

Revisiting Text-to-Image Evaluation with Gecko🦎: On Metrics, Prompts, and Human Ratings

Olivia Wiles^{*,1}, Chuhan Zhang^{*,1}, Isabela Albuquerque^{*,1}, Ivana Kajić¹, Su Wang¹, Emanuele Bugliarello¹, Yasumasa Onoe¹, Chris Knutsen¹, Cyrus Rashtchian², Jordi Pont-Tuset¹ and Aida Nematzadeh¹

^{*}Equal contributions, ¹Google DeepMind, ²Google Research

While text-to-image (T2I) generative models have become ubiquitous, they do not necessarily generate images that align with a given prompt. While previous work has evaluated T2I alignment by proposing metrics, benchmarks, and templates for collecting human judgements, the quality of these components is not systematically measured. Human-rated prompt sets are generally small and the reliability of the ratings—and thereby the prompt set used to compare models—is not evaluated. We address this gap by performing an extensive study evaluating auto-eval metrics and human templates. We provide three main contributions: (1) We introduce a comprehensive skills-based benchmark that can discriminate models across different human templates. This skills-based benchmark categorises prompts into sub-skills, allowing a practitioner to pinpoint not only which skills are challenging, but at what level of complexity a skill becomes challenging. (2) We gather human ratings across four templates and four T2I models for a total of >100K annotations. This allows us to understand where differences arise due to inherent ambiguity in the prompt and where they arise due to differences in metric and model quality. (3) Finally, we introduce a new QA-based auto-eval metric that is better correlated with human ratings than existing metrics for our new dataset, across different human templates, and on TIFA160. Data and code will be released at: https://github.com/google-deepmind/gecko_benchmark_t2i

1. Introduction

While text-to-image (T2I) models [51, 63, 3, 50] generate images of impressive quality, the resulting images do not necessarily capture all aspects of a given prompt correctly. See panel (b) of Fig. 1, for an example of a high quality generated image that does not fully capture the prompt, “a cartoon cat in a professor outfit, writing a book with the title “what if a cat wrote a book?””. Notably, reliably evaluating for the correspondence between the image and prompt, referred to as their alignment, is still an open problem requiring three steps: (1) creating a prompt set, (2) designing experiments to collect human judgements, and (3) developing metrics to measure image–text alignment. We discuss each step in turn and how we address shortcomings of previous work.

Prompt set. To thoroughly evaluate T2I models, the choice of the prompts is key, as it defines the various abilities or skills being evaluated; in the above example, a model is tested for alignment that requires action understanding, style understanding, and text rendering. Previous work has curated datasets by grouping prompts taken from existing vision–language datasets into high-level categories, such as reasoning [35]. However, the existing datasets do not provide coverage over a range of skills with various levels of complexity and might test for a few skills at the same time. We address these limitations by developing *Gecko2K*: a comprehensive skill-based benchmark consisting of two subsets, Gecko(R) and Gecko(S), where the prompts are tagged with skills listed in panel (a) of Fig. 1. Gecko(R) is curated by resampling from existing datasets as in [10] to obtain a more comprehensive distribution of skills (see Fig. 2). Gecko(S) is developed semi-automatically, where for each skill we manually design a set of sub-skills to probe models’ abilities in a fine-grained manner.

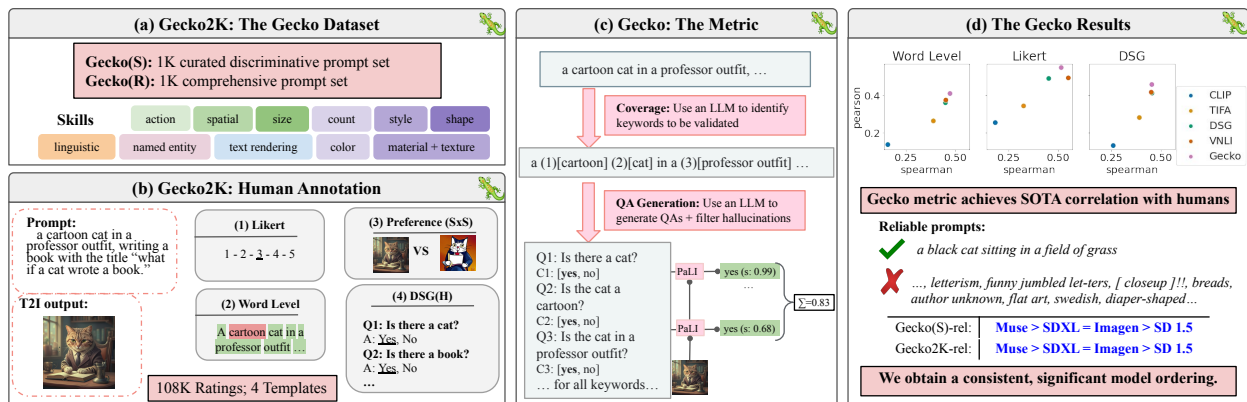


Figure 1 | **The Gecko Framework.** (a) Our Gecko2K dataset consists of two subsets covering a range of skills; Gecko(S) divides skills further into sub-skills. (b) In Gecko2K, we collect judgements for four human templates, resulting in ~108K annotations. (c) The Gecko metric, which includes better coverage for questions, NLI filtering, and better score aggregation. (d) Experiments: using Gecko2K, we compare metrics, T2I models, and the human annotation templates.

For example, as visualised in Fig. 3, for text rendering, it probes if models can generate text of varying length, numerical values, etc.

