



Rodent-Bench

Anonymous Author(s)

Affiliation

Address

email

Abstract

We present Rodent-Bench, a novel benchmark designed to evaluate the ability of Multimodal Large Language Models (MLLMs) to annotate rodent behaviour footage. We evaluate state-of-the-art MLLMs, including Gemini-2.5-Pro, Gemini-2.5-Flash and Qwen-VL-Max, using this benchmark and find that none of these models perform strongly enough to be used as an assistant for this task. Our benchmark encompasses diverse datasets spanning multiple behavioral paradigms including social interactions, grooming, scratching, and freezing behaviors, with videos ranging from 10 minutes to 35 minutes in length. We provide two benchmark versions to accommodate varying model capabilities and establish standardized evaluation metrics including second-wise accuracy, macro F1, mean average precision, mutual information, and Matthew's correlation coefficient. While some models show modest performance on certain datasets (notably grooming detection), overall results reveal significant challenges in temporal segmentation, handling extended video sequences, and distinguishing subtle behavioral states. Our analysis identifies key limitations in current MLLMs for scientific video annotation and provides insights for future model development. Rodent-Bench serves as a foundation for tracking progress toward reliable automated behavioral annotation in neuroscience research.

1 Introduction

Behavioral analysis is fundamental to neuroscience and biomedical research, yet manual annotation of animal behavior videos remains a time-consuming bottleneck that limits research scale and reproducibility (Sturm et al., 2020; Mathis & Mathis, 2020). While Multimodal Large Language Models (MLLMs) have shown impressive capabilities in vision-language tasks (Fu et al., 2024; Yin et al., 2024), their application to specialized scientific domains like behavioral analysis remains largely unexplored. MLLMs offer particular promise for scientific annotation tasks as they can potentially handle diverse behavioral paradigms through natural language instructions without requiring specialized model training for each new behavior or experimental setup.

Unlike general computer vision tasks, behavioral analysis requires models to identify subtle, context-dependent actions, maintain temporal coherence across extended sequences, and produce structured outputs aligned with ethological frameworks. Traditional computer vision approaches require training specialized models for each behavioral task, but MLLMs could streamline this process by accepting task descriptions in natural language and adapting to new behaviors without retraining. Existing video understanding benchmarks inadequately address these scientific requirements, creating a significant gap between current MLLM capabilities and practical research applications.

We present **Rodent-Bench-Short** and **Rodent-Bench-Long**, the first comprehensive benchmarks for evaluating MLLMs on rodent behavioral annotation tasks. Our benchmarks encompass diverse

37 datasets spanning multiple behavioral paradigms and provides standardized evaluation metrics to
 38 assess current model capabilities. We evaluate state-of-the-art MLLMs including Gemini-2.5-Pro,
 39 Gemini-2.5-Flash, and Qwen-VL-Max, revealing significant performance gaps that limit their utility
 40 as research assistants. While some of these models show fair performance on some datasets, our
 41 analysis identifies specific challenges in temporal segmentation, long video processing, and handling
 42 varied experimental conditions, providing insights for future improvements in scientific applications
 43 of multimodal models.

44 2 Related Work

45 The emergence of Multimodal Large Language Models (MLLMs) has opened new possibilities
 46 for video understanding tasks across diverse domains. Recent comprehensive benchmarks such as
 47 Video-MME (Fu et al., 2024) have evaluated state-of-the-art MLLMs including GPT-4 and Gemini
 48 on video analysis tasks, revealing significant challenges in temporal reasoning and long-form video
 49 understanding. Surveys on video understanding with large language models (Tang et al., 2025; Wang
 50 et al., 2024) highlight the emergent capabilities of these systems for multi-granularity reasoning,
 51 while identifying key limitations in handling long-form videos and maintaining alignment between
 52 visual and textual modalities. Despite these advances, the application of MLLMs to specialized
 53 scientific domains remains under-explored, with recent work suggesting significant potential for
 54 leveraging these models in natural science research (Yin et al., 2024; Testard et al., 2024).

55 Traditional animal behavior analysis has undergone significant transformation with the advent of
 56 deep learning and computer vision techniques. Deep learning-based behavioral analysis systems have
 57 demonstrated the ability to reach human-level accuracy in recognizing specific ethological behaviors
 58 (Sturm et al., 2020), with markerless pose estimation tools like DeepLabCut enabling robust tracking
 59 of individual body parts in freely moving rodents (Mathis & Mathis, 2020). Specialized tools such
 60 as DeepBehavior (Arac et al., 2019), ezTrack (Pennington et al., 2019), MoSeq (Wiltshko et al.,
 61 2015), SLEAP Pereira et al. (2022) and real-time behavior recognition systems (de Chaumont et al.,
 62 2022) have been developed specifically for automated analysis of animal behavior. However, these
 63 systems typically require task-specific training and lack the flexibility and generalization capabilities
 64 that modern MLLMs potentially offer. The specific application of MLLMs to behavioral annotation
 65 tasks in laboratory settings remains largely unexplored, representing a significant gap that our
 66 Rodent-Bench benchmark aims to address.

67 3 Rodent-Bench

68 We produced two benchmarks: **Rodent-Bench-Short**, with videos up to 10 minutes long; and
 69 **Rodent-Bench-Long**, with videos up to 35 minutes long. We created these two versions because
 70 current MLLMs have varying video length limitations—while models like Gemini can process videos
 71 up to 1 hour, others like Qwen-VL-Max are restricted to 10 minutes or less. This dual-benchmark
 72 approach ensures compatibility across all evaluated models and enables investigation of how video
 73 length affects annotation performance.

74 Along with this we suggest evaluation metrics. The task posed to the MLLM is to annotate each video,
 75 determining which of a fixed set of behaviours is occurring and to produce a JSON file segmenting
 76 each video into discrete non-overlapping time segments with behaviour labels.

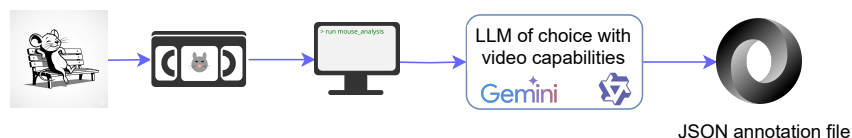


Figure 1: Workflow for annotating rodent videos.

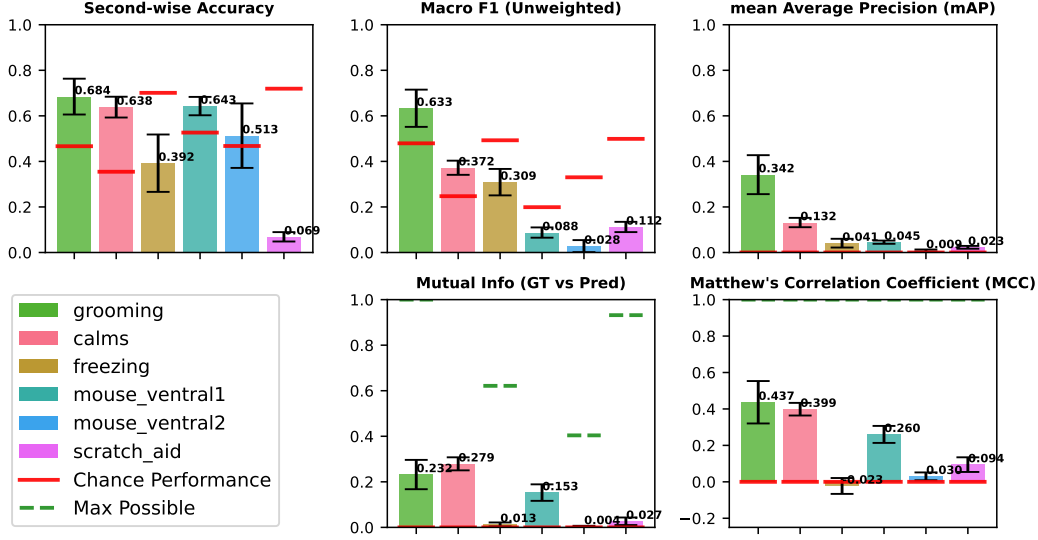


Figure 2: Performance metrics for Gemini-2.5-Pro across all datasets. Each metric shows substantial variation across behavioral paradigms, with the grooming detection dataset achieving the highest performance across most metrics. Social behaviors (CaMS21) show moderate performance, while challenging datasets like freezing and scratch detection exhibit poor performance approaching chance levels. Dashed lines indicate theoretical maximum performance where applicable. Error bars represent $2 \times$ standard error across videos within each dataset. The consistently low performance on certain datasets highlights the difficulty of fine-grained temporal behavioral annotation for current MLLMs.

3.1 Data Collection

We collected our data from several openly available datasets, as well as one private dataset which we now make freely available.

Caltech Rodent Social Interactions (CaMS21):

The CaMS21 dataset (Sun et al., 2021) is intended for multi-agent behaviour modelling. It consists of footage of multiple rodents interacting socially, with 6 million un-labelled frames and 1 million labelled frames. The labelled frames consist of both frame level behaviour and pose tracking annotations. For our purposes we use the first 25 labelled videos in the training set.

Rodent Grooming Detection Annotated Dataset:

The rodent grooming dataset (Geuther et al., 2021) was collected in order to train a neural network rodent grooming classifier. It consists of 1,253 video clips with 2,637,363 frames. Each frame is labelled “Grooming” or “Not Grooming”. We use the first 25 videos in the training set for our eval.

Mouse-Ventral 1&2: We use the Mouse-Ventral 1&2 subsets of the Deep Ethogram dataset (Bohnslav et al., 2021). These consists of 30 minute videos of a rodents shot from below, 16 for MV2 and 28 for MV1, the videos are annotated with behaviour labels. In the Mouse-Ventral1 subset the rodents are either “grooming”, “digging”, “scratching”, “licking” or “background” (neither scratching nor licking). In the Mouse-Ventral2 subset the rodents are either “scratching”, “licking” or “background”.

Scratch-AID: The Scratch-AID dataset (Yu et al., 2022) was collected to train a neural network CRNN rodent scratching classification model. The dataset consists of 40 videos of rodents shot from below, the rodents were injected with an itching agent causing them to scratch compulsively. The model trained especially for this task achieved 97.6% recall and 96.9% precision on previously unseen test videos.

Freezing: Our collaborators have given us access to nine videos of rodents displaying a “freezing” behaviour. This is a behavior distinct from resting, and is characterised by the ears being oriented towards the front indicating alert but immobile behavior (Blanchard & Blanchard, 1969). There are three types of videos, and three videos of each type: Low freezing, high freezing and extinction.

103 Extinction is a behavioral paradigm where the conditioned freezing response is gradually reduced
 104 through repeated exposure to the conditioned stimulus without the unconditioned stimulus.

105 This dataset is particularly important for evaluating MLLMs because freezing behavior presents
 106 challenges that traditional pose estimation approaches cannot address. While tools like DeepLabCut
 107 excel at tracking body parts and movements, they cannot distinguish between freezing (an active
 108 fear response) and other motionless states such as sleeping, resting, or general inactivity. These
 109 behaviors are not easily distinguishable when relying solely on pose or movement data. MLLMs,
 110 with their ability to integrate visual context, temporal patterns, and behavioral understanding, may
 111 offer advantages for this subtle but scientifically important distinction.

112 **Rodent-Bench-Short:** Some MLLMs will not accept long video files (30 minutes to an hour), so to
 113 evaluate these models we produce a shortened version of the dataset in which any file longer than 10
 114 minutes is shortened to that length. We evaluate all models on both datasets for comparison.

115 3.2 Metrics

116 **Second-wise accuracy:** We treat each second as a binary classification problem: is the behaviour
 117 in that second correctly classified or not. We then report the proportion of seconds in which the
 118 behaviour was correctly classified.

119 **Macro F1:** We calculate the F1 score for each class and average with no weighting.

