

Evaluation of Deep Learning Approaches for Pitch Scoring in Piano Practice and Performance

Juthakan Mekkoktanphira
School of Information Technology
King Mongkut's Institute of
Technology Ladkrabang
Bangkok, Thailand
64070181@kmitl.ac.th

Priyakorn Pangwapee
School of Information Technology
King Mongkut's Institute of
Technology Ladkrabang
Bangkok, Thailand
64070133@kmitl.ac.th

Nat Dilokthanakul
School of Information Technology
King Mongkut's Institute of
Technology Ladkrabang
Bangkok, Thailand
nat@it.kmitl.ac.th

Sirasit Lochanachit
School of Information Technology
King Mongkut's Institute of
Technology Ladkrabang
Bangkok, Thailand
sirasit@it.kmitl.ac.th

Praphan Pavarangkoon
School of Information Technology
King Mongkut's Institute of
Technology Ladkrabang
Bangkok, Thailand
Corresponding author: praphan.pa@kmitl.ac.th

Nont Kanungsukkasem
School of Information Technology
King Mongkut's Institute of
Technology Ladkrabang
Bangkok, Thailand
nont@it.kmitl.ac.th

Abstract—This study evaluates deep learning approaches for pitch scoring in piano practice and performance through two experiments. The first experiment compares Gated Recurrent Units (GRU) and Transformer architectures using datasets that include diverse musical elements such as pitch, rhythm, rest, and tempo. The results demonstrate that Transformers significantly outperform GRUs in terms of accuracy and robustness across all conditions. The second experiment investigates modifications to the Transformer model, specifically increasing the number of attention heads, to assess its impact on transcribing musical sequences of varying complexity. Overall, these experiments highlight the strengths and limitations of Transformer architectures, emphasizing their potential to advance music transcription tools for education and professional applications.

Index Terms—Pitch Detection, Piano Transcription, Deep Learning, Music Education.

I. INTRODUCTION

Accurate pitch assessment is a cornerstone of musical education, particularly in learning and mastering the piano. For students, the ability to receive precise, real-time feedback during practice is critical for identifying errors, improving technique, and building confidence. Traditionally, pitch evaluation has relied on manual observation by instructors or self-assessment by students [1]. While effective to some extent, these approaches are inherently limited by human fatigue, subjective interpretation, and the lack of consistent real-time feedback. As the demand for more efficient, objective, and scalable solutions grows, advancements in artificial intelligence (AI) and deep learning present transformative opportunities for automating this process.

Deep learning has revolutionized various domains, including image recognition, speech processing, and natural language understanding. Its application to music transcription and pitch scoring is an emerging area with immense potential. By

leveraging sequence-to-sequence modeling techniques, deep learning can address the complexities of musical data, such as polyphonic textures, temporal variations, and intricate dependencies. Among the many architectures, Gated Recurrent Units (GRU) and Transformers have emerged as prominent candidates for tasks involving sequential data.

GRUs, a type of recurrent neural network (RNN), are well-suited for capturing temporal dependencies due to their ability to mitigate the vanishing gradient problem. They have been widely applied in sequence modeling tasks, including pitch detection. However, GRUs often struggle with longer sequences and complex dependencies, which are intrinsic to music. On the other hand, Transformer architectures, which utilize self-attention mechanisms, are highly effective at capturing long-range dependencies and handling diverse input structures [2]. These features make Transformers particularly advantageous for music transcription tasks, where understanding intricate temporal and harmonic relationships is essential.

This study investigates the capabilities of GRU and Transformer models for pitch scoring in piano practice through two experiments. The first experiment conducts a comparative analysis of GRU and Transformer architectures using a dataset that encompasses diverse musical elements, including pitch, rhythm, rest, and tempo. The dataset is designed to reflect real-world variations in note sequence length, structure, and complexity. The results of this experiment establish a baseline for the performance of these models, highlighting the advantages of Transformers in terms of accuracy and robustness.

The second experiment focuses on enhancing the Transformer architecture by modifying the multi-head attention mechanism. Specifically, it examines the impact of increasing the number of attention heads on the model's ability to transcribe musical notes and sequences of varying complex-

ities. This experiment aims to uncover whether architectural modifications can further improve the Transformer’s performance, particularly in handling challenging scenarios with more intricate data structures.