Human judgement. Another key component is the human experiment design used to collect gold-standard data. One example is to ask the human raters to judge the image–text alignment on a fixed scale of 1 to 5. The design of human experiments, which we refer to as the *annotation template*, impacts the level of fine-grained information we can collect, *i.e.*, whether a specific word is depicted or not. It might also affect the quality and reliability of the collected data. However, various work uses different annotation templates, making it hard to compare results across papers. As a result, we perform a comprehensive study of four main annotation templates (shown in panel (b) of Fig. 1) and collect ratings of four T2I models on Gecko2K. We observe that, across templates, SDXL [47] is best on Gecko(R) and Muse [7] is best on Gecko(S). We find that under Gecko(S), we are able to *discriminate* between models—all templates agree on a model ordering with statistical significance. For Gecko(R), while general trends are similar across templates, results are noisier. This demonstrates the impact of the prompt set when comparing models. Finally, we find some prompts (*e.g.*, “*stunning city 4k, hyper detailed photograph*”) can be ambiguous, subjective, or hard to depict. As a result, we introduce a subset of *reliable prompts* where annotators agree across models and templates. Using this subset in Gecko(R), we observe a more consistent model ordering across templates, where the fine-grained annotation templates agree more.

Auto-eval metrics. Prior work comparing T2I automatic-evaluation (auto-eval) metrics [24, 49, 61, 29, 10] of alignment does so on small ground-truth datasets (with ~2K raw annotations [29, 69]), where a handful of prompts are rated using one annotation template (see Table 1). Using our large set of human ratings (>100K), we thoroughly compare these metrics. We also improve upon the question-answering (QA) framework in [29] (panel (c) in Fig. 1). Our metric, Gecko, achieves state-of-the-art (SOTA) results when compared to human data as a result of three main improvements: (1) enforcing that each word in a sentence is covered by a question, (2) filtering hallucinated questions that a large language model (LLM) generates, and (3) improved VQA scoring.

Fig. 1 summarises our main contributions: **(1) Gecko2K:** A comprehensive and discriminative benchmark with a better coverage of skills and fine-grained sub-skills to identify failures in T2I models and auto-eval metrics. **(2) A thorough analysis** on the impact of annotation templates in model and

	Likert	Word Level	DSG(H)	SxS	# prompts annotated	# anns # img	# anns	# Skills (All Categories)	# Sub-Skills
DSG1K [10]	✗	✗	✓	✗	1.06K	3	9.6K	11(13)	✗
DrawBench [51]	✗	✗	✗	✓	200	25	25K	7(11)	✗
PartiP. [63]	✗	✗	✗	✓	1.6K	5	16K	9(23)	✗
TIFA160 [29]	✓	✗	✗	✗	160	2	1.6K	8(12)	✗
PaintSkills [11]	✓	✗	✗	✗	150	5	2.25K	3(3)	✗
W-T2I [69]	✓	✗	✗	✗	200	3	2.4K	15(20)	9
HEIM [35]	✓	✗	✗	✗	708	~5.4	~150K	6(6)	✗
Gecko2K	✓	✓	✓	✓	2K	~13.5	~108K	12(12)	36
Gecko(R)	✓	✓	✓	✓	1K	~13.5	~54K	11(11)	✗
Gecko(S)	✓	✓	✓	✓	1K	~13.5	~54K	12(12)	36

Table 1 | **Comparison of annotated alignment datasets.** We report the amount of human annotation and skill division for each dataset. We can see that many datasets include only a handful of annotated prompts or a small number of annotations (anns) per image or overall. No dataset besides Gecko2K collects ratings across multiple different human annotation templates. We also include the number of skills and sub-skills in each dataset. Again, Gecko includes the most number of sub-skills, allowing for a fine-grained evaluation of metrics and models. When datasets do not include skills, we map their categories into skills/sub-skills as appropriate.

metric comparisons, which leads us to introduce *reliable prompts*. (3) *Gecko*: An auto-eval metric that is SOTA across our challenging benchmark (Gecko(R)/(S) and the four annotation templates) and TIFA160 [29].

2. Related Work

Benchmarking alignment in T2I models. Many different benchmarks have been proposed to holistically probe a broad range of model capabilities within T2I alignment. Early benchmarks are small scale and created alongside model development to perform side-by-side model comparisons [51, 63, 3]. Later work (e.g., TIFA [29], DSG1K [10] and HEIM [35]) focuses on creating holistic benchmarks by drawing from existing datasets (e.g., MSCOCO [40], Localized Narratives [48] and CountBench [45]) to evaluate a range of capabilities including counting, spatial relationships, and robustness. Other datasets focus on a specific challenge such as compositionality [30], contrastive reasoning [69], text rendering [57], reasoning [11] or spatial reasoning [19]. The Gecko2K benchmark is similar in spirit to TIFA and DSG1K in that it evaluates a set of skills. However, in addition to drawing from previous datasets—which may be biased or poorly representative of the challenges of a particular skill—we collate prompts across sub-skills for each skill to obtain a discriminative prompt set. Moreover, we gather human annotations across multiple templates and many prompts (see Table 1); as a result, we can reliably estimate image–text alignment and model performance.

Automatic metrics measuring T2I alignment. Image generation models initially compared a set of generations with a distribution from a validation set using Fréchet Inception Distance (FID) [25]. However, FID does not fully reflect human perception of image quality [54, 32] nor T2I alignment. Inspired by work in image captioning, a subsequent widely used auto-eval metric for T2I alignment is CLIPScore [24]. However, such metrics poorly capture finer-grained aspects of image–text understanding, such as complex syntax or objects [5, 64]. Motivated by work in natural language processing on evaluating the factuality of summaries using entailment or QA metrics [43, 34, 28], similar metrics have been devised for T2I alignment. VNLI [61] is an entailment metric which finetunes a pretrained vision-and-language (VLM) model [8] on a large dataset of image and text pairs. However, such a metric may not generalise to new settings and is not interpretable—one cannot diagnose why an

Sub-category	Example
Numbers / symbols	equation of " $3+4=7$ " etched into a rock
Length	a neon sign with the words " <i>the future is already here...</i> " reflected on a rainy street. (n=29)
Gibberish	graffiti made with bright pink paint on the concrete, saying " <i>fluff floop floof!</i> "
Typography	"i love you" written in serif font in grass

Table 2 | **Subcategories and corresponding motivations for the text rendering skill.**

alignment score is given. Visual question answering (VQA) methods such as TIFA [29], VQ² [61] and DSG [10] do not require task-specific finetuning and give an interpretable explanation for their score. These metrics create QA pairs which are then scored with a VLM given an image and aggregated into a single score. However, the performance of such methods is conditional on the behaviour of the underlying LLMs used for question generation, and VLMs used for answering questions. For example, LLMs might generate questions about entities that are not present in the prompt or may not cover all parts of the prompt. By enforcing that the LLM does indeed cover each keyword in the prompt and by filtering out hallucinated questions, we are able to substantially improve performance over other VQA-based metrics.