120 For each behavior class c :

$$\text{Precision}_c = \frac{TP_c}{TP_c + FP_c} \quad (1)$$

$$\text{Recall}_c = \frac{TP_c}{TP_c + FN_c} \quad (2)$$

$$F1_c = \frac{2 \cdot \text{Precision}_c \cdot \text{Recall}_c}{\text{Precision}_c + \text{Recall}_c} \quad (3)$$

121 Macro F1 is the unweighted average across all classes:

$$\text{Macro F1} = \frac{1}{|C|} \sum_{c \in C} F1_c \quad (4)$$

122 where $|C|$ is the number of behavior classes, TP_c is true positives for class c , FP_c is false positives,
 123 and FN_c is false negatives.

124 **mean Average Precision (mAP):** This is calculated by comparing predicted and ground truth
 125 behaviour segments across a range of IoU (Intersection over Union) thresholds (from 0.1 to 0.9)
 126 (Henderson & Ferrari, 2016). For each threshold, we match predicted segments to ground truth
 127 segments of the same behaviour if their IoU exceeds the threshold, counting true positives (TP), true
 128 negatives (TN), false positives (FP), and false negatives (FN). Precision and recall are computed at
 129 each threshold, and the average precision is accumulated as the sum of precision values weighted by
 130 the change in recall, this is to approximate the area under the precision-recall curve. The final mAP is
 131 the total of these values, providing a single metric that summarizes how well the predictions align
 132 with the ground truth across different levels of overlap.

133 **Mutual Information:** We calculate the mutual information between the ground-truth second-wise
 134 labels and the predicted labels.

135 **Matthew’s Correlation Coefficient (MCC):** This is a correlation coefficient between -1 and 1. It is
 136 calculated:

$$\text{MCC} = \frac{(TP \times TN) - (FP \times FN)}{\sqrt{(TP + FP)(TP + FN)(TN + FP)(TN + FN)}}$$

137 **Dataset label entropy weighted Matthew’s Correlation Coefficient:** To provide a singular score
 138 which takes into account differing datasets “difficulty”:

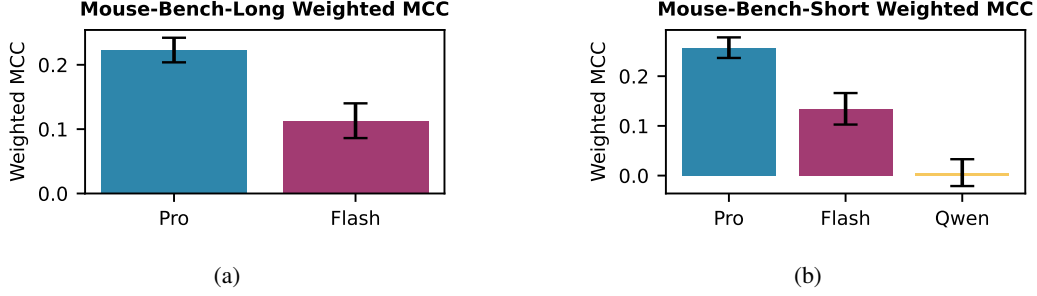


Figure 3: Weighted Matthew’s Correlation Coefficient (MCC) performance across models. (a) Rodent-Bench-Long: Gemini-2.5-Pro achieves the highest performance with lower variance compared to Gemini-2.5-Flash. (b) Rodent-Bench-Short: Similar performance hierarchy with Gemini-2.5-Pro outperforming Flash, while Qwen-VL-Max shows near-chance performance. Error bars represent $2 \times$ standard error across datasets. All models show modest performance levels, indicating substantial room for improvement in behavioral annotation tasks.

$$w_i = \frac{(H_i + \epsilon) \cdot T_i}{\sum_j (H_j + \epsilon) \cdot T_j}$$

Where, H_i is the entropy of dataset i , T_i is the duration of dataset i in seconds, $\epsilon = 10^{-8}$ is a small constant to avoid zero weights. This weighting scheme assigns higher weight to longer datasets with more diverse labels, which should be more challenging.

We prioritize mutual information, MCC, and mAP metrics for our primary analysis. Mutual information and MCC provide interpretable baselines, both equalling zero when predictions are statistically independent of ground truth labels, making chance-level performance easily identifiable across all datasets. The mAP metric is valuable for evaluating temporal segmentation quality, as it directly measures how well predicted behavioral segments align with ground truth boundaries across multiple IoU thresholds. In contrast, metrics like second-wise accuracy and macro F1 have dataset-dependent chance baselines that vary with class distributions and choice of random baseline strategy (uniform random vs. frequency-matched random prediction), complicating cross-dataset comparisons and performance interpretation.

4 Experiments

To provide an idea of how models currently perform we evaluate some of the available MLLM’s on this benchmark.

4.1 Experimental Setup

Models We evaluate our benchmark on three MLLMs. As of July 2025, a number of MLLMs which claim support for video actually just sample frames from the video at regular intervals and use these (e.g. Qwen-VL-Max). We use Gemini-2.5-Flash, Gemini-2.5-Pro and Qwen-VL-Max. Specifications of these models can be found in Appendix B.

Prompting Strategy Because our dataset is heterogenous we use different prompts for each of the sub-datasets, these appear in appendix G. Despite this we tried to use a similar prompt for each dataset in the interest of fairness. The same prompt is used for all models.

Computational Cost: For the complete benchmark, costs are approximately \$7 for Gemini-2.5-Pro, \$7 for Qwen-VL-Max, and under \$1 for Gemini-2.5-Flash.

4.2 Results

We evaluate three MLLMs on our benchmark. Further figures showing individual dataset performance for each model can be found in appendices D, E and F.

167 In figure 3a you can see that Gemini Pro outperforms Flash in both absolute performance, as well as
168 variability.

169 In 3b, evaluating on the shortened dataset, we see a similar comparison between Pro and Flash, while
170 Qwen-VL-Max performs no better than chance.

171 We notice that both Flash and Qwen models struggle with correct formatting, sometimes the wrong
172 key for certain segments (“end_long_time” instead of “end_time”) or in the case of Qwen-VL it will
173 simply stop partway through a segment, rendering the JSON file unreadable (without modifications).

174 We speculate that these models perform strongest on datasets with shorter videos on average, that
175 have clearly defined labels which depend only on behaviour (i.e do not require the rodent to be in
176 a particular position in the cage when acting for that label to apply), and have behaviours that last
177 at least a few seconds. They perform weakly on videos which have visual filters applied but that
178 require some degree of colour recognition for labelling (i.e “the feeding box is at the back of the
179 cage and is black”), or that are taken from a “non-standard” (i.e not front facing, from above or from
180 below) camera angle, or that have very short or ambiguous behaviours. For instance the freezing
181 dataset features behaviours shorter than a second, and the difference between the rodent “freezing”
182 and simply not moving is quite subtle.

183 5 Limitations

184 Our benchmark has several important limitations. First, ground-truth annotations were taken from
185 existing datasets with varying labelling schemes and quality standards, potentially containing inconsis-
186 tencies that affect evaluation reliability. Second, we lack human annotator baselines to contextualize
187 model performance—while current MLLMs perform poorly, we cannot determine how their accuracy
188 compares to average human annotators on these specific videos.

189 Third, our evaluation uses zero-shot inference without fine-tuning or extensive prompt optimization.
190 While this ensures fair comparison across models, it may underestimate achievable performance
191 through model adaptation or specialized prompting strategies. Additionally, dataset-specific prompts
192 introduce variability that could advantage certain models.

193 Finally, our evaluation is limited to three commercially available models with native video processing,
194 and the rapid pace of model development means newer capabilities may alter these findings. Despite
195 these limitations, Rodent-Bench provides a valuable initial assessment of current MLLM capabilities
196 for scientific behavioral annotation tasks.

197 6 Conclusion

198 We introduced Rodent-Bench, the first comprehensive benchmark for evaluating Multimodal Large
199 Language Models on scientific behavioral annotation tasks. Our evaluation of state-of-the-art
200 MLLMs—Gemini-2.5-Pro, Gemini-2.5-Flash, and Qwen-VL-Max—reveals that current models
201 perform substantially below the accuracy levels required for practical deployment as research assis-
202 tants in behavioral neuroscience.

203 While MLLMs showed modest success on certain datasets (notably grooming detection), perfor-
204 mance varied dramatically across behavioral paradigms. Models struggled particularly with subtle
205 temporal distinctions, brief behavioral episodes, and tasks requiring integration of spatial context
206 with behavioral understanding. The freezing behavior dataset exemplified these challenges, where
207 distinguishing between active freezing responses and passive inactivity proved difficult even for
208 advanced multimodal systems.

209 Our findings highlight several critical areas for improvement. First, enhanced temporal reasoning
210 capabilities are needed to handle the fine-grained segmentation required for behavioral analysis.
211 Second, models must develop better contextual understanding to distinguish between visually similar
212 but behaviorally distinct states. Finally, output formatting consistency remains a practical barrier,
213 with some models frequently producing malformed JSON responses that complicate automated
214 processing.

215 Despite current limitations, Rodent-Bench establishes a foundation for tracking progress in scientific
216 applications of multimodal AI. The benchmark’s diverse behavioral paradigms and standardized

evaluation framework provide a testbed for future model improvements. As MLLMs advance, their potential to democratize behavioral annotation, eliminating the need for specialized model training for each experimental paradigm, remains promising. Rodent-Bench will enable researchers to objectively assess when these models achieve the reliability threshold necessary for practical scientific deployment.

The gap between current capabilities and scientific requirements underscores the need for continued research at the intersection of multimodal AI and domain-specific applications. Our benchmark contributes to this effort by providing concrete evaluation targets and highlighting the unique challenges that scientific video understanding presents to current generation models.

References

- Arac, A., Zhao, P., Dobkin, B. H., Carmichael, S. T., and Golshani, P. Deepbehavior: A deep learning toolbox for automated analysis of animal and human behavior imaging data. *Frontiers in Systems Neuroscience*, 13:20, 2019.
- Blanchard, R. J. and Blanchard, D. C. Passive and active reactions to fear-eliciting stimuli. *Journal of comparative and physiological psychology*, 68(1p1):129, 1969.
- Bohnslav, J. P., Wimalasena, N. K., Clausen, K. J., Dai, Y. Y., Yarmolinsky, D. A., Cruz, T., Kashlan, A. D., Chiappe, M. E., Orefice, L. L., Woolf, C. J., and Harvey, C. D. Deepethogram, a machine learning pipeline for supervised behavior classification from raw pixels. *eLife*, 10:e63377, sep 2021. ISSN 2050-084X. doi: 10.7554/eLife.63377. URL <https://doi.org/10.7554/eLife.63377>.
- de Chaumont, F., Coura, R. D. S., Serreau, P., Cressant, A., Chabout, J., Granon, S., and Olivo-Marin, J.-C. Real-time analysis of the behaviour of groups of mice via a depth-sensing camera and machine learning. *bioRxiv*, pp. 2022–02, 2022.
- Fu, C., Dai, Y., Luo, Y., Li, L., Ren, S., Zhang, R., Wang, Z., Zhou, C., Shen, Y., Zhang, M., et al. Video-mme: The first-ever comprehensive evaluation benchmark of multi-modal llms in video analysis. *arXiv preprint arXiv:2405.21075*, 2024.
- Geuther, B. Q., Peer, A., He, H., Sabnis, G., Philip, V. M., and Kumar, V. Action detection using a neural network elucidates the genetics of mouse grooming behavior. *eLife*, 10:e63207, mar 2021. ISSN 2050-084X. doi: 10.7554/eLife.63207. URL <https://doi.org/10.7554/eLife.63207>.
- Henderson, P. and Ferrari, V. End-to-end training of object class detectors for mean average precision. In *Asian conference on computer vision*, pp. 198–213. Springer, 2016.
- Mathis, M. W. and Mathis, A. Deep learning tools for the measurement of animal behavior in neuroscience. *Current Opinion in Neurobiology*, 60:1–11, 2020.
- Pennington, Z. T., Dong, Z., Feng, Y., Vetere, L. M., Page-Harley, L., Shuman, T., and Cai, D. J. eztrack: An open-source video analysis pipeline for the investigation of animal behavior. *Scientific reports*, 9(1):19979, 2019.
- Pereira, T. D., Tabris, N., Matsliah, A., Turner, D. M., Li, J., Ravindranath, S., Papadoyannis, E. S., Normand, E., Deutsch, D. S., Wang, Z. Y., et al. Slep: A deep learning system for multi-animal pose tracking. *Nature methods*, 19(4):486–495, 2022.
- Sturm, G., Friede, T., Müller, S., Mathis, A., and Mathis, M. W. Deep learning-based behavioral analysis reaches human accuracy and is capable of outperforming commercial solutions. *Neuropsychopharmacology*, 45(11):1942–1952, 2020.
- Sun, J. J., Karigo, T., Chakraborty, D., Mohanty, S. P., Wild, B., Sun, Q., Chen, C., Anderson, D. J., Perona, P., Yue, Y., and Kennedy, A. The multi-agent behavior dataset: Mouse dyadic social interactions. *arXiv preprint arXiv:2104.02710*, 2021.
- Tang, Y., Bi, J., Xu, S., Song, L., Liang, S., Wang, T., Zhang, D., An, J., Lin, J., Zhu, R., Vosoughi, A., Huang, C., Zhang, Z., Liu, P., Feng, M., Zheng, F., Zhang, J., Luo, P., Luo, J., and Xu, C. Video understanding with large language models: A survey. *IEEE Transactions on Circuits and Systems for Video Technology*, 2025. doi: 10.1109/TCSVT.2025.3566695.

