

Comparison of Librarian and Patron Ratings of Synchronous Chat Interactions

Erin Elizabeth Owens and Kat Brooks

While virtual reference has become more critical during and after COVID-19, there remains a lack of current research in patron and librarian perceptions of the service. This study aims to compare librarian and patron ratings of chat interactions and highlight trends in what these ratings may suggest. Researchers collected randomized samples of patron rated chat transcripts from two large academic libraries. The transcripts were then blind reviewed according to a rubric based on the *RUSA Guidelines for Behavioral Performance of Reference and Information Service Providers*. Analysis of these ratings found discrepancies between patron and librarian perceptions of successful interactions. Patrons and librarians seemed to differ on their criteria for high or low ratings, the level of impact of time in interactions, and trends for overall perception of success. The lack of alignment between librarian and patron perceptions suggests areas for further research in how to improve chat services and patron experiences.

Introduction

After many libraries were pressed into mostly or entirely remote service by COVID-19 pandemic lockdowns, virtual channels became critical in maintaining reference services. For example, a study by Radford et al. (2022) found that 71% of the libraries they interviewed reported dramatic increases in chat reference encounters in the early stages of the pandemic. As libraries seek to launch, assess, or enhance chat reference services in any stage of maturity, we must be aware of whether we are meeting both patron expectations and our field's own professional expectations. This study seeks to compare patron assessments and librarian assessments of live chats to determine how well each group's expectations for this service are being met in practice. If librarians believe they are meeting their profession's expectations, but patron ratings do not seem to agree, determining the cause for this disagreement would be important. Meanwhile, librarians may be able to improve the patron experience in chat if evaluations by both parties can help to identify and address areas of professional practice requiring greater attention or new approaches.

Institutional Contexts

Sam Houston State University (SHSU), a member of the Texas State University System, is a large public university in Huntsville, Texas, with a Fall 2022 headcount enrollment of 21,480.

* Erin Elizabeth Owens, Professor, Associate Director of Library Public Services & Scholarly Communications Librarian, Newton Gresham Library, Sam Houston State University, eowens@shsu.edu; Kat Brooks, Assistant Professor, Collections Strategist, University of Tennessee Knoxville, kbrooks@utk.edu. ©2025 Erin Elizabeth Owens and Kat Brooks, Attribution-NonCommercial (<https://creativecommons.org/licenses/by-nc/4.0/>) CC BY-NC.

SHSU is Carnegie-classified as a Doctoral University: High Research Activity (R2). SHSU is also designated as a Hispanic-Serving Institution (HSI) and a Carnegie Community Engaged Campus, and it enrolls and graduates a higher-than-average number of first-generation college students. The Newton Gresham Library (NGL) at SHSU has offered live virtual chat services since 2004 and currently uses the LibChat platform from Springshare.

The University of Tennessee Knoxville (UTK), the flagship institution of the University of Tennessee System, is a large public university located in Knoxville, Tennessee, with a Fall 2022 headcount enrollment of 33,805. UTK is Carnegie-classified as a Doctoral University: Very High Research Activity (R1). The UTK Libraries have offered live virtual chat services since 2013 and currently use the LibChat platform from Springshare.

Literature Review

Live chat, once an uncertain tool in the suite of reference offerings, has become the focus of a wide number of research studies (Matteson, 2011). In the ten years following Matteson's review, chat has become an established mainstay of reference services in academic libraries. However, research on perceptions and assessment of the service has not evolved with the role of chat in 2021. Most relevant studies found were published between 2000 and 2010, when libraries were determining the value of chat and how to develop the pedagogies of the service (Desai & Graves, 2006; Smyth & MacKenzie, 2006; Arnold & Kaske, 2005; Hansen et al., 2009). Due to this lack of recent literature, this review includes some studies from ten years ago or older.

Our proposed study consists of two main elements: developing a rubric based on the *RUSA Guidelines for Behavioral Performance of Reference and Information Service Providers (RUSA Guidelines)* and applying that rubric to compare patron and librarian evaluation of chat interactions. Desai and Graves (2006, 2008; Graves & Desai, 2006) published several studies examining transcripts along with patron surveys, focusing on different elements of instruction in chat. Logan et al. (2019) analyzed exit surveys and transcripts to determine which behaviors correlated to patron dissatisfaction. Smyth and MacKenzie (2006) were able to highlight a disconnect in patron satisfaction and librarian assessment through their comparison study. Several studies used patron surveys and exit interviews to assess chat services without transcript comparisons (Foley, 2002; Lee, 2008; Neuhaus & Marsteller, 2002; Ruppel & Vecchione, 2012; Stoffel & Tucker, 2004). Hansen et al. (2009) surveyed both patron and librarian following reference interactions, finding provider pessimism to be a notable theme.

Great variation was found in rubrics used by researchers to assess chat. Many researchers developed their own rubric and coding schemas (Arnold & Kaske, 2005; Fuller & Dryden, 2015; Marsteller & Mizzy, 2003; Meert & Given, 2009; Pomerantz et al., 2006; Radford & Conaway, 2013; Butler & Byrd, 2016). A few used the READ Scale (Mawhinney & Hervieux, 2022; Mavodza, 2019; Cabaniss, 2015), while others built upon the ACRL Framework (Hervieux & Tummon, 2018) or used SERVQUAL (Gómez-Cruz, 2019). Of the studies that developed rubrics based on the *RUSA Guidelines*, most were concerned with the adherence of providers to the guidelines (Hughes, 2010; Maness et al., 2009; Van Duinkerken et al., 2009), while some used their rubric to assess provider skills (Keyes & Dworak, 2017; Ronan et al., 2006; Ward, 2004). Of all the research found, only two publications combined discussion of *RUSA Guidelines* with patron perception of chat interactions, both published more than ten years ago (Kwon & Gregory, 2007; Haynes, 2009).

Review of the relevant literature has highlighted a gap in current, methodologically similar, and extensive research concerned with the interactions of *RUSA Guidelines*, patron perceptions, and librarian perceptions of virtual chat interactions.

Aims

This study was guided by two primary research questions:

1. How well do librarians respond to live chat reference questions, in terms of professional guidelines from *RUSA*?
2. How do librarian ratings of chat responses compare to patron ratings of the same chats?

Methodology

We used the *RUSA Guidelines* to represent librarians' professional expectations for chat. A rubric based on these Guidelines, with an emphasis on the *Remote* aspect of the guidelines, was adapted for the current study from a rubric previously published by Cassidy et al. (2014). Because the original rubric was designed to evaluate SMS/text messaging rather than live synchronous chat, small modifications were made, mostly pertaining to the speed of response. The applied rubric is included in Appendix A.

With IRB approval from both institutions, each researcher downloaded their institution's LibChat transcripts that included patron ratings from August 1, 2019, to July 30, 2020. This period spans both pre- and post-COVID-19 pandemic; although the pandemic may have impacted the frequency and content of the chats, the researchers concluded that it should not impact the fundamental guidelines for library personnel behavior in responding to chats. The transcripts and associated metadata were cleaned to remove any personally identifying information, primarily patron names, identification numbers, email addresses, and phone numbers. Where applicable, identifying metadata fields were simply deleted from the dataset, but additionally the text of the transcript was carefully read and edited in context. Appendix B provides a data dictionary of the chat transcript data fields. One field worth explaining here in the methods is the *Patron Rating*, which users can optionally select after a chat ends; ratings are on a scale from 1 to 4, with the scores labeled as *Bad* (1), *So-so* (2), *Good* (3), and *Excellent* (4).