3. Gecko2K: The Gecko Benchmark

We curate a fine-grained skill-based benchmark with good coverage. If we do not have good coverage and include many prompts for one skill but very few for another, we would fail to highlight a model or a metric’s true abilities. We also consider the notion of sub-skills—for a given skill (e.g., *counting*), we curate a list of sub-skills (e.g., *simple modifier*: ‘1 cat’ vs *additive*: ‘1 cat and 3 dogs’). This notion of sub-skills is important as, without it, we may be testing a small, easy part of the distribution such as generating counts of 1-4 objects and not the full distribution that we care about. We introduce two new prompt sets within Gecko2K: Gecko(R) and Gecko(S). Gecko(R) extends the DSG1K [10] benchmark by using automatic tagging to improve the distribution of skills (see App. C for details). However, it does not consider sub-skills nor cover the full set of skills we are interested in. Moreover, a given prompt (as in DSG1K and similar datasets) may evaluate for many skills, making it difficult to diagnose if a specific skill is challenging or if generating multiple skills is the difficult aspect. As a result, we also introduce Gecko(S), a curated skill-based dataset with sub-skills. Finally, we gather a large number of annotations per image across four templates (see Table 1) but defer to Sec. 4 for details.

3.1. Gecko(R): Resampling Davidsonian Scene Graph Benchmark

The recent Davidsonian Scene Graph Benchmark (DSG1K, [10]) curates a list of prompts from existing datasets* but does not control for the coverage or the complexity degrees of a given skill. The authors randomly sample 100 prompts and limit the prompt length to 200 characters.

The resulting dataset is imbalanced in terms of the distribution of skills. Also, as T2I models take in longer and longer prompts, the dataset will not test models on that capability. We take a principled approach in creating Gecko(R) by resampling from the base datasets in DSG1K for better coverage

*TIFA[29], Stanford Paragraphs [33], Localized Narratives [48], CountBench [45], VRD [42], DiffusionDB [59], Midjourney [58], PoseScript [14], Whoops [4], DrawText-Creative [41].

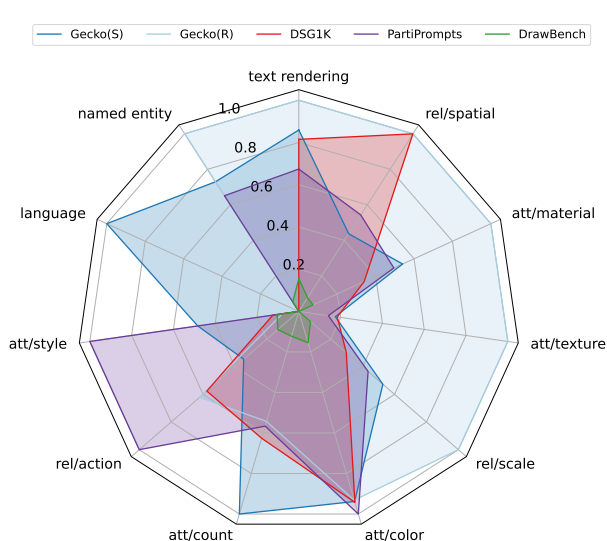


Figure 2 | Distribution of skills. We visualise the distribution of prompts across different skills for Gecko(S)/(R), DSG1K [10], PartiPrompts [63] and DrawBench [51]. We use automatic tagging and, for each skill, normalise by the maximum number of prompts in that skill over all datasets. For most skills, Gecko2K has the most number of prompts within that skill.

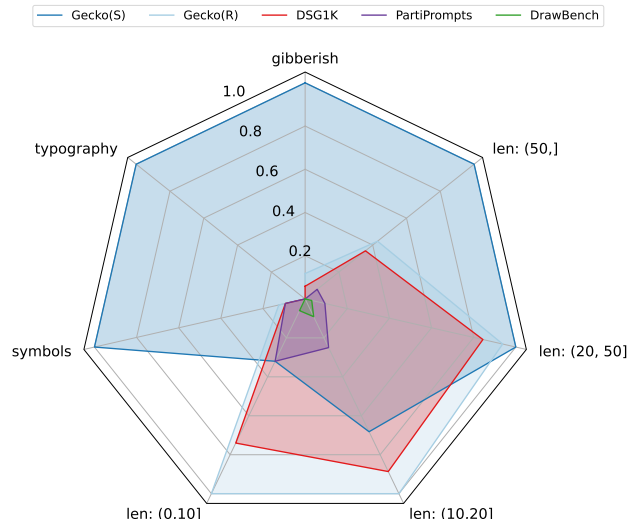


Figure 3 | TEXT RENDERING skill. We visualise the distribution of prompts across seven sub-skills explained in Table 2 (‘len: ...’ corresponds to bucketing different lengths of the text to be rendered). We normalise by the maximum number of prompts in the sub-skill (note that we only count unique texts to be rendered). The Gecko(S) dataset fills in much more of the distribution here than other datasets.

and lifting the length limit. After this process, there are 175 prompts longer than 200 characters and a maximum length of 570 characters. Also, this new dataset has better coverage over a variety of skills than the original DSG1K dataset (see Fig. 2).