265 Testard, C., Buch, A., Gauthier, J., Marvin, J., Yartsev, M., and Fairhall, A. Data science opportunities
266 of large language models for neuroscience and biomedicine. *Neuron*, 112(4):499–510, 2024.

267 Wang, J. et al. From seconds to hours: Reviewing multimodal large language models on comprehen-
268 sive long video understanding. *arXiv preprint arXiv:2409.18938*, 2024.

269 Wiltschko, A. B., Johnson, M. J., Iurilli, G., Peterson, R. E., Katon, J. M., Pashkovski, S. L., Abaira,
270 V. E., Adams, R. P., and Datta, S. R. Mapping sub-second structure in mouse behavior. *Neuron*, 88
271 (6):1121–1135, 2015.

272 Yin, S., Fu, C., Zhao, S., Li, K., Sun, X., Xu, T., and Chen, E. A survey on multimodal large language
273 models. *National Science Review*, 11(12):nwae403, 2024.

274 Yu, H., Xiong, J., Ye, A. Y., Cranfill, S. L., Cannonier, T., Gautam, M., Zhang, M., Bilal, R.,
275 Park, J.-E., Xue, Y., Polam, V., Vujovic, Z., Dai, D., Ong, W., Ip, J., Hsieh, A., Mimouni, N.,
276 Lozada, A., Sosale, M., Ahn, A., Ma, M., Ding, L., Arsuaga, J., and Luo, W. Scratch-aid,
277 a deep learning-based system for automatic detection of mouse scratching behavior with high
278 accuracy. *eLife*, 11:e84042, dec 2022. ISSN 2050-084X. doi: 10.7554/eLife.84042. URL
279 <https://doi.org/10.7554/eLife.84042>.

280 A Implementation Details

281 A.1 Model Access and Configuration

282 **Gemini Models:** We access Gemini-2.5-Pro and Gemini-2.5-Flash through Google’s Vertex AI
283 GenAI SDK. Videos are processed directly from Google Cloud Storage URIs using the native video
284 input capabilities. The response format is constrained to JSON using structured output schemas
285 specific to each dataset’s behavior categories.

286 **Qwen-VL-Max:** We access Qwen-VL-Max through Alibaba’s DashScope API using the OpenAI-
287 compatible interface.

288 A.2 Video Processing Pipeline

289 Each video is processed independently with dataset-specific prompts that include behavior definitions,
290 temporal annotation requirements, and output format specifications.

291 **Structured Output Schemas:** We define Pydantic models for each dataset’s behavior categories to
292 ensure consistent JSON output formatting. The following is the model for CaLMS21:

```
293 1 class RodentBehaviorSegment(BaseModel):  
294 2     segment_number: int = Field(..., description="Segment number in  
295     order")  
296 3     start_time: str = Field(..., description="Start time in MM:SS  
297     format")  
298 4     end_time: str = Field(..., description="End time in MM:SS format")  
299 5     behavior: str = Field(..., description="Behavior label (e.g.,  
300     attack, investigation, mount, other)")
```

301 A.3 Batch Processing Implementation

302 For Gemini models, we implement both individual and batch processing modes. Batch mode generates
303 JSONL files conforming to Gemini’s batch API requirements, uploads input files to Google Cloud
304 Storage, and monitors job completion through the batch API. This approach significantly reduces
305 API costs for large-scale evaluations while maintaining identical model configurations.

306 **Error Handling:** The system logs all API responses, including malformed outputs, to facilitate
307 debugging. For models producing incomplete JSON (particularly Qwen-VL-Max), we save raw
308 responses to text files for manual inspection. Output validation ensures all required fields are present
309 and temporal segments are non-overlapping.

310 **B Model Specifications**

311 We evaluate three state-of-the-art multimodal large language models with native video understanding
312 capabilities.

313 **B.1 Gemini-2.5-Pro**

314 **Key Specifications:**

- 315 • Maximum video length: 1 hour (without audio), 45 minutes (with audio)
- 316 • Context window: 1,048,576 tokens (input), Maximum 65,535 tokens (output)
- 317 • Maximum video file size: 2 GB

318 **B.2 Gemini-2.5-Flash**

319 **Key Specifications:**

- 320 • Maximum video length: 1 hour (without audio), 45 minutes (with audio)
- 321 • Context window: 1,048,576 tokens (input), Maximum 65,535 tokens (output)
- 322 • Maximum video file size: 2 GB

323 **B.3 Qwen-VL-Max**

324 Qwen-VL-Max is Alibaba Cloud’s most advanced vision-language model. Unlike the Gemini models
325 which process video natively, Qwen-VL extracts frames from video files for analysis, extracting one
326 frame every 0.5 seconds when using the OpenAI SDK.

327 **Key Specifications:**

- 328 • Maximum video length: 10 minutes (Qwen2.5-VL series)
- 329 • Context window: 129,024 input tokens, 8,192 output tokens
- 330 • Maximum video file size: 1 GB (via URL), 10 MB (Base64 encoded)
- 331 • Video processing: Frame extraction (no audio support)

332 **B.4 Model Selection Rationale**

333 We selected these models based on three criteria: (1) native video processing capabilities, (2)
334 availability through stable APIs for reproducible evaluation, and (3) demonstrated performance on
335 complex reasoning tasks.

336 **C Datasets**

337 We include screenshots from each dataset, demonstrating each behavior. We additionally include a
338 chart showing the proportion of each behavior in each dataset.

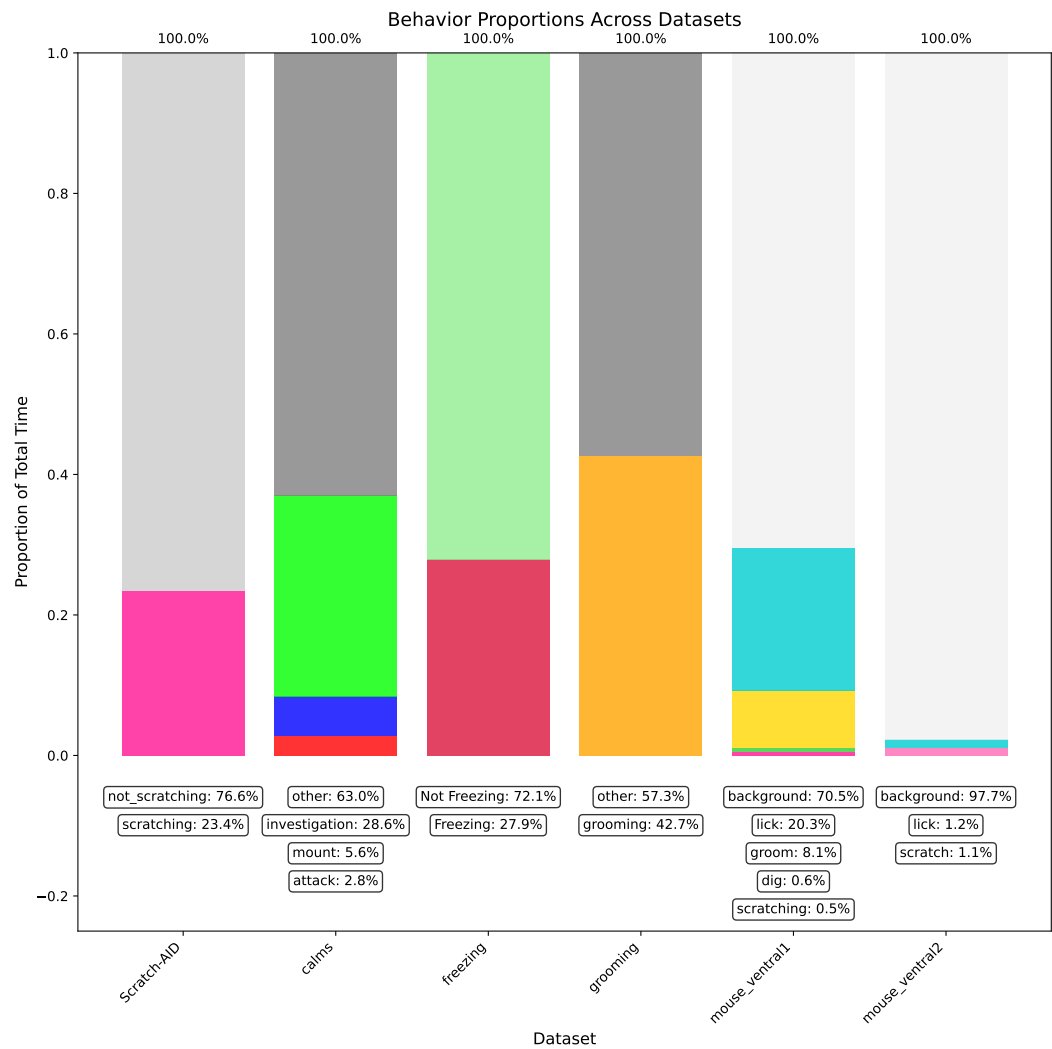


Figure 4: Behavior Proportions for each dataset.