A sample of 360 transcripts was taken from each institution after data cleaning; this sample amounted to almost 100% of SHSU's potential transcripts for the period and 20% of UTK's potential transcripts for the period. In order to take a random sampling of UTK's larger dataset, a *RAND()* function was inserted in a blank column of the data spreadsheet to generate a random number; records were sorted according to that random value, and the first 360 randomly ordered rows were selected for analysis.

To test interrater reliability, two sample transcripts were selected from each institutional dataset (four records total) and scored by both researchers using the rubric. Two weeks later, without referencing their first set of scores, both researchers scored the same sample records again. An average agreement of 84.4% between raters demonstrated acceptable interrater reliability. An average of 84.4% agreement was also found between each researcher's first and second sets of scores, indicating acceptable intra-rater reliability as well. After establishing reliability, each researcher proceeded to rate the first 25 records from the other institution to check for any issues or questions which might require clarification; no issues arose. Each researcher then completed ratings for half the chat transcripts from their own institution and half the transcripts from the other institution (one rater per chat). During the rating process,

a total of ten transcripts were excluded from the dataset as being insufficient for rating (e.g., one transcript simply represented a patron logging back in to say “thank you” after they had accidentally lost the chat connection). Finally, these rubric-based librarian scores were analyzed and compared to the patron ratings. Descriptive statistics were collected, and Pearson’s correlation tests were run as appropriate.

Results

After scoring and excluding insufficient transcripts, a total of 710 transcripts were found eligible for analysis, including 357 (50.3%) from Hodges and 353 (49.7%) from SHSU. Approximately 60% of these chats occurred before the COVID pandemic (defined for convenience by a date of March 15, 2020), including 212 chats from Hodges and 217 chats from SHSU. The remaining 40% of the chats analyzed occurred during the COVID pandemic up to the close of the data collection period (March 15, 2020, through July 30, 2020), including 145 chats from Hodges and 136 chats from SHSU.

Chat Characteristics

The length of time each patron spent waiting for an initial response to their chat was documented in the chat transcripts as “Wait Time” and was measured in seconds. The average wait time overall was 19.7 seconds, but this varied between institutions, with Hodges averaging 15.6 seconds, which was 8.2 seconds faster than SHSU’s slower average of 23.8 seconds. The extreme outliers for wait time were documented at 1707 seconds (Hodges) and 1562 seconds (SHSU)—about 28.5 minutes and 26 minutes, respectively. Even with these two outliers removed from the data, Hodges still averaged 8.6 seconds faster than SHSU (10.8 seconds versus 19.4 seconds, respectively).

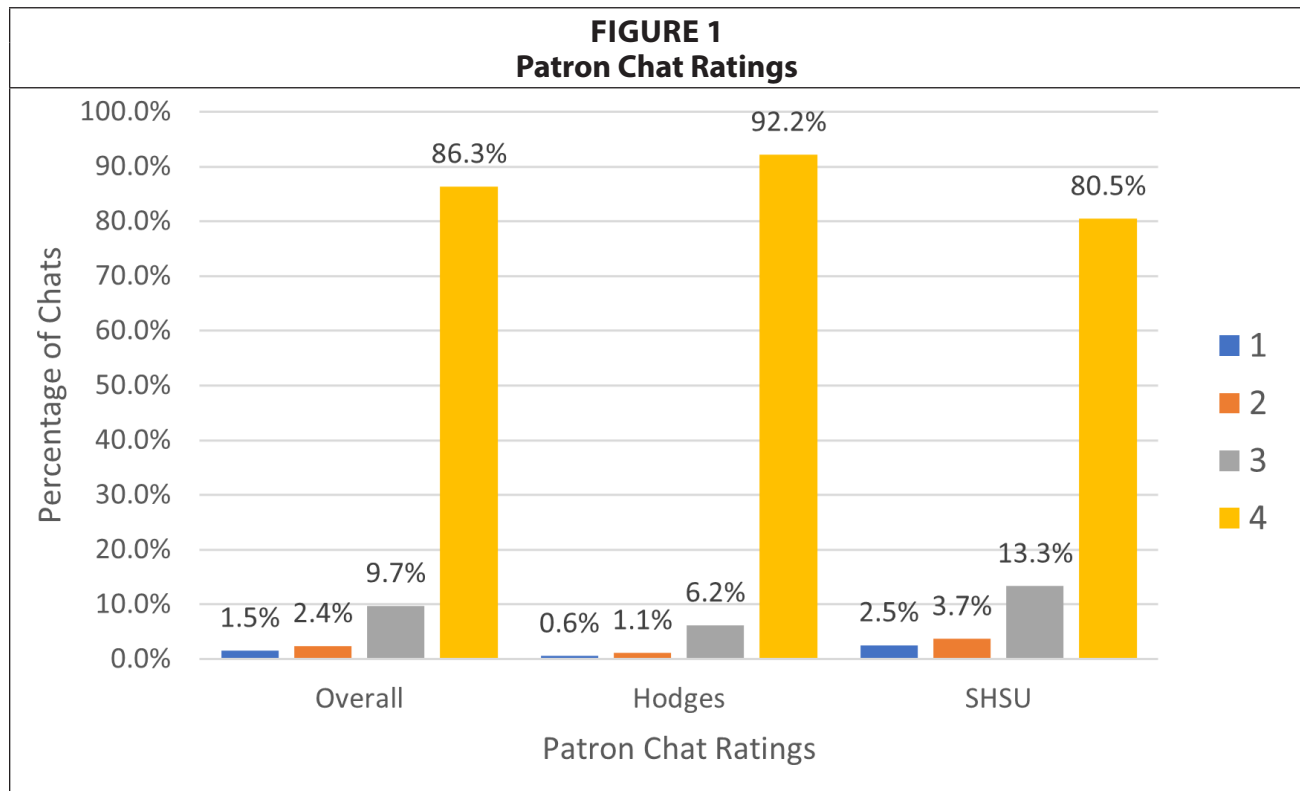
The duration of each chat in seconds was also included in the chat data. The average chat overall lasted 615.8 seconds (just over 10 minutes), but again the data showed variance between institutions. Hodges library personnel chatted longer on average than SHSU personnel, at 754.2 seconds (12.5 minutes) versus 475.8 seconds (just shy of 8 minutes), respectively. Along with longer average duration of chat, Hodges also exchanged a larger number of messages in the average chat: 14.8 messages, compared to just 10.8 messages on average at SHSU (the overall average number of messages was 12.8).

Statistical testing was conducted to determine whether these more mundane chat characteristics correlated to patron ratings. The wait time, duration, and message count all lacked statistically significant correlations to patron rating (Pearson’s correlation coefficients = -0.04, 0.04, and 0.06, respectively); patron ratings did not appear to be influenced by chats being either faster or lengthier. The month and weekday of the chat were also found to have no correlation to the average patron rating, which was always between 3.6 and 3.9; students did not seem more satisfied or more frustrated with library personnel’s chat performance at any particular time during the week, semester, or year. The same was true for patron ratings pre- and post-COVID: although individual chats occasionally earned poor scores at intermittent points in both periods, the overall average did not change appreciably, and a consistent proportion of chats earned each rating on the scale from 1 to 4. Pearson’s correlation coefficients for librarian scores also did not identify any relationships of strong significance, although librarian scores had a low positive correlation with both message count and duration and a very low negative correlation with wait time (-0.12, 0.24, and 0.24, respectively). In other words, librarian scores

were likely to be slightly higher when more messages were exchanged, when a chat lasted for a longer duration, or when the wait time was briefer.