While resampling improves the distribution of skills, it has the following shortcomings. Due to the limitations of automatic tagging, it does not include all skills we wish to explore (e.g., language). It also does not include sub-skills: most text rendering prompts do not focus on numerical text, non-English phrases (e.g., Gibberish), the typography, or longer text (see Fig. 3). Finally, automatic tagging can be error prone.

3.2. Gecko(S): A Controlled and Diagnostic Prompt Set

The aim of Gecko(S) is to generate prompts in a controllable manner for skills that are not well represented in previous work. We divide skills into sub-skills to diversify the difficulty and content of prompts. We take inspiration from psychology literature where possible (e.g., colour perception) and known limitations of current text-to-image models to devise these sub-skills. We first explain how we curate these prompts before discussing the final set of skills and sub-skills.

Curating a controlled set of prompts with an LLM. To generate a set of prompts semi-automatically, we use an LLM. We first decide on the sub-skills we wish to test for. For example, for text rendering, we may want to test for (1) English vs Gibberish to evaluate the model’s ability to generate uncommon words, and (2) the length of the text to be generated. We then create a template which conditions the generation on these properties. Note that as we can generate as much data as desired, we can define a distribution over the properties and control the number of examples generated for each sub-skill. Finally, we run the LLM and manually validate that the prompts are reasonable, fluent, and match

the conditioning variables (e.g., the prompt has the right length and is Gibberish / English). A sample template is given in App. C.3.

Gecko(S) make up. Using this approach and also some manual curation, we focus on twelve skills falling into five categories (Fig. 7): (1) NAMED ENTITIES; (2) TEXT RENDERING; (3) LANGUAGE/LINGUISTIC COMPLEXITY; (4) RELATIONAL: ACTION, SPATIAL, SCALE; (5) ATTRIBUTES: COLOR, COUNT, SURFACES (TEXTURE/MATERIAL), SHAPE, STYLE. An overview is given in Fig. 7. We give more details here about how we break down a given skill, TEXT RENDERING, into sub-skills but note that we do the same for all categories; details are in App. G. For TEXT RENDERING, we consider the sub-skills given in Table 2. Using this approach, we get better coverage over the given sub-skills than other datasets (including Gecko(R)) as shown in Fig. 3 and we find results are consistent across different annotation templates.

4. Collecting Human Judgements

When evaluating T2I models, existing work uses different *annotation templates* for collecting human judgements. While human annotation is the gold-standard for evaluating generative models, previous work shows the design of a template and background of annotators can significantly impact the results [12]. For example, a common approach is to use a Likert scale [39] where annotators evaluate the goodness of an image with a fixed scale of discrete ratings. However, previous work [10, 29, 35] uses different designs (e.g., scales from 1 to 5 vs 1 to 3), so results cannot be directly compared. We examine differences between templates and how the choice of a template [10, 38] impacts the outcomes when comparing four models: SD1.5 [50], SDXL [46], Muse [7], and Imagen [51]. (Imagen and Muse are variants based on the original models, but trained on internal data sources.) We consider three *absolute* comparison templates (Likert, Word Level, and DSG(H)) which evaluate models individually, and a template for *relative comparison* (SxS). The templates are defined below and a high-level visualisation of each template is in panel (b) of Fig. 1. More details are in App. E.1.

4.1. Annotation Templates

Likert scale. We follow the template of [10] and collect human judgements using a 5-point Likert scale by asking the annotators “How consistent is the image with the prompt?” where *consistency* is defined as how well the image matches the text description. Annotators are asked to choose a rating from the given scale, where 1 represents *inconsistent* and 5 *consistent*, or a sixth *Unsure* option for cases where the text prompt is not clear. Choosing this template enables us to compare our results with previous work, but does not provide fine-grained, word-level alignment information. Moreover, while Likert provides a simple and fast way to collect data, challenges such as defining each rating especially when used without textual description (e.g., what 2 refers to in terms of image–text consistency), can lead to subjective and biased scores [23, 37].

Word-level alignment (WL). To collect word-level alignment annotations, we use the template of [38] and define an overall image–text alignment score using the word-level information. Given a text–image pair, raters are asked to annotate each word in the prompt as *Aligned*, *Unsure*, or *Not aligned*. Note that for each text–image pair under the evaluation, the number of effective annotations a rater must perform is equal to the number of words in the text prompt. Although potentially more time consuming than the Likert template,[†] We compute a score for each prompt–image pair per rater by aggregating the annotations given to each word. A final score is then obtained by averaging the scores of 3 raters. More details can be found in App. E.1.

[†]We find that annotators spend ~30s more to rate a prompt–image pair with WL than Likert.

DSG(H). We also use the annotation template of [10] that asks the raters to answer a series of questions for a given image, where the questions are generated automatically for the given text prompt as discussed in [10].

In addition, raters can mark a question as *Invalid* in case a question contradicts another one. The total number of *Invalid* ratings per evaluated generative model is given in App. E.1. Annotators could also rate a question as *Unsure*, in cases where they do not know the answer or find the question subjective or not answerable based on the given information.

For a given prompt, the number of annotations a rater must complete is given by the number of questions. We calculate an overall score for an image–prompt pair by aggregating the answers across all questions, and then averaging this number across raters to obtain the final score.

Side-by-side (SxS). We consider a template in which pairs of images are directly compared. The annotators see two images from two models side-by-side and are asked to choose the image that is *better aligned* with the prompt or select *Unsure*. We obtain a score for each comparison by computing the majority voting across all 3 ratings. In case there is a tie, we assign *Unsure* to the final score of an image–prompt pair.

Data Collection Details. We recruited participants ($N = 40$) through a crowd-sourcing pool. The full details of our study design, including compensation rates, were reviewed by our institution’s independent ethical review committee. All participants provided informed consent prior to completing tasks and were reimbursed for their time. Considering all four templates, both Gecko subsets, and the four evaluated generative models, approximately 108K answers were collected, totalling 2675 hours of evaluation.