339 **C.1 Rodent-Bench**

Table 1: Timing Statistics by Dataset

Dataset	Average Time (Mins)	Minimum Time (Mins)	Maximum Time (Mins)	Total Time (Mins)
CalMS21	4.41	1.02	11.87	110.22
Freezing	14.35	4.01	32.67	129.17
Grooming	1.32	0.34	5.99	32.88
Rodent Ventral 1	8.33	8.28	8.33	233.25
Rodent Ventral 2	29.97	29.97	29.97	479.57
Scratch-AID	20.00	20.00	20.00	300.03

340 **C.2 Rodent-Bench Short**

Table 2: Timing Statistics by Dataset

Dataset	Average Time (Mins)	Minimum Time (Mins)	Maximum Time (Mins)	Total Time (Mins)
CalMS21	4.30	1.02	9.98	107.57
Freezing	9.04	4.01	9.98	81.33
Grooming	2.87	0.44	9.98	71.68
Rodent Ventral 1	6.57	0.34	8.33	183.89
Rodent Ventral 2	9.26	8.33	9.98	148.19
Scratch-AID	9.98	9.98	9.99	149.77

Dataset: calms

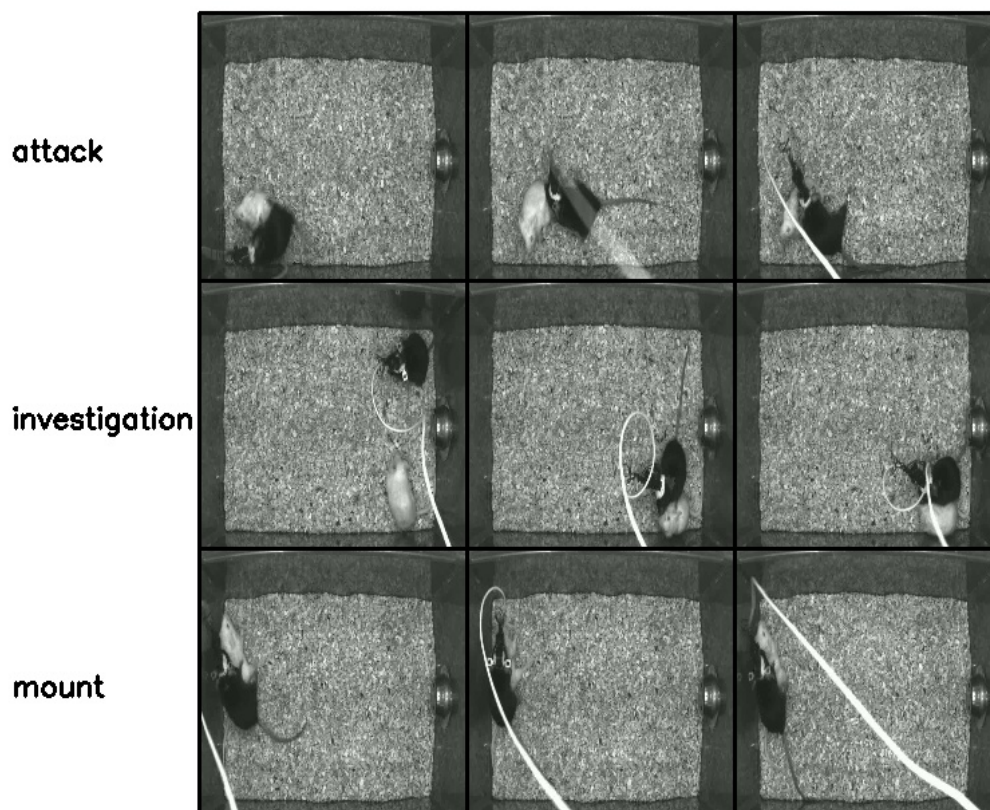


Figure 5: CaLMS21 Behaviors

Dataset: grooming

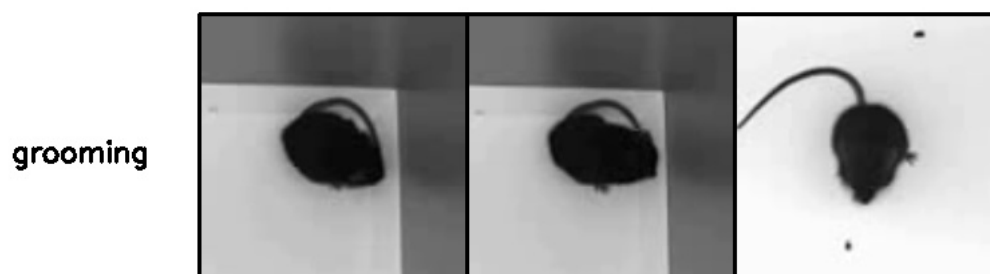


Figure 6: Rodent Grooming Behaviors

Dataset: mouse_ventral1

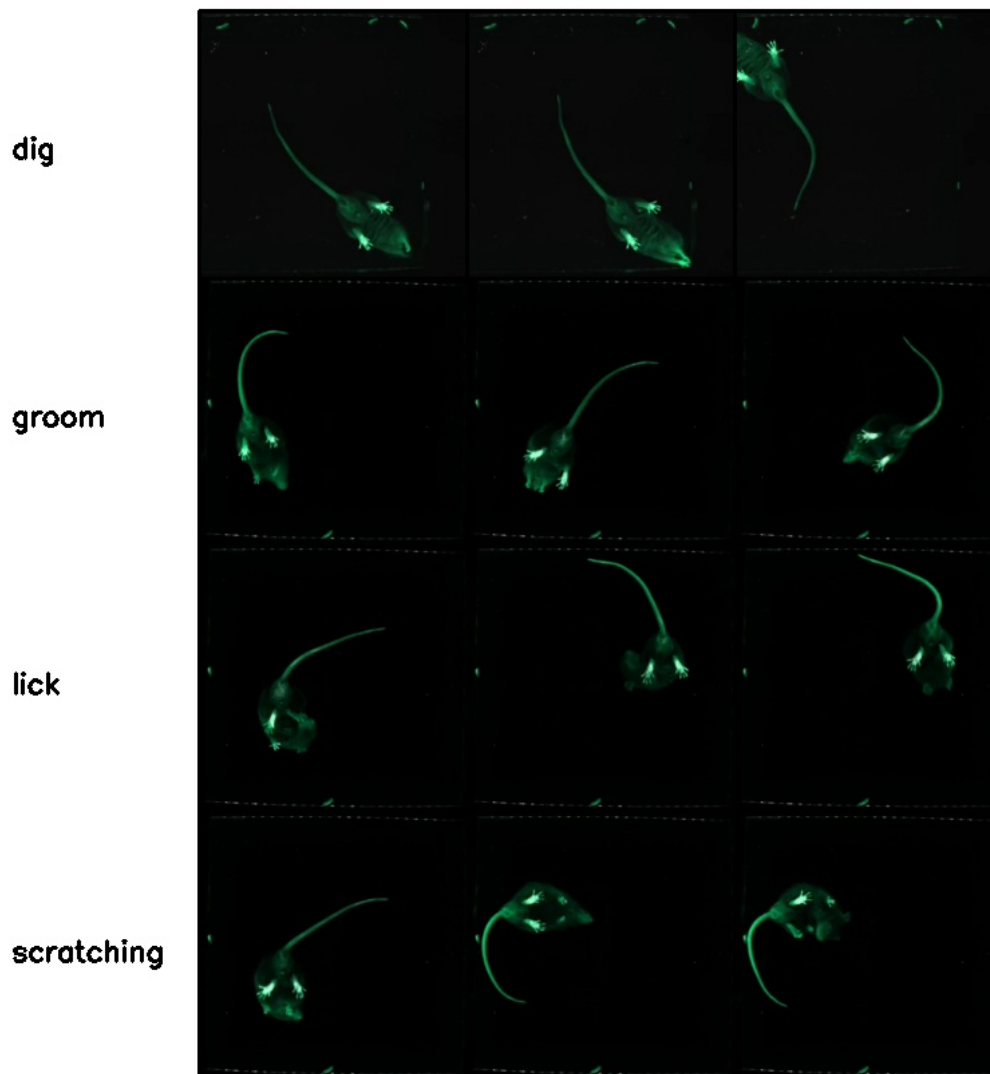


Figure 7: Mouse-Ventral 1 Behaviors

Dataset: mouse_ventral2

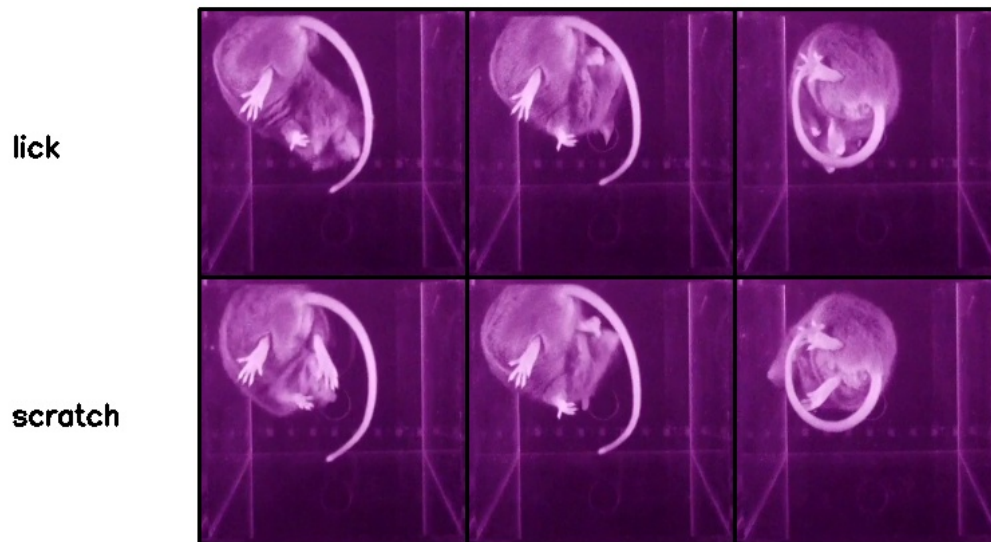


Figure 8: Mouse-Ventral 2 Behaviors

344 **C.6 Scratch-AID**

Dataset: Scratch-AID

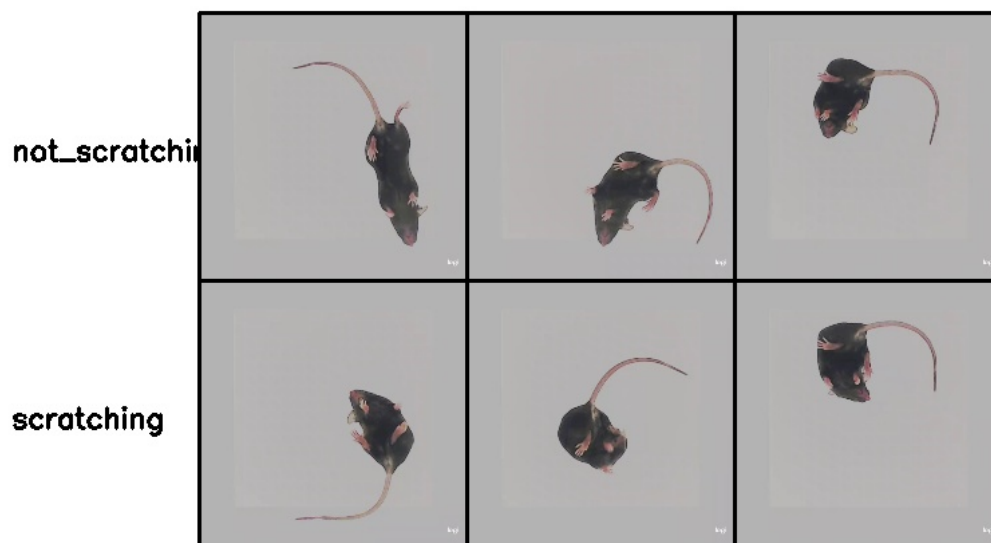


Figure 9: Scratch-AID Behaviors

345 **C.7 Freezing**

Dataset: freezing

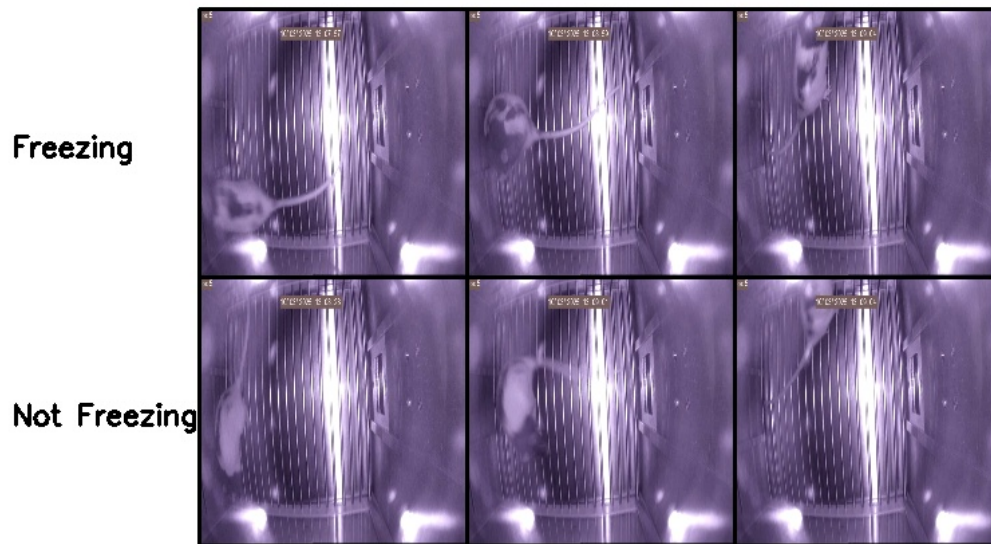


Figure 10: Freezing Behaviors

346 D Gemini-Pro Results

347 D.1 Rodent-Bench

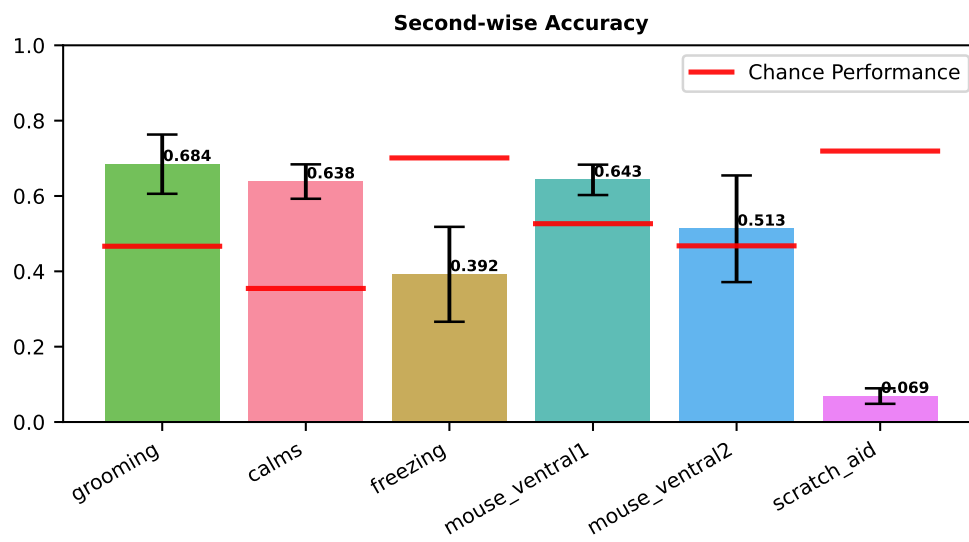


Figure 11: Per second accuracy

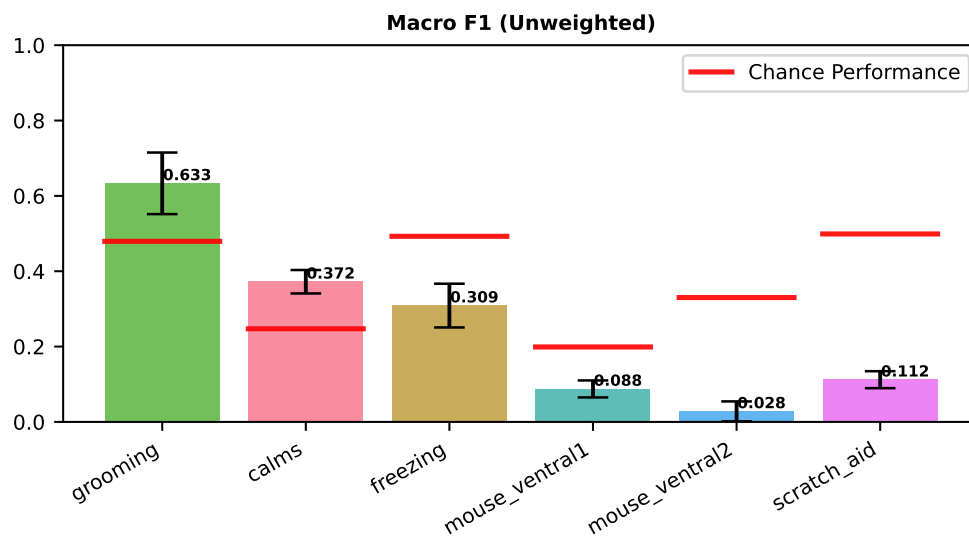


Figure 12: Macro F1 score

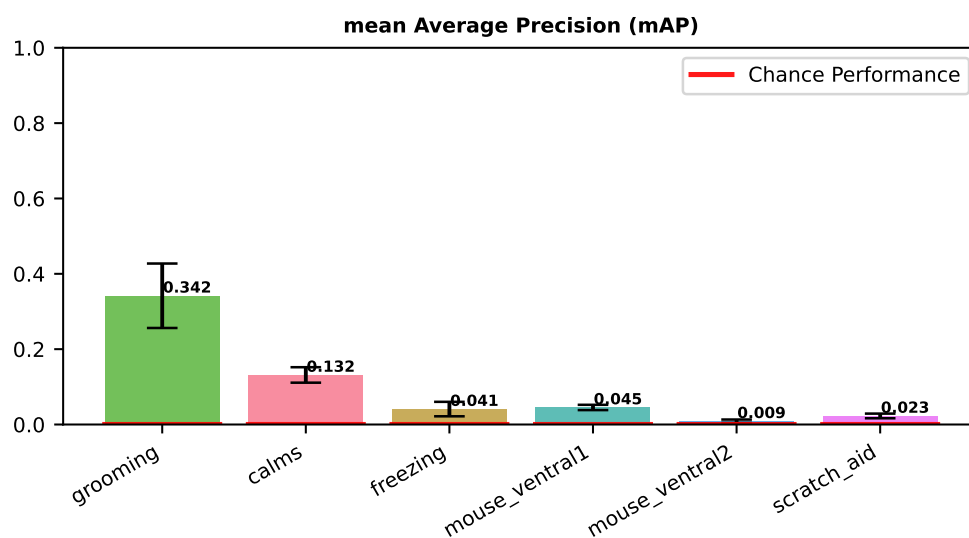


Figure 13: Per second accuracy

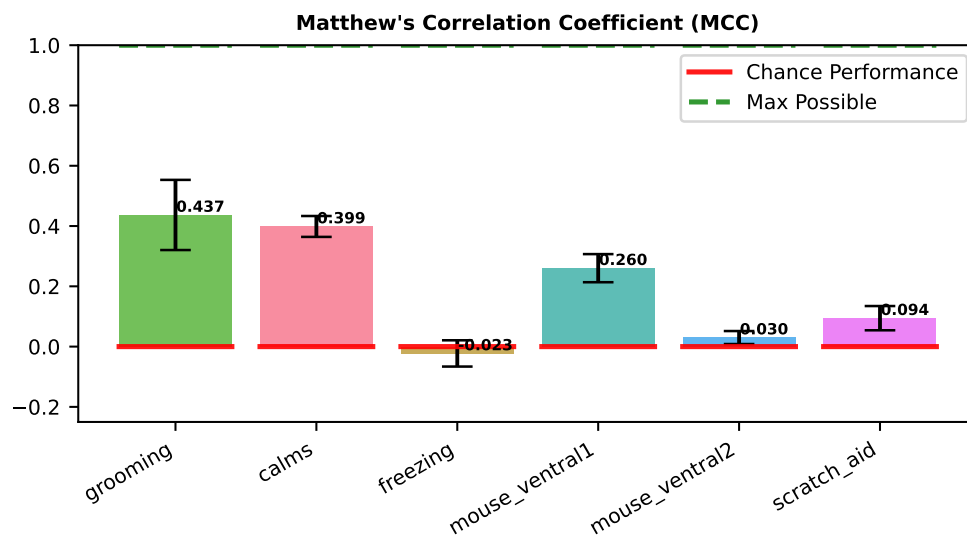


Figure 14: Matthew's Correlation Coefficient

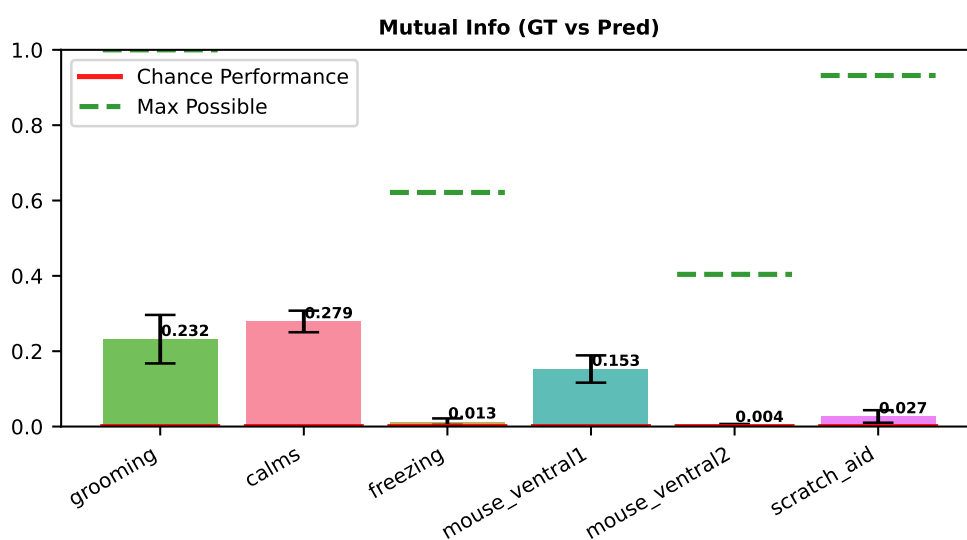


Figure 15: Mutual information between ground truth and predictions

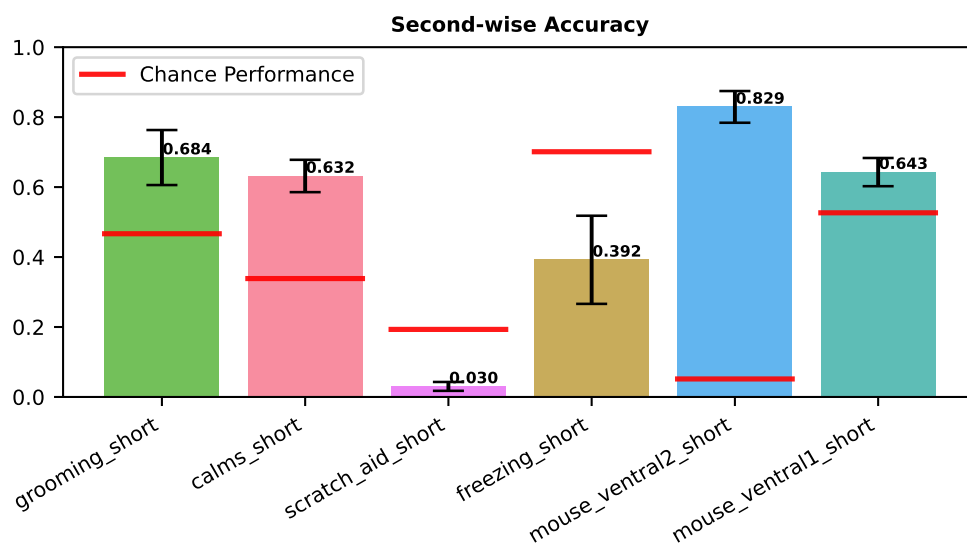