Patron Ratings

The average patron rating overall was 3.8 (on a scale of 1 to 4). The average at Hodges was slightly higher (3.9), while the SHSU average was slightly lower (3.7). The median and mode scores were 4, both overall and for each institution. Both overall and at each separate institution, the patrons who chose to submit a rating for their chats overwhelmingly rated them at the highest score of 4 (see Figure 1).

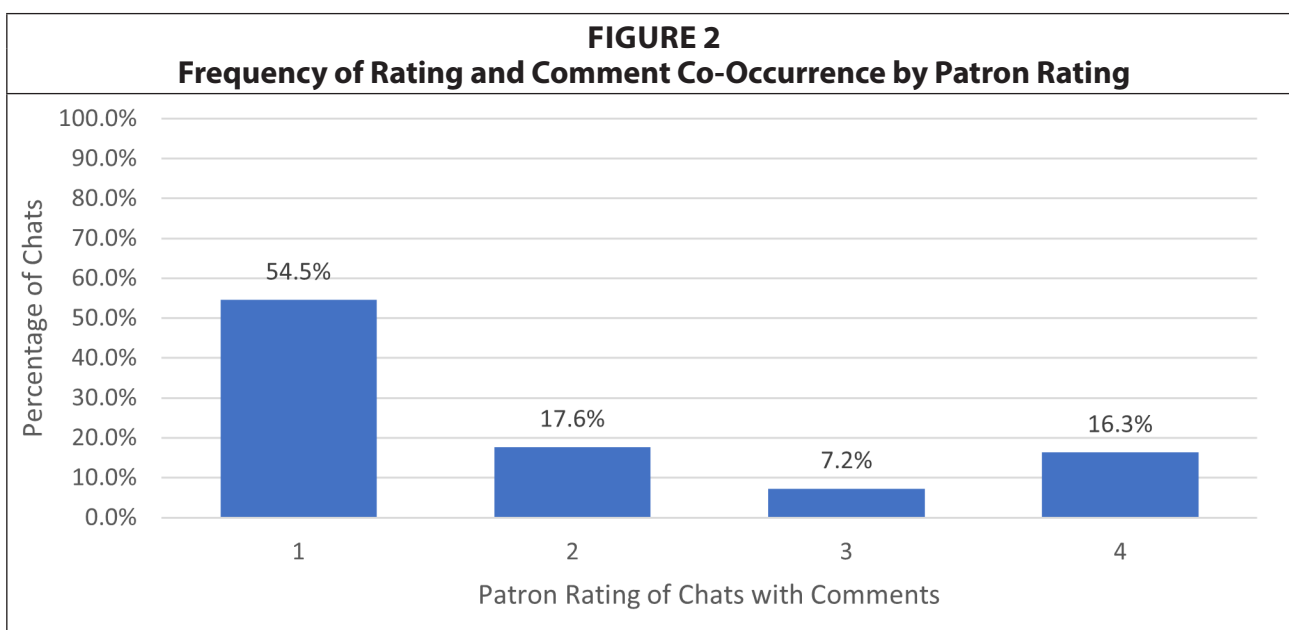


Only eleven chats received a patron rating of 1 or *Bad*. These chats were examined qualitatively to look for themes of behavior that resulted in patron dissatisfaction strong enough to warrant leaving a negative rating. As it turns out, patrons have straightforward expectations: they want library personnel to respond to their chat and answer their question. Three out of eleven chats earning the lowest patron rating were characterized by the library personnel either never responding to the chat at all, or by indicating that they would be right back and then never returning to the chat. At the same time, five of these eleven chats shared the common theme that the patron felt their question was never fully answered or their problem never fully resolved. In one chat which illustrates a variation of the *unresponsive* theme, a patron chatting during off-hours reached a library student worker, who redirected their research question to an email address being monitored by a librarian. The patron left no comment, but their negative rating—left well before they would have received a satisfactory or dissatisfactory answer to their email—may reflect discontent with being encouraged to use a less immediate communication method rather than receive a prompt, real-time reply.

Occasionally, however, a patron's poor rating expressed a frustration which, although possibly warranted, had nothing to do with the chat service at all: in one instance, the patron's low rating was accompanied by a comment expressing their disagreement with a specific pandemic-related service limitation. In another example, a patron was appropriately directed to the campus bookstore for the answer to their question; however, as indicated in their post-chat comment, the bookstore phone was not answered and the voicemail inbox was full, so they were unable to leave a message.

Rating and Comment Co-Occurrence

Overall, students were very unlikely to leave a comment at all when rating a chat; only 16.1% (n = 114) of patron-rated chats analyzed included a comment. Among these, students who rated their chat poorly were more than three times as likely to leave a comment than those who gave a middling or high rating: 54.5% of patrons who assigned a chat rating of 1 also left a comment, compared to just 16.3% of patrons who assigned a chat rating of 4 (see Figure 2).

