

On Designing Collaboration in a Mixed-Methods Scenario. Reflecting Quantitative Drama Analytics (QuaDrama)

Janis Pagel, Benjamin Krautter, Melanie Andresen, Marcus Willand, Nils Reiter

Abstract ‘Quantitative Drama Analytics’ (QuaDrama) is a mixed-methods project that brings together researchers from modern German literature and computational linguistics, thus belonging to the field of computational literary studies. The goal of the project is to define, annotate, automatically detect, and quantitatively analyze different dramatic character types in German-language plays. To this end, we extract textual and structural properties from plays and investigate their distribution among literary characters, such as Romeo and Juliet. One of the key decisions in any mixed-methods project is to agree on a specific collaboration workflow between the different parties, yet this decision is rarely made deliberately and explicitly. In QuaDrama, the collaboration is based on three pillars. (i) Text annotation is used to both clarify concepts and enrich data required in machine learning. (ii) The different project parts perform close and frequent personal collaboration, particularly in the beginning of the project. In addition, quantitative analyses are made by researchers from both disciplines by using generic programming tools (thus avoiding the need for custom graphical user interfaces). (iii) Quantitative and qualitative research methods go hand in hand and are closely integrated. We discuss the reasoning behind these pillars in this article and provide some recommendations for projects with a similar setup.

Introduction

This article focuses on organization and collaboration in ‘Quantitative Drama Analytics’ (QuaDrama), a project between researchers from literary studies and computational linguistics (CL).¹ QuaDrama aims to extend the possibilities for large- and small-scale analysis of plays by using quantitative and formal methods, while focusing on dramatic characters as objects of investigation.

1 The scientific contributions of the project can be found on its web page: <https://quadrma.github.io>.

One of the biggest challenges in this collaboration is that research questions and presumptions from literary studies cannot be put into a computer directly, as they are usually context-dependent, example-based, and not aimed at quantification. In other words, they are defined in a less formalised way than would be needed for computational processing.² This gap between notions in literary studies and technical possibilities needs to be bridged by means of operationalization: Concepts, questions and phenomena need to be made measurable.³ This step involves conceptual decisions and, therefore, it is more than a simple translation process. In addition, the results that are generated using algorithmic methods need to be 're-mapped' into the field of literary studies. The challenge of operationalization has two directions: From the humanities' conceptual realm into the formal realm and back. We postulate that both directions are interpretative processes. An algorithm formalizing the detection of a type of text is an interpretation of this type of text; mapping quantitative results (that are generated, e.g., in the form of bar charts) to humanities findings requires not only an interpretation of the numbers and their visual representation, but also simultaneous consideration of relevant contexts. These processes can fail for numerous reasons. We argue, however, that many of these reasons do not lie in technical or conceptual issues, but are rather organizational, pragmatic, or communication problems.

Operationalization, in this sense, can be thought of as being hierarchical: A target concept is operationalized in terms of more basic concepts that, in turn, are operationalized in even more basic concepts. Take, for instance, the 'protagonist' or 'main character' of a play. Operationalization of this concept consists of a definition of properties that are relevant for detecting a protagonist, such as the importance of the character for the plot, its relations to other characters or the themes they talk about. These properties themselves need to be operationalized: Themes in character speech could be operationalized through topic modeling or word field analysis, character relations could be operationalized through social networks. This way, an operationalization hierarchy is established, thus connecting the target concept to the textual surface.

Due to this hierarchical structure, operationalizations also provide information on a more fine-grained level: investigating the distribution of properties 'below' the

2 Cf. Jan Christoph Meister, "Computerphilologie vs. 'Digital Text Studies'", in: Christine Grond-Rigler and Wolfgang Straub (eds.): *Literatur und Digitalisierung* (Berlin, Boston: De Gruyter, 2013), 294.

3 See Axel Pichler and Nils Reiter, "Reflektierte Textanalyse", in: Nils Reiter, Axel Pichler, and Jonas Kuhn (eds.), *Reflektierte Algorithmische Textanalyse* (Berlin, Boston: De Gruyter, 2020), 45–47 and Axel Pichler, and Nils Reiter, "Zur Operationalisierung literaturwissenschaftlicher Begriffe in der algorithmischen Textanalyse: Eine Annäherung über Norbert Altenhofers hermeneutischer Modellinterpretation von Kleists *Das Erdbeben in Chili*", *Journal of Literary Theory* 15, no. 1–2 (2021): 4–7.

actual target concept may be just as revealing as looking at the target concept itself. Instead of only knowing that a character is a protagonist, we are *also* aware of the (operationalized) underlying property values and can compare them:⁴ the number of words a character utters, the number of relations they have in a copresence network,⁵ or the distribution of different topics they speak about. This feature-based description is more specific and tries to be more transparent than the holistic and interpretative notion of ‘protagonist’ as it is used in literary studies; as a result, it provides new opportunities to pose research questions.⁶

Summarizing our experiences in QuaDrama, we argue that digital humanities (DH) projects work better if the interdisciplinary collaboration is thought through in detail and constantly evaluated as well as adapted. We, therefore, describe the three pillars that our collaboration relies on and give insight into our reasoning behind them. In the next section, we provide an overview of the project itself and its research agenda. We then discuss the three pillars individually: (i) annotation as a means to enable interdisciplinary discussion, (ii) our collaboration model, i.e., the setup of our tools and workflows, and (iii) a close integration of quantitative and qualitative research. Finally, we provide conclusions.

Quantitative Drama Analytics (QuaDrama)

The project QuaDrama develops methods for the quantitative analysis of dramatic texts. The methods are integrated into a research framework that combines the computing power of digital tools in CL with established qualitative methods of literary

-
- 4 For instance, in Lessing's *Minna von Barnhelm* (1767) Minna's maid Franziska is the play's character who is most frequently present on stage: she appears and speaks in 32 of the 56 scenes. The title character Minna, however, speaks only in 25 scenes. Still, Franziska would not be regarded as the protagonist of the play. A quantitative indication for this is Franziska speaking a lot less than Minna. Franziska's character speech comprises only 3,904 tokens, while Minna's comprises 7,347. See also Manfred Pfister, *The Theory and Analysis of Drama*, transl. John Halliday (Cambridge et al.: Cambridge University Press, 1988), 165–66.
 - 5 Copresence networks are networks based on literary characters that are on stage at the same time. Such social networks have many different applications (see Fazli Can, Tansel Özyer and Faruk Polat (eds.), *State of the Art Applications of Social Network Analysis* (Cham: Springer, 2014)), but their use is not without pitfalls. See Katharina A. Zweig, *Network Analysis Literacy. A Practical Approach to the Analysis of Networks* (Wien: Springer, 2016). One might expect protagonists to have a central position in the network, interacting with many of the other characters.
 - 6 It may not be directly obvious that exactness is achievable or even desirable when applied to humanities concepts. We argue that this depends on the specific context in which such concepts are used. Concepts that are primarily descriptive need to be as exact as possible, while the same does not necessarily hold true for interpretations of literary texts.

studies. In this way, we gain a better understanding of drama history, its contexts, and the interpretation of individual plays. While structure-oriented drama research is straightforward, and has been conducted extensively, e.g., in the form of network analysis going back to the early 1970s,⁷ QuaDraMA combines the analysis of the structure of the plays with that of their content.

