN, absichtlich so, denn man definiert nicht etwas, was man dann verbieten will. Für B enthält es Aspekte des Identifizierens, Zielens und Zerstörens eines Gegners am Boden.	
3	Was ist Ihre Stellungnahme zu autonomen Waffensystemen?
Es sollte ein Verbot für alle Technologien für AW geben, die als zerstörend und gefährlich angesehen werden, bevor ein Rüstungswettrennen oder ein Entwicklungssprung unsere Kriegsführung beginnt. Es ist möglicherweise die dritte Revolution in der Kriegsführung. Es wäre ein Entwicklungssprung in der Effizienz und Geschwindigkeit, mit der die andere Seite getötet würde. Es gäbe viel Kollateralschaden und wäre nicht, wie Leute denken mögen, eine besonnenere und sauberere Kriegsführung. Es würde potenziell das Töten noch mehr industrialisieren, was ja bereits ein Teil der Geschichte der Kriegsführung ist. Es würde schrecklich destabilisieren, da Kriegsführung früher auf die Wirtschaft angewiesen war (eine große Armee und viel Gerätschaft zu erhalten). Diese Waffentypen sind potenziell eher billig und brauchen nicht viel Wartung. Man brauchte	

früher Tausende von Leuten, um Krieg zu führen und mit AW braucht man nur einen Programmierer. Das wird die aktuelle politische Ordnung ins Wanken bringen.

Es wird wahrscheinlich die Hemmschwelle einen Krieg zu erklären, senken und uns vom Kämpfen in einem Krieg und dem physischen Akt des Tötens distanzieren. Auf einer mehr moralischen und persönlichen Ebene war Krieg immer der letzte Ausweg und hochpersönlich, blutig und gefährlich. Wenn wir so tun, als wäre er das nicht, werden wir wahrscheinlich öfter Krieg führen und das wäre etwas, das Politiker rechtfertigen müssen, wenn Leute in Leichensäcken zurückkommen.

Aber dieses Argument ist nicht immer anwendbar. Diese Argumente sind für heutige autonome Waffensysteme anwendbar, aber in 50 Jahren können andere Argumente relevanter sein, wenn die Technologie weiter entwickelt wurde. Zurzeit entspricht die Technologie dem Menschenrecht nicht und kann nicht zwischen Kämpfern und Nicht-Kämpfern unterscheiden. Irgendwann in der Zukunft werden wir Systeme haben, die präziser und fähiger sind, Gesetze zu befolgen, aber dann kommen andere Argumente auf, wie die größere Effizienz und das Absenken der Hemmschwellen, um Krieg zu führen.

Zurzeit ist es unethisch, AW einzusetzen, aber es ist nicht schwarz und weiß, wie die Einführung von Minen uns gezeigt hat. Dies führt uns zu den Gesetzen und Regulierungen, die ihren Einsatz verbieten, und eine Mine kann als eine dumme autonome Waffe angesehen werden, die nicht zwischen Menschen diskriminiert.

4	Welche drei Fotos automatischer Waffensysteme [in dieser Version der Masterarbeit sind die drei Fotos wegen möglicher Urheberrechtsprobleme mit Worten beschrieben] möchten Sie in die Diskussion
---	---

	aufnehmen*?
1 Schwarm von sechs Drohnen 2 Sea Hunter (autonomes unbemanntes Oberflächenfahrzeug) 3 URAN-9 (unbemanntes Ketten-Bodenkampffahrzeug)	
5	Können Sie für jedes [beschriebenes] Bild erklären, warum Sie es gewählt haben, welchen Wert es für sie innehält und warum es wichtig ist*?
2.3.3.1Viele Drohnen Es beweist, dass wir über eine ganz simple Technologie reden, die bereits vorhanden ist. Sie wird es sein, die in einer Konfliktsituation auf uns zu schwärmen wird, man wird sich nicht verteidigen können. Es sind die Waffen des Terrors und werden benutzt, um Zivilisten und andere Personengruppen zu terrorisieren, und sie sind relativ billig. Man kann mit diesen schlicht aussehenden Drohnen viel anrichten.	
2.3.3.2 Sea hunter Dies wurde gewählt, weil wir über sämtliche Sphären von Kriegsführung sprechen. Also werden auf Land, der See, in der Luft, überall wo man es sich denken kann AW eingesetzt. Es ist ein interessanter Fall, wie schnell die Wettrüstung eintritt. Das wurde bei der Veröffentlichung des offenen Briefes nicht gewusst, aber seitdem passt es auf weitere militärische Prototypen von AW. Und es ist interessant, dass sie, als sie sie gebaut haben, behaupteten, dass sie keine Waffen tragen würden, sondern für Minenräumung und Erkennung von U-Booten benutzt würden. Und noch vor einigen Monaten gaben sie zu, dass sie auch Waffen darauf anbringen würden. Es beweist wieder einmal, dass wir uns auf einem fragwürdigen Weg befinden und die Tatsache, dass eine Wettrüstung vor sich geht. Man braucht keine Leute oder Einrichtungen, um ein Schiff zu führen und könnte viel mehr in sehr viel weniger Raum erreichen und sehr viel weiter kommen. Sie bauen jetzt von Solarenergie angetriebene Schiffe, die jahrelang auf See	

funktionieren können. Man braucht keine Leute, die zu ernähren sind. Man kann den militärischen Vorteil dieser Arten von Technologie deutlich verstehen, und die Armeen werden versucht sein, dies auszunutzen.

2.3.3.3 Uran-9

Der russische autonome Panzer, der beweist, dass viele Spieler, China, Russland, Großbritannien, Israel sowie auch die Vereinigten Staaten deutlich um die Wette rüsten. Das russische Militär hat Aufträge von Millionen Dollar in ein autonomes Waffensystem gesteckt, und es stellt ein ganz anderes Umfeld, als das des Schlachtfeldes dar. Man kann überall sehen, dass, überall wo wir Kriege führen, ein autonomes Waffensystem eingesetzt wird, weil es offensichtliche militärische Vorteile hat. Wenn man sich diesen Panzer ansieht, weiß man, dass wir nicht zu futuristisch denken und dass es mit den Terminator-Robotern nichts zu tun hat. B begreift, dass Russland und China ebenfalls AW entwickeln, da sie wissen, dass die Vereinigten Staaten 18 Billionen Dollar in ihr Budget einschließen, um die nächste Generation für hauptsächlich AW zu entwickeln. Das Problem ist, dass diese Technologie nur mit eigenen autonomen Waffen verteidigt werden kann, da niemand die Reaktionszeit kennt, oder die Fähigkeit hat, rund um die Uhr zu verteidigen, folglich wäre es erforderlich, weil man sich sonst nicht verteidigen könnte. Natürlich ist es nicht überraschend, dass eine Wettrüstung stattfindet.

6	Haben Sie noch weitere Kommentare zu diesem Interview?
---	--

Einige weitere Kommentare wurden von B während des Interviews gemacht:

Es ist wahnsinnig kurzsichtig von Politikern und dem Militär zu denken, dass man eine Technologie taktisch führen kann. Dass haben wir mit der Wasserstoffbombe oder Atombombe

nicht gekonnt. Man kann den Geist nicht zurück in die Flasche bannen.

Es wäre schwer zu bewerten, wie sich die Technologie bewähren wird, da wir über höchst komplizierte Systeme reden, die in einem ungeordneten und komplizierten Umfeld zusammenkommen. Es überrascht nicht, dass wir vom Kriegsnebel sprechen, und zu denken, dass wir gewährleisten können, wie die Systeme sich verhalten werden und dass wir sie beobachten und bewerten können, geht weit über das, was technisch möglich ist, hinaus. Es wird nie möglich sein, deshalb denkt er, dass es unverantwortlich ist, diese Technologie ins Feld zu führen. Roboter können in der Industrie sehr von Hilfe sein, da es sich um ein kontrolliertes Umfeld handelt und die Menschen nicht im Wege sind, aber der letzte Ort, um Roboter einzusetzen, ist ein chaotisches Schlachtfeld, wo wir keinerlei Erwartungen haben, wie sie sich verhalten werden. Wir wissen nicht, wie wir ethische Regelungen in die Roboter einbauen können, aber wir müssen daran arbeiten, damit autonome Fahrzeuge sicher auf unseren Straßen fahren können. Er meinte, dass das Schlachtfeld so ziemlich der letzte Ort wäre, um dies auszuarbeiten, aber kurioserweise wird es der erste sein, in dem diese Technologie ausprobiert wird.

B möchte ebenfalls darauf hinweisen, dass AI enorme Vorteile für das Militär mit sich bringt. Die amerikanische Armee war stets der größte Investor in KI-Forschung und das ist gut, weil Leute nicht Leben oder Glieder riskieren brauchen, um ein Minenfeld zu räumen. Dies ist ein perfekter Job für einen Roboter, dem, wenn er einen Fehler macht und zersprengt wird, niemand nachtrauert. Das gilt auch für Versorgung in Konfliktzonen mit autonomen Fahrzeugen. Es gibt viele Anwendungen von KI im militärischen Bereich, die AW nicht einbeziehen, welche eine endgültige Entscheidung über Leben oder Tod fällen. B hat keine großen moralischen

Einsprüche gegen Verteidigungssysteme, wie z.B. Phalanx, die in klar definierten Bereichen arbeiten, so dass sie konzipiert sind, Leben zu retten, nicht zu nehmen, aber selbstverständlich kann jede Technologie neue Positionen einnehmen. Es würde schwerfallen, ein moralisches Argument aufzustellen, dass die Systeme, die bereits vorhanden sind, von Marineschiffen entfernt werden und sie so angriffsfähiger machen sollten. Es geht vor allem um den Umfang der Einsätze und die Bereiche, in denen sie sich befinden. Dies ist eine ganz andere Szene, als die einer Drohne, die aus Verteidigungsgründen rund um die Uhr eingesetzt wird.

B ist auch wegen anderer Akteure besorgt, denen wir nicht wie unserem Militär vertrauen können, Akteure wie Terroristen, Schurkenstaaten und Leute, die kein Problem damit haben, ethische Sicherheitsbarrieren abzuschalten oder dort hinein zu hacken, um Schaden anzurichten. Folglich, selbst wenn wir diese Systeme so bauen könnten, dass sie sich richtig verhalten, damit wir damit Kriege führen könnten und man nicht in sie eindringen kann, ist die Welt immer noch ein sehr unsicherer Ort.