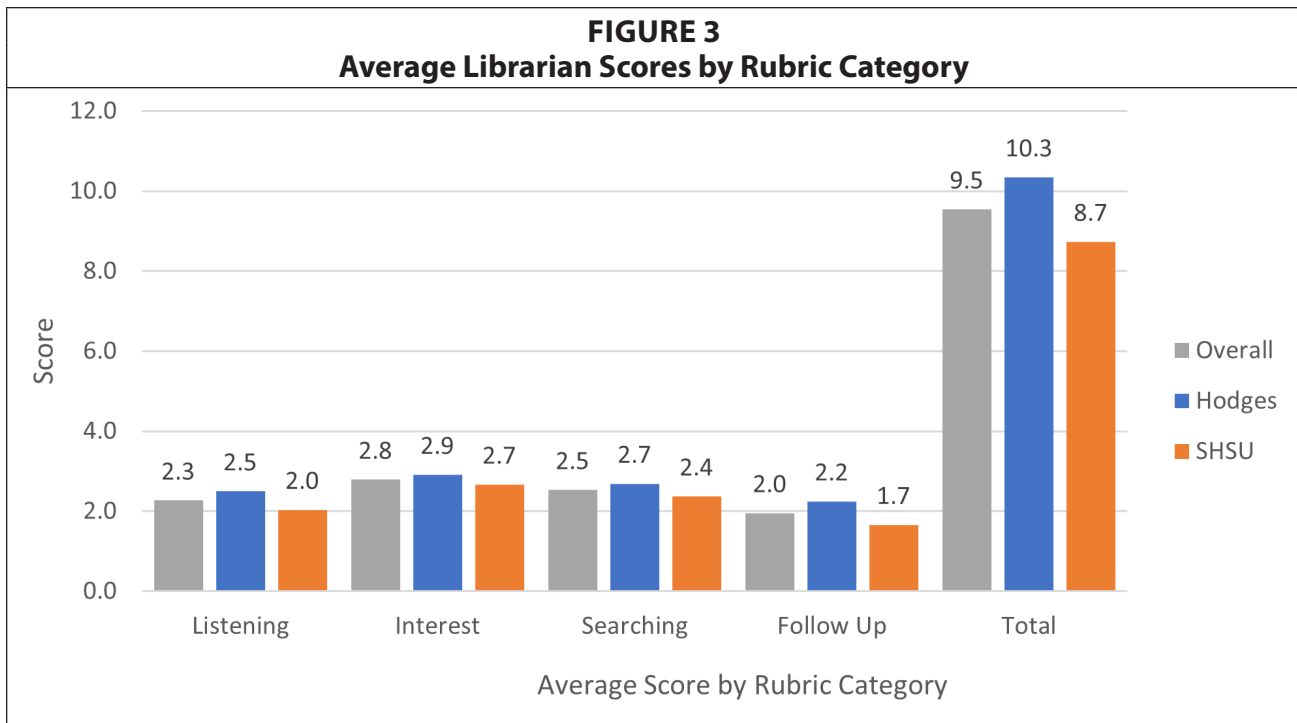


Librarian Scores

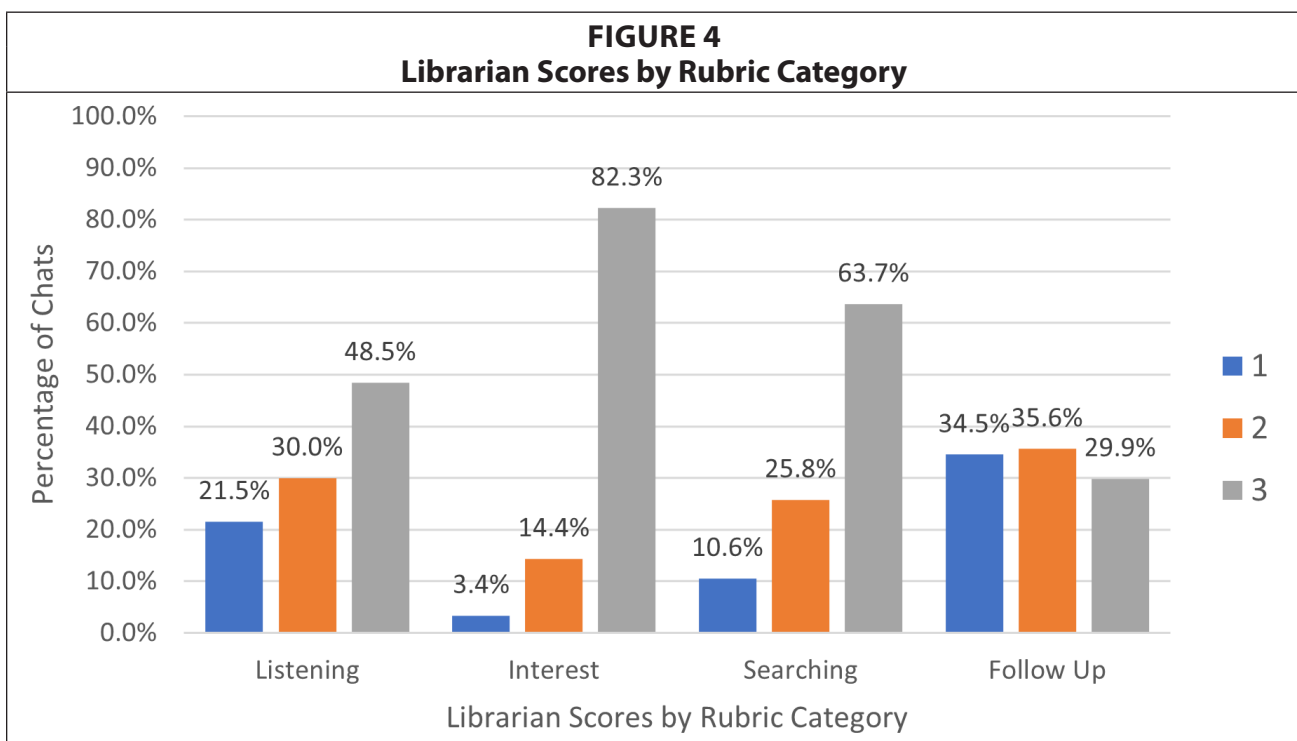
The librarian scores, determined through a rubric-based evaluation, reflected more nuance than the patron ratings. Each chat was scored on a scale from 1 to 3 (*Beginning*, *Developing*, or *Accomplished*) in the areas of *Listening*, *Interest*, *Searching* and *Follow Up*, corresponding to areas in the *Remote* guidelines of the *RUSA Guidelines*. Average scores in each area were very similar overall and for each institution, though scores for chats from Hodges were consistently a fraction higher than the scores of chats from SHSU (see Figure 3). Overall, the average total score was 9.5, while the institutional averages were 10.3 at Hodges and 8.7 at SHSU.

Some of the most frequently seen reasons for librarian scores being lower in a given rubric category included:

- Interest: abrupt language lacking pleasantries and empathy; failure to clarify vague queries.

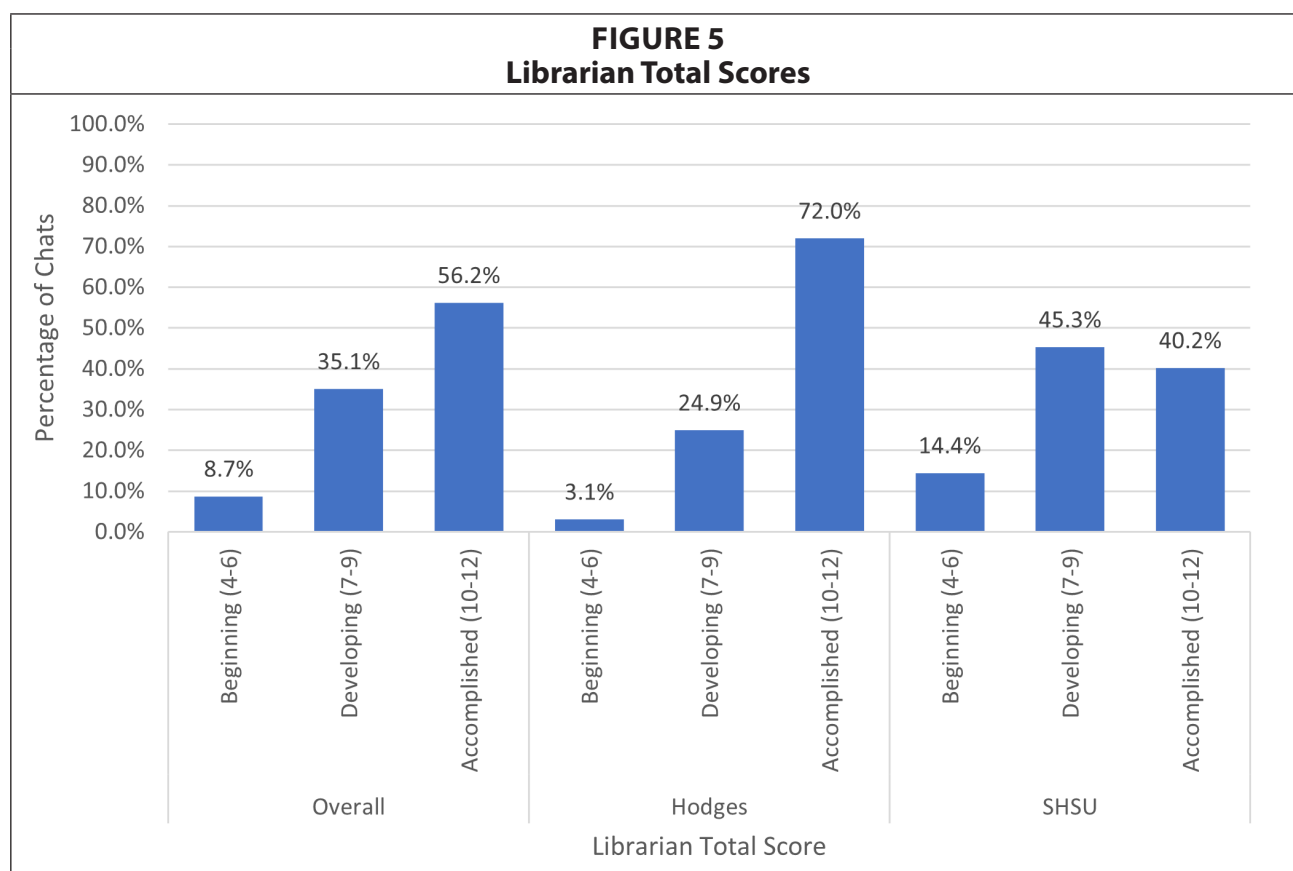


- Listening/Inquiring: slow response to initial query; failure to maintain regular contact while searching.
- Searching: failure to be transparent about where and how they searched; failure to provide appropriate links or contact information for patrons to move forward easily.
- Follow Up: failure to end chat politely; failure to encourage patron to return for further help.