The corpus we are investigating is the German Drama Corpus (GerDraCor)⁸ which comprises more than 590 German plays, mainly from 1730 to 1930. Our core method to generate insight is a statistical view on the literary data, properly interpreted with the background knowledge that has been established in literary studies, and, if possible, used as a backdrop for traditional close reading. The project, thus, is ‘mixed-methods’ in the sense that quantitative and hermeneutic methods are intertwined and used side by side, consequently allowing us to evaluate their results in comparison. As not all relevant textual or structural properties are directly accessible on the text surface, we employ methods from computational linguistics—based on machine learning or hand-crafted rule sets—to assign relevant properties such as themes characters talk about on a large scale. In addition to this use as a pre-processing step, machine-learning techniques are used to gain insight into the target categories, e.g., the different character types in a play. This is done by inspecting important features in a classification system or by doing an in-depth error analysis of automatic systems.

In comparison to traditional literary studies, quantitative methods add a new layer of analysis to the spectrum of hermeneutical reading. Although a pluralism of methods does exist in literary studies, these methods are to a great extent qualitative and derive from theories focusing on different entities (author, text, or reader) as their main reference for interpretation. For the practices of literary studies, quantitative methods and the emerging data-based text knowledge can support the understanding of single literary texts and, on a larger scale, literary history, including the formation of literary movements, genres, or conditions of production and reception.

Computational linguistics, the D-discipline in this DH-collaboration, is in general concerned with detecting linguistic properties in texts (and other forms of human language expressions). CL has made tremendous progress in the past 10 years, mostly due to the availability of large data sets, efficient hardware, and machine-learning algorithms that make use of improved computing power (in particular, artificial neural networks). At the same time, most of the progress in CL is made on

7 See Solomon Marcus, *Mathematische Poetik*, transl. Edith Mândroiu (Frankfurt a.M.: Athenäum Verlag, 1973), 287–370.

8 Frank Fischer et al., “Programmable Corpora. Die digitale Literaturwissenschaft zwischen Forschung und Infrastruktur am Beispiel von DraCor”, in *Abstracts of DHd*, ed. Patrick Sahle (Frankfurt, Mainz, 2019), 194–197, <https://doi.org/10.5281/zenodo.2596095>.

English newspaper texts, because ‘news’ is considered to be the default domain and English a world language.⁹ While applying CL tools on other textual domains, the drop in their performance can be quite significant,¹⁰ and other languages are much less supported. While the progress is impressive and the result of hard work, developing CL methods only on newspaper texts leads to models that lack linguistic properties because they simply do not exist in this genre. For the processing of (German) dramatic texts, this applies to the way plays are layered into uttered tokens, stage directions and other kinds of ‘paratext’ (segmentation markers, speaker designations, etc.). Crucially, these three layers cannot be considered as three separate texts and processed independently, because a number of linguistic relations can be discerned between them. The excerpt below, taken from Lessing’s *Emilia Galotti* (1772), gives an example for this phenomenon. The underlined words are part of a coreference chain¹¹ that starts within the Prince’s utterance, but continues within the stage direction associated with Conti, and is then referenced by Conti in his utterance. It shows a linguistic continuity that cuts across two different textual layers, also revealing that a proper linguistic modeling of these texts needs to take the different layers and the specificities of the text genre into account.

As we hope has become clear, the research questions in QuaDrama require a close integration of literary studies and computational linguistics, and therefore a close collaboration between researchers of these disciplines. The next sections describe guiding principles that have made this collaboration constructive for us.

[Prince and Conti discuss two paintings that Conti painted]

Prince. ‘Tis true, but why did you not bring it a month sooner? Lay it aside. What is the other?

Conti. (*taking it up and holding it still reversed*). It is also a female portrait.¹²

-
- 9 The focus on a single domain has enabled the recent boost in prediction quality because it keeps one ‘variable’ controlled while doing experiments. Another reason why news (and, increasingly, social media texts) are in focus is that the amount of text in need of processing increases daily and will continue to do so. This justifies a large investment of resources, because every increase in performance will pay off in the long run. The situation is very different if the text corpus is limited and cannot be expected to grow in the future.
- 10 See Nils Reiter, *Discovering Structural Similarities in Narrative Texts using Event Alignment Algorithms* (PhD diss., Heidelberg University, 2014), <https://doi.org/10.11588/heidok.00017042> and Benedikt Adelman et al., “Evaluating Part-of-Speech and Morphological Tagging for Humanities’ Interpretation”, in: Andrew U. Frank et al. (eds.): *Proceedings of the Second Workshop on Corpus-Based Research in the Humanities* (Wien, 2018), 5–14, <https://www.oeaw.ac.at/fileadmin/subsites/academiaecorpora/PDF/CRH2.pdf>.
- 11 In coreference annotation, every mention of an entity (e.g., person, object, abstract concept) is marked in the text, be it a proper name, a noun phrase or a pronoun.
- 12 Taken from Lessing’s *Emilia Galotti* (1772). The excerpt shows a coreference chain (underlined) that crosses the text layers (transl. by Ernest Bell, 1878).

Collaboration Model

1. Text Annotation

In a mixed-methods scenario, many different traditions, methods, and backgrounds meet and need to interact in a meaningful and efficient way. For QuaDrama, it means finding a way to collaborate between researchers from modern German literary studies and CL. One pillar of our solution is to use manual *annotation* as a layer of communication¹³ that allows us to talk about data, concepts, theories, and, of course, elements of literary texts such as character types.

Annotation is the process of enriching data with overt information in the form of markup, which makes underlying, implicit information¹⁴ explicit.¹⁵ In CL, annotation and the use of annotated data has a long tradition.¹⁶ E.g., computational linguists will mark all occurrences of words in a text and assign a specific part of speech to each word. These markups have (at least) two benefits: (i) The created data can be used to inform and train automatic models, and (ii) the process of annotation discloses difficult cases for assigning markup (e.g., deciding which part of speech a certain word should receive) and thus sparks discussions about the underlying theory and guidelines used to annotate the phenomenon in question. It relates to the notion of operationalization that we have outlined above. In order to create sensible annotations, the target concepts need to be operationalized. At the same time, annotation is one possible option to make the target concepts measurable. This problem of interdependence can be solved by iteratively annotating data and including the insights from the annotations into new or refined operationalizations which in turn inform new annotation decisions.

13 An observation that we have made is that in order to establish a healthy communication between different disciplines, a chicken-and-egg problem needs to be solved: On the one hand, annotations are a gateway for developing a common language but, on the other, they are a realm for which a common language is a prerequisite. This problem can be solved by first creating annotations 'blindly' and then using them to establish mutual viewpoints. In this sense, annotations are a concrete object for the two parties to refer to while discussing phenomena and assumptions about the texts.

14 Implicit information can be located relatively overtly on the text surface, e.g., when appearances and exits of characters are annotated in a play. Usually, however, implicit information is not directly part of the text, but requires theoretical assumptions which provide categories that can then be mapped to textual properties, e.g., when the sentiments of character speech are annotated.

15 See James Pustejovsky and Amber Stubbs, *Natural Language Annotation for Machine Learning: A Guide to Corpus-Building for Applications* (Sebastopol/Boston/Farnham: O'Reilly Media, 2012).