By conducting these two experiments, this paper aims to provide a comprehensive evaluation of deep learning models for automated music transcription. The results contribute to the understanding of the strengths and limitations of GRU and Transformer architectures, offering insights into their potential applications in educational and professional settings. Moreover, this research underscores the importance of model optimization and architectural tuning for advancing music transcription technologies, paving the way for innovations that enhance learning experiences and support professional musicians.

II. RELATED WORKS

Recent advancements in artificial intelligence and deep learning have significantly advanced the field of music transcription and pitch detection. These technologies have facilitated the development of automated systems that analyze musical sequences, recognize pitch, and transcribe performances with high accuracy. This section reviews key contributions and existing research in this domain, emphasizing their relevance to pitch scoring and transcription for piano practice and performance.

The Transformer architecture, introduced by Vaswani et al. [3], has become a foundational model in sequence-to-sequence tasks. Its self-attention mechanism enables it to effectively capture long-range dependencies, a critical feature for processing complex musical sequences. Transformers have shown superior performance across various domains, including natural language processing and audio signal analysis, making them a promising candidate for music transcription applications. On the other hand, GRUs, proposed by Cho et al. [4], address the vanishing gradient problem in traditional RNNs. GRUs are efficient in learning sequential dependencies and have been successfully applied in numerous tasks requiring temporal data analysis, including pitch detection.

Music transcription often involves capturing pitch information in polyphonic settings. Wang et al. [5] proposed a harmonic structure-based neural network that uses convolutional layers to capture local frequency patterns and harmonic features for robust pitch detection. Their research demonstrated significant improvements in detecting active pitches, especially in piano music, where overlapping frequencies pose a challenge. Similarly, Hawthorne et al. [6] explored the use of sequence-to-sequence models for piano transcription, integrating onset and frame detection to enhance accuracy. Their models, trained on datasets such as MAESTRO and MAPS, established a benchmark for music transcription tasks.

Meanwhile, multitask learning frameworks, such as the MT3 model introduced by Gardner et al. [7], have advanced multitrack music transcription. By leveraging the Transformer architecture, MT3 processes raw audio spectrograms and transcribes multiple instruments simultaneously, showcasing the

scalability of Transformer-based systems in handling complex musical data.

The application of deep learning to music transcription has extended beyond academic research into practical use cases. For example, models integrating convolutional and recurrent architectures have been used for real-time pitch tracking, offering immediate feedback for musicians. Hawthorne et al. [8] further refined the onset and frame detection framework to produce symbolic representations of piano performances, enabling accurate transcription and evaluation of complex pieces. These contributions underscore the importance of combining robust architectures with effective data representations to achieve state-of-the-art results in music transcription.

III. EXPERIMENTS

This study evaluates the performance of GRU and Transformer models for pitch scoring in piano practice through two complementary experiments. The first experiment establishes a comparative baseline for the GRU and Transformer architectures, while the second experiment investigates the impact of modifying the Transformer architecture by varying the number of attention heads.

A. Experiment 1: Comparative Evaluation of GRU and Transformer Models

1) *Objective:* The first experiment aims to compare the performance of GRU and Transformer architectures in processing musical sequences. This evaluation focuses on how well each model handles diverse musical elements, including pitch, rhythm, rest, and tempo, across varying sequence lengths and structures.

2) *Dataset:* The dataset for this experiment was designed to reflect real-world variations in music. It consists of sequences ranging from 3 to 10 notes, incorporating both uniform and non-uniform rest lengths, varying tempos, and complex rhythms. Figure 1(a) illustrates an example of a sequence with four half notes, without any rests, highlighting the model’s ability to handle continuous notes seamlessly. Figure 1(b) shows a sequence of four half notes with equal rests between each note, demonstrating the model’s capability to maintain uniform spacing. Finally, Figure 1(c) depicts a sequence of four notes with diverse rhythms, where rests are not consistently placed and differ in length, showcasing the model’s proficiency in managing irregular rhythmic patterns and non-uniform rest distributions. The dataset was divided into training (80%), validation (10%), and test (10%) sets, ensuring a balanced representation across all variables.

3) *Methodology:* Both GRU and Transformer models were trained on identical datasets. MIDI files were converted to WAV format and processed into mel spectrograms using Short-Time Fourier Transform (STFT) with the following parameters: 1024 FFT components, a hop size of 512, 1024 mel bands, and a 16,000 Hz sample rate. Label preprocessing included one-hot encoding and the addition of Start of Sequence (SOS) and End of Sequence (EOS) tokens to ensure proper sequence alignment.