Figure 16: Per second accuracy

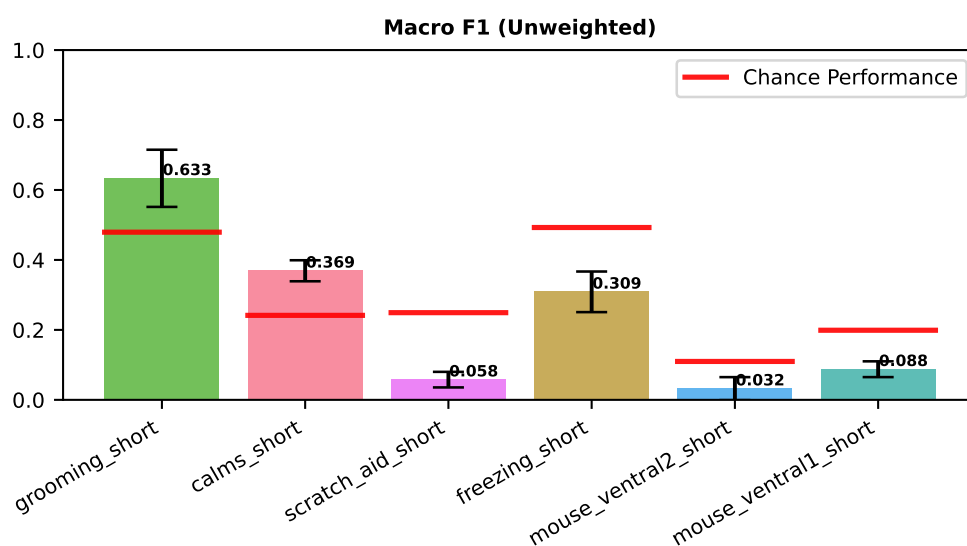


Figure 17: Macro F1 score

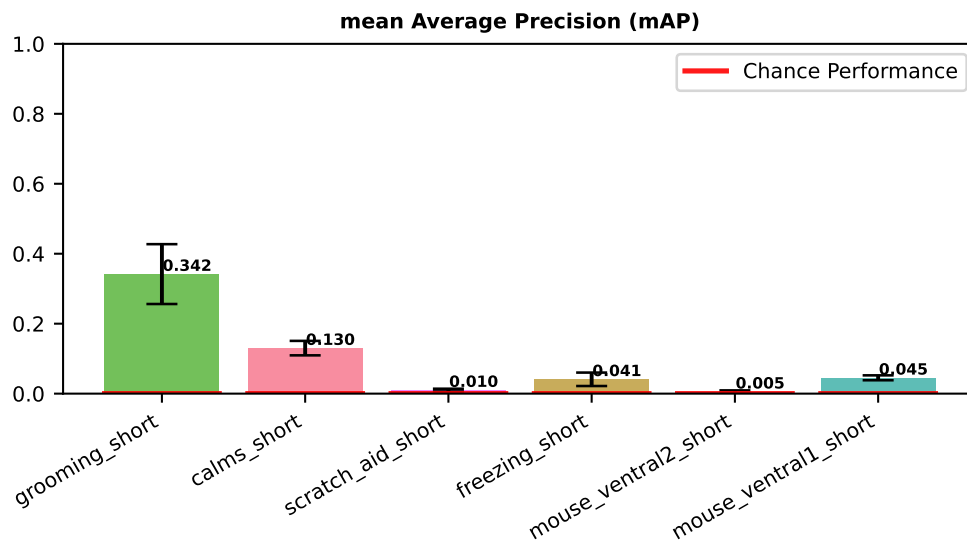


Figure 18: Per second accuracy

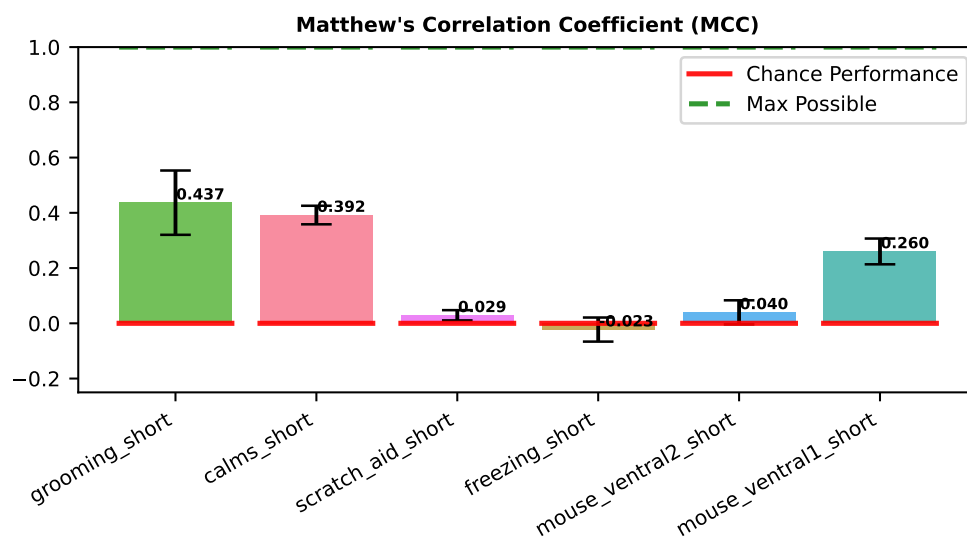


Figure 19: Matthew's Correlation Coefficient

349 **E Gemini-Flash Results**

350 **E.1 Rodent-Bench**

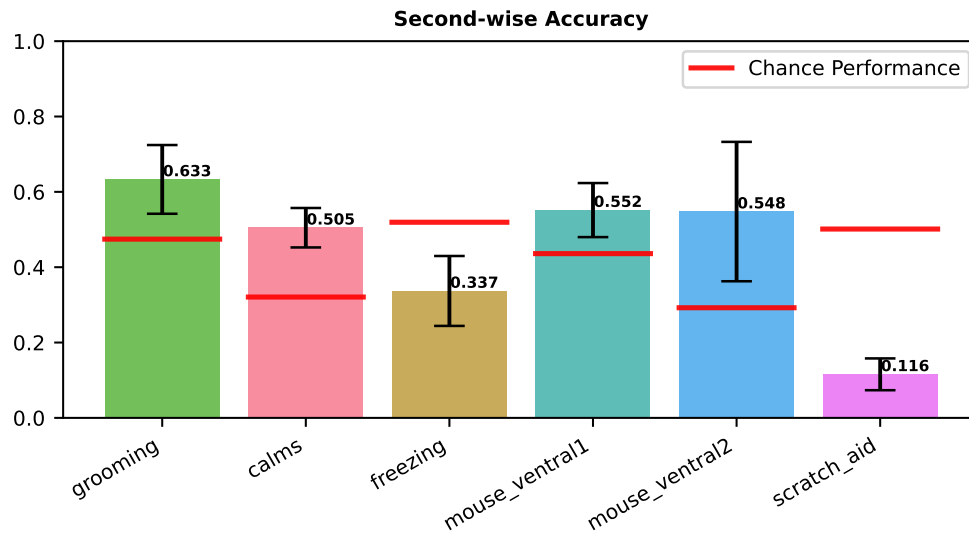


Figure 20: Per second accuracy

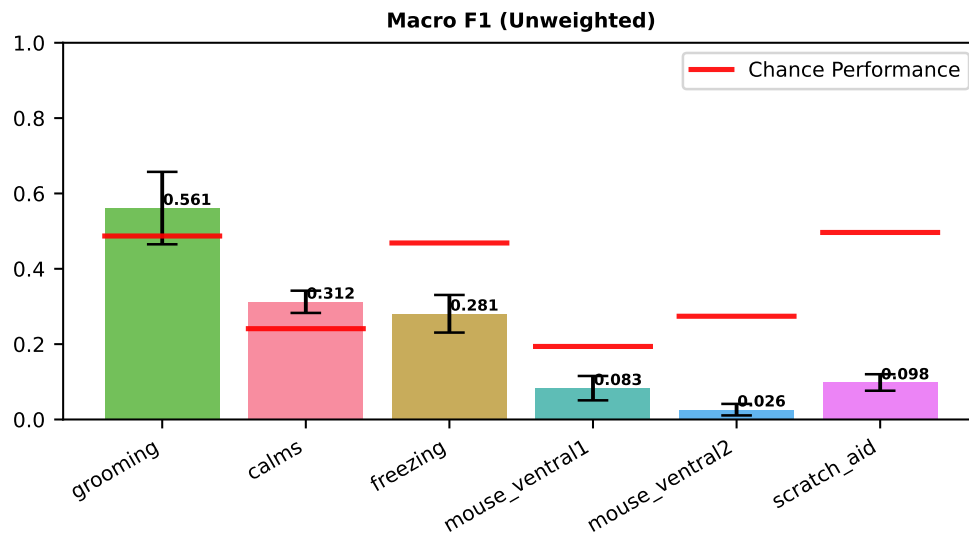


Figure 21: Macro F1 score

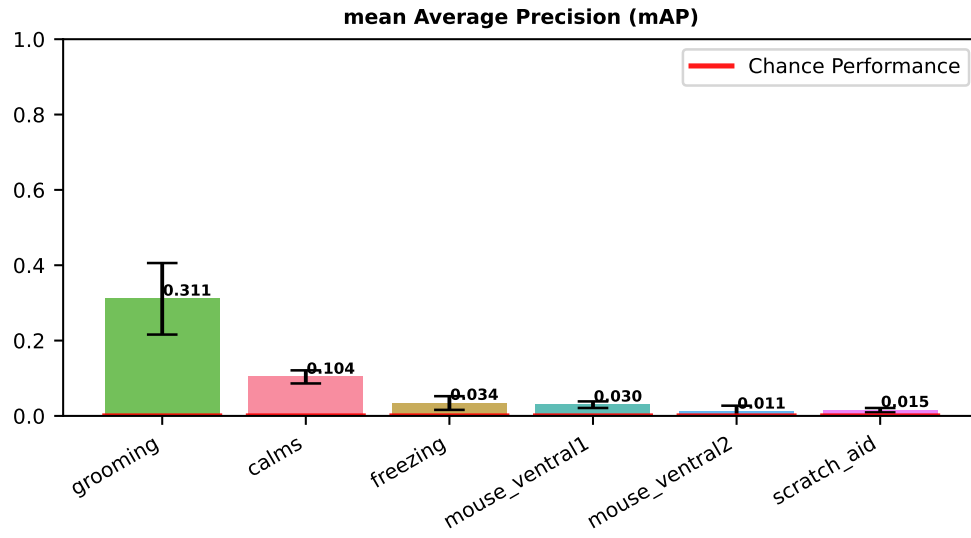


Figure 22: Per second accuracy

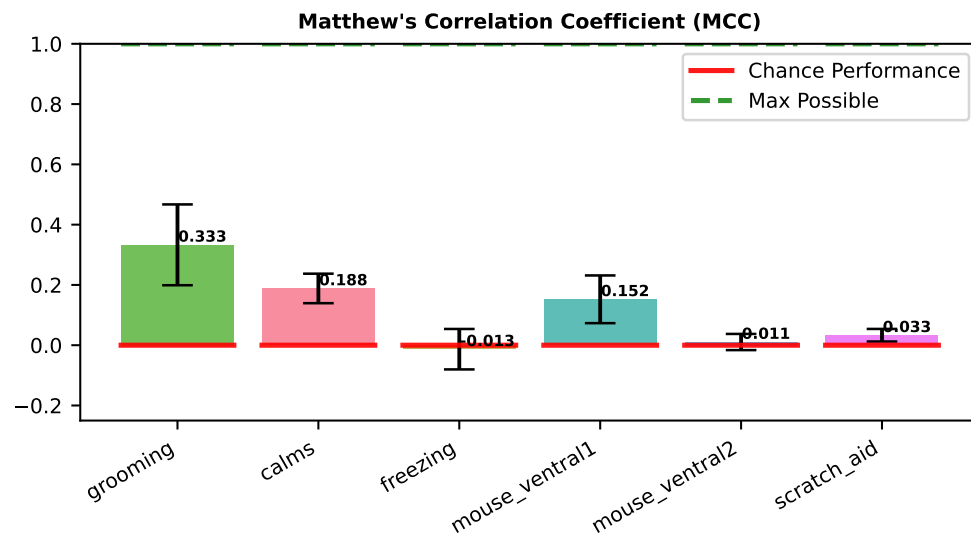


Figure 23: Matthew's Correlation Coefficient

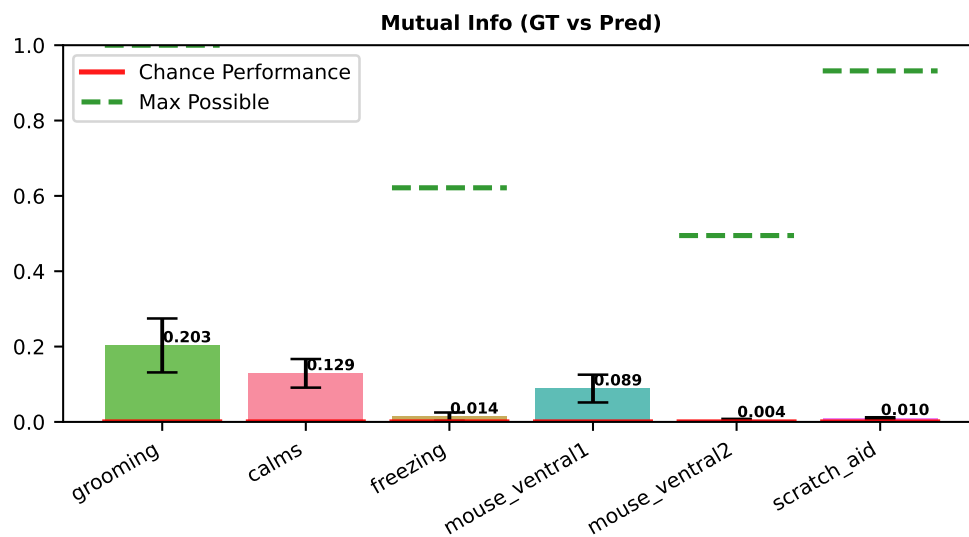


Figure 24: Mutual information between ground truth and predictions

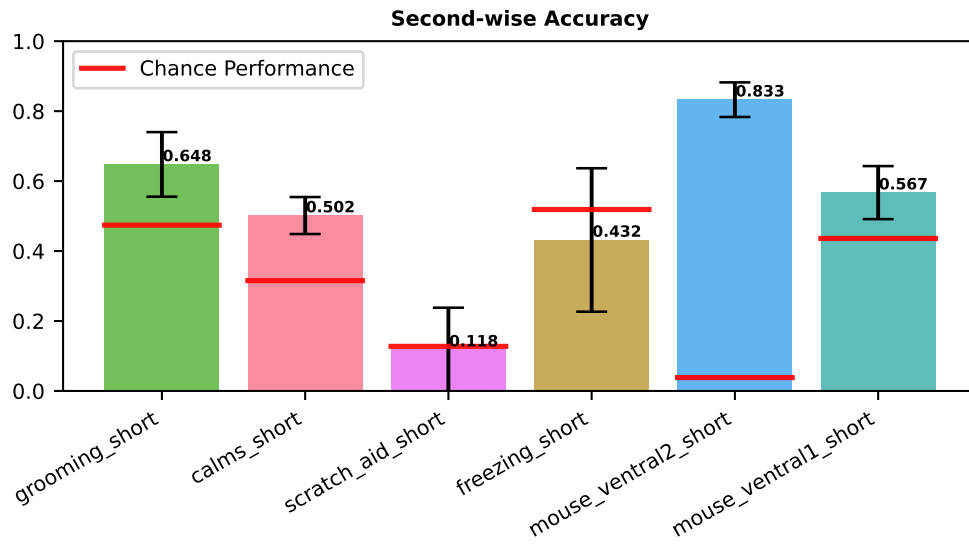


Figure 25: Per second accuracy

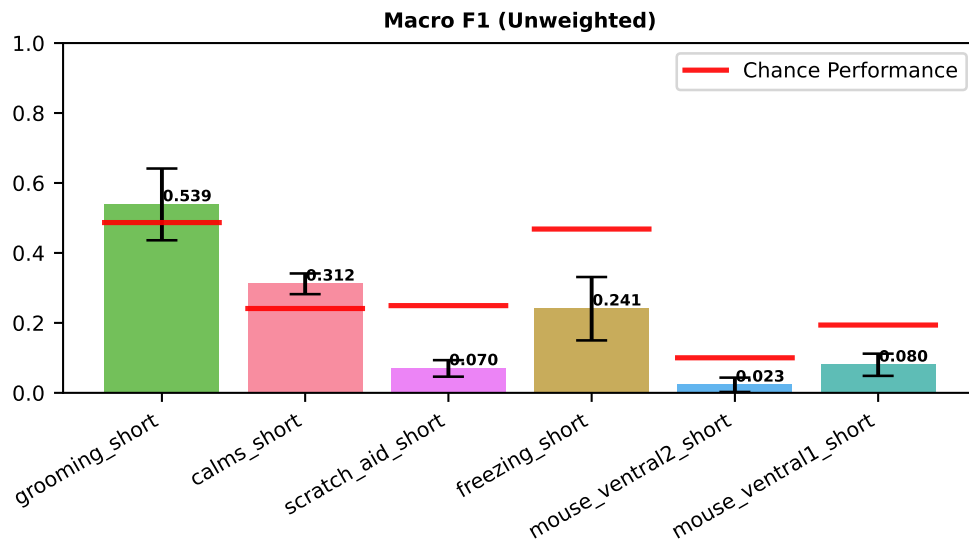


Figure 26: Macro F1 score

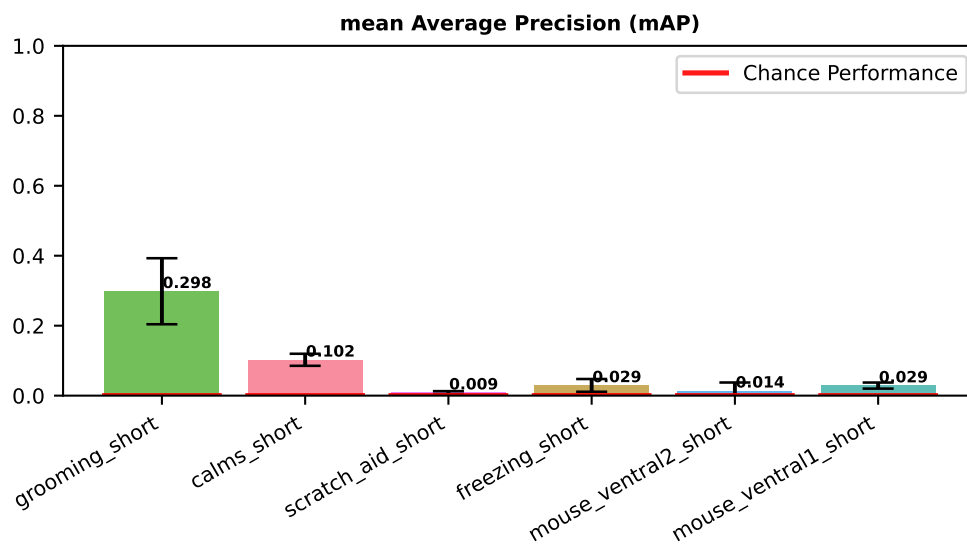


Figure 27: Per second accuracy

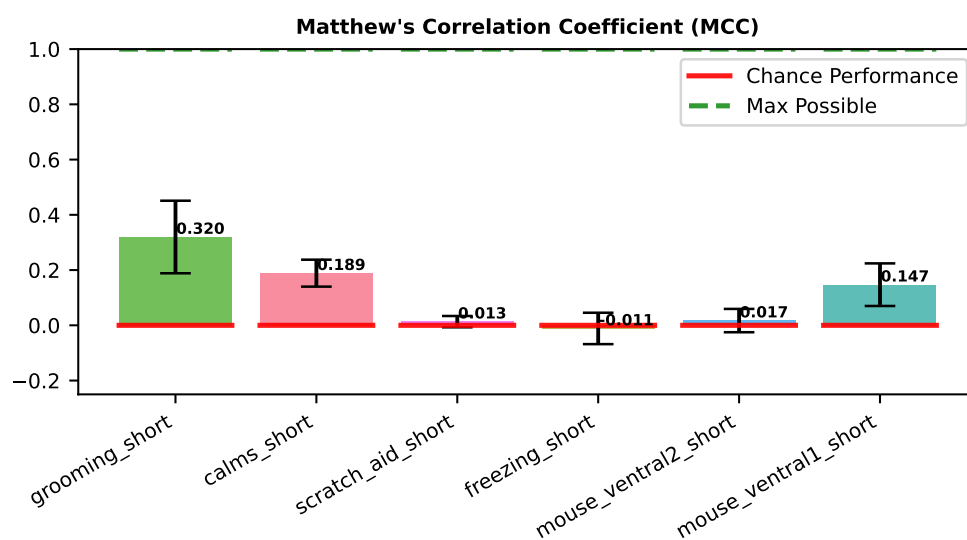


Figure 28: Matthew's Correlation Coefficient

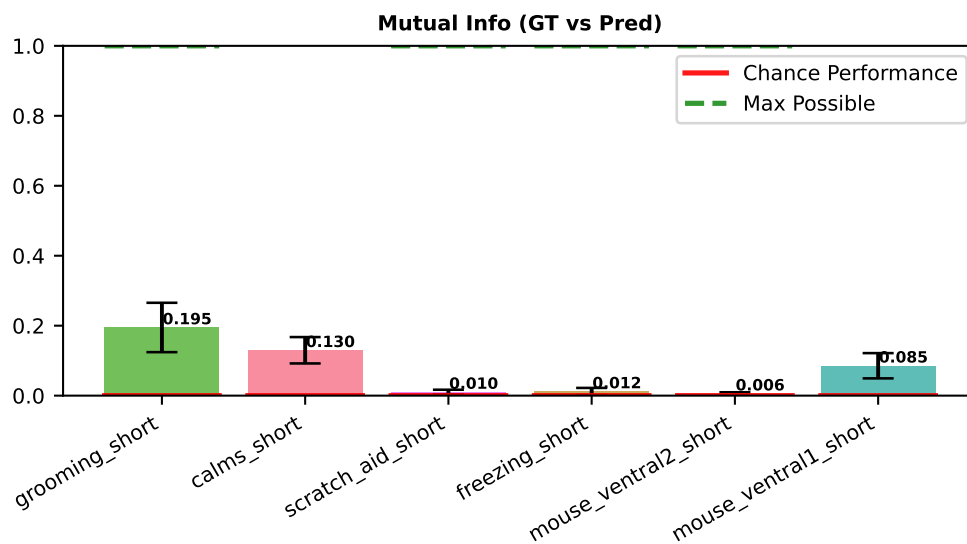


Figure 29: Mutual information between ground truth and predictions

352 **F Qwen-VL-Max Results**

353 Because the Qwen-VL-Max can't ingest videos longer than 10 minutes we only have results for
354 Rodent-Bench-Short.

355 **F.1 Rodent-Bench-Short**

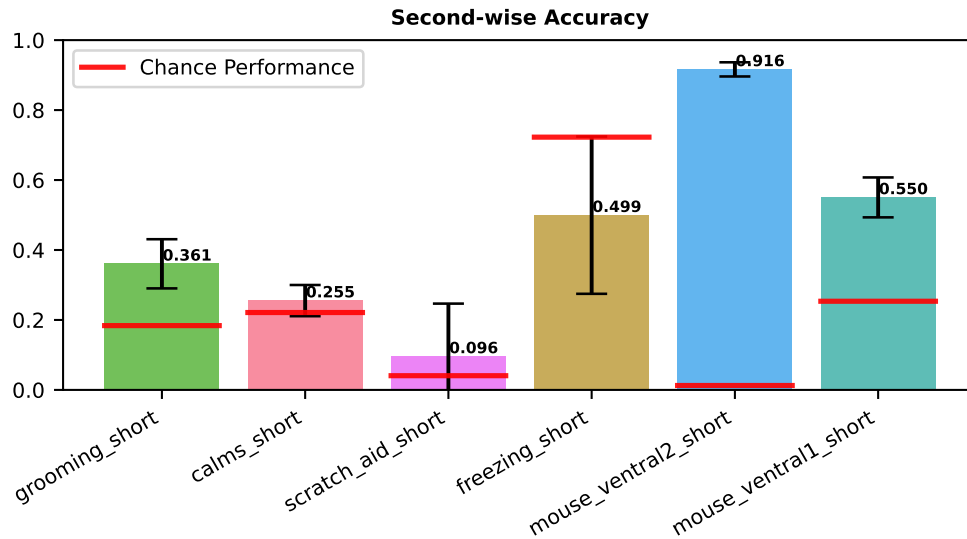


Figure 30: Per second accuracy

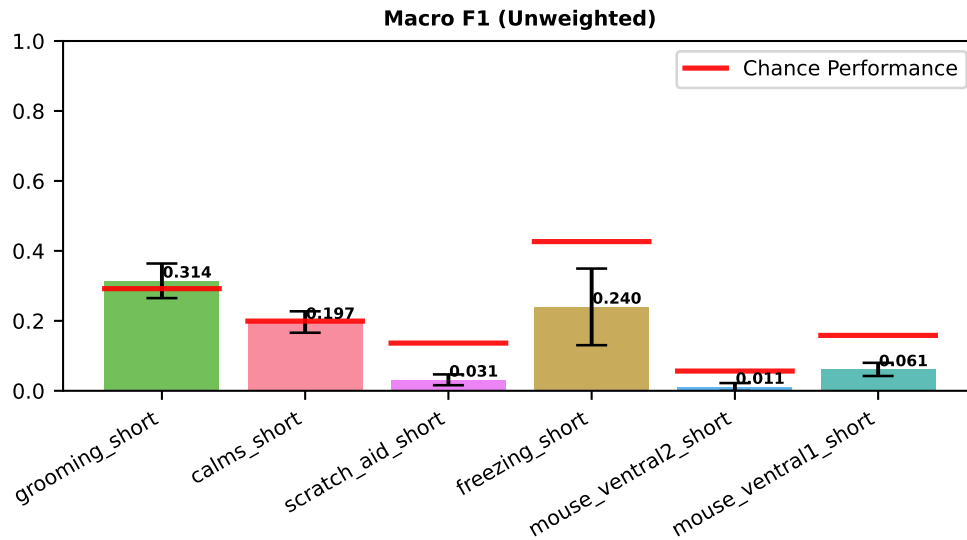


Figure 31: Macro F1 score

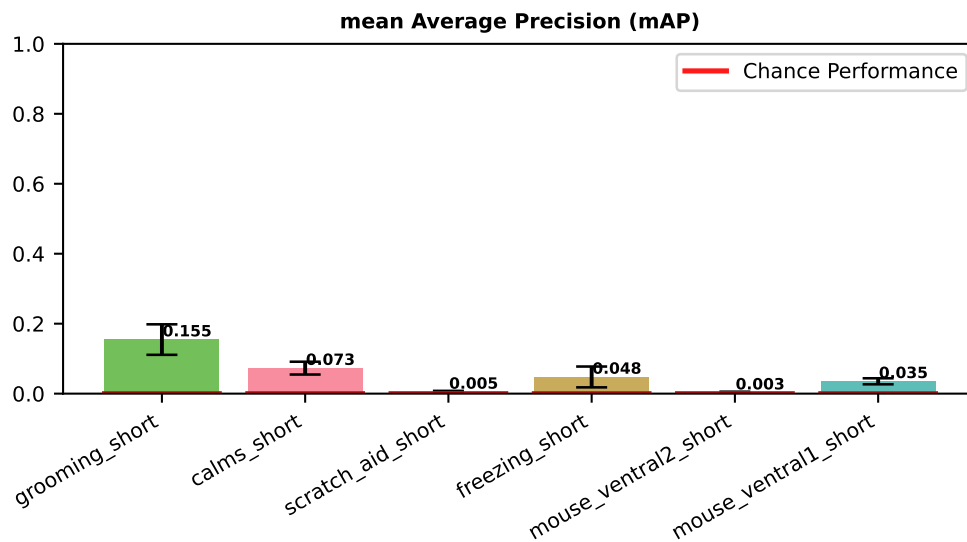


Figure 32: Per second accuracy

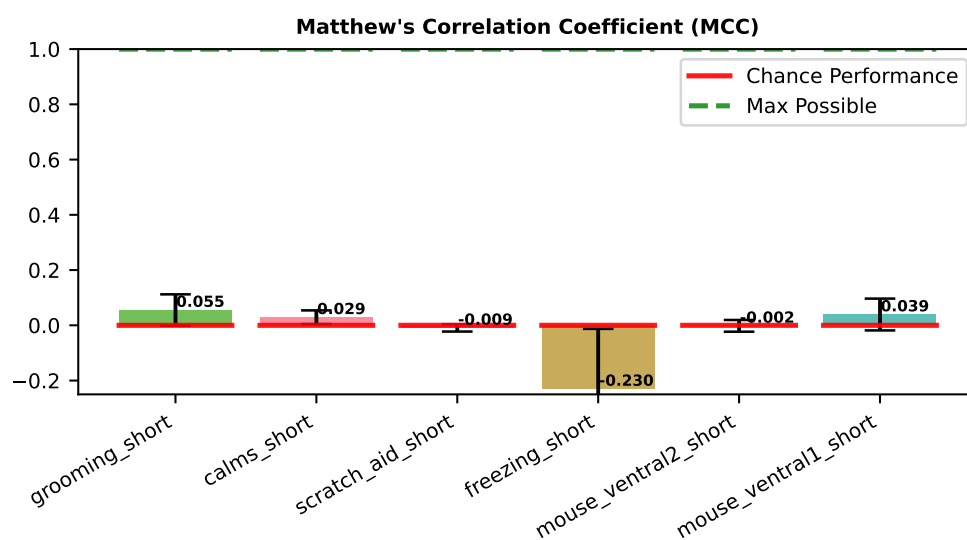


Figure 33: Matthew's Correlation Coefficient

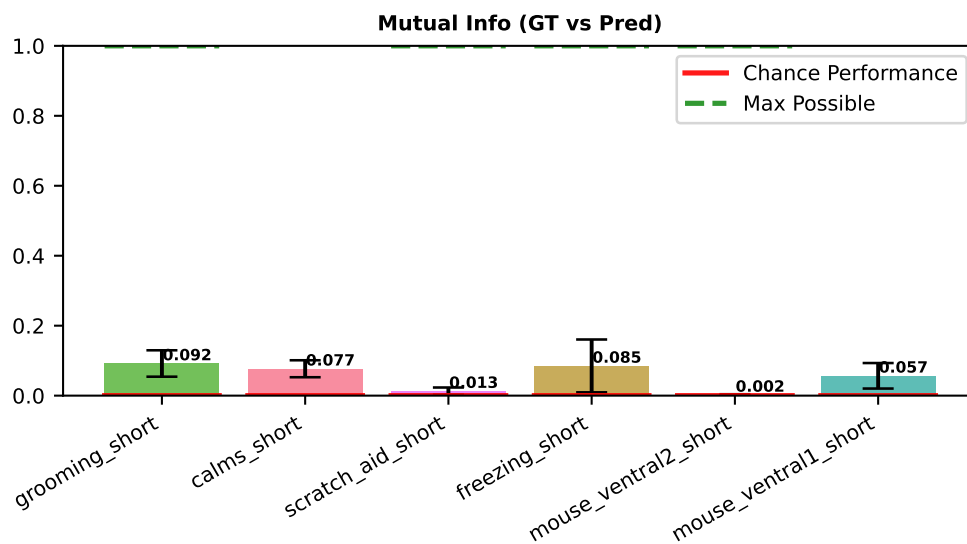


Figure 34: Mutual information between ground truth and predictions

356 G Prompt Templates

357 We provide the complete prompt templates used for each dataset. All prompts follow a consistent