16 See Eduard Hovy and Julia Lavid, "Towards a Science of Corpus Annotation: A New Methodological Challenge for Corpus Linguistics", *International Journal of Translation Studies* 22, no. 1 (2010): 13–36.

In our project, we realize this iterative process in the following, partially overlapping steps: Firstly, we use the process of annotation to condense the knowledge of literary studies about certain aspects of dramatic theory, and thus make it more concrete. This approach of concretization is used in numerous other DH projects.¹⁷ It can also be exemplified with our efforts to annotate protagonists of plays.¹⁸ We collected theories and characteristics of what it means to be a play's protagonist drawing on both primary sources (e.g., Aristotle's *Poetics*) and secondary literature. Next, we condensed the given descriptive observations or normative instructions into an annotation framework, such that independent annotators could mark characters of plays by using our guidelines. In doing so, we were able to establish which properties protagonists should share most of the time. Based on the annotated data we can start to investigate the different properties that can be individually identified in a given text. How much do protagonists talk? Who do they talk to? What do they talk about? Deriving from general and abstract conceptualizations, the annotation allows us to directly compare characters that we have identified as protagonists. Finally, the resulting data, i.e., the information about which character in a play can be classified as a protagonist by which features, can be used to conduct further research, e.g., by automatically classifying protagonists and determining which properties of a character are most useful for classification decisions. Furthermore, it is possible to use the annotations as input for other classification tasks, such as identifying character types.¹⁹

Text annotations can be represented digitally in a multitude of ways. We opted for two main formats that allow us to cover the needs of the two research areas in QuaDrama: TEI and CoNLL. The TEI²⁰ format is an XML²¹ standard that describes a rich catalog of categories for encoding and marking surface and semantic aspects of texts. It is popular in the DH and computational literary studies community and as such perfectly suited our purpose. A well-known format in CL that is used to encode

17 See Evelyn Gius and Janina Jacke. "The Hermeneutic Profit of Annotation: On Preventing and Fostering Disagreement in Literary Analysis", *International Journal of Humanities and Arts Computing* 11, no. 2 (2017): 233–254, <https://doi.org/10.3366/ijhac.2017.0194> and Nils Reiter, Marcus Willand and Evelyn Gius, "A Shared Task for the Digital Humanities Chapter 1: Introduction to Annotation, Narrative Levels and Shared Tasks", *Cultural Analytics* (2019), <https://culturalanalytics.org/article/11192>

18 See Benjamin Krautter et al., "Titelhelden und Protagonisten – interpretierbare Figurenklassifikation in deutschsprachigen Dramen", *LitLab Pamphlets* 7 (2018), https://www.digitalhumanitiescooperation.de/wp-content/uploads/2018/12/po7_krautter_et_al.pdf.

19 See Krautter et al., "Titelhelden und Protagonisten"; Janis Pagel et al., "Annotation als flexibel einsetzbare Methode", in: Nils Reiter, Axel Pichler, and Jonas Kuhn (eds.): *Reflektierte Algorithmische Textanalyse* (Berlin, Boston: De Gruyter, 2020), 125–141.

20 Text Encoding Initiative, <https://tei-c.org/>.

21 Extensible Markup Language.

linguistic annotations is the CoNLL²² format. It has been used for many shared tasks and constitutes a column-based format where each row represents a token and the columns represent linguistic information such as parts of speech, lemma, or syntactic information. As we can convert information between these formats and other formats if needed, we ensure that our annotations can be used by a variety of different communities and for different purposes.

One type of annotation that poses unique challenges is that of coreference (see Figure 1 for an example). The annotation of coreference is often performed on short texts such as newspaper or Wikipedia articles. Dramatic texts, however, are much longer and contain far more entities, which makes existing workflows for annotating coreference unapplicable. Thus, we created our own annotation tool, CorefAnnotator,²³ that allowed us to manually perform coreference annotation on longer texts and to display the texts in a suitable format. We also established our own guidelines and a specific workflow.²⁴ The annotation tool CorefAnnotator supports the import of dramatic texts directly from TEI/XML and can export both to the aforementioned CoNLL format and TEI/XML. Internally, CorefAnnotator uses the Apache UIMA framework to represent annotations.²⁵

Coreference annotation also reveals ambiguities and uncertainties with respect to literary interpretation. Oftentimes, several annotations are valid since the text itself offers multiple possible readings. These cases are to be distinguished from ‘false’ ambiguities that are caused by annotation mistakes or improper guidelines.²⁶ In the case of valid ambiguities, we ask annotators to mark the different possibilities.

We have published a first version of a corpus with coreference annotations on dramatic texts,²⁷ called *GerDraCor-Coref*. The corpus release allows us to document a milestone within the project and enables further subsequent research, as the corpus can now be cited and used as a resource for various investigations, such as personality traits of characters or the function of abstract entities in drama. As mentioned

22 Named after the ‘Conference on Computational Natural Language Learning’, which first used it.

23 <https://github.com/nilsreiter/CorefAnnotator>.

24 See Ina Rösiger, Sarah Schulz and Nils Reiter, “Towards Coreference for Literary Text: Analyzing Domain-Specific Phenomena”, in: *Proceedings of the Second Joint SIGHUM Workshop on Computational Linguistics for Cultural Heritage, Social Sciences, Humanities and Literature* (Santa Fe, 2018), 129–38, <http://aclweb.org/anthology/W18-4515>.

25 For a detailed description, see Nils Reiter, “CorefAnnotator – a New Annotation Tool for Entity References”, in: *Abstracts of EADH: Data in the Digital Humanities* (Galway, 2018), <https://doi.org/10.18419/opus-10144>.

26 See Gius and Jacke, “The Hermeneutic Profit of Annotation”.

27 Janis Pagel, and Nils Reiter, “GerDraCor-Coref: A Coreference Corpus for Dramatic Texts in German”, in: *Proceedings of the 12th Language Resources and Evaluation Conference (LREC)* (Marseille, 2020), 55–64, <https://www.aclweb.org/anthology/2020.lrec-1.7.pdf>.

before, plays contain much longer coreference chains on average, which we can now quantify in our corpus (328.8 mentions in a coreference chain on average, compared to 14.4 in newspaper data²⁸). This shows the value that annotations as a stand-alone task have while simultaneously fostering further research.

2. Workflow

Organizing an interdisciplinary collaboration on a daily basis is never easy, as all DH-experienced researchers are aware of. The collaboration model in QuaDrama consists of a technical and a non-technical part, and both have been carefully and consciously developed. This section discusses the core ideas of our collaboration model, namely to a) focus on programming libraries instead of tools, b) modularize the technical workflow such that its steps have clearly defined interfaces, and c) provide external tutorials and workshops.

In early stages, we envisaged an interactive analysis tool with a graphical user interface and made some experiments in that direction. It quickly turned out to be unfeasible: The significant development time between ‘having an idea for an analysis’ and ‘inspecting the results of this analysis’ made an explorative approach quite difficult and, at the same time, prohibited serious research on the technical side because the computational linguists in the project would be busy developing and refining the user interface all the time. For the scholars in the humanities, this would also lead to fundamental dependence: They would never be able to conduct analyses on their own and therefore would always need help from a ‘technician’. This experience led us to drop the tool idea entirely. Instead, we focused on programming libraries.