(a) Four half notes without rests.



(b) Four half notes with equal rests between each note.



(c) Four notes with varying rhythms and different rests.

Fig. 1: Examples of note sequences demonstrating different rest patterns and rhythms.

The GRU model was configured with two hidden layers of 512 units each, while the Transformer model utilized three encoder and decoder layers, single-head attention, and a feedforward network with 512 hidden units. Both models were trained using the Adam optimizer, and their performance was evaluated using exact score matching, which considers a sequence correct only if all predicted notes match the ground truth.

B. Experiment 2: Impact of Attention Heads in Transformer Architecture

1) *Objective*: The second experiment investigates how increasing the number of attention heads in the Transformer architecture affects its ability to transcribe musical sequences. This experiment focuses on both simpler and more complex scenarios to determine whether architectural modifications enhance performance.

2) *Dataset*: The same dataset from Experiment 1 was used, allowing direct comparisons between the experiments. It encompasses diverse musical elements, including uniform and non-uniform rests, varying tempos, and rhythms, with note sequences ranging from 3 to 10 notes.

3) *Methodology*: The Transformer model was modified to use 4, 8, and 16 attention heads while keeping all other hyperparameters constant. The preprocessing pipeline remained the same as in Experiment 1, ensuring consistency in data handling and evaluation criteria. Performance was assessed using exact score matching, emphasizing transcription accuracy.

C. Evaluation Metrics

Both experiments utilized exact score matching as the primary evaluation metric. This rigorous metric evaluates a model's ability to transcribe musical sequences by requiring all notes in a sequence to match the ground truth for a correct

score. This ensures a stringent and consistent assessment across both experiments.

IV. RESULTS AND DISCUSSION

The experiments conducted aimed to evaluate the performance of Transformer and GRU-based sequence-to-sequence models for musical note transcription, focusing on their robustness and accuracy across diverse musical scenarios. Additionally, a second experiment studied the impact of modifying the multi-head attention mechanism in Transformer models. Below, the results of both experiments are presented, including all 10 tables, and their implications are discussed.

A. Experiment 1: Comparative Performance of GRU and Transformer Models

1) *Results*: The first experiment clearly demonstrate the superior performance of the Transformer model compared to the GRU-based sequence-to-sequence model across a range of transcription scenarios. For sequences involving three consecutive notes under simpler conditions, such as consistent tempo and rhythm, the Transformer achieved an exact score of up to 0.995, as shown in Table II. In contrast, Table I reveals that while the GRU model performed exceptionally well in structured conditions, achieving a perfect exact score of 1.00 for sequences with equal rests, its accuracy significantly declined in scenarios with greater complexity, such as varying tempos and rhythms, with scores dropping as low as 0.35 in cases without rests.

When analyzing sequences of 3 to 10 notes, the performance gap between the two models becomes even more pronounced. As shown in Table III, the GRU model's performance deteriorates sharply in scenarios involving complex variations, with exact scores plummeting to as low as 0.005 under conditions of different tempos and rhythms, especially when dealing with unequal rests not placed between every note. Conversely, Table IV highlights the Transformer model's ability to maintain a high degree of accuracy, achieving scores of up to 0.85 in structured scenarios and consistently outperforming the GRU model even in the most challenging conditions, where its lowest score of 0.555 far exceeded the GRU model's performance.

Overall, these results highlight a consistent trend: while the GRU model shows acceptable performance in controlled and simple settings, its accuracy diminishes substantially as the complexity of the musical patterns increases. On the other hand, the Transformer model demonstrates remarkable robustness, retaining high accuracy even under diverse and intricate conditions.

2) *Discussion*: The findings emphasize the advantages of the Transformer architecture, particularly its ability to handle complex musical transcription scenarios. The self-attention mechanism inherent to the Transformer plays a critical role in capturing long-range dependencies and temporal relationships, which are essential for accurately transcribing sequences with variable rhythms, tempos, and rests. This capability is reflected

TABLE I: Exact scores of the GRU-based sequence-to-sequence model for 3 consecutive notes.

	Same Tempo, Same Rhythm	Different Tempo, Same Rhythm	Same Tempo, Different Rhythm	Different Tempo, Different Rhythm
No Rest	0.995	0.905	0.495	0.35
Equal Rest Between Notes	1.00	0.965	0.735	0.47
Unequal Rest Between Notes	0.845	0.65	0.37	0.29
Equal Rest Not Between Every Note	0.945	0.67	0.365	0.33
Unequal Rest Not Between Every Note	0.8	0.655	0.38	0.285

TABLE II: Exact scores of the Transformer model for 3 consecutive notes.