4.2. Comparing Annotation Templates

We first evaluate the reliability of each absolute comparison template by measuring inter-annotator agreement—does the choice of the template impact the quality of the data we can collect?

We compute Krippendorff’s α [22] for each generative model and template and report the results in Table 3. Krippendorff’s α assumes values between -1 and 1, with 1 indicating perfect agreement and 0 chance level [65]. As our goal is evaluating agreement between annotators, in this experiment we keep all *Unsure* ratings, but for the following experiments we pre-process the annotations by removing all such ratings so as to only reflect cases where annotators were confident. We observe that all the templates yield agreement above chance levels for all generative models, with $\alpha > 0.5$, except for the Likert–SD1.5 pair for Gecko(R). We also observe that for Gecko(R), WL is a more reliable template but for synthetic prompts, DSG(H) has higher agreement. We then investigate which templates agree/disagree more with each other by computing Spearman’s rank and Pearson correlations between scores from each template. Results are in App. E.2 and show that there is an overall high correlation between all templates on both Gecko(S)/(R) and the strongest observed correlation among all models and metrics is between the finer-grained templates, WL and DSG(H), on Gecko(S).

Reliable Prompts. Given the human ratings we have collected, we can examine how *reliable* a prompt is—the extent to which the annotators agree in their ratings regardless of the annotation template and the model used to generate images. To get a set of *reliable prompts*, for each model–template pair, we select the prompts for which inter-rater *disagreement*[‡] is below 50% the maximum *disagreement* observed across all prompts for that model–template pair. We then repeat this step

[‡]Defined as the variance across the scores given by each rater in the case of Likert, and the average variance across words and questions scores in the case of WL and DSG(H), respectively.

Gen. model	Inter annotator agreement						Scores					
	Gecko(R)			Gecko(S)			Gecko(R)			Gecko(S)		
	WL	Likert	DSG(H)	WL	Likert	DSG(H)	WL	Likert	DSG(H)	WL	Likert	DSG(H)
Imagen	0.81	0.64	0.68	0.72	0.57	0.75	0.74 $\pm$ 0.30	0.60 $\pm$ 0.22	0.84 $\pm$ 0.18	0.80 $\pm$ 0.24	0.59 $\pm$ 0.20	0.78 $\pm$ 0.23
Muse	0.82	0.78	0.72	0.69	0.58	0.72	0.84 $\pm$ 0.24	0.61 $\pm$ 0.25	0.83 $\pm$ 0.22	0.88 $\pm$ 0.18	0.63 $\pm$ 0.21	0.84 $\pm$ 0.21
SDXL	0.75	0.76	0.57	0.67	0.56	0.70	0.87 $\pm$ 0.19	0.68 $\pm$ 0.22	0.86 $\pm$ 0.16	0.80 $\pm$ 0.23	0.60 $\pm$ 0.21	0.79 $\pm$ 0.22
SD1.5	0.66	0.36	0.66	0.69	0.59	0.74	0.86 $\pm$ 0.16	0.67 $\pm$ 0.22	0.76 $\pm$ 0.23	0.61 $\pm$ 0.33	0.49 $\pm$ 0.21	0.68 $\pm$ 0.27

Table 3 | **Inter-annotator agreement and ratings for all models and templates.** We measure inter-annotator agreement for each human evaluation template with Krippendorff’s α . Higher values indicate better agreement. We also show the mean and standard deviation for the scores of all templates after mapping the ratings to the $[0, 1]$ interval, with 1 indicating perfect alignment.

for all model–template pairs, and consider the intersection of prompts across these settings as the *reliable prompts*. We further refine this set by removing instances for which all ratings from the Likert template are *Unsure*. This procedure yields a total of 531 and 725 *reliable prompts* for Gecko(R) and Gecko(S), respectively. As we only consider ratings from the absolute comparison templates when finding *reliable prompts*, we investigate how the SxS template agreement (*i.e.*, Krippendorff’s α) changes given this set. This experiment validates our notion of reliability by evaluating if it is transferable to other templates. For Gecko(R), we find that using reliable prompts increases the average Krippendorff’s α from 0.45 to 0.47, and for Gecko(S) from 0.49 to 0.54, showing that both sets of reliable prompts generalise to other templates. Detailed results are in App. E.2.

SxS comparison. To compare the SxS template with the absolute comparison ones, we compute the accuracy obtained by each absolute template when predicting the preferred model given by SxS annotations with the reliable subsets of both Gecko(R) and Gecko(S). For Gecko(R), we find that DSG(H) presented the best accuracy in 4 out of 6 evaluations, followed by Likert which was the most accurate absolute template in the remaining 2 evaluations. On the other hand, for Gecko(S), we find that WL and Likert predict SxS judgements with higher accuracy, each one being the most accurate template in 3 out of 6 cases, indicating that, overall, Likert scores were able to better predict results of SxS comparisons. The complete set of results is in App. E.2.

4.3. Comparing I2T Models

We first compare the T2I models with the human ratings collected using the three absolute comparison templates as they presented overall higher inter annotator agreement. In Table 3, we report the average and standard deviation of all ratings. We verify the significance of outcomes by performing the Wilcoxon signed-rank test for all pairs of models with $p < 0.001$. Where results indicate the null-hypothesis is rejected (*i.e.*, the distribution of ratings is significantly different), we can say there is enough evidence that one model is better than another. To determine which model is best, we compare the mean values of their scores. Fig. 4 presents the results for the comparisons between all pairs of models with reliable (Gecko(*)-rel) prompt subsets. When considering significance in Fig. 4, we see that Muse is not worse than any of the contenders across all templates and prompt sets; we determine it is the best overall model given human judgement. Also, SD1.5 is worse or the same as other models except for Imagen when evaluated with Gecko(R)-rel.