358 structure: role definition, task description, available behavior labels, formatting requirements, and
359 JSON output schema.

360 G.1 CalMS21 Social Behaviors

```
361 1 You are a Rodent Behavior Labeler specializing in rodent social
362   behavior.
363 2 Your task is to analyze a video of rodents and segment it into periods
364   of distinct behaviors.
365 3
366 4 Available behavior labels:
367 5 - attack - when the black rodent is attacking another rodent
368 6 - investigation - when the black rodent is investigating another
369   rodent
370 7 - mount - when the black rodent is mounting another rodent
371 8 - other - when the black rodent is doing something else
372 9
373 10 Important:
374 11 You must use ONLY the labels listed above. Do not create new labels or
375   modify existing ones.
376 12
377 13 Start your analysis from the start of the video and continue until the
378   end of the video.
379 14
380 15 For each segment, provide:
381 16 - segment number (in order)
382 17 - start and end time in MM:SS format
383 18 - behavior label (must be one of the above labels)
384 19
385 20 Your response must be in JSON format with the following structure:
386 21 {
387 22   "segments": [
388 23     {
389 24       "start_time": MM:SS,
390 25       "end_time": MM:SS,
391 26       "behavior": "behavior_label",
392 27       "segment_number": INTEGER,
393 28     },
394 29     ...
395 30   ]
396 31 }
```

397 G.2 Scratch-AID

```
398 1 You are a Rodent Behavior Labeler specializing in telling when a
399   rodent is scratching.
400 2 Your task is to analyze a video of rodents and segment it into periods
401   of distinct behaviors.
402 3
403 4 Available behavior labels:
404 5 - scratching - when the rodent is scratching, usually with the hind
405   legs
406 6 - not scratching - when the rodent is not scratching.