	Same Tempo, Same Rhythm	Different Tempo, Same Rhythm	Same Tempo, Different Rhythm	Different Tempo, Different Rhythm
No Rest	0.985	0.985	0.885	0.645
Equal Rest Between Notes	0.995	0.96	0.975	0.94
Unequal Rest Between Notes	0.94	0.94	0.94	0.84
Equal Rest Not Between Every Note	0.94	0.96	0.935	0.83
Unequal Rest Not Between Every Note	0.95	0.91	0.855	0.845

TABLE III: Exact scores of the GRU-based sequence-to-sequence model for 3-10 notes.

	Same Tempo, Same Rhythm	Different Tempo, Same Rhythm	Same Tempo, Different Rhythm	Different Tempo, Different Rhythm
No Rest	0.385	0.07	0.00	0.005
Equal Rest Between Notes	0.22	0.03	0.025	0.005
Unequal Rest Between Notes	0.05	0.015	0.005	0.005
Equal Rest Not Between Every Note	0.04	0.00	0.00	0.005
Unequal Rest Not Between Every Note	0.02	0.015	0.01	0.015

TABLE IV: Exact scores of the Transformer model for 3-10 notes.

	Same Tempo, Same Rhythm	Different Tempo, Same Rhythm	Same Tempo, Different Rhythm	Different Tempo, Different Rhythm
No Rest	0.85	0.7	0.395	0.075
Equal Rest Between Notes	0.17	0.665	0.57	0.27
Unequal Rest Between Notes	0.665	0.4	0.37	0.26
Equal Rest Not Between Every Note	0.685	0.535	0.335	0.265
Unequal Rest Not Between Every Note	0.555	0.35	0.175	0.19

in the model’s consistent performance across all experimental conditions, as seen in Tables II and IV.

In contrast, the GRU-based sequence-to-sequence model, despite performing well in structured scenarios such as sequences with equal rests and consistent rhythms, struggles significantly with more complex patterns. The substantial decline in its accuracy, as evidenced in Tables I and III, indicates that the GRU model’s reliance on sequential processing limits its ability to generalize to diverse and sparse input conditions. This limitation becomes particularly apparent as the number of notes in the sequence increases or when rests and rhythms vary significantly.

These results have important implications for the selection of models in musical transcription tasks. The superior performance of the Transformer across varying conditions underscores its suitability for real-world transcription scenarios, where variations in tempo, rhythm, and rests are inevitable. Furthermore, the robustness of the Transformer highlights its potential for broader applications in music analysis and processing.

Looking ahead, future research could explore the integration of GRU and Transformer architectures to leverage their respective strengths. Hybrid models might address the GRU’s limitations in capturing long-range dependencies while

preserving its strengths in sequential processing. Additionally, further investigation into optimizing hyperparameters for specific datasets and transcription tasks could provide deeper insights into improving model performance.

In conclusion, the experiments confirm that the Transformer model offers significant advantages over the GRU-based model, particularly in handling complex musical structures. Its consistent performance and adaptability across diverse conditions make it a promising choice for deep learning-based piano transcription applications.

B. Experiment 2: Impact of Multi-Head Attention in Transformer Models

1) *Results:* The second experiment examined the impact of varying the number of attention heads in the Transformer architecture, specifically 4, 8, and 16 heads, on the transcription performance. Results for three consecutive notes are shown in Tables V, VI, and VII, while Tables VIII, IX, and X present the findings for sequences ranging from 3 to 10 notes. The results reveal that while all configurations performed well under simpler conditions—such as sequences with consistent tempo and rhythm—the introduction of greater complexity, such as variations in tempo and rhythm, highlighted notable differences in model performance.

For sequences with three consecutive notes, the model with 4 heads consistently achieved the highest scores across most scenarios, particularly under complex conditions involving different tempos and rhythms, as shown in Table V. The performance advantage was less pronounced for simpler cases, where all configurations demonstrated comparable accuracy. Tables VI and VII show that increasing the number of heads to 8 or 16 did not yield consistent improvements and, in some cases, resulted in diminished performance.