Fig. 4 results also shed light on the agreement between templates: for Gecko(S)-rel all templates agree with each other, and for Gecko2K-rel all templates agree or do not contradict each other (*i.e.*, one is significant, the other not).

Summary. We see that our discriminative prompt set, Gecko(S), results in consistent model rankings,

Imagen	Muse			SDXL			SD1.5			SDXL			SD1.5			SD1.5		
	WL	L	D(H)	WL	L	D(H)	WL	L	D(H)	WL	L	D(H)	WL	L	D(H)	WL	L	D(H)
	G(S)-rel	<	<	<	=	=	=	>	>	>	>	>	>	>	>	>	>	>
	G(R)-rel	<	<	=	<	<	>	<	<	>	<	<	>	<	<	>	<	<
G2K-rel		<	<	<	=	<	=	>	>	>	>	>	>	>	>	>	>	>
		<	<	<	=	<	=	>	>	>	>	>	>	>	>	>	>	>

Figure 4 | **Comparing models using human annotations.** We compare model rankings on the reliable subsets of Gecko(S) (G(S)-rel), Gecko(R) (G(R)-rel) and both subsets (G2K-rel). Each grid represents a comparison between two models. Entries in the grid depict results for WL, Likert (L), and DSG(H) (D(H)) scores. The > sign indicates the left-side model is better, worse (<), or not significantly different (=) than the model on the top.

independent of the template used. For prompts that test for aspects other than image–text alignment, *i.e.*, Gecko(R), the choice of the annotation template impacts the ordering of the models, but the best and worst models are mostly consistent across templates. In terms of templates, the finer-grained ones (WL and DSG(H)) have better inter-annotator agreement in comparison to Likert and SxS. We also notice that when comparing model rankings with Gecko(R)-rel, DSG(H) is the template that mostly disagrees with the others. This finding, along with the higher inter-annotator agreement shown by the WL template, makes it the overall best choice for evaluating image–text alignment on Gecko2K.

5. The Gecko Metric

Question answering (QA) based metrics have the advantage of attributing failures to specific questions as opposed to giving a single, uninterpretable score (as in VNLI [61] and CLIP [49]), and benefit from the increasing capabilities of LLMs and VQA models. The Gecko metric improves upon QA [29, 10, 61] metrics by enforcing coverage of the prompt and reducing hallucinations. A standard QA setup (from [29]) consists of three steps: (1) QA generation: prompting an LLM to generate a set of binary question-answer pairs $\{Q_i, A_i\}_{i=1}^N$ on a given T2I text description T . (2) VQA assessment: employing a VQA model to predict answer $\{A'_i\}_{i=1}^N$ for the generated questions given the generated image I . (3) Scoring: computing the alignment score by assessing the VQA accuracy using Eq. (1).

$$\text{Alignment}(T, I) = \frac{1}{N} \sum_{i=1}^N \mathbb{1}[A'_i = A_i]. \quad (1)$$

However, this pipeline exhibits several limitations. First, controlling the quantity of QA pairs generated, especially from lengthy sentences, poses a challenge. In such cases, the prompted LLM often selectively generates questions for specific segments of the text while overlooking others. Second, LLMs are prone to producing hallucinations [2, 20], leading to the generation of low-quality, unreliable QA pairs. Furthermore, the reliance on binary judgement—strictly matching A'_i and A_i without considering the predicted probability of A'_i —overlooks the inherent uncertainty in the predictions.

While TIFA[29] and DSG[10] have proposed some solutions to overcome some of the limitations highlighted above, such as using a QA model to verify the generated QAs or breaking the prompt into atomic parts, the effectiveness of these solutions remains limited, especially on more complicated text prompts. Here we propose a simpler, but more robust method that consistently outperforms on various benchmarks.

Coverage. To improve the coverage of questions over the text sentence T , we split the QA generation

Metrics	FT	Gecko(R)				Gecko(S)				Gecko(R)-Rel				Gecko(S)-Rel			
		WL	Likert	DSG(H)	SxS	WL	Likert	DSG(H)	SxS	WL	Likert	DSG(H)	SxS	WL	Likert	DSG(H)	SxS
		SpearmanR			Acc	SpearmanR			Acc	SpearmanR			Acc	SpearmanR			Acc
CLIP	✗	0.14	0.16	0.13	54.4	0.25	0.18	0.26	67.2	0.11	0.09	0.08	59.7	0.24	0.19	0.25	67.1
PyramidCLIP	✗	0.26	0.27	0.26	64.3	0.22	0.25	0.23	70.7	0.26	0.26	0.23	65.8	0.21	0.26	0.22	71.0
TIFA	✗	0.26	0.34	0.28	41.7	0.39	0.32	0.39	53.2	0.34	0.37	0.33	47.3	0.41	0.36	0.40	52.8
DSG	✗	0.35	0.47	0.42	49.6	0.45	0.45	0.45	58.1	0.47	0.50	0.50	53.9	0.48	0.47	0.48	56.9
Gecko	✗	0.41	0.55	0.46	62.1	0.47	0.52	0.45	74.6	0.52	0.58	0.53	71.3	0.52	0.54	0.49	75.2
VNLI	✓	0.37	0.49	0.42	54.4	0.45	0.55	0.45	72.7	0.49	0.57	0.46	65.6	0.50	0.61	0.48	72.7

Table 4 | **Correlation between VQA-based, contrastive, and fine-tuned (FT) auto-eval metrics and human ratings across annotation templates on Gecko2K and Gecko2K-Rel.** Gecko scores highest in 13/14 conditions, even out-performing the fine-tuned VNLI metric.

into two distinct steps. We first prompt the LLM to index the visually groundable words in the sentence. For example, the sentence ‘A red colored dog.’ is transformed into ‘A {1}[red colored]{2}[dog].’ Subsequently, using the text with annotated keywords $\{w'_i\}_{i=1}^N$ as input, we prompt the LLM again to generate a QA pair $\{q_i, a_i\}$ for each word labelled $\{w'_i\}$ in an iterative manner (see App. D for the prompting details). This two-step process ensures a more comprehensive and controllable QA generation process, particularly for complex or detailed text descriptions.