A programming library can be described as a collection of functions that work well together for a clearly defined and limited functional area.²⁹ To use them, one needs to write programming code to execute the functions and do something with the results. Most of the analysis functions we are using on dramatic texts are collected in the library ‘DramaAnalysis’³⁰, developed for the programming language called ‘R’. The library is available on the standard repository for R packages (Comprehensive R Archive Network, CRAN) which makes the installation quite easy. In

28 Since newspaper texts are usually much shorter than plays, we have also normalized the values by text length, but are still getting higher values for plays on average (0.0089 for plays, 0.0013 for newspaper data).

29 For instance, the Python library ‘scikit-learn’ provides functions for machine learning, the Java library ‘Apache Commons CSV’ provides functions for handling CSV files etc. Programming libraries are an essential part of programming and complex programs may use hundreds or thousands of libraries.

30 <http://doi.org/10.5281/zenodo.597509>.

addition to our own project, the package is also used by others to analyze and visualize dramatic texts.³¹

Functions provided by DramaAnalysis can easily be used in RStudio, an integrated development environment (IDE) and a data science platform for the R programming language. Basic knowledge of the programming language R is, however, required to use it: The main code that calls the functions of the package and sets its parameters still needs to be written by the user. We believe that it is an acceptable requirement since it is a skill that will benefit the non-programmers beyond the project, making them more independent in the long run.³² To soften the start, the package provides documentation with examples, a report function to easily generate a comprehensive report for a single play, and a 100-page tutorial with examples, guides and interpretation hints.³³

With respect to visualization, this setup allows one to make use of the broad visualization capabilities of R. As R is oriented towards data analysis and statistics, the visualization methods we use are bar charts, box and scatter plots, and node-link diagrams for networks. As they are generated on the fly using R code, they can be manipulated easily, offering a simple way of interactive visualization by adapting the code directly.

This shift from tools to libraries reflects an adjustment in our ‘mental’ setup: In contrast to a common, seemingly natural division of labor, we do not see literary scholars primarily as users and computational linguists as developers. In our experience, it amounts to collaboration which is equally beneficial for both parties.

Modularization is another concept that promotes independence among researchers. It refers to the idea that a software project is subdivided into modules, each with their own goal and ability to operate largely independent of another. The interaction between the modules is regulated by a clear definition of what the input and output of each module are, so that one module can further process the output of another. The technical setup in QuaDrama consists of two main modules, reflecting the research interests and goals of the involved parties. The first module includes the language processing and is mostly the working field of the computational linguists. The second is the subsequent quantitative analysis of dramatic texts and is mostly done by the literary scholars. Both modules and their interaction will be briefly described in the following.

31 See Hanno Ehrlicher et al., “Gattungspoetik in quantitativer Sicht. Das Werk Calderón de la Barca und die zeitgenössischen Dramenpoetiken”, *LitLab Pamphlet* 9 (2020): 1–29. https://www.digitalhumanitiescooperation.de/wp-content/uploads/2020/04/p09_ehrlicher_lehmann_al_de.pdf.

32 In contrast, a project-specific tool will quickly be disbanded once the project is finished and to continue to be familiar with it will be of little use.

33 <https://quadrma.github.io/DramaAnalysis/tutorial/3/>.

Our corpus is encoded in TEI-XML and provides a machine-readable structuring (in acts/scenes) as well as speaker designations and separable stage directions. In our pipeline, the texts are processed by the custom-built DramaNLP library³⁴ which is built on top of Apache UIMA.³⁵ All the information in the original XML files is initially mapped to UIMA data structures. The processing pipeline automatically adds linguistic information such as parts of speech, lemmas, syntactic parse trees, and coreference chains. As the latter is further developed within the project, its pipeline component is updated from time to time. The final component of DramaNLP is an export of the data into several comma-separated values (CSV) files which are stored in a Git repository on GitHub, containing information about the segmentation of the plays, the spoken tokens, the tokens in stage directions, the characters, and meta-data about the play. The use of Git provides a straightforward versioning mechanism, such that earlier data versions and experiments based on them can be reproduced easily. The CSV files serve as an interface to the second module comprising the quantitative analysis with R.

The second part of the technical setup consists of the R library DramaAnalysis which reads the CSV files and offers analysis and visualization functions. In order to make access and navigation of the package's capabilities easier, we have created a tutorial which interested scholars can work through. Starting from loading a specified play (with the function `loadDrama()`), one can investigate themes in character speech by using word fields (with the function `dictionaryStatistics()`), access character statistics (function: `characterStatistics()`), and access utterance statistics (`utteranceStatistics()`). Also, one can inspect configurations (`configuration()`), that is, co-occurrence matrices),³⁶ get information on active and passive presence of characters (`presence()`),³⁷ calculate different metrics that measure the rate of stage personnel exchange over scene boundaries (`hamming()`, `scenicDifference()`)³⁸ and calculate a general frequency matrix of words that can, for instance, be used in the stylometry library 'stylo'.³⁹

34 <http://doi.org/10.5281/zenodo.597866>; <https://github.com/quadrada/DramaNLP>.

35 UIMA is a processing library that provides data structures and efficient indexing for handling annotations on textual data. <https://uima.apache.org>.

36 See Pfister, *The Theory and Analysis of Drama*, 171–176.

37 See Marcus Willand et al., "Passive Präsenz tragischer Hauptfiguren im Drama", in: Christof Schöch (ed.): *Abstracts of DHd* (Paderborn, 2020), 177–181, <https://zenodo.org/record/3666690>.

38 See Frank Fischer et al., "Network Dynamics, Plot Analysis. Approaching the Progressive Structuration of Literary Texts", in *Digital Humanities 2017. Conference Abstracts* (Montréal, 2017), 437–441, <https://dh2017.adho.org/abstracts/DH2017-abstracts.pdf>.

39 Maciej Eder, Jan Rybicki and Mike Kestemont, "Stylometry with R: A Package for Computational Text Analysis", *R Journal* 8, no. 1 (2016): 107–21. <https://doi.org/10.32614/RJ-2016-007>.

This modular structure of the project ensures that project partners can, for the most part, work independently of each other as long as the interface remains stable, and thus makes smooth collaboration possible.

In addition, we consider dissemination activities geared towards the ‘outside’ to be a part of the collaboration model that has worked well in our project. These activities are mostly workshops and tutorials held on various occasions.⁴⁰ While such activities are obviously beneficial for an external view on a project, they also serve an important internal goal: While preparing such an activity, the team is forced to condense the important aspects into a short period of time and to restructure it in a way that is useful for others. Thus, the preparation fosters the uncovering of misunderstandings and misconceptions, which is very useful for the project itself. Finally, such milestones also encourage everyone to provide proper documentation and code structures.

3. Integration of Quantitative and Qualitative Analysis

The following section discloses our reasoning in combining qualitative and quantitative methods for the analysis of drama. This is the third pillar of our collaboration and it ensures that our two disciplines work together, instead of next to each other, to benefit from each other's insights. The employment of quantitative methods has been hotly debated in literary studies and we, therefore, start by providing a brief overview of the ongoing methodological debate on distant and scalable reading. Thereafter, we describe three examples of our research that illustrate the integration of quantitative and qualitative methods in QuaDrama.