For sequences ranging from 3 to 10 notes, the trend became even more apparent. Tables VIII, IX, and X demonstrate that the 4-head configuration consistently outperformed both the 8-head and 16-head configurations across all experimental conditions. Notably, in the most challenging scenarios, such as sequences with unequal rests not between every note and varying tempos and rhythms, the 4-head model exhibited a clear performance advantage. This suggests that increasing the number of attention heads does not necessarily translate to better performance in complex transcription tasks, especially when handling intricate patterns and longer sequences.

2) *Discussion:* The findings highlight the nuanced role of attention heads in Transformer-based transcription models. While adding more attention heads increases the model's representational capacity, this also introduces additional computational overhead and potential overfitting, particularly when applied to datasets with significant variability in tempo, rhythm, and rests. The superior performance of the 4-head configuration suggests that it achieves an optimal balance between computational efficiency and model complexity, enabling it to effectively capture the temporal and sequential dependencies in the data without becoming overwhelmed by extraneous patterns or noise.

In simpler transcription scenarios, such as those involving consistent tempos and rhythms, the differences in performance among the 4-head, 8-head, and 16-head configurations were minimal. This indicates that the additional capacity provided by more attention heads is underutilized in such cases. However, as the complexity of the sequences increased, the 4-head model demonstrated greater robustness, likely due to its ability to focus on critical dependencies without being hindered by excessive parameterization.

C. Overall Implications

The experiments collectively underscore the importance of carefully tuning the Transformer model's hyperparameters to match the complexity of the transcription task. While the self-attention mechanism inherently equips the Transformer with the capacity to handle intricate temporal relationships, the number of attention heads plays a pivotal role in determining the model's effectiveness and efficiency. The results suggest that for musical transcription tasks involving sequences with varying tempos, rhythms, and rests, a smaller number of attention heads, such as 4, provides the best trade-off between accuracy and computational demands.

D. Limitations and Future Work

This study focused on single-octave datasets with controlled variations in tempo, rhythm, and rests. While this provided a controlled environment for evaluating the effects of attention head configurations, it limits the generalizability of the findings to real-world transcription scenarios involving polyphonic music, multiple octaves, and diverse musical styles. Future research should explore the applicability of these results to broader and more complex datasets. Additionally, the integration of hybrid architectures that combine the strengths of GRU and Transformer models could further enhance transcription accuracy and efficiency. Exploring dynamic adjustment mechanisms for attention head configurations based on the complexity of the input data may also yield promising results.

V. CONCLUSION

This study evaluated GRU and Transformer-based models for musical note transcription, highlighting the Transformer's superior ability to handle complex temporal dependencies and diverse musical scenarios. While both models performed well in simpler conditions, the Transformer consistently outperformed the GRU in complex datasets. Additionally, increasing the Transformer's attention heads did not consistently improve performance, with the 4-head configuration offering the best balance of efficiency and accuracy. These findings underscore the Transformer's suitability for scalable and adaptable music transcription tools, particularly in educational and professional contexts. Future research should explore broader datasets, polyphonic music, and hybrid architectures to further enhance transcription accuracy and efficiency.

REFERENCES

- [1] K. Zhukov, "Evaluating new approaches to teaching of sight-reading skills to advanced pianists," *Music Education Research*, vol. 16, no. 1, pp. 70–87, August 2013.
- [2] A. Dosovitskiy, L. Beyer, A. Kolesnikov, D. Weissenborn, X. Zhai, T. Unterthiner, M. Dehghani, M. Minderer, G. Heigold, S. Gelly and J. Uszkoreit, and N. Houlsby, "An Image is Worth 16x16 Words: Transformers for Image Recognition at Scale," in *Proc. the 9th International Conference on Learning Representations (ICLR)*, May 2021.
- [3] A. Vaswani, N. Shazeer, N. Parmar, J. Uszkoreit, L. Jones, A. N. Gomez, L. Kaiser, and I. Polosukhin, "Attention is all you need," in *Proc. the 31st International Conference on Neural Information Processing Systems (NIPS'17)*, pp. 6000–6010, December 2017.
- [4] K. Cho, B. V. Merriënboer, C. Gulcehre, D. Bahdanau, F. Bougares, H. Schwenk, and Y. Bengio, "Learning Phrase Representations using RNN Encoder–Decoder for Statistical Machine Translation," in *Proc. the 2014 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pp. 1724–1734, October 2014.
- [5] X. Wang, L. Liu and Q. Shi, "Harmonic Structure-Based Neural Network Model for Music Pitch Detection," in *Proc. the 19th IEEE International Conference on Machine Learning and Applications (ICMLA)*, pp. 87–92, December 2020.
- [6] C. Hawthorne, E. Elsen, J. Song, A. Roberts, I. Simon, C. Raffel, J. Engel, S. Oore, and D. Eck, "Onsets and Frames: Dual-Objective Piano Transcription," in *Proc. the 19th International Society for Music Information Retrieval Conference (ISMIR)*, pp. 50–57, September 2018.
- [7] J. Gardner, I. Simon, E. Manilow, C. Hawthorne, and J. Engel, "MT3: Multi-Task Multitrack Music Transcription," in *Proc. the 10th International Conference on Learning Representations (ICLR)*, April 2022.
- [8] C. Hawthorne, I. Simon, R. Swavely, E. Manilow, and J. Engel, "Sequence-to-Sequence Piano Transcription with Transformers," in *Proc. the 22nd International Society for Music Information Retrieval Conference (ISMIR)*, pp. 246–253, November 2021.