Name: C

Datum: 10-03-2017

Dauer: 25:25

No.	Frage
1	Warum beschäftigen Sie sich mit dem Thema Autonome Waffensysteme?
Ich stellte diese Frage nicht, wegen der begrenzten Zeit (30 Min.) für das Interview.	
2	Wie würden Sie autonome Waffen definieren?
Die Definition des Verteidigungsministeriums ist, dass diese Waffe selbst das Ziel auswählen und auch angreifen kann. Sie können zwei verschiedene Dinge vertreten. Erstens fügte C hinzu, dass die Waffe nicht zuvor von einem menschlichen Operateur designiert wurde. Zweitens, und dies ist immer stärker die Definition der Kampagne „Stop Killer Robots“, kann man AW als negativ definieren, da es Waffensysteme ohne bedeutungsvolle menschliche Kontrolle sind. Und das hilft uns immer noch nicht weiter, da wir irgendwann ja etwas definieren müssen. Ob wir nun definieren, was bedeutungsvolle menschliche Kontrolle ist, oder was man wählen und einsetzen soll, ist unwichtig, aber es ist tatsächlich eine harte Nuss zu knacken, da wir wissen, was es im Extremfall bedeutet, aber wo die Grenze gezogen werden soll ist wirklich schwierig.	
3	Was ist Ihre Stellungnahme zu autonomen Waffensystemen?
Seine persönliche Einstellung zu AW ist, dass es schwer ist zu entscheiden, was wir in Bezug auf AW tun sollen, wenn wir nicht einmal wissen, wie sie aussehen. Wenn Sie mir sagen, welche Waffensysteme autonom sind und welche Waffen inbegriffen oder ausgeschlossen sind, dann kann ich Ihnen	

sagen, was ich über sie denke. Aber die Technologie ist erst aufkeimend und die Vagheit der Definition so groß, dass es schwer ist, Ja oder Nein dazu zu sagen. Die Kategorie ist potenziell so groß, und wir wissen nicht, in welche Richtung die Technologie geht.	
4	Welche drei Fotos automatischer Waffensysteme [in dieser Version der Masterarbeit sind die drei Fotos wegen möglicher Urheberrechtsprobleme mit Worten beschrieben] möchten Sie in die Diskussion aufnehmen*?
1. Unbemanntes Raupen-Bodenfahrzeug 2. MQ-9 (Reaper) (Unbemanntes Luftfahrzeug) 3. Nahaufnahme des Roboters der „ <i>Ban the Killer Robot Campaign</i> “ (Kampagne gegen Killer-Roboter) vor dem Big Ben in London.	
5	Können Sie für jedes [beschriebenes] Bild erklären, warum Sie es gewählt haben, welchen Wert es für sie innehält und warum es wichtig ist*?
1. Landroboter Dies stellt die Zukunft der militärischen Robotertechnik da, die die Leute wirklich verwirklicht sehen oder besprechen möchten, und es ist die übliche Form der Robotertechnik, der wir in Wahrheit auf dem Schlachtfeld begegnen werden. Eines der interessantesten Dinge in der Debatte zu AW ist, dass man fast eine ganze Schlacht von Analogien hat, wo Gruppen wie die Kampagne <i>Stop Killer Robots</i> sich auf anthropomorphe Maschinen wie den Terminator konzentrieren, die in ein Haus marschieren, um zu entscheiden, ob die Person darin ein legaler Gegner oder nicht ist, z.B. ein Kind. B und ein weiterer Autor reden oft über Luftwaffen- und Marine-Szenarien, wo das Schlachtfeld ein wenig übersichtlicher ist und reden mehr über Waffenplattformen statt über marschierender Roboter. Das erste Bild stellt eine funktionalere Form da, wo solche	

Roboter auf dem Schlachtfeld wahrscheinlicher sind, als dass Terminator Typen herumlaufen. Diese Systeme werden schon jetzt als ferngesteuerte Systeme benutzt.

7. MQ-9 (Reaper)

Der MQ-9 ist interessant, da viele Bedenken zu AW mit dem Einsatz von Drohnen begannen, und die Leute machen sich Sorgen, dass Drohnen einen Krieg zu führen zu leicht machen, und das ist gefährlich. Aber man könnte Drohnen nicht wirklich aufhalten, denn sie sind bereits auf dem Schlachtfeld und werden bereits eingesetzt. Sein früheres Werk zeigte, dass Drohnen sich rapide vermehren, und das könnte nicht aufgehalten werden. Aber vielleicht kann man etwas dagegen tun, was als Nächstes kommt. Und eine Menge Gruppen versuchen, von ihrem Standpunkt aus, wie wir sicherstellen können, dass dies der richtige Weg ist und sie machen sich Sorgen darum, dass Drohnen für sich selbst entscheiden können.

Dies ist auch für die Frage der Entscheidungstreffung relevant und wenn die Rede davon ist, dass ein Roboter eine eigene Entscheidung trifft, statt im Vorherin programmiert zu sein. Besonders wie eng und zentralisiert seine Programmierung ist. Wenn die Leute negativ über Roboter reden, meinen sie oft, dass Roboter Entscheidungen treffen und die Leute sind weniger besorgt, wenn Menschen die Roboter anweisen, da es dann immer noch Menschen sind, die die Entscheidungen treffen.

Es ist faszinierend, dass es Menschen gibt, die glauben dass man ein Recht hat, von einer Person getötet zu werden, während die Menschen sich gegenseitig alle möglichen schrecklichen Dinge antun, während man weiß, wenn man tot ist, ist man tot. Es ist interessant zu sehen, wie schnell dies zur Moralphilosophie ansteigt, fast mehr als alles andere.

8. Kampagne gegen Killer-Roboter

Es ist interessant, dass eine Gruppe von Leuten mit Erfahrung in Kampagnen zu Landminen, Streumunition und anderen weniger zentralen Technologien nicht wissen, wie Militärs arbeiten. Viele Bedenken in Bezug auf AW drehen sich darum, welche AW wir haben und was sie bewirken können, aber wir wissen, dass die Integration von Autonomie in Militärsysteme im Allgemeinen stattfindet, doch ob die Waffensysteme autonom sind ist eine andere Frage. In der Tat ist die Integration von Autonomie in Waffensysteme unabwendbar. Es ist interessant, dass die Kampagne dies erfasst hat, besonders weil es eine andere militärische Technologie ist, als die, die sie zuvor aufgegriffen haben. Für C ist es eher, als ob sie versuchten den Panzer oder das U-Boot zu verbieten. Wenn die Killerroboter-kampagne recht hat, dann wollen die Militärs sie und wenn sie nicht recht hat und diese Waffen keine große Sache ist, dann ist es eigentlich nicht besonders wichtig. Es ist eine interessante Gruppe und sie versuchen ernsthaft, menschliches Leid zu reduzieren, aber sie benutzen das Drehbuch für ganz andere Siege und C denkt, dass dies eine andere Situation ist, weil die Technologie noch ungewiss ist und wie weiträumig die potenzielle Kategorie ist, weil niemand weiß, wie wichtig die Kategorie möglicherweise für das Militär ist.

6	Haben Sie noch weitere Kommentare zu diesem Interview?
---	--

Dieses sind einige weitere Punkte, die während des Interviews angesprochen wurden.

Auf meine Frage, dass nicht viele empirische Untersuchungen durchgeführt worden sind, antwortete C, dass er einige Untersuchungen mit schnellen Experimenten vorgenommen hatte, Im Allgemeinen stellen Meinungsforscher eine Liste mit Fragen zur Einstellung auf und präsentiert sie in einem anderen Kontext, um ihre Ansichten in diesen Szenarien kennenzuler-

nen und dann versuchen einzuschätzen, was sie empfinden, aber diese Arbeit ist auf die USA beschränkt.

Von Interesse ist auch, dass die Politik sehr interessant ist und bis zu einem gewissen Ausmaß deshalb, weil wir in einer Welt leben, in der wir die Überlegenheit der amerikanischen militärischen Technologie für selbstverständlich ansehen. Stellen Sie sich eine Welt vor, in der China AW einsetzt und die USA und EU-Länder haben keine, dann hörte man nächste Woche im Kongress von einer AW-Lücke und Fragen würden gestellt wie: „Warum hat die USA keine Waffensysteme wie die Chinesen?“ Dies führt zu der Frage, die nach Cs Meinung von der Kampagne nicht berücksichtigt wurde, nämlich was passiert, wenn ein undemokratisches Land diese Systeme einsetzt und sie nicht nur Nischenwaffen sind, sondern wichtig für die allgemeinen militärischen Einsätze und sie einem tatsächlich einen Vorteil auf dem Schlachtfeld gibt? „Was würde die Welt darauf wohl antworten?“ Die braucht nicht unbedingt negativ ausfallen, aber das Problem ist sehr kompliziert.

Die Frage ist, ist es ein Wettrüsten, bei dem man Angst hat, dass jemand eine Technologie erwirbt und man ist in direkter Konkurrenz mit ihm, oder dass sie sich so schnell vermehrt, weil ihr Erwerb so einfach ist. Bei einem Wettrüsten geht es um Politik, da man einen Grund für das Wettrüsten benötigt, und wenn eines mit AW stattfindet, ist der Grund, dass es schwer ist, wissen was ein AW darstellt. Was ist der Unterschied zwischen einen MQ-9 Reaper und einen autonomen Reaper. Es ist die Software, nicht die Hardware, und man kann sie nicht sehen. Die Ungewissheit, ob das System des Gegners eine Drohne oder ein AW ist, könnte ein möglicher Grund sein, der zu einem Wettrüsten führen würde, da man sich nicht sicher sein kann, dass die andere Seite keine AW erwirbt. Es sei denn, man könnte in ihr Waffensystem hacken, aber wer würde so etwas zulassen?

Name: Person D

Datum: 19-04-2017

Dauer: 55:18 Min

No.	Fragen
1	Wie beschäftigen Sie sich mit dem Thema Autonome Waffen?
Diese Frage habe ich nicht gestellt, weil ich sie den anderen Befragten auch nicht gestellt habe.	
2	Wie würden Sie autonome Waffensysteme definieren?
Es ist ein Waffensystem, das Gewalt anwendet und diese Gewalt an sich ist tödlich, aber das ist keine notwendige Folgerung. Dies ist keine zwingende Notwendigkeit, aber eine häufige Ansicht in einem militärischen Kontext, obwohl diese Vermutungen bei der Diskussion der Fotos aufkommen könnten. Für die Anwendung in naher Zukunft, ist folgendes zu erwarten. Autonomie ist für mich, dass sie fähig sind, selbständig ein Ziel zu wählen, um ein Ziel in einem gewissen Kontext mit einer vagen Aufgabe zu wählen. Es wird nicht speziellen Koordinaten zugesendet oder zu einem speziellen Fahrzeug, sondern eher: „gehe in diesen Bereich und setze alle Fahrzeuge, die sich feindlich verhalten, außer Gefecht.“	
3	Was ist Ihre Stellungnahme zu autonomen Waffensystemen?
D würde sich dem gerne widersetzen, aber meine Analyse zeigt, dass das nicht möglich ist, also wäre es das nächstbeste, dass wir sie so verantwortungsbewusst wie möglich einsetzen und so das Risiko, dass Situationen aus dem Ruder laufen, reduzieren. Ich kann diese Einstellung basiert auf nachstehenden Ideen erklären: 1. Auf der zivilen Seite floriert ein starkes geschäftliches Interesse daran, KI zu entwickeln, um viel Geld zu	

verdienen.

2. Gleichzeitig ist die Entscheidungstreffung im militärischen Kontext von Wichtigkeit. Früher war Informationssammlung das Problem, aber heute ist die Verarbeitung der Informationen das Problem. Das bedeutet, das menschliche Entscheidungstreffung langsam zum schwächsten Glied wird, was dazu führen wird, das es stattdessen dem System überlassen wird.
3. Und drittens ist es möglich, dass eine übermenschliche Intelligenz aufkommt. D sieht keinen technischen Grund, warum das unmöglich sein sollte. Wenn dem so ist, können wir nicht voraussehen, was das bedeuten wird. Dies, kombiniert mit tödlicher Gewalt, ist keine Situation, die wir als Menschheit wünschen.