In the early 2000s, ‘distant reading’, a term first coined by Franco Moretti in his essay *Conjectures on World Literature* (2000), sparked a methodological controversy about the analysis of literary history. Moretti, in the light of his diagnosis that literary history should not solely rely on its canonical fracture,⁴¹ but also on what Mar-

40 Tutorial at the Heidelberg summer school on quantitative drama analytics, 2019; Tutorial in a research seminar at Tübingen University, 2020; Operationalization workshop at the German DH conference (DHd), 2020 and 2022; Tutorial in a research seminar at the University of Bochum (2021); Operationalization workshop at the international DH conference, 2022.

41 The literary canon is a compilation of texts that are recognized as exemplary and are, therefore, considered particularly worthy of remembrance. Consequently, the canon is highly selective and excludes major parts of published literature. Moretti, for instance, projects that the literary canon of nineteenth-century British novels would consist ‘of less than one per cent of the novels that were actually published’. Franco Moretti, “Graphs, Maps, Trees: Abstract Models for Literary History”, *New Left Review* 24 (2003): 67. For researchers such as Matthew Jockers, relying on the literary canon to study literary history is not the way forward: ‘Just as we would not expect an economist to generate sound theories about the economy by studying a few consumers or a few businesses, literary scholars cannot be content to read literary history from a canon of a few authors or even several hundred texts’. Matthew L.

garet Cohen called the “great unread”,⁴² polemically suggested embarking on “second hand” criticism instead of close reading of literature.⁴³

“the trouble with close reading (in all of its incarnations, from the new criticism to deconstruction) is that it necessarily depends on an extremely small canon. This may have become an unconscious and invisible premise by now, but it is an iron one nonetheless: you invest so much in individual texts *only* if you think that very few of them really matter. Otherwise, it doesn't make sense. And if you want to look beyond the canon [...] close reading will not do it.”⁴⁴

His proposition to do ‘distant reading’ back then was based on mostly two ideas that were rooted in his philosophical conviction of a materialistic theory of history: (i) reading secondary literature instead of literature itself, i.e., analyzing other researchers’ analyses, and (ii) activating research networks with experts for different languages, genres and time spans.⁴⁵ His ambitious endeavor to analyze world literature, in the end, depends on merging as much existing knowledge as possible for as many literary texts as possible. Thus, the focus shifts from understanding and interpreting a single literary text to the explanatory power of patterns to explain literary conventions and their historical development.⁴⁶ While doing ‘distant reading’, “the reality of the text undergoes a process of deliberate reduction and abstraction”⁴⁷ for the bigger picture to emerge: ‘it allows you to focus on units that are much smaller or much larger than the text: devices, themes, tropes—or genres and systems.’⁴⁸

Back in 2000, Moretti neither referred to computational approaches, nor did he specifically mention quantitative methods to analyze literature. Later, in an essay called *Graphs, Maps, Trees* (2003), which became part of an essay collection with the same title, Moretti decidedly called for the construction of abstract models imported from theories of other disciplines, namely quantitative history, geography and evolutionary theory.⁴⁹ The resulting graphs, maps and trees have drawn a first connec-

Jockers, *Macroanalysis. Digital Methods and Literary History* (Urbana/Chicago/Springfield: University of Illinois Press, 2013), 9.

42 Margaret Cohen, *The Sentimental Education of the Novel* (Princeton: Princeton University Press, 1999), 23 and Margaret Cohen, “Narratology in the Archive of Literature”, *Representations* 108, no. 1 (2009): 59.

43 Franco Moretti, “Conjectures on World Literature”, *New Left Review* 1 (2000): 57.

44 Ibid.

45 See Ibid., 58–60.

46 See Franco Moretti, *Graphs, Maps, Trees: Abstract Models for Literary History* (London/New York: Verso, 2005), 91.

47 Ibid., 1.

48 Moretti, “Conjectures on World Literature”, 57.

49 See Moretti, “Graphs, Maps, Trees”, 2003, 67.

tion of 'distant reading' to computational methods that has deepened remarkably in the following years, mainly since 2010 when the Stanford Literary Lab was formed.⁵⁰

Nowadays, 'distant reading' is pervasively used as a metaphor for computational quantitative corpus analyses.⁵¹ However, its polemical introduction more than 20 years ago and its preference for reductionist models have since induced many critics to fundamentally object against analyzing art by calculating,⁵² against the lack of innovative and new insights compared to the massive efforts needed,⁵³ and against the nature of computational methods from a perspective of intellectual work ethics: quantitative analysis of art was perceived as solely exploratory, as 'mere play'.⁵⁴ All these criticisms can be understood as central problems of the integration of quantitative and qualitative methods.⁵⁵

Martin Mueller, a proponent of computational methods praising the "second-order query potential of digital surrogates",⁵⁶ is also one of the critics of Moretti's 'distant reading' idea. For Mueller, 'distant reading' did not "express adequately the powers that new technologies bring to the old business of reading", as the term

"implicitly set[s] the 'digital' into an unwelcome opposition to some other—a trend explicitly supported by the term 'Digital Humanities' or its short form DH,

50 See Ted Underwood, "The Stanford Literary Lab's Story", last modified February 11, 2017, <https://www.publicbooks.org/the-stanford-literary-labs-narrative/https://www.publicbooks.org/the-stanford-literary-labs-narrative/>.

51 See Ted Underwood, "A Genealogy of Distant Reading", *Digital Humanities Quarterly* 11, no. 2 (2017): § 7–11 and Thomas Weitin, Thomas Gilli and Nico Kunkel, "Auslegen und Ausrechnen. Zum Verhältnis hermeneutischer und quantitativer Verfahren in den Literaturwissenschaften", *Zeitschrift für Literaturwissenschaft und Linguistik* 46 (2016): 104. Peer Trilcke and Frank Fischer point out that Moretti, in his recent publications, no longer uses the term 'distant reading'. Instead, he would refer to 'computational criticism' or 'digital humanities'. See Peer Trilcke, and Frank Fischer, "Fernlesen mit Foucault? Überlegungen zur Praxis des *distant reading* und zur Operationalisierung von Foucaults Diskursanalyse", *Le foucauldien* 2, no. 1 (2016): 13. <http://doi.org/10.16995/lefou.15>.

52 See Lamping, Dieter: "Distant Reading", aus der Nähe betrachtet: Zu Franco Morettis überschätzter Aufsatzsammlung". *Literaturkritik.de*, no. 9 (2016). <https://literaturkritik.de/id/22506#biblio>.

53 See Kathryn Schulz, "Distant Reading: To Uncover the True Nature of Literature, a Scholar Says, Don't Read the Books", *New York Times*, June 26, 2011, B14.

54 Stanley Fish, "Mind Your p's and b's: The Digital Humanities and Interpretation", *New York Times*, last modified January 23, 2012, https://opinionator.blogs.nytimes.com/2012/01/23/mind-your-ps-and-bs-the-digital-humanities-and-interpretation/?_r=0.

55 Cf. Fotis Jannidis, "Digitale Geisteswissenschaften: Offene Fragen – schöne Aussichten", *Zeitschrift für Medien- und Kulturforschung* 10, no. 1 (2019): 63–70.