NLI filtering. For filtering out the hallucinated QA pairs, inspired by previous work in NLP [43, 34], we employ an Natural Language Inference (NLI) model [27] model for measuring the factual consistency between the text T and QA pairs $\{Q_i, A_i\}$. QA pairs with a consistency score lower than a pre-defined threshold r are removed, ensuring that only the factually aligned ones are retained.

VQA score normalisation. The final improvement we make is how to aggregate scores from the VQA model. The motivation is that the VQA model can predict a very similar score for two answers. If we simply take the max, then we lose this notion of uncertainty reflected in the scores of the model. If the negative log likelihood of answer A'_i is s_i and the correct answer is A'_a with score s_a , we normalise the scores as follows:

$$Alignment(T, I) = \frac{1}{N} \sum_{i=1}^N \frac{s_a}{\sum_i s_i}. \quad (2)$$

6. Experiments on Auto-Eval Metrics

To compare auto-eval metrics, we perform two sets of experiments: first, we compare how various metrics correlate with human judgements—whether their predicted alignment scores match those of the raters for each prompt. We demonstrate that the proposed Gecko metric consistently performs best in Sec. 6.2 for Gecko(S)/(R) and analyse results per category for Gecko(S). Additionally, we validate the utility of each component.

Finally, in Sec. 6.3, we validate that metrics give the same model ordering as human judgements and find that Gecko is the same or better than other metrics at this task.

6.1. Experimental Setup

Metrics Evaluated. We benchmark three types of auto-eval metrics on Gecko2K: (1) contrastive models: CLIP[49] and PyramidCLIP [16]; (2) NLI models: VNLI [61]; (3) QA-based methods: TIFA [29], DSG [10] and our proposed metric Gecko.

Pre-trained Models. We use a ViT-B/32 [15] model for CLIP and ViT-B/16 [15] for PyramidCLIP

Skill (subskill): Prompt:	lang/complexity (negation) A bridge with no cars on it.				Shape: (hierarchical) The number 0 made of smaller circles			
								
	Imagen	Muse	SDXL	SD1.5	Imagen	Muse	SDXL	SD1.5
WL:	1.	1.	1.	1.	1.	1.	0.	0.67
Likert:	1.	1.	1.	0.87	1.	0.87	0.2	0.67
DSG(H):	0.5	0.5	0.5	0.67	0.92	1.	0.	0.89
Gecko:	0.96	0.94	0.93	0.91	0.9	0.95	0.55	0.75
DSG:	0.25	0.25	0.25	0.25	1.	1.	0.	0.
VNLI:	0.4	0.42	0.32	0.30	0.36	0.79	0.24	0.32

Figure 6 | **Visualisations of scores from different auto-eval metrics.** We show the image generations by the four generative models given two prompts from Gecko(S), with the alignment scores estimated by human annotators and auto-eval metrics.

in our experiments. For all the VQA-based models, we use PaLM-2 [1] as the LLM and PaLI [8] as the VQA model for fair comparison. When evaluating the Gecko metric, we use a T5-11B NLI model from [27] and set the threshold r at 0.005. This threshold was determined by examining QA pairs with NLI probability scores below 0.05. We observed that QAs with scores below 0.005 are typically hallucinations. We re-use the original prompts from TIFA for generating QAs, and add coverage notation to their selected texts as described in Sec. 5.

Finally, some baseline models are trained with a maximum text input length L , where $L_{\text{CLIP}} = 77$ and $L_{\text{VNLI}} = 82$, thus unable to process longer descriptions; for these models, we only take the first L tokens from the texts as input.

6.2. Correlation with Human Judgement

We evaluate all the auto-eval metrics on the Gecko2K benchmarks, by calculating correlations with their scores and those obtained from human ratings across four annotation templates as discussed in Sec. 4.1. We compute Pearson and Spearman Ranked correlation—the former is a stricter measure of linear correlation, while the latter assesses the ranking agreement between paired data points, making it more robust to outliers and non-linear relationships.

Gecko	0.65	0.54	0.52	0.45	0.50	0.67	0.59	0.40	0.71	0.60	0.63	0.27
DSG	0.53	0.48	0.55	0.41	0.42	0.64	0.63	0.28	0.64	0.61	0.58	0.08
VNLI	0.71	0.54	0.62	0.45	0.55	0.49	0.43	0.46	0.66	0.50	0.63	0.46
	color	count	shape	style	texture material	lang complexity	lang compositional	named entity	action	scale	spatial	text

Figure 5 | **Per skill results of Gecko, DSG and VNLI.** We visualise the Likert correlation for each skill. Where p-values are > 0.05 , we colour the square black. Results for other metrics are in App. G.

Results on Gecko(R). We first compare the auto-eval metrics (CLIP, TIFA, DSG, Gecko) that do not rely on fine-tuning (results and findings from additional auto-eval metrics are in Sec. G.3.) As shown in Tab. 4, the Gecko metric outperforms the others and shows a higher correlation with most of the human annotation templates. As expected, contrastive metrics (CLIP variants) are worse overall than QA-based metrics, as they only measure coarse-grained alignment at the sentence level. However, it is worth noting that a better CLIP model (*e.g.*, PyramidCLIP) can be a very strong baseline on SxS accuracy, showing its superior ability in doing pair-wise comparison.