Examining each rubric category in turn, librarians found *Interest* to be the strongest area of performance, with 82.3% of chats receiving a score of 3/*Accomplished* and only 3.4% of chats receiving a score of 1/*Beginning*. *Searching* also scored well, with 63.7% of chats at a 3 and only 10.6% of chats at a 1. The category of *Listening* showed more tepid performance: slightly less than half of chats (48.5%) scored a 3, while nearly a quarter (21.5%) scored a 1. Finally, performance in the *Follow Up* category showed the most room for improvement, with only 29.9% of chats scoring a 3 and 34.5% scoring a 1. Figure 4 shows the full details of how each category received librarian scores 1 through 3.

All told, the librarians assigned 143 chats the highest possible score of 12, including 121 chats from Hodges and only 22 chats from SHSU. At the other end of the spectrum, the librarians assigned just six chats, all from SHSU, the lowest possible score of 4. Overall, more than half of chats (56.2%) were scored as *Accomplished*; by institution, 72.0% of Hodges' chats were *Accomplished*, while only 40.2% of SHSU's chats scored at this level. Instead, the largest proportion of SHSU's chats scored in the *Developing* range (see Figure 5).

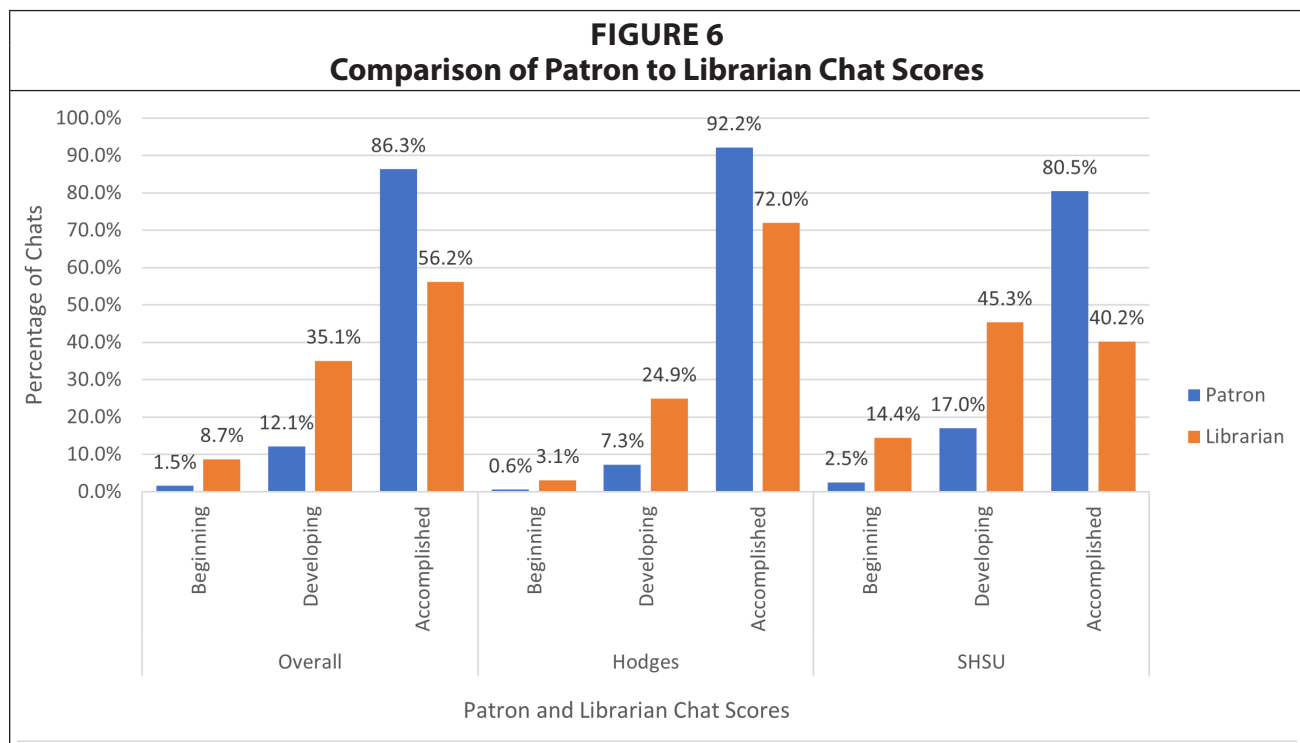


Unresponsiveness from library personnel earned poor librarian scores, just as this behavior earned poor patron ratings. Of the six transcripts which earned the lowest possible scores from the librarian evaluators, four did poorly because of simple failure to respond to a query. The remaining two chats were evaluated as simply being poor reference interactions: library personnel responses were abrupt and lacking in detail expected for clarity. For example, one patron asked, "What would be a good book to research educating the youth about protesting?" The library personnel monitoring chat stated, without any greetings, "There might be

some good articles on it using the main search engine on the page. As for books I'm not quite sure there would be anything recent in book form," and that was it. They asked no clarifying questions about the vague topic: was the patron's interest limited to "recent" discussion, or would they have been equally interested in books from youth protests of the 1960s? The library personnel also did not clarify the context of the need/use, such as whether the book format was specifically required by a class assignment, and they failed to ask other similar questions which would be expected in a strong reference interview. They did not walk the patron through a search attempt or any possible results. They did not even specify what they meant by "the main search engine," or on which page it could be found, which might have been clarified with a hyperlink. Curiously, the patron in this instance rated the chat a 4/4, compared to the librarian score of 4/12. Perhaps this was a more advanced user, who readily understood the search recommendation, and perhaps the unverified understanding of their query was accurate enough that they were satisfied with the result of the answer. In any case, patrons and librarians clearly share some common expectations and priorities, while others may differ significantly. We will examine such comparisons further in the next section.

Comparing Patron Ratings to Librarian Scores

The patron rating scale of 1 to 4 did not perfectly correspond to the librarian scoring rubric, which could yield a total score between 4 and 12. Rather than tear down and recreate a well-tested rubric, for purposes of comparison, the researchers decided to equate a patron rating of 1 to librarian scores 4 to 6, or *Beginning*. Patron ratings of 2 and 3 equated to librarian scores of 7 to 9, or *Developing*. Finally, patron ratings of 4 equated to librarian scores of 10 to 12, or *Accomplished*. Overall, patrons were more likely to consider a chat *Accomplished* than were librarians (86.3% versus 56.2%). Conversely, librarians were more likely to rate chats as *Beginning* than were patrons (8.7% versus 1.5%). Figure 6 shows the comparisons of patron and librarian scores, both overall and for each institution.