Vom Standpunkt der Hoffnung aus, dass wir autonome Waffensysteme nicht entwickeln und einsetzen werden, ist es nicht möglich, da die ersten beiden Punkt sicherstellen, dass wir diese Situation erwirken. „Die schönsten Blumen blühen dem Klippenrand am nächsten.“ Also, das Beste, das wir tun können, ist sicherzustellen, dass wir nicht über den Klippenrand stürzen. Er ist ziemlich pessimistisch, dass wir dies kontrollieren können, aber wir müssen es versuchen.

Um es zu kontrollieren, könnte man ein System benutzen, zum Beispiel in der Design-Phase, um sie in unsere erwünschte Welt einzubetten und zu verhüten, dass diese Technologien ihre eigene Version davon kreieren. Dies steht ganz im Vordergrund der Entwicklung, aber während des Einsatzes müsste man stark darüber nachdenken, wie man das Zielverfahren konzipieren würde. Je fähiger die Selbstbestimmung dieser Systeme wird, je mehr müssen wir die Zielsetzungsparameter spezifizieren, und das kann anders aussehen, als wir es uns zurzeit vorstellen.

Und wir sollten versuchen, ein internationales Abkommen

darüber zu schließen, und das ist das Schwerste, denn wenn man es nicht befolgt, hat man einen Riesenvorteil. Es ist das typische Gefangenendilemma. Wir müssen darüber nachdenken, und da es keine leicht zu beantwortende Frage ist, hat D keine Antwort darauf. Aber es ist offensichtlich, dass die Geschäftsparteien hier eingeschlossen sein sollten.

Auf meine Anfrage, seine erste Aussage bezüglich seinem Widerspruch gegen diese Technologie zu erklären, entgegnete er, dass autonome Waffensysteme sich insofern von konventionellen Waffensystemen ohne Autonomie unterscheiden, dass letztere nicht ganze Gattungen vernichten können, obgleich das rein theoretisch mit Atomwaffen ja möglich ist, aber sie sind zu teuer im Bau und AW sind das nicht und stellen somit ein höheres Risiko da. Folglich kann das Verfahren der Konstruktion nuklearer Waffen kontrolliert werden, und es gibt sie seit 60 Jahren. Es hat uns erlaubt, sich an sie zu gewöhnen. Die Frage ist, ob wir diese Möglichkeit mit den AW und ihrer rapiden Entwicklung haben. Hinzu kommt, dass Waffensysteme mit künstlicher Intelligenz einige Risiken innehaben, während KI Ärzte in ihrer Diagnosestellung unterstützt, was zum Guten führen soll. AW kombinieren ein hoch intelligentes System mit einem feindlichen Vorhaben, folglich ist D besorgt, dass AW ein größeres Risiko darstellen. Autonome Waffensysteme sind nicht das einzige, um die sich D Sorgen macht Synthetische Biologie ist ebenfalls ein hoch riskantes Feld. Also, Technologien mit einer niedrigen Hemmschwelle aber potenziell unkontrollierbaren Risiken, bereiten der Menschheit Kopfschmerzen.

4	Welche drei Fotos automatischer Waffensysteme [in dieser Version der Masterarbeit sind die drei Fotos wegen möglicher Urheberrechtsprobleme mit Worten beschrieben] möchten Sie in die Diskussion aufnehmen*?
---	---

1 Screenshot des Films Ex Machina. 2 Bild eines Rettungsroboters, der einen menschenähnlichen Dummy trägt, bei einem Roboter-Wettbewerb zur Hilfeleistung für Umstehende. 3. Roboterhand und menschliche Hand, deren Zeigefinger einander berühren.	
5	Können Sie für jedes [beschriebenes] Bild erklären, warum Sie es gewählt haben, welchen Wert es für sie innehält und warum es wichtig ist*?
Alle Fotos sind verknüpft. 1. Das Interessante in dem Film Ex Machina ist, dass die für diesen Film entwickelte Technologie auch die Überlebensinstinkte der Menschheit geerbt hat (oder sie einprogrammiert werden) und somit Maßnahmen trifft, um daran hindern zu sterben aber sie auch keinen Deut darum gibt, was anderen passiert. Es stellt die Tatsache da, dass man mögliche Ereignisse berücksichtigen muss, wenn man diese Art von Technologie konstruiert und was falsch laufen könnte. wenn man es nicht tut. Da wir über ein System mit einem gewissen Niveau der Selbstbestimmung reden, könnten wir die Dinge übersehen, die sie sich selbst beigebracht haben. Dies sind 2 Elemente: erstens, unabsichtlich außer Kontrolle zu geraten und zweitens, absichtlich gegen die Grenzen zu stoßen, d.h. man kann diese Grenze überschreiten. 2. Das zweite Bild ist von einem Ausschreiben für Rettungsroboter, der bei Katastrophen eingesetzt werden kann. Dieser Roboter schadet niemanden, aber er tut auch nichts. Folglich ist das Dilemma, dass man eine Kombination von sowohl 1 als auch 2 haben will, aber dafür muss man den Kontext gut kennen und über eine Art „guten Willen“ verfügen. D deutet an, dass KI bessere Entscheidungen als ein Mensch treffen sollte. Er gibt ein	

	Beispiel von KI, der befohlen wird, so viele von Hand geschriebene Notizen wie möglich zu machen, und zu einem bestimmten Zeitpunkt entwickelt sie sich in eine Super-KI und entscheidet innerhalb von Sekunden, was die beste Methode ist, die ganze Erde in Notizen zu wandeln, alle Menschen auf der Erde töten, da dies eine Möglichkeit ist, so viele von Hand geschriebenen Notizen wie möglich zu erzeugen. Auftrag ausgeführt, aber nicht auf die Weise, die wir geplant hatten. Also muss man darüber während der Design-Phase nachdenken.
3.	Doch wir müssen auch über die Kooperation zwischen Mensch und KI nachdenken (obwohl dies nicht in den frühen KI-Systemen geschieht), aber in der Zukunft werden wir KI-Systeme haben, die es tun. Die aktuellen KI-Systeme wie Harpyie haben wenig Kontext-Bewusstsein, aber wenn wir Roboter in Landeinsätzen einsetzen, müssen wir mit ihnen kooperieren, und Zielsetzung wird überaus wichtig werden. Foto 3 stellt diese Kooperation da. Dies bedeutet, dass das System wissen muss, was wir meinen, aber wir müssen unsere Leute auch anlernen, mit diesen Systemtypen zusammenzuarbeiten und ihnen die Ziele angeben. Je mehr die Systeme grenzenlose Möglichkeiten haben, desto mehr werden sie über das Feststecken von Grenzen zu diesen Aufgaben nachdenken müssen. D gibt das Beispiel, dass ein Zauberspruch gut wäre, aber man muss ihn auch verstehen. Wir müssen also unser Personal darin schulen, Aufgaben mit festen Grenzen zu stellen.
6	Haben Sie noch weitere Kommentare zu diesem Interview?
	Auf meine Frage, was D darüber denkt, einen Menschen in den Entscheidungsprozess einzuschließen, antwortete er, dass es davon abhängt, wie grenzenlos das System ist. Erstens, je mehr Freiheit es hat, desto mehr muss man in der Design-Phase

darüber nachdenken, wie man dem System gestattet, guten Willen zu entwickeln. Zweitens, da man diese Systeme nicht eingrenzen kann, muss man einen Weg finden, das Risiko einzuschränken.

D sieht diese Entwicklung in Phasen. Phase 1 ist die großflächige Einführung von AW mit begrenzten Fähigkeiten. In Phase 2 erhält man Mensch-Maschine-Teams, und sie kann als eine Zwischenphase angesehen werden, in der Leute lernen, mit der Technologie zu kooperieren. Doch in der Phase wird es gefährlich und riskant, da wir dann mit unabhängigen Systemen konkurrieren und in Phase 4 wird dies zum Terminator oder Skynet-Beispiel, die ein Beispiel für eine KI sind, die sich als ernsthaft böse entwickelt hat, doch D denkt, dass dies nicht nächstliegend ist und KI sich nicht als böse herausstellt, sondern es wird eher der Fall sein, dass wir die KI durch kleine Abweichungen nicht mit der richtigen Denkweise ausgerüstet haben und somit immer noch die Welt zerstören. Er meint auch, dass es keine scharfe Grenze zwischen Phase 3 und 4 gibt, aber dass man das erst merkt, wenn es zu spät ist.

Name: Person E

Datum: 17-04-2017

Dauer: 22:19

No.	Frage
1	Warum beschäftigen Sie sich mit dem Thema Autonome Waffensysteme?
Ich stellte diese Frage nicht, wegen der begrenzten Zeit (30 Min.) für das Interview.	
2	Sie, Wie würden Sie autonome Waffensysteme definieren?
Autonome Waffensysteme sind ein Waffensystem, dem ich einen Befehl geben kann zu gehen, es findet den Weg, wählt selbst sein Ziel (basiert auf den von Menschen erstellten Parametern), basiert auf dem Bild des Umfeldes das es selbst kreiert. Also wählen, finden und selbstständig handeln sind die drei Komponenten.	
3	Was ist Ihre Stellungnahme zu autonomen Waffensystemen?
Positiv, ich bin offen für diese Art von Technologie. Also, positiv, aber wir müssen darüber nachdenken, wie wir sie mit Besonnenheit einsetzen. Wir nehmen sie nicht einfach achtlos in die Organisation auf ('niet zomaar naar binnen gooien').	
4	Welche drei Fotos automatischer Waffensysteme [in dieser Version der Masterarbeit sind die drei Fotos wegen möglicher Urheberrechtsprobleme mit Worten beschrieben] möchten Sie in die Diskussion aufnehmen*?
1 Harpunen-Marschflugkörper, von einem Marineschiff aus abgefeuert 2 Schwarm von 19 Drohnen 3 Roboter im Future-Style mit einer großen Handwaffe.	