56 Martin, Mueller, "Digital Shakespeare, or towards a Literary Informatics", *Shakespeare* 4, no. 3 (2008): 290.

which puts phenomena into the ghetto of an acronym that makes its practitioners feel good about themselves but allows the rest of the humanities to ignore them.”⁵⁷

Drawing on experiences from Google Earth and its possibility to seamlessly zoom in and out, Mueller instead proposes the term “Scalable Reading as a happy synthesis of ‘close’ and ‘distant reading’”.⁵⁸ For the study of literature, he suggests, in other words, mixing qualitative and quantitative methods, and thus combining the benefits of different methodological approaches. In contrast to distant reading, scalable reading promises infinite scaling which forms a connection from individual texts to huge text corpora, and also allows a change of scale, i.e., the methodological leap from qualitative interpretation to quantitative analysis—and vice versa.⁵⁹ This change of scale, however, is in no way as seamless as in Google Earth.⁶⁰ Rather, it requires well reflected operationalization to bridge the methodological gap.

Although the promise of infinite scalability is problematic, we argue that a reflected combination of qualitative and quantitative methods for the analysis of literature, as Mueller intends with his term ‘scalable reading’, can make for highly instructive and hypothesis-driven research; as numbers are in need of interpretation: they demand and “enable new and powerful ways of shuttling between ‘text’ and ‘context’”.⁶¹ Hereafter, we will give three cursory examples of QuaDramA research that demonstrate the potential of a mixed-methods setup—or, in Mueller’s words, the potential of ‘scalable reading’—not only for bigger corpus analyses, but also for the understanding and interpretation of single literary texts. The examples try to show how qualitative and quantitative methods can build upon each other.

Our first example concentrates on Heinrich von Kleist’s play *Die Familie Schroffenstein* (1803) and showcases a hypothesis-driven approach. With the establishment of the bourgeois family, authors analytically took gender differences into account for their writings, which are, e.g., important for the bourgeois tragedy.⁶² Kleist’s letters,

57 Martin, Mueller, “Scalable Reading”, last modified April 26, 2020, <https://web.archive.org/web/20211201185120/https://sites.northwestern.edu/scalablereading/2020/04/26/scalable-reading/>.

58 Ibid.

59 Benjamin Krautter, and Marcus Willand, “Close, distant, scalable. Skalierende Textpraktiken in der Literaturwissenschaft und den Digital Humanities”, in: Carlos Spoerhase, Steffen Siegel and Nikolaus Wegmann (eds.), *Ästhetik der Skalierung* (Hamburg: Felix Meiner Verlag, 2020), 86–87.

60 Cf. Thomas Weitin, *Digitale Literaturgeschichte. Eine Versuchsreihe mit sieben Experimenten* (Berlin: J.B. Metzler, 2021), 116.

61 Mueller, “Scalable Reading”.

62 See Sigrid Lange, *Die Utopie des Weiblichen im Drama Goethes, Schillers und Kleists* (Frankfurt a.M.: Peter Lang, 1993), 138 and Ulrike Vedder, “Biblische Muster und ihre Spielräume in Kleists Familien- und Geschlechterordnungen”, *Kleist-Jahrbuch* (2018): 87.

especially the remarks in his letters to Marie von Kleist, suggest that for him, gender might have been one of the more important poetological categories.⁶³ This hypothesis finds additional weight when one looks at the final act of *Die Familie Schroffenstein*. When the loving couple Agnes and Ottokar—they come from two hostile houses of the same family—change their clothes in the light of the danger that their nearing fathers are posing, metaphorically they also change their gender. To verify the hypothesis through supplementary quantitative analysis, we employed five word fields (love, family, reason, religion, and war) to investigate the semantics of the characters' speech. The results show that the speech parts of the male and female characters do differ. The values for male and female characters show gender-specific patterns regarding the word fields.⁶⁴ Ottokar and Agnes, however, do not fit in these patterns shown by the other characters. Instead, their values seem to be somewhat mirrored: When one looks at the distributions of the word fields, one sees that Ottokar rather speaks like the female characters, Agnes rather like the male ones. Thus, the semantics of Agnes' and Ottokar's character speech correspond to the lovers' final external identification when they are changing their clothes. In the end, this leads to the tragic death of both of them.⁶⁵

The second example derives from the annotation of coreference chains in German plays, in this case Friedrich Schiller's *Die Räuber* (1781): When one looks at the chains of Franz Moor, he is brother and antagonist of Karl Moor, the head of a band of robbers, instructive findings can be made. Compared to the other main characters, Franz' utterances feature a lot of references to other characters.⁶⁶ It comes as no surprise, as his father and his brother are omnipresent in his thoughts. However, starting with the fifth act of the play, in which Franz will commit suicide, these references diminish, and his utterances become more self-centered. As soon as father, brother and the main female character Amalia disappear as references, i.e., Franz' thought process no longer targets the perceived injustice of nature that he projected onto other characters, his own horrible intrigues collapse.⁶⁷

While the first two examples showcase the use of qualitative and quantitative methods for the analysis and interpretation of single literary texts, our third exam-

63 See Marcus Willand and Nils Reiter, "Geschlecht und Gattung. Digitale Analysen von Kleists Familie Schroffenstein", *Kleist-Jahrbuch* (2017): 188, https://doi.org/10.1007/978-3-476-04516-4_16.

64 See *Ibid.*, 187.

65 See *Ibid.*, 189–190.

66 See Benjamin Krautter and Marcus Willand, "Vermessene Figuren – Karl und Franz Moor im quantitativen Vergleich", in: Peter-André Alt and Stefanie Hundehöge (eds.): *Schillers Feste der Rhetorik* (Berlin, Boston: De Gruyter, 2021), 107–138.

67 See *Ibid.*

ple is geared towards a bigger corpus of dramas.⁶⁸ Based on literary history and theory, we operationalize dramatic character types in German-language drama. Applying Moretti's original idea of 'distant reading', i.e., the extensive study of secondary literature, we identified characters that belonged to one of the three types: 'schemer', 'virtuous daughter' and 'tender father'. Furthermore, we annotated those properties of the investigated characters that were referenced in the secondary literature. In total, we created a data set of 257 characters and their annotated properties.

These annotations are in turn used for the automatic classification of character types. The experiments show that the selected character types emerge as definable and automatically identifiable subsets within a bigger population. These results lay the groundwork for further, more extensive literary historical analyses of dramatic character types: When do character types emerge, when do they disappear? How do different character types interact within the plays?

All three examples combine quantitative and qualitative approaches for the analysis of dramatic texts. They reveal that the computational counting of surface text elements, such as character mentions or word frequencies, can provide insights into the meaning of a text, especially when it comes to counting the elements that mostly remain invisible to human readers. The examples are also to be understood as a statement that interdisciplinary mixed-methods approaches in digital humanities can fruitfully expand existing research discussions on different humanities disciplines such as literary studies.

Conclusions

In this article, we have argued for a consciously designed collaboration in mixed-methods projects. Interdisciplinary collaboration does not happen by itself, in particular in the digital humanities, where the underlying research assumptions, workflows, and interests can be quite different between D and H. In our project, the joint annotation of humanities concepts, technical workflows that ensure independence of researchers and a close integration of quantitative and qualitative findings are the core principles of successful collaboration.