Discussion

According to the comparison between patron and librarian evaluations of live chats, both patrons and librarians seem generally satisfied with chat interactions, with patrons ranking most chats 4 out of 4 and the librarian evaluators assigning an average overall score of 9.5 out of 12. However, a librarian's score for a chat does not necessarily predict the patron's rating; the discrepancies between their assessments yield some valuable discussion points.

In a few noteworthy instances, the librarian evaluators scored a transcript highly (e.g., 10 out of 12 or even 12 out of 12) while the patron scored the same chat very poorly (e.g., 1 out of 4 or 2 out of 4, respectively). These discrepancies provide an interesting opportunity to investigate areas where librarian and patron expectations may be out of sync. In five out of eight chats with such mismatched scores, the library personnel fulfilled all the behavioral expectations of a good reference interaction, but the patrons were nevertheless dissatisfied with the services or collections available from the library. Such cases are difficult to eliminate, since the behavior of an individual professional may never fully satisfy a patron who simply wants the library to own more or different resources, or to provide services beyond what is appropriate (e.g., doing a research assignment for a student).

In two other instances, the librarian exhibited generally strong behaviors, except that the patron ultimately felt their question was not fully answered or that their problem was not resolved. One final chat, however, illustrates a very different kind of case. A patron requests "academically acceptable" sources on a topic, and the library personnel assists obligingly. They use welcoming language and ask clarifying questions about information needs and any constraints of a class assignment. When they suggest searching the library's discovery layer, they take the time to explain what it is and why the student would be better off searching there (as opposed to web search engines). They both explain and link to a sample search they formulated as a starting point for the student. By all accounts, this chat hits all the behavioral high points, and it earned a librarian score of 10 out of 12. However, the patron stopped responding to the library personnel before the end of the chat and subsequently rated the chat only 2 out of 4. Although no patron comment accompanies this rating, one possible explanation is that the patron became overwhelmed with the unfamiliar nature of the information presented and the speed of its presentation; perhaps the library personnel could have checked on the patron's understanding at more frequent intervals during the chat, making sure they were absorbing and understanding the guidance. On the other hand, perhaps the patron had hoped to have specific sources named for them, as opposed to instructions for searching. Whatever the patron may have been expecting, they didn't seem to feel that it was delivered.

The findings also suggest that librarians' professional expectations go beyond patron expectations in certain areas, such as cordiality of language and transparency about the search process. Patrons are more concerned about actually getting the answer they want, and they are the least satisfied when they do not perceive this has happened. While librarians routinely lowered a chat's *Interest* score based on abrupt language that lacked pleasantries or empathy, patron ratings indicated little concern about such cordiality, as long as the abrupt responses were prompt and constructive. This may reflect the ever-widening gap between expectations of digital versus in-person conversation. It may also result in part from increased interaction with chatbots on commercial websites, not to mention search engines, which fulfill requests dispassionately. Some of the initial questions posed in the chat even resemble the efficient,

incomplete sentences often fed into a search engine, such as: Recalling a recall; Help finding a book; Posting a flyer; SRDS; What is Interlibrary Loan?

Perhaps today's university patrons don't place as much priority on seeing textual evidence that they are speaking to a human. For that matter, as artificial intelligence text generators like ChatGPT become ever more skilled and ever more integrated into other online technologies, that which constitutes such "evidence" of humanity is perhaps less clear-cut. When even a machine can engage in a polite two-way discussion, does it really matter whether your partner in a dialogue is human, if they satisfy your requirements of the conversation? On the other hand, quite a few of the initial questions that launched chat interactions began (and ended) with indicators of establishing human rapport (e.g., Hello!; Heyy!; Good afternoon; Thank you!). This suggests that many patrons are aware and respectful of the human connection being initiated. And once a two-way dialogus has been established in a chat, patrons overwhelmingly conversed in a polite and appreciative manner. Even if some of them are satisfied by a chat that lacks such courteous "fluff" on the part of library personnel, continuing to strive for a warm human communication style is still probably the best policy for library chat providers, as long as that warm communication is balanced with efficiency and accuracy.

Related to this topic of warm human communication, the researchers discovered that *Follow Up* was particularly difficult to gauge via transcripts from LibChat. The transcripts from the study period did not record any details regarding which chat participant ended or disconnected from the chat or the timestamp at which the disconnection occurred. Therefore, it may often look as though the library personnel failed to appropriately end the interaction, when the patron may in fact have abruptly left the chat before an appropriate closing could be made. As a result of this limitation in the data, *Follow Up* scores may be lower than warranted. In the digital environment, adherence to RUSA's general *Follow Up* guideline of "Takes care not to end the reference interview prematurely," is not truly under the library staff member's control in the same way that it is during face-to-face interactions, and this factor should be taken into consideration by libraries undertaking assessment of their own performance in chat transcripts. The LibChat platform has since updated this aspect of transcript recording, so newer data analysis may be able to assess chat disconnection with more nuance.

Limitations and Further Research

The average patron chat ratings were consistently lower at SHSU than at Hodges. Similarly, librarian evaluators scored chats from Hodges higher in every rubric category, compared to chats from SHSU. This suggests that substantive differences may exist between the two libraries in terms of how library personnel are trained to provide reference via chat, how such training is refreshed over time and/or monitored for quality control, and the general norms in each library for engaging in reference interactions. While the current study did not gather details to compare these aspects of practice, the findings suggest that future research delving into this may provide insight for improving patron satisfaction with virtual chat reference.

Conclusions

In exploring the relationship between how librarians and patrons rate virtual chat interactions, this study highlighted both similarities and variances in how well the expectations of chat participants were met. While patrons were most satisfied when they got the answer they were seeking, librarians put a greater emphasis on professional expectations of cordiality

and customer service. Unresponsiveness was an issue for both, earning poor ratings from librarians and patrons. Overall, perceptions of librarians and patrons were often out of sync, which suggests there is room for further research in this area. These discrepancies provide an opportunity to better understand and serve patrons as chat reference evolves in a modern landscape.

Author Contributions

Erin Elizabeth Owens: Conceptualization (equal), Methodology (equal), Data Curation (equal), Investigation (equal); Formal analysis; Visualization; Writing: original draft (Introduction, Methodology, Results, Discussion, Limitations); Writing: review and editing (equal); **Kat Brooks:** Conceptualization (equal), Methodology (equal), Data Curation (equal), Investigation (equal); Writing: original draft (Abstract, Literature review, Conclusions); Writing: review and editing (equal).