5	Können Sie für jedes [beschriebenes] Bild erklären, warum Sie es gewählt haben, welchen Wert es für sie innehält und warum es wichtig ist*?
1.	Dies ist eines der derzeit existierenden Systeme, das von der Marine seit über 30 Jahren benutzt wurde. Es ist ein Schiff, das Torpedos auf U-Boote abschießt. Sobald Sie diese Waffe abschießen, wählt sie autonom ihr Ziel und greift an. Man kann es nicht überschalten. Es heißt schießen und vergessen. Ein ähnliches System ist die Harpunenkanone, ein Lenkflugkörper, der gegen gegnerische Schiffe auf der Wasseroberfläche eingesetzt wird. Diese Waffensysteme suchen sich ihre eigenen Ziele und wenn sie sie finden, könnte es ein anderes sein, als das, auf das Sie gezielt haben. Es ist ein System, das mit größter Vorsicht einzusetzen ist, aber es ist eine Art von Autonomie. Vom Standpunkt der Marine aus, sind diese gelenkten Flugkörper Kontrollsysteme, die mit einem gewissen Satz Parameter eingestellt werden können, z.B. Höhe, Geschwindigkeit und gewissen Identifizierungsmustern, und sie sind unterwegs. Es befindet sich kein Mensch im Entscheidungsprozess, und wir benutzen diese Systeme seit über 20 Jahren. Diese Systeme können autonom eingesetzt werden, aber da wir diese Systeme einsetzen, bauen wir immer eine Funktion ein, die einen Menschen in den Entscheidungsprozess einschließt. Ein Mensch muss dem System grünes Licht geben, bevor es sein Ziel vernichten kann. Dies ist ein Prozess zentriert auf diese autonome Funktion, und es gibt sogar einen Hardware-Schalter, der bedient werden muss, bevor das System das Ziel vernichten kann. Eine Art von Hemmschwellenschalter, aber sie können das System auch so einstellen, dass Sie diese Überschaltung nicht benötigen. Diese Systeme wurden für Einsätze bei hoher Luftbedrohung entwickelt, in denen

jede Sekunde zählt. Dies ist bereits ein gewisser Grad der Autonomie, mit dem Gefühl, dass man zu jeder Zeit eingreifen kann. Die Tatsache, dass es möglich ist, zu jeder Zeit einzugreifen, ist für die Marine sehr wichtig. Dies steht im Gegensatz zum Torpedo, der einem nicht immer gestattet, einzugreifen, will sagen, man muss sich absolut sicher sein, bevor man ihn einsetzt. Für beide Waffensysteme hat die Marine Einsatz-Prozedere erstellt.

2. Der nächste Schritt für das Verteidigungsministerium sind die ferngesteuerten Drohnen, die wir in unsere Organisation aufzunehmen versuchen, aber die Entwicklung in der Technologie geht schneller vor sich, als dass wir ihr zurzeit folgen können, und wir bemühen uns, uns diesem Tempo anzupassen. Schwärmer sind die Systeme, die in naher Zukunft verwendet werden. Diese Systeme werden eingesetzt und werden ihren Auftrag ganz alleine ausführen, aber wie sie das tun werden, das haben wir noch nicht im Griff. Die funktionieren mit einem ziemlich hohen Level der Selbstbestimmung. Zurzeit geht es um gemeinsamen Zugriff auf Informationen und Sensortechnologie, aber der nächste Schritt ist ein fühlergesteuerter Angriff mit einem Waffensystem. Die Technologie ist derzeit verfügbar. Zurzeit entwickelt das Ministerium einen Entwicklungsplan für ferngesteuerte Luftwaffensysteme, aber der Sprung, vollständig autonome Waffensysteme zu bauen, ist zu groß, da dies etwas ist, dass wir, falls wir es haben wollen, noch nicht innerhalb des Verteidigungsministeriums besprochen haben. Dies ist zurzeit eine ethische Barriere im Verteidigungsministerium. Das Thema unbemannter und autonomer Systeme ist Teil der aktuellen Vision für die Streitkräfte, dass von der Hoofd Directie Beleid (vom Defensiestaf) konzipiert wird, und damit beginnt auch die

Diskussion über die AW. Das Verteidigungsministerium wird diese Waffentypen nur erwerben, nachdem das Thema AW in unserer Gesellschaft debattiert worden ist.

Alle Fachkräfte der Armee, Marine und Luftwaffe sagen aus, dass Autonomie unterwegs ist, aber wir wollen immer die Kontrolle behalten. Wir wollen sie immer zurückrufen können, wenn sich die Situation geändert hat. Oder, dass wir sie ggf. neu programmieren können. Es ist typisch für militärische Einsätze, dass sich Situationen rapide ändern und man seine Pläne entsprechend umwerfen muss. Besonders, wenn Sie irgendwo arbeiten, wo es noch nicht zum bewaffneten Konflikt gekommen ist und Sie im Rahmen der Verhaltensregeln arbeiten. Sie müssen vorgreifen und schnell reagieren können, und Sie möchten keine Situation erhalten, in der Sie ihr System in der Luft eingesetzt haben, ein gewisses Ziel anzugreifen, aber wenn sich die taktische oder operative Situation ändert, möchten Sie in der Lage sein, es zurückzurufen. Wird diese Anforderung in die Systeme eingebaut, dann würden wir als Verteidigungsministerium diese Systemart gerne einsetzen.

3. Foto 3 stellt das Bild dar, das die Zivilbevölkerung von Killerrobotern hat, ein System, das nicht kontrolliert werden kann und selbstständig denkt und handelt. Leute stehen ihm misstrauisch gegenüber, und es stellt die Trennlinie dar, zwischen dem, was wir zurzeit einsetzen können, und welche Technologie wir uns ersparen möchten.

5 Haben Sie noch weitere Kommentare zu diesem Interview?

E betonte, dass er bereits seit Jahren mit AW gearbeitet hat (wie Bild 1 darstellt) und dass es eine gewisse Kontrollebene gibt, die man in seine Prozedere und Einsatzführung

(bedrijfsvoering') einbettet. Besonders, wenn eine neue Technologie sich noch nicht bewährt hat, dann müssen Sie mehr Sicherheiten dieser Art darin einbauen. E glaubt, dass wir in diese Richtung gehen und dass wir das auch sollten, denn es würde uns kostbare Zeit sparen und unsere Gegner tun es auch, folglich muss man dem folgen. Wenn ein Schwarm Drohnen auf Sie zukommt, dann müssen Sie sie bekämpfen können.

Name: Person F

Datum: 08-05-2017

Dauer: 53:44 Min

No.	Frage
1	Warum beschäftigen Sie sich mit dem Thema Autonome Waffensysteme?
Diese Frage habe ich nicht gestellt, weil ich nur wenig Zeit hatte, das Interview zu führen (30 Minuten).	
2	Wie würden Sie autonome Waffensysteme definieren?
F akzeptiert die ICRAC-Definition, die Autonomie in den kritischen Funktionen der Zielsetzung und des Abschießens der Waffen befürwortet. Es ist nicht nur die Wahl des Ziels als eine Funktion sondern die Kommunikation, die Waffe auf jemanden zu richten ist auch eine anthropologische und soziologische Handlung. Im formellen Sinne, wählt es ein Ziel oder bestimmt, dass jemand ein gültiges Ziel ist. Und in der nächsten Phase gibt es die Waffe frei und schießt sie ab. Wenn man in diesen beiden Funktionen Autonomie hat, dann ist es definitiv ein autonomes Waffensystem. Die Frage, ob man so ein System oder ein anderes besitzt und bis zu welchen Grad man menschliche Kontrolle hat, ist eine größere Herausforderung. Und was bedeutungsvolle menschliche Kontrolle über diese Funktionen eigentlich heißt, steht immer noch zur Debatte. Die Grundlage für menschliche Kontrolle vom technischen Standpunkt aus bedeutet, dass der Mensch effektiv die Kontrolle hat, er kann intervenieren oder etwas zu Boden bringen oder das Verhalten eines Systems ändern. Bedeutungslose Kontrolle ist, wenn die Person zum gedankenlosen Automaten einer Maschine wird, also wenn man in einem Raum mit einem Knopf sitzt, und jedes Mal,	

wenn ein rotes Licht aufleuchtet, man ihn drücken muss. In gewissem Sinn hat man natürliche kausale Kontrolle über den Knopf, aber man ist sich keiner wirklich wichtigen Bedeutung bewusst, ob ein gültiges Ziel festgestellt wurde, da eben nur ein rotes Licht aufleuchtet und man diesem Licht zutraut, dass es korrekte Ziele produziert und einen gültigen Angriff auf dieses Ziel autorisiert. Aber man hat in Wirklichkeit kein kontextuelles oder situationsbezogenes Bewusstsein oder Verständnis oder unabhängiges Mittel, die Legalität oder Moral eines Ziels zu beurteilen. Im Effekt bedeutet es eben, dass man keine bedeutungsvolle Kontrolle hat, sondern nur kausale Kontrolle.

Im Vergleich dazu, damit Töten eine bedeutungsvolle Handlung ist, muss ein Vorhaben da sein, oder ein militärischer Zweck, ein Ziel zu vernichten. Und ganz einfach weil etwas ein legales militärisches Ziel ist, sollte man es angreifen. Es könnte sein, dass man nicht angreift, weil es teuer ist, oder man seine eigene Stellung verrät, es gibt alle möglichen Rücksichtnahmen, dass man auch bei gültigen Zielen überlegen muss, ob man angreift oder nicht, und das wird dann Teil einer menschlichen Maßnahme, und wenn man sie automatisiert, wird sie bedeutungslos und eigentlich zu einer arbiträren gerichtlichen Hinrichtung. Man hat keine sinnvolle Weise der Bewertungsmöglichkeit, ob ein Ziel legal oder sinnvoll ist.

3	Was ist Ihre Stellungnahme zu autonomen Waffensystemen?
---	---

F ist gegen sie und denkt, man sollte sie verbieten. Allerdings könnte man ein Verbot als eine positive Anforderung als bedeutungsvolle menschliche Kontrolle formulieren. Es ist wirklich schwer, über die ICRC-Definition hinaus zu definieren, die er als vager als bedeutungsvolle menschliche Kontrolle empfindet, oder man könnte sagen, man braucht bedeutungsvolle menschliche Kontrolle für diese beiden

kritischen Funktionen. Er meint, dass wenn bedeutungsvolle menschliche Kontrolle verstanden wird, dann wird es eine Anforderung für alle Waffensysteme, man definiert es für eine bestimmte Klasse von Waffensystemen, aber man muss gewährleisten, dass es für alle Waffensysteme in Frage kommt. Den Vergleich, den er ziehen möchte, ist der des unnötigen Leidens und der überflüssigen Verwundung, nicht dass dies die gleiche Definition ist, aber sie hat eine ähnliche Funktion, ist auch irgendwie unbestimmt darin, was unnötiges Leiden und überflüssige Verwundung eigentlich ist. Es gab auch große Debatten darüber, als diese Sprache in der St. Petersburg Konvention eingeführt wurde, und Juristen legten es beiseite, aber es wird nun so verstanden, dass es zusätzlicher Schaden ist, wenn sie keiner militärischen Funktion dienen, außer Leiden und Schmerz zuzufügen. Wenn dies auch unbestimmte Attribute sind, zeigen sie doch auf, ob wir bedeutungsvolle Kontrolle über ein System haben. Kann ich es aufhalten? Kann ich in seine Handlung eingreifen? Kann ich voraussehen, was es tun wird? Kann ich für mich selber beurteilen, ob die gewählten Ziele wirklich legal und moralisch verfechtbar sind? Denn wenn ich das nicht kann, dann haben wir ein Problem mit diesem Waffensystem. Die Systeme mit begrenzter Kontrolle sind gefährlich, und die globale Gesellschaft muss ausarbeiten, was wir mit ihnen anfangen sollen.