As discussed in the introduction, we regard operationalization as one of the key challenges not only for our collaboration, but also for digital humanities in general. From our perspective, *good* operationalization of a literary studies concept is one that is both technically feasible and, at the same time, worth arguing about within

68 See Benjamin Krautter et al., "[E]in Vater, dächte ich, ist doch immer ein Vater". Figurentypen und ihre Operationalisierung", in *Zeitschrift für digitale Geisteswissenschaften* (2020), http://dx.doi.org/10.17175/2020_007.

the literary studies community: while it may not be consensual, the literary studies community accepts it (and its results) as a contribution to the discussions in the field.⁶⁹ This is hard to achieve, as operationalization requires a lot of conceptual work depending on the degree of abstraction of the target categories. While measuring certain properties of a text is straightforward, e.g., how many words a character utters in a play, connecting these values to established concepts, thereby fostering the understanding of a specific text, the understanding of a writer's *œuvre*, or literary historical relations, requires sound theoretical reflection. Thus, filling the gap between 'concepts and measurement'⁷⁰ and then between measurement and meaning is one of the major challenges a mixed-methods project must tackle.

Digital humanities practitioners often refer to their discipline as a 'big tent' allowing for a multitude of methods, questions, objects, and approaches.⁷¹ The disciplinary status together with its relation to the variety of 'origin disciplines' in the humanities is a subject of much debate. The digital humanities are an invaluable community that unites researchers employing digital, quantitative, and algorithmic methods to research questions from the humanities. The exploration of their potential should not be stopped solely because of disciplinary boundaries or because these methods still live in the shadows of many humanities disciplines. While development and discussion of digital methods are well accommodated in the digital humanities, once they are established, we envisage bringing the results of applying these methods back into the general field of the humanities. In our opinion, mixed-methods projects that focus and reflect on their collaboration in the light of their research questions and used methods are in a prime position to do so.

Bibliography

- Adelmann, Benedikt, Melanie Andresen, Wolfgang Menzel, and Heike Zinsmeister. "Evaluating Part-of-Speech and Morphological Tagging for Humanities' Interpretation". In *Proceedings of the Second Workshop on Corpus-Based Research in the Humanities*, edited by Andrew U. Frank, Christine Ivanovic, Francesco Mambrini, Marco Passarotti, and Caroline Sporleder, 5–14. Wien, 2018. <https://www.oeaw.ac.at/fileadmin/subsites/academiaecorpora/PDF/CRH2.pdf>.
- Can, Fazli, Tansel Özyer, and Faruk Polat (eds.). *State of the Art Applications of Social Network Analysis*. Cham: Springer, 2014.

69 The German word *satisfaktionsfähig* describes our intention well: The term originates in the context of duels and describes dualists as worthy of a duel.

70 Franco Moretti, "'Operationalizing': Or, the Function of Measurement in Modern Literary Theory," *Pamphlets of the Stanford Literary Lab* 6 (2013): 1.

71 Cf. Patrik Svensson, "Beyond the Big Tent", in: Matthew K. Gold (ed.): *Debates in the Digital Humanities* (Minneapolis, London: University of Minnesota Press, 2012), 36–49.

- Cohen, Margaret. "Narratology in the Archive of Literature". In *Representations* 108, 1 (2009): 51–75.
- Cohen, Margaret. *The Sentimental Education of the Novel*. Princeton: Princeton University Press, 1999.
- Eder, Maciej, Jan Rybicki, and Mike Kestemont. "Stylometry with R: A Package for Computational Text Analysis". *R Journal* 8, 1 (2016): 107–121. <https://doi.org/10.32614/RJ-2016-007>.
- Ehrlicher, Hanno, Jörg Lehmann, Nils Reiter, and Marcus Willand. "Gattungspoetik in quantitativer Sicht. Das Werk Calderón de la Barca und die zeitgenössischen Dramenpoetiken". In *LitLab Pamphlet* 9 (2020): 1–29. https://www.digitalhumanitiescooperation.de/wp-content/uploads/2020/04/p09_ehrlicher_lehmann_al_de.pdf.
- Fischer, Frank, Mathias Göbel, Dario Kampkaspar, Christopher Kittel, and Peer Trilcke. "Network Dynamics, Plot Analysis. Approaching the Progressive Structuration of Literary Texts". In *Digital Humanities 2017. Conference Abstracts*, 437–441. Montréal, 2017. <https://dh2017.adho.org/abstracts/DH2017-abstracts.pdf>.
- Fischer, Frank, Ingo Börner, Mathias Göbel, Angelika Hechtel, Christopher Kittel, Carsten Milling, and Peer Trilcke. "Programmable Corpora. Die digitale Literaturwissenschaft zwischen Forschung und Infrastruktur am Beispiel von DraCor". In *Abstracts of DHd*, edited by Patrick Sahle, 194–197. Frankfurt, Mainz, 2019. <https://doi.org/10.5281/zenodo.2596095>.
- Fish, Stanley. "Mind Your p's and b's: The Digital Humanities and Interpretation". In *New York Times*. Last modified January 23, 2012. https://opinionator.blogs.nytimes.com/2012/01/23/mind-your-ps-and-bs-the-digital-humanities-and-interpretation/?_r=0.
- Gius, Evelyn, and Janina Jacke. "The Hermeneutic Profit of Annotation: On Preventing and Fostering Disagreement in Literary Analysis". In *International Journal of Humanities and Arts Computing* 11, 2 (2017): 233–254. <https://doi.org/10.3366/ijhac.2017.0194>.
- Hovy, Eduard, and Julia Lavid. "Towards a Science of Corpus Annotation: A New Methodological Challenge for Corpus Linguistics". In *International Journal of Translation Studies* 22, 1 (2010): 13–36.
- Jannidis, Fotis. "Digitale Geisteswissenschaften: Offene Fragen – schöne Aussichten". In *Zeitschrift für Medien- und Kulturforschung* 10, 1 (2019): 63–70.
- Jockers, Matthew L. *Macroanalysis. Digital Methods and Literary History*. Urbana, Chicago, Springfield: University of Illinois Press, 2013.
- Krautter, Benjamin, Janis Pagel, Nils Reiter, and Marcus Willand. "Titelhelden und Protagonisten – interpretierbare Figurenklassifikation in deutschsprachigen