References

- Arnold, J., & Kaske, N. (2005). Evaluating the quality of a chat service. *portal: Libraries & the Academy*, 5(2), 177–93. <https://doi.org/10.1353/pla.2005.0017>
- Butler, K., & Byrd, J. (2016). Research consultation assessment: Perceptions of students and librarians. *The Journal of Academic Librarianship*, 42(1), 83–86. <https://doi.org/10.1016/j.acalib.2015.10.011>
- Desai, C. M., & Graves, S. (2006). Instruction via instant messaging: What's happening?, *The Electronic Library*, 24(2), 174–89. <https://doi.org/10.1108/02640470610660369>
- Desai, C. M., & Graves, S. J. (2008). Cyberspace or face-to-face: The teachable moment and changing reference mediums. *Reference & User Services Quarterly*, 47(3), 242–55. <https://www.jstor.org/stable/20864890>
- Foley, M. (2002). Instant messaging reference in an academic library: A case study. *College & Research Libraries*, 63(1), 36–45. <https://doi.org/10.5860/crl.63.1.36>
- Fuller, K., & Dryden, N. H. (2015). Chat reference analysis to determine accuracy and staffing needs at one academic library. *Internet Reference Services Quarterly* 3(4), 163. <https://doi.org/10.1080/10875301.2015.1106999>
- Gómez-Cruz, M. E. (2019). Electronic reference services: A quality and satisfaction evaluation. *Reference Services Review*, 47(2), 118–133. <https://doi.org/10.1108/RSR-07-2018-0057>
- Graves, S. J., & Desai, C. M. (2006). Instruction via chat reference: Does co-browse help? *Reference Services Review*, 34(3), 340–57. <https://doi.org/10.1108/00907320610685300>
- Hansen, D., Johnson, M., Norton, E., & McDonough, A. (2009). Virtual provider pessimism: Analysing instant messaging reference encounters with the pair perception comparison method. *Information Research*, 14(4), 9–9. <https://informationr.net/ir/14-4/paper416.html>
- Haynes, W. (2009). ASSISTing you online: Creating positive student experiences at the University of Wolverhampton. *SCONUL Focus*, 46(Summer), 86–90. <https://www.sconul.ac.uk/publication/assisting-you-online-creating-positive-student-experiences-at-the-university-of>
- Hervieux, S., & Tummon, N. (2018). Let's chat: The art of virtual reference instruction. *Reference Services Review*, 46(4), 529–542. <https://doi.org/10.1108/rsr-07-2018-0060>
- Hughes, A. M. (2010). Adherence to RUSA's Guidelines for Virtual Reference Services is below expected in academic libraries. *Evidence Based Library & Information Practice*, 5(4), 105–7. <https://doi.org/10.18438/B8JP6W>
- Keyes, K., and Dworak, E. (2017). Staffing chat reference with undergraduate student assistants at an academic library: A standards-based assessment. *The Journal of Academic Librarianship*, 43(6), 469–78. <https://doi.org/10.1016/j.acalib.2017.09.001>
- Kwon, N., & Gregory, V. L. (2007). The effects of librarians' behavioral performance on user satisfaction in chat reference services. *Reference and User Services Quarterly*, 47(2), 137–148. <http://dx.doi.org/10.5860/rusq.47n2.137>
- Lee, L. S. (2008). Reference services for students studying by distance: A comparative study of the attitudes distance students have towards phone, email and chat reference services. *New Zealand Library & Information Management Journal*, 51(1), 6–21.
- Logan, J., Barrett, K., & Pagotto, S. (2019). Dissatisfaction in chat reference users: A transcript analysis study. *College & Research Libraries*, 80(7), 925–44. <https://doi.org/10.5860/crl.80.7.925>
- Mawhinney, T., & Hervieux, S. (2022). Dissonance between perceptions and use of virtual reference methods. *College & Research Libraries*, 83(3), 503. <https://doi.org/10.5860/crl.83.3.503>

- Maness, J. M., Naper, S., & Chaudhuri, J. (2009). The good, the bad, but mostly the ugly: Adherence to RUSA Guidelines during encounters with inappropriate behavior online. *Reference & User Services Quarterly*, 49(2), 151–62. <https://doi.org/10.5860/rusq.49n2.151>
- Marsteller, M. R., & Mizzy, D. (2003). Exploring the synchronous digital reference interaction for query types, question negotiation, and patron response. *Internet Reference Services Quarterly*, 8(1-2), 149–65. https://doi.org/10.1300/J136v08n01_13
- Matteson, M. L., Salamon, J., & Brewster, L. (2011). A systematic review of research on live chat service. *Reference & User Services Quarterly*, 51(2), 172–90. <https://doi.org/10.5860/rusq.51n2.172>
- Mavodza, J. (2019). Interpreting library chat reference service transactions. *Reference Librarian*, 60(2), 122–33. <https://doi.org/10.1080/02763877.2019.1572571>
- Meert, D. L., & Given, L. M. (2009). Measuring quality in chat reference consortia: A comparative analysis of responses to users' queries. *College & Research Libraries*, 70(1), 71–84. <https://doi.org/10.5860/0700071>
- Neuhaus, P., & Marsteller, M. R. (2002). Chat reference at Carnegie Mellon University. *Public Services Quarterly*, 1(2), 29–41. https://doi.org/10.1300/J295v01n02_04
- Pomerantz, J., Luo, L., & McClure, C. R. (2006). Peer review of chat reference transcripts: Approaches and strategies. *Library and Information Science Research*, 28(1), 24–48. <https://doi.org/10.1016/j.lisr.2005.11.004>
- Radford, M. L., Costello, L., & Montague, K. E. (2022). Death of social encounters: Investigating COVID-19's initial impact on virtual reference services in academic libraries. *Journal of the Association for Information Science and Technology*, 73(11), 1594–1607. <https://doi.org/10.1002/asi.24698>
- Radford, M. L., & Connaway, L. S. (2013). Not dead yet! A longitudinal study of query type and ready reference accuracy in live chat and IM reference. *Library & Information Science Research*, 35(1), 2–13. <https://doi.org/10.1016/j.lisr.2012.08.001>
- Ronan, J., Reakes, P., & Ochoa, M. (2006). Application of reference guidelines in chat reference interactions: A study of online reference skills. *College & Undergraduate Libraries*, 13(4), 3–23. https://doi.org/10.1300/J106v13n04_02
- Ruppel, M., & Vecchione, A. (2012). 'It's research made easier!' SMS and chat reference perceptions. *Reference Services Review*, 40(3), 423–48. <https://doi.org/10.1108/00907321211254689>
- Smyth, J. B., MacKenzie, J. C. (2006). Comparing virtual reference exit survey results and transcript analysis: A model for service evaluation. *Public Services Quarterly*, 2(2–3), 85–105. https://doi.org/10.1300/J295v02n02_07
- Stoffel, B., & Tucker, T. (2004). E-mail and chat reference: Assessing patron satisfaction. *Reference Services Review*, 32(2), 120–40. <https://doi.org/10.1108/00907320410537649>
- Van Duinkerken, W., Stephens, J., & Macdonald, K. I. (2009). The chat reference interview: Seeking evidence based on RUSA's guidelines. *New Library World*, 110(3–4), 107–121. <https://doi.org/10.1108/03074800910941310>
- Ward, D. (2004). Measuring the completeness of reference transactions in online chats: Results of an unobtrusive study. *Reference & User Services Quarterly*, 44(1), 46–56. <https://www.jstor.org/stable/20864287>

Appendix A

Chat Assessment Rubric

Rubric Purpose

The purpose of this rubric is to provide measurable criteria to assess the chat reference skills of library personnel in a selected set of chat transcripts. Results of this rubric are intended to be used as a teaching/training tool to communicate expectations and give informative feedback. The assessment goal is to improve the performance of library personnel in the area of chat reference services.

Rubric Credits

This rubric was adapted in 2021, by Erin Owens and Kat Brooks, from the rubric published in Cassidy, E. D., Colmenares, A., & Martinez, M. (2014). So text me—maybe: A rubric assessment of librarian behavior in SMS reference services. *Reference and User Services Quarterly* 53(4), 300–312. [doi:10.5860/rusq.53n4.300](https://doi.org/10.5860/rusq.53n4.300).

	Accomplished – 3	Developing – 2	Beginning – 1
Listening/ Inquiring The reference interview is the heart of the reference transaction and is crucial to the success of the process. The librarian must be effective in identifying the patron's information needs and must do so in a manner that keeps patrons at ease. Strong listening and questioning skills are necessary for a positive interaction.	1. Communicates in a clearly receptive/ cordial/ encouraging manner 2. Uses open-ended questioning techniques if appropriate to encourage the patron to expand on the request or present additional information. Some examples of such questions include: – Please tell me more about your topic. – What additional information can you give me? – How much information do you need?	1. Communicates in a receptive/cordial/ encouraging manner 2. Does not use open-ended questioning techniques even when appropriate to encourage the patron to expand on the request or present additional information.	1. Communicates in an abrupt manner 2. Does not use open-ended questioning techniques even when appropriate to encourage the patron to expand on the request or present additional information.