Aktuelle autonome Systeme, wie Goalkeeper auf Schiffen, müssen diese Anforderungen erfüllen. Und Staaten, die diese Art von Systemen einsetzen, müssen begründen, wie sie diese Anforderungen für ballistische oder abwehrende Systeme wie Goalkeeper, Phalanx oder Patriot erfüllen, da viele dieser Systeme herankommende Flugkörper zu Boden bringen und das System, das F persönlich erlebt hat, hatte eine sehr kurze Zeit und wenig Spielraum in Bezug auf wie lange es autonom zielgerichtet funktionieren konnte, nämlich weniger als eine

Minute (20 oder 30 Sekunden), aber man kann immer noch argumentieren, dass da Menschen ihre Hände an diese Systeme legen, die entscheiden, wann sie den Systemen gestatten, abzuschießen. F hat nicht vor, dass das ICRAAC diese Systemtypen eliminiert, aber sie sind trotzdem gefährlich, da sie immer noch den Beschuss eigener Truppen verursachen. F gibt ein Beispiel an, laut dem 2 Flugzeuge des Patriotensystems im Irak heruntergeschossen wurden, da ihr Transponder unsachgemäß reagierte, und danach wurde der Patriot technisch überarbeitet, sodass der Operateur positive Autorisierung für den Abschuss vom System geben muss, statt dass es eigenmächtig automatisch schießt. Dies erhöht die psychologische Belastung, sodass die Leute nicht bereit sind, den Knopf zu drücken, es sei denn sie sind sich völlig sicher, da man dieses Vorurteil gegen Automatisierung hat, wo das System bestimmt, dass das Ziel sicher ist. Nur dann sind Leute willens, das System handeln zu lassen. Es gibt auf jeden Fall eine bedeutungsvollere menschliche Kontrolle für die neue Patriotensplattform als für die frühere.	
4	Welche vier Fotos automatischer Waffensysteme [in dieser Version der Masterarbeit sind die vier Fotos wegen möglicher Urheberrechtsprobleme mit Worten beschrieben] möchten Sie in die Diskussion aufnehmen*?
1. Terminator mit Metallskelett 2. Predator (Unbemanntes Luftfahrzeug) 3. Bild des Roboters der „Ban the Killer Robot Campaign“ vor dem britischen Parlament in London. 4. Dreibeiniger Stuhl vor der UNO in Genf, Schweiz.	
5	Können Sie für jedes [beschriebenes] Bild erklären, warum Sie es gewählt haben, welchen Wert es für sie innehält und warum es wichtig ist*?
1. Terminator	

In den ersten Jahren der Kampagne, wurde das metallene Skelett des Terminators in jedem Artikel über die Kampagne dargestellt, interessant vom Standpunkt der Medien und sozialen Aktivisten aus. Einerseits ist es ein beeindruckendes Bild, dass das öffentliche Bewusstsein stimuliert und die Aufmerksamkeit der Leute erregt, da es Killerroboter darstellt.

Der Terminator ist wirklich sehr interessant, wenn man sich den Film anschaut, da er für ein bestimmtes Ziel eingesetzt wird; in diesem Sinn ist er also kein autonomes Waffensystem, obwohl er viele andere Leute auf seinem Marsch umbringt, somit gibt es einige autonome Zielentscheidungen. Doch die eigentliche Direktive ist Sara Conner, da sie das eigentliche Ziel darstellt. Und es gibt viele andere Elemente zu dem Film, die die Ängste der Leute stimulieren, wie Roboter, die sich ihrer selbst bewusst werden, sich gegen die Menschheit wenden und solche Sachen, die nicht in die Kampagne gehören. Es sind mehr die Zwischenstufen, dumme autonome Waffensysteme, die in den nächsten 5 bis 20 Jahren eingesetzt werden können, die nicht über anspruchsvolle Gedankensysteme verfügen werden, sondern auf einen kleinen Satz Kriterien achten, die auf ihre Sensordaten begründet sind und ihre Waffen auslösen, was negative Konsequenzen für die Zivilbevölkerung haben kann und alles Mögliche sonst. Das Image des Terminators erfasst nicht wirklich einen Wert, aber es erregt Aufmerksamkeit und das war zu Anfang der Kampagne sehr wichtig. Es war nicht Furcht wegen des Films, sondern echte Furcht davor, wie diese Waffensysteme funktionieren werden.

2. Predator

Das zweite, woran Leute denken, sind Predator- und Reaper-Drohnen, ferngesteuerte Roboter, folglich sind sie nicht wirklich autonom und passen somit nicht richtig ins Bild. Es sind realistischere Bilder, die im weitesten Sinne die Vorläufer zu autonomen Waffensysteme sind, weil man die Zielsetzung

und das Auslösen von Drohnen automatisieren könnte, und sie würden zu autonomen Waffensystemen, und es gibt sogar Entwicklungen der nächsten Generation Drohnen, von denen viele sagen, dass sie autonom Ziele setzen können, zum Beispiel die X47-B, die Taranis und der Neuron, heimlich suchende Drohnen mit Düsenantrieb, die viele Waffen mit sich führen, die wirklich für im Konflikt stehenden Luftraum konzipiert sind, im Gegensatz zum Predator und Reaper, die mit Propeller angetrieben werden und langsam sind, und für sie muss man die Kontrolle über den Luftraum haben. In einem Luftraum, wo der Gegner alle Flugkörper hat oder Kommunikationen sperren kann, sind diese Waffentypen nicht geeignet. Die nächste Generation Drohnen kann in Räumen funktionieren, wo keine Kommunikation herrscht, ein Argument, das für Autonomie stimmt, weil man keine Kommunikationsverbindung benötigt, die unterbrochen werden kann, sodass man die Kontrolle über den Flugkörper verliert. Für diese Bereiche müssen sie autonom sein. Die Frage ist, wie man sie zum Ziel bringt. Es könnte wie bei vorhandenen Lenkflugkörpern vorgehen, bei denen man eine Liste von GPS-Koordinaten bekommt, sie suchen diese Koordinaten und bombardieren sie; in diesem Fall bestimmen Menschen diese Ziele, oder der Flugkörper hat bis zu einem gewissen Grad eine eigene Fähigkeit des Suchens und Wählens eines Ziels von Sensordaten basiert auf Zufall.

Die Bilder dieses neuen Systems sehen sehr unheimlich aus, vor allem die der Taranis, mit Blitzen und Stapeln von Sprengköpfen davor. Sie stellen dar, was uns Besorgnis wegen der nächsten Generation autonomer Waffensysteme erregen sollte. Diese Besorgnisse würden darin liegen, dass sie eine Fähigkeit der autonomen Auswahl von Zielen haben. Es ist ein Langstreckenkampfflugzeug, das seine eigenen Ziele auswählt. Es wird zu einer Frage von Legalität und Bedienungskontrolle.

Je nach Konzept dieser Versuchssysteme fallen sie entweder unter die autonomen Waffensysteme oder den nicht von Menschen vorgewählten Zielen und ihre Funktionalität ist noch nicht richtig entwickelt, und weil alles streng geheim behandelt wird, wissen wir nicht, wie sie eigentlich funktionieren.

Der zugrundeliegende Fall ist die Menschenwürde, will sagen, dass die Entscheidung zu töten, von einem Menschen und nicht von einer Maschine ergriffen werden soll. Es ist in die Bedeutsamkeit des Tötens eingebunden. Ein anderer Mensch bestimmt, dass du ein feindlicher Gegner bist, aufgrund deiner militärischen Präsenz auf dem Schlachtfeld, deiner Uniform und deiner Teilnahme an einen bewaffneten Konflikt und solche Kriterien. Das ist ein legaler und moralischer Akt der Tötung, und die Möglichkeit eines würdevollen Todes ist vorhanden. Ist es nur eine Maschine, die einem Algorithmus gehorcht und versucht, einige Kriterien, die die Legalität bestimmen, zu erfüllen, dann ist es immer noch keine menschliche Beurteilung und F findet, dass es keine Garantie für die Menschenwürde ist. Es kann richtig oder falsch vom rechtlichen Standpunkt aus sein, einer Art von Wahrscheinlichkeitsfunktion, aber ultimativ gibt es keinen Plan und keinen Sinn in diesem System, es kann nicht bestimmen, ob eine militärische Notwendigkeit vorliegt. Und da es keine moralische oder rechtliche Beurteilung machen kann, weil es kein rechtliches oder moralisches Agens ist, so kann sein Töten definitiv nicht würdig sein. Es kann gemäß des internationalen Rechts utilitär oder praktisch sein, aber von der moralischen Perspektive aus mangelt es an Würde.

Zu meiner Frage, ob in gewissen Situationen Roboter bessere Entscheidungen treffen als Menschen, weil sie nicht ermüden, nicht emotional werden und nicht Rache suchen, antwortete F, dass dies zwar ein wichtiger Punkt ist, aber dass

wir Systeme entwerfen sollten, die Menschen helfen, bessere Entscheidungen zu treffen und insgesamt Fehler zu reduzieren, doch dass sie all diese Elemente der Bedeutung und Schaffen von Bedeutung auch beibehalten sollten und so dem Menschen helfen, sicherzustellen, dass dies ein moralischer, legaler und bedeutungsvoller Prozess ist, und die Maschine zu benutzen, um zu versuchen, diese Art von Fehlern zu reduzieren und gesteuerte Vorschläge zu erhalten. Das wird Menschen helfen, bessere Entscheidungen zu treffen, wenn sie müde sind oder unter Stress stehen.

3. Freundliche Roboter

Das dritte Bild, ist das der Kampagne, in der der Roboter vor dem britischen Parlament steht. Dies ist ein Bild eines freundlichen Roboters, da F erklärt, dass er Roboter eigentlich gerne mag und dass es militärische Aufgaben gibt, die von Robotern erledigt werden können, aber er betont, dass alle Aufgaben, die moralische und legale Urteilsfällungen erfordern, von Menschen ausgeführt werden müssen.

4. „Dreibeiniger Stuhl“ als Symbol der Landminen-Kampagne.

Ein weiteres gutes Image der Kampagne ist das des sogenannten „Broken Chair“ vor dem UN-Sitz in Genf, der die Landminen-Kampagne symbolisiert. Zu Beginn der Kampagne (Verbot des Killerroboters), versuchten wir, wie bei der Landminen-Kampagne das gleiche zu tun und zu sagen, dass dies extrem unterschiedslose Waffen sind, mit riesigen Auswirkungen auf die Zivilbevölkerung. Das ist äußerst unbehaglich. Wenn man es vergleicht, können Killerroboter als Landminen, die sich bewegen und Ziele wählen, gelten. Es ist eine andere Art, darüber nachzudenken, wie unheimlich sie sind, da sie wirklich nicht viel gescheiter als Landminen sind und die Welt auch nicht besser verstehen, als eine Landmine es tut. Sie sind nur feiner ausgeklügelt, weil sie navigieren und

Daten über die Welt abtasten und ihre Ziele wählen können. Folglich werden sie einerseits nicht dermaßen blind handeln, und man kann den Mangel an Unterscheidungsfähigkeit im Laufe der Zeit mit Technologien verbessern, aber es gibt eine fundamentale Verschiebung, da man die Art der tödlichen Maßnahme und was sie überhaupt rechtfertigen kann, geändert hat, wenn man diese Killerroboter einsetzt. Aber das trifft nicht auf das Militär und die Kriegsführung zu, sondern auf die gesamte Gesellschaft, da wir viele Aufgaben, die früher unter menschlicher Kontrolle standen, automatisieren, und natürlich machen Menschen viele Fehler in Entscheidungstreffungen, es gibt Vorurteile und alle möglichen Neigungen usw. ... die angesprochen werden müssen, aber oft konzipieren wir Automatisierung und geben vor, dass diese Vorurteile nicht länger vorhanden sind, da es sich um ein technisch konzipiertes System handelt, und es dürfte keine Fehler machen. Aber diese Systeme sind kompliziert, und wir wissen nicht, wie sie sich verhalten oder was sie in unerwarteten Situationen tun werden und man könnte einen ganzen Satz neuer Themen zur Sprache bringen, im Hinblick auf die Technik und Probleme, und man würde immer noch keine sozialen Fragen beantworten.