- Dramen". In *LitLab Pamphlets* 7 (2018): 1–56. https://www.digitalhumanitiescooperation.de/wp-content/uploads/2018/12/p07_krautter_et_al.pdf.
- Krautter, Benjamin, Janis Pagel, Nils Reiter, and Marcus Willand. "[E]in Vater, dächte ich, ist doch immer ein Vater." Figurentypen und ihre Operationalisierung". In *Zeitschrift für digitale Geisteswissenschaften* (2020). http://dx.doi.org/10.17175/2020_007.
- Krautter, Benjamin, and Marcus Willand. "Close, distant, scalable. Skalierende Textpraktiken in der Literaturwissenschaft und den Digital Humanities". In *Ästhetik der Skalierung*, edited by Carlos Spoerhase, Steffen Siegel, and Nikolaus Wegmann, 77–97. Hamburg: Felix Meiner Verlag, 2020.
- Krautter, Benjamin, and Marcus Willand. "Vermessene Figuren – Karl und Franz Moor im quantitativen Vergleich". In *Schillers Feste der Rhetorik*, edited by Peter-André Alt, and Stefanie Hundehage, 107–138. Berlin, Boston: De Gruyter, 2021.
- Lamping, Dieter. "'Distant Reading', aus der Nähe betrachtet: Zu Franco Morettis überschätzter Aufsatzsammlung". In *Literaturkritik.de*, no. 9 (2016). <https://literaturkritik.de/id/22506#biblio>.
- Lange, Sigrid. *Die Utopie des Weiblichen im Drama Goethes, Schillers und Kleists*. Frankfurt a.M.: Peter Lang, 1993.
- Marcus, Solomon. *Mathematische Poetik*. Translated by Edith Mändroiu. Frankfurt a.M.: Athenäum Verlag, 1973.
- Meister, Jan Christoph. "Computerphilologie vs. 'Digital Text Studies'". In *Literatur und Digitalisierung*, edited by Christine Grond-Rigler and Wolfgang Straub, 267–296. Berlin, Boston: De Gruyter, 2013.
- Moretti, Franco. "Conjectures on World Literature". In *New Left Review* 1 (2000): 54–68.
- Moretti, Franco. "Graphs, Maps, Trees: Abstract Models for Literary History". In *New Left Review* 24 (2003): 67–93.
- Moretti, Franco. *Graphs, Maps, Trees: Abstract Models for Literary History*. London and New York: Verso, 2005.
- Moretti, Franco. "'Operationalizing': Or, the Function of Measurement in Modern Literary Theory". In *Pamphlets of the Stanford Literary Lab* 6 (2013): 1–13.
- Mueller, Martin. "Digital Shakespeare, or towards a Literary Informatics". *Shakespeare* 4, 3 (2008): 284–301. <https://doi.org/10.1080/17450910802295179>
- Mueller, Martin. "Scalable Reading". Last modified April 26, 2020. <https://web.archive.org/web/20211201185120/https://sites.northwestern.edu/scalablereading/2020/04/26/scalable-reading/>.
- Pagel, Janis, and Nils Reiter. "GerDraCor-Coref: A Coreference Corpus for Dramatic Texts in German". In *Proceedings of the 12th Language Resources and Evaluation Conference (LREC)*, 55–64. Marseille, 2020. <https://www.aclweb.org/anthology/2020.lrec-1.7.pdf>

- Pagel, Janis, Nils Reiter, Ina Rösiger, and Sarah Schulz. "Annotation als flexibel einsetzbare Methode". In *Reflektierte Algorithmische Textanalyse*, edited by Nils Reiter, Axel Pichler, and Jonas Kuhn, 125–141. Berlin, Boston: De Gruyter, 2020.
- Pfister, Manfred. *The Theory and Analysis of Drama*. Cambridge et al.: Cambridge University Press, 1988.
- Pichler, Axel, and Nils Reiter. "Reflektierte Textanalyse". In *Reflektierte Algorithmische Textanalyse*, edited by Nils Reiter, Axel Pichler, and Jonas Kuhn, 43–59. Berlin, Boston: De Gruyter, 2020.
- Pichler, Axel, and Nils Reiter. "Zur Operationalisierung literaturwissenschaftlicher Begriffe in der algorithmischen Textanalyse: Eine Annäherung über Norbert Altenhofers hermeneutischer Modellinterpretation von Kleists *Das Erdbeben in Chili*". In *Journal of Literary Theory* 15, 1–2 (2021): 1–29.
- Pustejovsky, James, and Amber Stubbs. *Natural Language Annotation for Machine Learning: A Guide to Corpus-Building for Applications*. Sebastopol, Boston, Farnham: O'Reilly Media, 2012.
- Reiter, Nils. Discovering Structural Similarities in Narrative Texts using Event Alignment Algorithms. PhD diss., Heidelberg University, 2014. <https://doi.org/10.11588/heidok.00017042>.
- Reiter, Nils. "CorefAnnotator – a New Annotation Tool for Entity References". In *Abstracts of EADH: Data in the Digital Humanities*. Galway, 2018. <https://doi.org/10.18419/opus-10144>.
- Reiter, Nils, Marcus Willand, and Evelyn Gius. "A Shared Task for the Digital Humanities Chapter 1: Introduction to Annotation, Narrative Levels and Shared Tasks". In *Cultural Analytics* (2019). <https://culturalanalytics.org/article/11192>.
- Rösiger, Ina, Sarah Schulz, and Nils Reiter. "Towards Coreference for Literary Text: Analyzing Domain-Specific Phenomena". In *Proceedings of the Second Joint SIGHUM Workshop on Computational Linguistics for Cultural Heritage, Social Sciences, Humanities and Literature*, 129–38. Santa Fe, 2018. <http://aclweb.org/anthology/W18-4515>.
- Schulz, Kathryn. "Distant Reading: To Uncover the True Nature of Literature, a Scholar Says, Don't Read the Books". *New York Times*, June 26, 2011, B14.
- Svensson, Patrick. "Beyond the Big Tent". In *Debates in the Digital Humanities*, edited by Matthew K. Gold, 36–49. Minneapolis, London: University of Minnesota Press, 2012.
- Trilcke, Peer, and Frank Fischer. "Fernlesen mit Foucault? Überlegungen zur Praxis des *distantreading* und zur Operationalisierung von Foucaults Diskursanalyse". In *Le foucauldien* 2, 1 (2016): 1–18. <http://doi.org/10.16995/lefou.15>.
- Underwood, Ted. "A Genealogy of Distant Reading". In *Digital Humanities Quarterly* 11, 2 (2017): § 1–44. <http://www.digitalhumanities.org/dhq/vol/11/2/000317/000317.html>.

- Underwood, Ted. "The Stanford Literary Lab's Story". Last modified February 11, 2017. <https://www.publicbooks.org/the-stanford-literary-labs-narrative/>.
- Vedder, Ulrike. "Biblische Muster und ihre Spielräume in Kleists Familien- und Geschlechterordnungen". *Kleist-Jahrbuch* (2018): 87–99.
- Weitin, Thomas. *Digitale Literaturgeschichte. Eine Versuchsreihe mit sieben Experimenten*. Berlin: J.B. Metzler, 2021.
- Weitin, Thomas, Thomas Gilli, and Nico Kunkel. "Auslegen und Ausrechnen. Zum Verhältnis hermeneutischer und quantitativer Verfahren in den Literaturwissenschaften". In *Zeitschrift für Literaturwissenschaft und Linguistik* 46 (2016): 103–115.
- Willand, Marcus, Benjamin Krautter, Janis Pagel, and Nils Reiter. "Passive Präsenz tragischer Hauptfiguren im Drama". In *Abstracts of DHD*, edited by Christof Schöch, 177–181. Paderborn, 2020. <https://zenodo.org/record/3666690>.
- Willand, Marcus, and Nils Reiter. "Geschlecht und Gattung. Digitale Analysen von Kleists Familie Schrockenstein". In *Kleist-Jahrbuch* (2017): 177–95. https://doi.org/10.1007/978-3-476-04516-4_16.
- Zweig, Katharina A. *Network Analysis Literacy. A Practical Approach to the Analysis of Networks*. Wien: Springer, 2016.