407 7
408 8 Important:
409 9 You must use ONLY the labels listed above. Do not create new labels or
410   modify existing ones.
411 10
412 11 Start your analysis from the start of the video and continue until the
413   end of the video.
```

```

414|2 The video is of a rodent and is taken from below.
415|3
416|4 For each segment, provide:
417|5 - start and end time in MM:SS format
418|6 - segment number (in order)
419|7
420|8 Your response must be in JSON format with the following structure:
421|9 {
422|10     "segments": [
423|11         {
424|12             "segment_number": INTEGER,
425|13             "start_time": MM:SS,
426|14             "end_time": MM:SS,
427|15             "behavior": "behavior_label",
428|16         },
429|17         ...
430|18     ]
431|19 }

```

432 G.3 Rodent Grooming Detection

```

433|1 You are a Rodent Behavior Labeler specializing in identifying grooming
434|2 behaviors in rodents.
435|3 Your task is to analyze a video of rodents and segment it into periods
436|4 of distinct behaviors.
437|5 The video shows a rodent from above.
438|6
439|7 Available behavior labels:
440|8 - grooming - when the rodent is actively grooming itself (e.g.,
441|9     licking, scratching, cleaning fur)
442|10 - other - when the rodent is not grooming (e.g., walking, exploring,
443|11     resting)
444|12
445|13 Important:
446|14 You must use ONLY the labels listed above. Do not create new labels or
447|15 modify existing ones.
448|16 Grooming behaviors are characterized by:
449|17 - Repetitive movements of paws over the face or body
450|18 - Licking of fur or paws
451|19 - Scratching with hind legs
452|20 - Cleaning of specific body parts
453|21
454|22 Start your analysis from the start of the video and continue until the
455|23 end of the video.
456|24
457|25 For each segment, provide:
458|26 - segment number (in order)
459|27 - start and end time in MM:SS format
460|28 - behavior label (must be one of the above labels)
461|29
462|30 Your response must be in JSON format with the following structure:
463|31 {
464|32     "segments": [
465|33         {
466|34             "segment_number": INTEGER,
467|35             "start_time": MM:SS,
468|36             "end_time": MM:SS,
469|37             "behavior": "behavior_label",
470|38         },
471|39         ...
472|40     ]
473|41 }

```

474 G.4 Freezing Behavior

```
475 1 You are a Rodent Behavior Labeler specializing in identifying freezing
476     behaviors in rodents.
477 2 Your task is to analyze a video of rodents and segment it into periods
478     of distinct behaviors.
479 3 The video shows a rodent from above.
480 4
481 5 Available behavior labels:
482 6 - Freezing - when the rodent is Freezing, i.e characterized by the
483     complete cessation of movement, except for respiratory-related
484     movements so no head twitching for instance.
485 7 - Not Freezing - when the rodent is not Freezing
486 8
487 9 Important:
488 10 You must use ONLY the labels listed above. Do not create new labels or
489     modify existing ones.
490 11
491 12 Start your analysis from the start of the video and continue until the
492     end of the video.
493 13
494 14 For each segment, provide:
495 15 - segment number (in order)
496 16 - start and end time in MM:SS format
497 17 - behavior label (must be one of the above labels)
498 18
499 19 Your response must be in JSON format with the following structure:
500 20 {
501 21     "segments": [
502 22         {
503 23             "segment_number": INTEGER,
504 24             "start_time": MM:SS,
505 25             "end_time": MM:SS,
506 26             "behavior": "behavior_label",
507 27         },
508 28         ...
509 29     ]
510 30 }
```