Um einen anderen Vergleich mit Landminen zu finden, gibt es eine gewisse Art von Voraussehbarkeit, doch andererseits auch eine offensichtliche Unberechenbarkeit. Landminen können gerechtfertigt werden, weil sie Truppen während eines Krieges aufhalten, aber 10 Jahre später, wenn jemand sein Land bebauen möchte, können Landminen immer noch explodieren, also kann man die Zukunft einfach nicht kontrollieren. Dies ist das inhärente Problem mit autonomen Waffensystemen, von denen man weiß, dass es losgeht, Ziele sucht und sie bombardiert. Das erscheint eine vernünftige und begrenzte Denkweise, dass es nur ganz spezielle Radarsignale

zerstört, aber wie akkurat ist das? Wie viele Dinge sehen wie das Radarsignal aus? Da sind andere Dinge, die Radarsignale emittieren und damit verwechselt werden können, oder Leute senden Signale aus, die diese Radarsignale imitieren, um eine der Raketen auf andere Dinge zu lenken. Nachdem man es über eine lange Zeit oder über einen breiten Bereich eingesetzt hat, weiß man wirklich nicht, was es zum Schluss tun wird. Man hat eine Ahnung von dem, was das System tun soll, aber man kann es danach nicht länger kontrollieren. Und dann begibt man sich in gefährliche Bereiche, weil es nur um diese Risikoerwartungen geht und wir nicht wirklich wissen, wie man das messen kann und diese Systeme immer komplizierter werden, sodass es immer schwieriger wird, sie einzuschätzen, ob als Bediener, der vorhersehen möchte, was zu tun ist, oder als Gericht bzw. eine Organisation, die jemanden zur Rechenschaft dafür zieht, dass er/sie etwas Schreckliches ausgelöst hat, statt tat nur das Vorhaben gehabt zu haben, etwas Schreckliches zu tun. Also kann man ihn/sie nicht wirklich für Kriegsverbrechen verurteilen, doch stattdessen ist es Fahrlässigkeit, da er/sie nicht wirklich wusste, auf was er/sie sich einließ, aufgrund einiger Umstände von außen, die man nicht voraussehen konnte oder die man nicht kannte. Folglich beginnt unser ganzes System des Verstehens von Rechenschaft und Verantwortlichkeit zu zerbrechen und erzeugt einen tollen Ausweg für Leute, die Schlechtes wollen und dafür zur Rechenschaft gezogen werden, und es belastet Leute, die zwangsverpflichtet wurden oder dem Militär beitraten, als sie 18 waren und zu dem Zeitpunkt nicht wussten, was sie tun würden und deshalb nicht dafür verantwortlich gemacht werden können.

6	Haben Sie noch weitere Kommentare zu diesem Interview?
---	--

Zu meiner Frag an F, ob er fähig wäre, ein Verbot für diese Art

von Technologie durchzusetzen, antwortete er: Es gibt einige Dinge, die man über internationales Recht und Verbote berücksichtigen sollte. Die Durchsetzung ist eines dieser Elemente, aber es ist nicht das einzige Element. Man kann es mit Mord vergleichen. Leute werden sich immer gegenseitig ermorden, sollten wir also Mord nicht gesetzlich verbieten? Nein, wir haben ihn als Straftat eingestuft, weil er böse ist und Leute ihn immer noch begehen, und damit muss man fertig werden. Es geht um die Etablierung einer Norm auf internationaler Ebene, auf der alle Länder dieser Erde zusammentreffen und vereinbaren, dass es grundsätzlich falsch ist, ein autonomes Waffensystem einzusetzen. Demzufolge, was wären die Folgen, wenn Leute dies übertreten. Die meisten Abkommen verfügen über keine Nachweismethoden und ausdrückliche punitive Maßnahmen und solche Sachen, aber es geht darum, Länder zu beschämen, Sanktionen zu erlassen und um internationale Konsequenzen für Länder, die die Abkommen übertreten.

Nichtstaatliche Akteure befinden sich in einem ganz anderen Bereich, und wir haben bereits chemische Waffen verboten, ein guter Vergleich, dass Länder, selbst Länder, die dieses Abkommen nicht unterschrieben, beobachtet und sanktioniert haben und ein Nachspiel erleben müssen, dafür, dass sie chemische Waffen benutzt haben. Es wirkt sich auch auf die Industrie aus, in dem Sinn, dass große militärische Unternehmer nicht aktiv chemische Waffen entwickeln und große chemische Unternehmen keine Riesenmenge von Vorläufern produzieren, die im Rahmen dieser Konvention verboten sind. Man wird Ähnliches mit autonomen Waffensystemen erleben. Natürlich kann man autonome Waffensysteme basiert auf Kram aus Baumärkten zusammenbauen, aber es wird kein schweres militärisches Waffensystem sein, das ganze Städte und Fabriken zerstören

kann. Und Leute produzieren bereits Sprengfallen und Bomben und solche Dinge, und man kann das nicht wirklich aufhalten, aber man kann versuchen, den Zugang zu explosivem Material und Schlüsseltechnologien zu regulieren. Man kann vermeiden, dass sich autonome Waffensysteme von großen Staaten auf kleine Staaten vermehren, die sie nicht mit Einschränkung benutzen oder ihre Software modifizieren und zum Schluss auf nichtstaatliche Akteure und Terroristen erweitern werden, die Zugang zu Systemen mit ernsthaften Fähigkeiten haben und nicht nur ein Spielzeug mit einer angeschnallten Bombe sind. Dies wird folglich mit einem internationalen Abkommen abgeschwächt. Aber es gibt nichts, was die Polizei oder Staaten aufhalten wird, es gegen ihre eigenen Bürger zu benutzen, deshalb braucht man nationale Gesetze und Normen, um den Einsatz dieser Systeme für die Polizei auszuschließen, um zu vermeiden, dass sie sie gegen Demonstranten anwenden. Ich werde wahrscheinlich dem Zivilrecht und dem Strafrecht für Privatpersonen unterstehen, wenn ich diese Systeme bauen würde, aber es wäre gut anzugeben, dass es verboten ist, autonomen Feuerleitanlagen Waffen zuzufügen.

Auf meine Frage, ob dies Kriterien der Transparenz bedeuten würde,
antwortete F: Da Sie die Bewertung hochentwickelter militärischer Systeme ernstnehmen, ist ein erforderlicher Teil für die Demonstration bedeutungsvoller menschlicher Kontrolle die Darstellung der Kommando- und Datenspeicherungskette, sodass wenn jemand erwartet, dass Ihr System vollständig autonom ist, Sie mithilfe des Datenflusses beweisen können, dass eine Zielentscheidung tatsächlich von Menschen getroffen wurden bzw. dass Menschen das Ziel in diesem speziellen Fall für jede Art von Waffensystem genehmigt haben. Es ist leicht, ein anspruchsvolles Informationsverarbeitungssystem zu haben,

die Genehmigung, es zu bearbeiten, ist etwas anderes. Der aktuelle Standard ist, dass Sie diese Disziplin ihrem eigenen Militär auferlegen, sodass Sie nicht unbedingt dementsprechend handeln müssen, aber Sie verfügen über diese Prozedere, sodass Sie Offiziere und Soldaten, wie sich selber auch, für einen Bruch des internationalen Rechts vor das Kriegsgericht stellen können. Sie werden von den Gegnern nicht erwarten, dass sie ein Verfahren dafür einleiten. Sie selbst bewerten Ihre Waffensysteme, folglich ist es ganz und gar nicht offensichtlich, und Sie werden diese Bewertungen nicht veröffentlichen. Aber die Anforderung ist, dass Sie diese Daten für den Fall, das etwas geschieht, speichern, sodass Sie sie zur Hand haben.

Auf meine Frage, dass die Kampagne in 2015 und 2016 stark gefördert wurde, aber dass es nun aussieht, als ob sie verlangsamt und oberer die neuesten Informationen dazu habe, antwortete F: Das Problem ist, dass die Zeit zwischen den VN-Treffen und den Schlagzeilen irgendwo im Dezember lag, dass die Diskussion vom informellen Treffen der Experten auf das formelle GGE erhöht wurde, was einen Schritt vor den Vertragsverhandlungen darstellt. Sie planten 2 Wochen für dieses Treffen, aber sie hatten in diesem Jahr eine Menge Probleme mit der Finanzierung, da sie das Budget-Modell der VN geändert haben, das vorschreibt, alle Beiträge bevor die Treffen abgehalten werden, zu bezahlen, doch während der letzten 50 Jahre gab es ein anderes System, gemäß dem man die Treffen einfach abhielt und wartete, bis die Leute ihre Beiträge später bezahlten. Es steht also zur Debatte, ob die zweite Woche mit den Treffen eigentlich stattfindet, wenn Länder ihre Beiträge nicht vollständig leisten. Das Treffen im August wird stattfinden, und hoffentlich gibt es ein zweites im November, in Verbindung mit ihrem Jahrestreffen. Sie sind der Hoffnung, dass sich etwas aus der GGE entwickelt, aber eine Woche ist

wirklich nicht lange genug und sie hatten auf drei Wochen gehofft, nicht 2 Wochen, die vielleicht auf eine Woche reduziert werden.

Im Allgemeinen hatte die Kampagne einen guten Start, im Vergleich mit den meisten anderen Abrüstungskampagnen, einschließlich Landminen und Streubomben, die tatsächlich 10 Jahre brauchten, bis wir da waren, wo wir jetzt in drei oder vier Jahren sind. Also, diese Art von Schwerkraft beizubehalten wäre sehr beeindruckend gewesen, und es wurde von Staaten herausgezögert, die wohl bereit waren, es zu besprechen, aber weniger interessiert, ihre Einstellungen wirklich zu formalisieren. Es gibt alle möglichen Art und Weisen des Hinhaltens, aber die VN formalisiert ihre Stellung dazu auf jeden Fall, vor allem, da die 3000/09-Richtlinie im Pentagon gerade jetzt, nach 5 Jahren, neu bearbeitet wird. Und sie hoffen, dass die Staaten es gründlicher in der GGE debattieren werden und sich auf die Sprache einigen, da es leicht ist, zu sagen, es gäbe keine Definition von Autonomie und autonomen Waffensystemen, und dass wir wissen, was bedeutungsvolle menschliche Kontrolle ist, aber es liegt an ihnen, zu entscheiden, ob sie diese Begriffe rechtfertigen können, wie gerne sie es auch möchten. Doch für das Abkommen muss der politische Konsens darüber was es für Definitionen sein sollen, gefunden werden. Der ICRC legt sich sehr ins Zeug, um die Definition autonomer Waffensysteme einzuengen und dann ist es eine Frage, bedeutungsvolle menschliche Kontrolle auszuarbeiten und welche Art von Sprache für ein Abkommen die Richtige sein würde, nicht zu freizügig und nicht zu beengend. Es ist eine schwierige Balance und mehr ein politisches als ein technisches Problem. F hofft ziemlich stark, dass sie vorankommen.