	Accomplished – 3	Developing – 2	Beginning – 1
	3. Uses closed questions if appropriate to refine the search query. Some examples of clarifying questions are: – What types of information do you need (books, articles, etc.)? – Do you need current or historical information?	3. Uses closed questions to refine the search query. Some examples of clarifying questions are: – What types of information do you need (books, articles, etc.)? – Do you need current or historical information?	3. Does not use closed questions to refine the search query.
Interest A successful librarian must demonstrate a high degree of interest in the reference transaction. While not every query will contain stimulating intellectual challenges, the librarian should be interested in each patron's information need and should be committed to providing the most effective assistance. Librarians who demonstrate a high level of interest in the inquiries of their patrons will generate a higher level of satisfaction among users.	1. An automatic response acknowledges user questions submitted outside of library operation hours (hours during which the library is open) 2. Provided a very timely initial response (wait time for chat opening) 3. Patron assisted in a very timely manner (overall time) 4. Maintained regular contact	1. No automatic response acknowledges user questions submitted outside of library operation hours (hours during which the library is open) 2. Provided a somewhat timely initial response (wait time for chat opening) 3. Patron assisted in a somewhat timely manner (overall time) 4. Maintained regular contact	1. No automatic response acknowledges user questions submitted outside of library operation hours (hours during which the library is open) 2. Did not provide a timely initial response (wait time for chat opening) 3. Patron not assisted in a timely manner (overall time) 4. Did not maintain regular contact

	Accomplished – 3	Developing – 2	Beginning – 1
Searching The search process is the portion of the transaction in which behavior and accuracy intersect. Without an effective search, not only is the desired information unlikely to be found, but patrons may become discouraged as well. Yet many of the aspects of searching that lead to accurate results are still dependent on the behavior of the librarian.	- Names the sources to be used, when appropriate. - Works with the patron to narrow or broaden the topic when too little or too much information is identified. - Recognizes when to refer the patron to a more appropriate guide, database, library, librarian, or other resource. - Offers detailed search paths or links/URLs to needed electronic resources. Excessively long links have been converted to a shorter link (for example, using Tiny.URL) - If appropriate, detailed directions to physical resources are given, for example – Call #s and Floor #s – Room #s	- Names the sources to be used, when appropriate. - Indicates that the patron needs to narrow or broaden the topic when too little or too much information is identified. - Recognizes when to refer the patron to a more appropriate guide, database, library, librarian, or other resource when appropriate - Offers detailed search paths or links/URLs to needed electronic resources. - If appropriate, general directions to physical resources are given, for example—either call #s or floor #s, but not both.	- Does not name the sources to be used when appropriate. - Does not work with the patron to narrow or broaden the topic when too little or too much information is identified. - Does not refer the patron to a more appropriate guide, database, library, librarian, or other resource when appropriate - Does not offer detailed search paths or links/URLs to needed electronic resources. - Even if appropriate, directions to physical resources are not given.
Rubric Notes for Searching: - Librarian answers that were clearly inaccurate to the scoring group received the “Beginning” (1) score. - Closed-ended questions that required little or no searching on the part of the librarian received the “Accomplished” (3) rating.			

Follow Up The reference transaction does not end when the librarian leaves the patrons. The librarian is responsible for determining if the patrons are satisfied with the results of the search, and is also responsible for referring the patrons to other sources, even when those sources are not available in the local library	1. Offers to answer more questions or asks the patron if they need help with anything else. 2. Encourages the patron to return if they have further questions by making a statement such as—"if you don't find what you are looking for, please come back and we'll try something else" or similar. 3. Makes the patron aware of other reference services, if appropriate (email, instant chat, phone, etc.) 4. Makes arrangements, when appropriate, with the patron to research a question even after the reference transaction has been completed. 5. Refers the patron to other sources or institutions when the query cannot be answered to the satisfaction of the patron. 6. Takes care not to end the reference interview prematurely.	1. Does not offer to answer more questions or ask the patron if they need help with anything else. 2. Does not encourage the patron to return if they have further questions. 3. Makes the patron aware of other reference services, if appropriate (email, instant chat, phone, etc.) 4. Does not make arrangements, when appropriate, with the patron to research a question even after the reference transaction has been completed. 5. Does not refer the patron to other sources or institutions when the query cannot be answered to the satisfaction of the patron. 6. Takes care not to end the reference interview prematurely.	1. Does not offer to answer more questions or ask the patron if they need help with anything else. 2. Does not encourage the patron to return if they have further questions. 3. Does not make the patron aware of other reference services even when appropriate (email, instant chat, phone, etc.) 4. Does not make arrangements, when appropriate, with the patron to research a question even after the reference transaction has been completed. 5. Does not refer the patron to other sources or institutions when the query cannot be answered to the satisfaction of the patron. 6. Ends the reference interview prematurely, before answering or addressing all parts of a question.
---	---	---	--

Appendix B

Data Dictionary

Field from LibChat	Definition	Disposition for Study
Chat ID	Unique numerical identifier assigned by the system to each chat	Kept
Name	Patron's name as input (may be a pseudonym)	Deleted
Contact Information	Patron's email or phone number as entered manually (if prompted in the library's chat setup)	Deleted
IP	IP address of the chat patron's device	Deleted
Browser	Name and version of the chat patron's internet browser	Deleted
Operating System	Name and version of the chat patron's device operating system	Deleted
User Agent	Any details available about the patron's user agent, such as browser type	Deleted
Referrer	URL/Web address from which the chat patron initiated a chat with the library	Deleted
Widget	LibChat system name of the specific widget used by the chat patron to initiate a chat A library may create multiple widgets that send chats to different library departments or otherwise have different configured behavior	Deleted
Department	The department, as defined in the LibChat setup, that received the chat	Deleted
Answerer	Username of the library personnel member who answered the chat	Deleted
Timestamp	The precise date and time at which a chat was initiated by a patron	Kept
Wait Time	Length of time in seconds that a patron waited for an initial chat response from library personnel	Kept
Duration	Length of time in seconds from library personnel's first response until a chat is ended	Kept
Screensharing	Indicates whether screensharing was used during a chat; Valid values: None; Yes	Kept
Rating (0-4)	Optional patron rating of a chat interaction after a chat has ended; Valid values: numerals 1 through 4 Scores correspond to labels in patron display: Bad (1), So-so (2), Good (3), Excellent (4)	Kept
Comment	Optional patron open-ended comments accompanying a chat rating	Kept; deidentified
User Field 1	Customizable field; not in use	Deleted
User Field 2	Customizable field; not in use	Deleted
User Field 3	Customizable field; not in use	Deleted

Field from LibChat	Definition	Disposition for Study
Initial Question	Initial question entered by patron when initiating a chat	Kept; deidentified
Transfer History	Details of the chat being transferred between operators (if applicable)	Deleted
Message Count	The total number of messages exchanged in the chat	Kept
Internal Note	Any notes added to the chat by library personnel after a chat ends	Kept; deidentified
Transcript	The complete text of the conversation	Kept; deidentified
Tags	Any tags applied to the chat by library personnel, eg, for searchability or other system uses	Kept