TABLE V: Exact scores for 3 consecutive notes (4 heads).

	Same Tempo, Same Rhythm	Different Tempo, Same Rhythm	Same Tempo, Different Rhythm	Different Tempo, Different Rhythm
No Rest	0.975	1.00	0.99	0.95
Equal Rest Between Notes	0.95	1.00	0.98	0.995
Unequal Rest Between Notes	0.99	0.975	0.94	0.865
Equal Rest Not Between Every Note	0.995	0.98	0.93	0.935
Unequal Rest Not Between Every Note	0.985	0.96	0.96	0.86

TABLE VI: Exact scores for 3 consecutive notes (8 heads).

	Same Tempo, Same Rhythm	Different Tempo, Same Rhythm	Same Tempo, Different Rhythm	Different Tempo, Different Rhythm
No Rest	0.99	0.965	1.00	0.95
Equal Rest Between Notes	0.95	0.995	0.98	0.99
Unequal Rest Between Notes	0.995	0.995	0.99	0.95
Equal Rest Not Between Every Note	0.995	0.99	0.97	0.935
Unequal Rest Not Between Every Note	1.00	0.99	0.855	0.9

TABLE VII: Exact scores for 3 consecutive notes (16 heads).

	Same Tempo, Same Rhythm	Different Tempo, Same Rhythm	Same Tempo, Different Rhythm	Different Tempo, Different Rhythm
No Rest	0.995	0.99	0.91	0.96
Equal Rest Between Notes	1.00	0.97	1.00	0.995
Unequal Rest Between Notes	0.99	0.97	0.925	0.95
Equal Rest Not Between Every Note	1.00	1.00	0.91	0.91
Unequal Rest Not Between Every Note	0.965	0.98	0.96	0.94

TABLE VIII: Exact scores for 3-10 notes (4 heads).

	Same Tempo, Same Rhythm	Different Tempo, Same Rhythm	Same Tempo, Different Rhythm	Different Tempo, Different Rhythm
No Rest	0.76	0.675	0.28	0.17
Equal Rest Between Notes	0.16	0.725	0.41	0.235
Unequal Rest Between Notes	0.62	0.475	0.235	0.255
Equal Rest Not Between Every Note	0.64	0.465	0.275	0.135
Unequal Rest Not Between Every Note	0.58	0.52	0.26	0.215

TABLE IX: Exact scores for 3-10 notes (8 heads).

	Same Tempo, Same Rhythm	Different Tempo, Same Rhythm	Same Tempo, Different Rhythm	Different Tempo, Different Rhythm
No Rest	0.75	0.565	0.13	0.155
Equal Rest Between Notes	0.125	0.68	0.445	0.385
Unequal Rest Between Notes	0.485	0.355	0.34	0.205
Equal Rest Not Between Every Note	0.65	0.39	0.2	0.155
Unequal Rest Not Between Every Note	0.52	0.395	0.28	0.275

TABLE X: Exact scores for 3-10 notes (16 heads).

	Same Tempo, Same Rhythm	Different Tempo, Same Rhythm	Same Tempo, Different Rhythm	Different Tempo, Different Rhythm
No Rest	0.68	0.19	0.165	0.125
Equal Rest Between Notes	0.175	0.56	0.34	0.1
Unequal Rest Between Notes	0.39	0.135	0.125	0.13
Equal Rest Not Between Every Note	0.56	0.43	0.205	0.195
Unequal Rest Not Between Every Note	0.4	0.125	0.205	0.175