Ein Grund, dass die CCW sich so einsetzte, war, dass sie im letzten Jahrzehnt wenig getan haben, und was sie getan hatten,

war irgendwie ungenügend, sodass sie versuchten, es wett zu machen und das ist eine gute Sache. Wenn es bald geschieht, wird es vor allem eine präventive Art von Abkommen sein, aber die vorhandenen Waffensysteme müssen gemäß der jetzt erforderlichen bedeutungsvollen menschlichen Kontrolle neu untersucht werden, doch die meisten Waffensysteme im Einsatz würden diese Erklärung erfüllen und solche, die das nicht können, sollten auf keinen Fall benutzt werden. Es ist leichter es jetzt präventiv zu tun, als in 10 Jahren, wenn Staaten voll-autonome Drohnen und U-Boote entwickelt haben und empfinden, dass sie für ihre Sicherheit notwendig sind, und dann wäre es unmöglich diese Art von Abkommen zu beschließen.

Anhang G. Codierungsnotizen

Wissenschaftlicher Mitarbeiter 1

Beim Codieren von Werten, wird der Text codiert, um die Werte⁸, Einstellungen⁹ und Ansichten¹⁰ (Saldaña, 2015) zu codieren. Diese Konzepte sind als voreingestellte Codes verzeichnet.

List of pre-set codes:

Concept	Colour
Value	Blue
Belief	Green
Attitude	Yellow

List of emergent codes:

Concept	Colour
Definition	Purple

TIPPS:

- Eine voreingestellte Liste braucht nur 10 Codes kann aber auch bis zu 40 oder 50 Codes haben. Wir empfehlen, nicht zu viele Codes zu erstellen, da die codierende Person

⁸ Wert: Die Bedeutung, die wir uns selbst, einer anderen Person, eines Gegenstands oder einer Idee beimessen.

⁹ Einstellung: Die Art, wie wir über uns selbst, eine andere Person, einen Gegenstand oder eine Idee nachdenken

¹⁰ Ansichten/Überzeugung: Der Teil des Systems, der unsere Werte und Einstellungen sowie unsere persönlichen Kenntnisse, Erfahrungen, Meinungen, Vorurteile und andere interpretierende Auffassungen der sozialen Welt einschließen.

überfordert werden oder Fehler im Codierungsprozess machen kann, wenn es zu viele gibt.

- Als Faustregel für das Codieren gilt, die Codes den Daten anzupassen, statt die Daten Ihren Codes.
- Das Erstellen von Notizen während des Codierprozesses ist für grounded und a Priori-Methoden gleich. Qualitative Forschung ist inhärent reflexiv; da die forschenden Personen tiefer in ihr Thema eindringen, ist es wichtig, ihre eigenen Denkprozesse mithilfe von reflektierenden oder methodologischen Notizen festzuhalten, da dies vielleicht ihre eigene subjektive Interpretation der Daten hervorheben könnte. Es ist äußerst wichtig, gleich zu Anfang der Forschung mit den Notizen zu beginnen. Egal welche Art von Notiz produziert wird, es ist wichtig, dass der Prozess kritisches Denken und Produktivität in der Forschung anregt. So etwas zu tun, erleichtert die Analyse und zusammenhängende Analysen, während das Projekt fortgeführt wird. Notizen können für das Ausarbeiten der Forschungsaktivitäten und Entdeckung der Bedeutung der Daten benutzt werden, sie erhalten den Schwung in der Forschung, Engagement und Eröffnungskommunikation.

MORE INFO:

http://onlineqda.hud.ac.uk/Intro_QDA/how_what_to_code.php

https://researchrundowns.files.wordpress.com/2009/07/rrqualcodinganalysis_7_19_09.pdf

http://programeval.ucdavis.edu/documents/Tips_Tools_18_2012.pdf

<https://www.slideshare.net/kontorphilip/qualitative-analysis-coding-and-categorizing>

WEITERE NOTIZEN:

Vor dem Start:

Es ist mein Vorhaben, die Werte-Codierung zu benutzen (Saldaña, 2015), aber Saldaña (2015) definiert den Werte-Begriff anders als Friedman und Kahn Jr (2003). Ich frage mich außerdem, ob der Unterschied zwischen Wert, Einstellung und Überzeugung wirklich deutlich ist, aber ich werde sehen, wie es auslaufen wird. Vielleicht sollte ein weiteres Code-Wort für eine „Definition“ hinzugefügt werden?

Während der Codierung:

- a. Ich bemerke, dass eine Einstellung oftmals ein Verb einschließt, oder so würde ich es jedenfalls interpretieren.
- b. Eine Überzeugung schließt oft eine Ansicht ein.
- c. Ich merke, dass ich oftmals an die Listen der Werte aus meiner Fachliteraturuntersuchung denke, wenn ich einen Begriff hervorheben möchte, zum Beispiel „Reflexivität“. >> Dies kann auf ein Vorurteil als Rechercheur meinerseits weisen, folglich wäre es gut, wenn ich meine Werte-Liste mit dem zweiten Rechercheur überprüfe.
- d. Wenn ein Begriff der befragten Person nicht gefällt, z.B. Determinismus, markiere ich ihn als einen Wert.
- e. Ist etwas viel Erwähntes oftmals wichtig für die befragte Person? Und kann es deshalb als ein Wert interpretiert werden? Zum Beispiel, „Distanz zwischen Mensch und Maschine“ (siehe j).
- f. Beim Codieren gehe ich manchmal zurück, um zu sehen, wie ich den gleichen Ausdruck klassifiziert habe (z.B. „zu unheimlich“) und kopiere die Codierung.
- g. Ich füge einen zusätzlichen Code für „Definition“ hinzu und füge es dann zum Code-Buch hinzu.

- h. Ich habe eine Ansicht einer befragten Person als Überzeugung interpretiert, da es Teil der Definition des Begriffs „Überzeugung“ von Saldaña (2015) ist.
- i. Nach dem Codieren einer Frage eines Interviews, lese ich die Frage noch einmal durch und prüfe meine Codierung.
- j. Wenn ein Begriff mehr als zweimal benutzt wurde, z.B. der Begriff „defensiv“ in Professor Walshes Interview, dann habe ich ihn als einen Wert markiert.
- k. Außerdem markierte ich einen Begriff, der besonders hervorgehoben wurde, z.B. wurde „Vorteile“ In Prof. Walshes Interview markiert.
- l. Danach codierte ich ein neues Interview, ich prüfte die vorherigen, um zu sehen, ob ich konsistent geblieben war.
- m. Wenn etwas als wichtig erörtert wird, z.B. im Interview mit Ad über die Wichtigkeit, dass die der Marine zu jederzeit die Möglichkeit hat, einzugreifen, dann habe ich das als einen Wert codiert.
- n. Nach dem Codieren der letzten drei Interviews, habe ich die ersten drei nicht noch einmal überprüft. Vor allem wegen Mangel an Zeit, aber auch, weil Codierung leichter für mich geworden war.

Nach dem Codieren:

Ich habe alle 6 Interviews an den zweiten Rechercheur geschickt und einige Hintergrundinformationen zu den Codierungsinterviews hinzugefügt. Außerdem fügte ich das Codebuch ohne meine Kommentare hinzu. Wir besprachen die Aufgabe über Skype, um sicherzustellen, dass sie sie verstand.

Codebuch des Korreferenten 2

Hiermit gebe ich die ordentlichen Richtlinien von Mutmaßungen, Regeln oder Schlussfolgerung, die halfen, die Aussagen in Richtung Wert, Einstellung oder Überzeugung zu wenden. Beachten Sie, dass diese Listen während dieses Codierungsverfahrens entwickelt und erweitert wurden, deshalb wurde es zu einem iterativen Verfahren. In anderen Worten, in späteren Interviews habe ich stets über die Richtlinien, die das Codierungsverfahren in früheren Interviews unterstützt haben, nachgedacht und darauf aufgebaut.

Wert: Wenn Aussagen etwas von großer Wichtigkeit wiedergeben, z.B. „wir müssen...“, „müssen unbedingt“, „wichtig“.

- wenn man eine starke Beurteilungen bedenkt, wie „sicherstellen“, „beunruhigend“, „Bedürfnis“, „Menschheit“.
- Wenn sie großartige oder hoch bewertete Wörter enthalten, z.B. „sicherstellen“, „menschliches Leiden“, „unvermeidbar“.
- Wenn etwas als sehr wichtig oder äußerst zweckmäßig wiedergegeben wird, „Notwendigkeit“, „Reflexivität“.
- Wenn es ein wichtiger Teil der Geschichte oder Diskussion oder Darstellung ist, z.B. „Bedürfnis“, „Moralität“.
- Wenn etwas als wertvoll oder als zu berücksichtigen erachtet wird, z.B. „rechtfertigen“, „Menschenrecht“.

Einstellung: Wenn Aussagen Trends oder Richtungen angeben, z.B. „nicht nur...“, „auch...“.

- Wenn persönliche Ansichten oder Präferenzen angesprochen werden, z.B. „ist nicht wichtig“, „wirklich schwierig“, „man möchte“.

- Wenn relative Positionen betrachtet werden, z.B. „ist nicht wichtig“, „wirklich schwierig“, „Unbestimmtheit“.
- Wenn etwas das Urteil beeinflussen könnte, z.B. „nicht für...“, „mehr als“, „wir möchten“.
- Wenn es aus Worten besteht, die Ansichten darstellen oder wiederholt werden, z.B. „gegen“, „vage“, „schrecklich“.
- Wenn man eine Art von Einstellung wiedergibt, z.B. „absichtlich getan“, „sollte sein“, „furchtbar“, „verführt“.

Überzeugung: Wenn Aussagen auf Mutmaßungen basiert sind oder etwas, das als wahr anerkannt ist, betrachtet.

- Wenn basiert auf (Future of Life-Institute) Erwartungen oder Mutmaßungen, z.B. „wir könnten...“, „es wird...“.
- Wenn Ungewisses besprochen (Future of Life-Institute) wird oder die Perspektiven anderer Leute vermutet werden.
- Wenn persönliches Wissen/persönliche Erfahrungen wiedergegeben werden, z.B. „wird geschehen“, „weiß nicht“.
- Wenn erhobene Fragen, Spekulationen oder Mutmaßungen betrachtet werden, z.B. „sie haben keine“, „sie werden“.
- Wenn interpretierende Aussichten oder die Zukunft betrachtet werden, z.B. „sie sind nur“, „sie werden“.

Anhang H. Ergebnisse des Codierungsprozesses

Interviews	Werte		
	<i>Rechercheur</i>	<i>Korreferent(in)</i>	<i>Ähnliche Werte</i>
1. Interview (Person E)	- stets eingreifen - Kontrolle haben - sie zurückrufen - abrufen - kann nicht kontrolliert werden - Kontrollebene	- Wir müssen daran denken, wie wir sie mit Vorsicht benutzen. - Die Tatsache, dass es möglich ist, zu jeder Zeit einzugreifen ist sehr wichtig für die Marine, - aber wir wollen immer die Kontrolle haben. Wir wollen immer fähig sein, sie zurückzurufen, wenn die Situation sich geändert hat. - Oder dass wir fähig sind, sie neu einzusetzen, falls erforderlich. - man möchte fähig sein, sie zurückzurufen - müssen mehr von dieser Art Sicherheiten darin einbauen - muss mitmachen.	- jederzeit eingreifen - Kontrolle haben - sie zurückrufen

2. Interview (Person D)	- außer Kontrolle - Kontrolle - Kontrolle - Risiko - kontrolliert - riskieren - Gutes zu tun - Risiko - unkontrollierbares Risiko - verliert unabsichtlich die Kontrolle - Grenzen - Guter Wille - abgesteckte Grenzen - Aufgaben mit Grenzen - grenzenlos - Guter Wille - Risiko - riskant	- sicherstellen, dass wir sie so verantwortlich wie möglich benutzen, - das Risiko, die Kontrolle über Situationen zu verlieren, minimieren - im militärischen Kontext, ist die Schnelligkeit der Entscheidungsfung wichtig. - ist keine Situation, die wir als Menschheit wollen. - sicherstellen, dass wir nicht über den Klippenrand stürzen - verhindern, dass diese Technologien ihre eigene Vision kreieren - Technologien mit niedriger Hemmschwelle, aber dem Potenzial unkontrollierbarer Risiken sind beunruhigend für die Menschheit. - mögliche Ereignisse erwägen, wenn diese Art Technologie konstruiert wird und was schief gehen könnte, wenn man es nicht täte.	- Risiken minimieren - Kontrolle haben - Guter Wille - abgesteckte Grenzen
------------------------------------	--	--	---

		- brauchen einen starken Sinn des Zusammenhangs und eine Art von gutem Willen. - Auftrag ausgeführt, aber nicht auf die Weise, die wir geplant hatten. - müssen mit ihnen kooperieren, und Zielsetzung wird äußerst wichtig. - System muss wissen, was wir meinen. - müssen unsere Leute schulen, mit diesen Systemtypen zusammen-zuarbeiten und ihnen die Ziele vermitteln - müssen unser Personal darin schulen, Aufgaben mit festen Grenzen zu stellen. - guter Wille - Risiken eingrenzen. - KI mit der richtigen Denkweise bauen und immer noch die Welt zerstören	
--	--	---	--

3. <i>Interview (Person C)</i>	- Entscheidungen entscheiden funktionale Form - Entscheidungen - Entscheidung - Entscheidung - Entscheidungen - Entscheidung - man hat das Recht, von einem Menschen getötet zu werden - menschliches Leiden	- wie stellen wir sicher, dass dies die Grenze ist - man hat das Recht, [von einer Person] getötet zu werden, - Einbau von Autonomie in ein Waffensystem ist unabwendbar - menschliches Leiden zu reduzieren, - die Politik ist sehr interessant - gibt es Ihnen einen Vorsprung auf dem Schlachtfeld? - kann nicht sicher sein, dass die andere Seite keine AWs erwirbt.	- das Recht, von einem Menschen getötet zu werden - menschliches Leiden
4. <i>Interview (Person A)</i>	- Ethischer Rahmen - wir dürfen keine Kontrolle verlieren - menschliche Kontrolle - Distanz - bedeutungsvolle menschliche Kontrolle - menschliche Rechenschaftspflicht - bedeutungsvolle menschliche Kontrolle	- der ethische Rahmen ist sehr wichtig - eine Waffe ohne menschliche Kontrolle bei kritischen Funktionen, z.B. Wahl des Ziels, müsste verboten werden. - bedeutungsvolle menschliche Kontrolle - menschliche Rechenschaftspflicht - menschliche Kontrolle sollte Angewendet werden, wenn ein Ziel gewählt und außer Gefecht	- Ethischer Rahmen - Bedeutungsvolle menschliche Kontrolle - Menschliche Rechenschaftspflicht - Voraussehbarkeit - Reflexivität - Distanz - Unbesiegbarkeit - Determinismus - Unabwendbarkeit - Menschen müssen über Leben und Tod entscheiden können

- bedeutungsvolle menschliche Kontrolle - Distanz - Distanz - Distanz - Distanz - Mangel an Rechenschaftspflicht [und] Unvorhersehbarkeit - es interessiert uns nicht - Reflexivität - Distanz - Unbesiegbarkeit - Determinismus - Unabwendbarkeit - Menschen müssen über Leben und Tod entscheiden. - Menschen müssen in der Entscheidung, Leben zu nehmen, eingeschlossen sein - Menschen müssen Kontrolle haben, über das was eine Maschine tut	gesetzt wird - aktuelle Waffensysteme müssten in Bezug auf Anwendung von menschlicher Kontrolle geprüft werden, - zweckmäßig für etwas konstruiert - erzeugt Distanz oder reduziert Schaden - Furcht, mangelnde Rechenschaftspflicht und Unvorhersehbarkeit - Reflexivität - Distanz - Unbesiegbarkeit - Determinismus - Unabwendbarkeit - menschliches Bedürfnis, über Tod und Leben zu entscheiden - menschliches Bedürfnis, die Kontrolle darüber zu haben, was eine Maschine innerhalb gewisser Parameter anstellt ('Kader'). - Wie menschliche Kontrolle gewährleistet wird
--	--

5. <i>Interview</i> <i>(Person F)</i>	- bedeutungsvolle menschliche Kontrolle - bedeutungsvolle menschliche Kontrolle - bedeutungsvolle Kontrolle - Plan - bedeutungsvoll - bedeutungsvoll - bedeutungsvoll - bedeutungsvolle menschliche Kontrolle - bedeutungsvolle menschliche Kontrolle - bedeutungsvolle menschliche Kontrolle - oberflächliche Verletzungen - unnötiges Leiden - oberflächliche Verletzungen - bedeutungsvolle Kontrolle - Ausrichtung der Automatisierung - bedeutungsvolle menschliche Kontrolle - Kontrolle - Menschenwürde - Bedeutung - würdig	- effektive menschliche Kontrolle - kontext- oder situationsbedingtes Wissen oder Verständnis - bewerten der Legalität oder Moralität - muss und hat die Absicht, oder einen militärischen Zweck, um ein Ziel außer Gefecht zu setzen - bedeutungsvolle menschliche Kontrolle - bedeutungsvolle menschliche Kontrolle - bedeutungsvolle menschliche Kontrolle - Unnötiges Leiden oberflächliche Verletzungen - legale und moralisch gültig - hochentwickelte begründete Systeme - Kontrolle verlieren - legale und aktive Kontrolle - Menschenwürde - legaler und moralischer Akt des Tötens - würdiges Töten - militärisches Bedürfnis - moralische oder rechtliche Beurteilung	- bedeutungsvolle menschliche Kontrolle - unnötiges Leiden - oberflächliche Verletzungen - Menschenwürde - würdiges Töten - Voraussehbarkeit - Kontrolle - Rechenschaftspflicht - Verantwortlichkeit
--	--	--	--

- Würde - Absicht - Bedeutung - würdig - bedeutungsvoll - menschliche Kontrolle - Ausrichtung - Ausrichtungen - Voraussehbarkeit - Unvorhersehbarkeit - Kontrolle - Kontrolle - rechenschaftspflichtig - Rechenschaftspflicht - Verantwortlichkeit - rechenschaftspflichtig - verantwortlich - bedeutungsvolle menschliche Kontrolle - rechenschaftspflichtig - bedeutungsvolle menschliche Kontrolle - bedeutungsvolle menschliche	- stellt sicher, dass [...] moralischer, richterlicher und bedeutungsvoller Prozess - Voraussehbarkeit - Kontrolle - rechenschaftspflichtig - Rechenschaftspflicht und Verantwortlichkeit - Gesetze und Normen - Befehlskette - Autorität - transparent - diese Daten speichern	
---	--	--

	Kontrolle - bedeutungsvolles menschliche Kontrolle - rechenschaftspflichtig - bedeutungsvolle menschliche Kontrolle - bedeutungsvolle menschliche Kontrolle - bedeutungsvolle menschliche Kontrolle		
6. <i>Interview</i> <i>(Person B)</i>	- Verteidigung - verteidigt - verteidigen - verteidigen - unverantwortlich - Vorteile - machtlos - Schaden	- rechtfertigen - Menschenrecht - Effizienz - Barrieren - Gesetze und Regeln - dagegen verteidigen - militärischer Vorteil - sich selbst verteidigen - Genauigkeit - bewerten - sich verhalten - sich auf eine bestimmte Weise verhalten - überprüfen und bewerten - Umwelt - Erwartung, wie es sich verhalten wird - ethische Regelungen	- Verteidigung - Schaden

		- gut definierte Spannen - Schaden - sich so verhalten, wie wir es gerne möchten - man kann nicht darin eindringen	
--	--	---	--

Basiert auf den Interviews und dem Vergleich der Werte-Codierung des Rechercheurs und dem zweiten Korreferenten wurde die folgende Liste der besonderen 23 Werte abgeleitet:

[fly leaf – intentionally left blank]

[reverse of rear fly leaf – intentionally blank]

[inside of back cover – intentionally blank]

In den letzten fünf Jahren findet in der Gesellschaft und bei der Konvention über bestimmte herkömmliche Waffen (CCW) bei den Vereinten Nationen eine Debatte über den Einsatz autonomer Waffen statt. Es wurde jedoch wenig empirische Forschung betrieben, die Aufschluss darüber gibt, wie autonome Waffen von der Öffentlichkeit und dem Militär wahrgenommen werden. In diesem Beitrag beschreiben wir eine empirische Studie, um einen Einblick in 1) die Wahrnehmung autonomer Waffen durch das Militär und die Öffentlichkeit zu erhalten und 2) welche moralischen Werte die Menschen für wichtig halten, wenn autonome Waffen in naher Zukunft eingesetzt werden. Als Forschungsansatz haben wir in unserer Studie die Value-Sensitive-Design (VSD)-Methode verwendet. VSD ist ein dreiteiliger Ansatz, der es ermöglicht, menschliche Werte während des gesamten Designprozesses der Technologie zu berücksichtigen. Es ist ein iterativer Prozess zur *konzeptionellen, empirischen und technologischen* Untersuchung menschlicher Werte, die mit dem Design verbunden sind. Unsere Ergebnisse deuten darauf hin, dass Militärangehörige und Zivilisten, die im niederländischen Verteidigungsministerium arbeiten, einer autonomen Waffe mehr Handlungskompetenz (d.h. die Fähigkeit zu denken und zu planen) zuschreiben als einer vom Menschen betriebenen Drohne. Wir haben auch festgestellt, dass es gerade in der Wahrnehmung der Werte Menschenwürde und Angst Gemeinsamkeiten zwischen militärischen und gesellschaftlichen Gruppen in der Debatte über autonome Waffen gibt. Diese beiden Werte, die im Diskurs oft erwähnt werden, sollten daher bei der Diskussion über die Ethik des Einsatzes autonomer Waffen unbedingt berücksichtigt werden

Diese Arbeit wurde im Rahmen des jährlichen Wettbewerbs der EuroISME um die beste Studienarbeit mit dem ersten Preis des Jahres 2018 ausgezeichnet. Für Informationen über den Wettbewerb besuchen Sie bitte www.euroisme.eu