DSG is the second best metric, except on SxS where it ranks third. It outperforms TIFA by a clear margin but falls behind Gecko. Finally, we compare Gecko with VNLI, our supervised baseline as the VNLI model is fine-tuned for text-image alignment on a mixed dataset containing COCO (which is used in DSG(R)), while other metrics are not fine-tuned. Nonetheless, Gecko still shows a much

Metrics	Gecko(R)						Gecko(S)					
	WL		Likert		DSG(H)		WL		Likert		DSG(H)	
	Pearson	Spearman-R	Pearson	Spearman-R	Pearson	Spearman-R	Pearson	Spearman-R	Pearson	Spearman-R	Pearson	Spearman-R
TIFA baseline	0.21	0.26	0.32	0.34	0.25	0.28	0.39	0.39	0.32	0.32	0.39	0.39
+ coverage	0.28	0.28	0.34	0.36	0.32	0.32	0.41	0.39	0.33	0.33	0.40	0.39
+ VQA score norm	0.32	0.32	0.42	0.44	0.37	0.38	0.43	0.42	0.37	0.40	0.41	0.42
+ NLI filtering	0.38	0.41	0.51	0.55	0.42	0.46	0.46	0.47	0.48	0.52	0.46	0.45

Table 5 | **Ablation on proposed Gecko metric on Gecko2K.** We ablate the effectiveness of the three proposed improvements by adding them to the TIFA baseline one by one. They all bring higher correlation with human judgement across the board on Gecko2K.

higher correlation with WL, Likert and DSG(H).

Results on Gecko(S). As shown in Tab. 4, we observe a similar pattern to Gecko(R) on our synthetic subset, Gecko(S): the Gecko metric has the highest correlation with human ratings; DSG is the second best metric, except for the SxS template where CLIP outperforms it. Given a template, the correlation scores are generally higher on Gecko(S) which is not surprising at this benchmark focuses on measuring alignment. To better understand the differences between auto-eval metrics/annotation templates with respect to various skills, we visualise a breakdown of skills in Gecko(S) in Fig. 5 and App. F.1, G.4. Different metrics have different strengths: for example, we see that while Gecko and VNLI metrics are consistently good across skills, the Gecko metric is better on more complex language, DSG is better on compositional prompts, and VNLI is better on text rendering. We visualise examples in Fig. 6. For the negation example, we can see that VNLI and DSG mistakenly think none of the images are aligned. We can also see that the reason DSG(H) gives inconsistent results with WL/Likert here is that the question generation is confused by the negation (asking if there *are* cars as opposed to *no* cars). VNLI and DSG perform better on the shape prompt but VNLI scores Imagen incorrectly and DSG gives hard results (0 or 1) and so is not able to capture subtler differences in the human ratings.

Results on Gecko2K-Rel. We also include the correlation of metrics on reliable prompts in Tab. 4. We can see that on this reliable subset, the Gecko metric shows even better results by outperforming other metrics on everything. Furthermore, we observe that when moving from Gecko2K to Gecko2K-Rel, all the correlation scores between the QA-based metrics and human ratings increase, while those of the CLIP variants drop. This shows that the relatively high correlations from CLIP in the full-set has some noise, and QA-based metrics align more closely with human judgement.

Ablation on proposed Gecko metric. We ablate the impact of the three key improvements we propose, namely, QA generation with coverage, linear normalisation of predicted probabilities, and NLI filtering on QA. Starting from our baseline TIFA, we include the improvements one at a time to see the benefits brought. The results in Tab. 5 uniformly demonstrate a positive impact from each improvement. Notably, NLI filtering brings the largest boost among the three, underscoring the limitation of LLMs in reliably generating high-quality and accurate QA pairs.

Results on TIFA160. We compare the Gecko metric with other metrics on TIFA160 [29]. It is a set of 160 text-image pairs, each with two Likert ratings. In Tab. 6, we list the results of other metrics reported in [29, 10], and compare them with Gecko as well as our re-implementation of TIFA and DSG. Gecko has the highest correlation with Likert scores, with an average correlation 0.07 higher than that of DSG, when using the same QA and VQA models. This shows that the power of our proposed metric is from the method itself, not from the advance of models used.

6.3. Model-ordering Evaluation

We next examine if each auto-eval metric can predict how two models are ordered according to the human ratings. We use the results obtained by comparing T2I models with scores from all absolute comparison templates reported in Fig. 4 and consider the majority relation assigned to a model pair as the ground truth: if two of the templates find model 1 is better than model 2, then we assume model 1 > model 2. Given the ground-truth relation between a model pair, we then count the number of times an auto-eval metric successfully found the same relation, in terms of both significance and direction. In addition, we count the number of times an auto-eval metric captured the ground-truth relationship following the most common practice in the literature: ordering models by solely comparing the mean values of scores regardless of statistical significance. In App. G.2, we show the results for all auto-eval metrics on both Gecko(R) and Gecko(S). Our results highlight that the number of significant successful comparisons is lower than comparisons with mean values. All metrics except for DSG and CLIP get the same number of successful comparisons for both criteria. Note that Gecko not only is within the metrics that best capture pairwise model orderings as per the human eval ground-truth values, but Gecko scores are also the best correlated with the human scores themselves (as we can see in Table 4).

Metrics	QA	VQA	Spearman’s ρ	Kendall’s τ
ROUGE-L	N/A	N/A	0.33	0.25
METEOR			0.34	0.27
SPICE			0.33	0.23
CLIP			0.33	0.23
TIFA	GPT-3	BLIP-2	0.56	0.44
	GPT-3	MPLUG	0.60	0.47
	PALM	PaLI	0.43	0.32
DSG	PALM	PaLI	0.57	0.46
Gecko	PALM	PaLI	0.64	0.50

Table 6 | **Comparing different metrics by their correlation with human Likert ratings on TIFA160.** The Gecko metric outperforms the others by a significant margin.

7. Discussion and Conclusions

We took stock of auto-eval of T2I models and the three main components: the benchmark, the annotation templates, and the auto-eval metric. We introduced the *Gecko2K* benchmark, gathered ratings across four templates and introduced a notion of *reliable* prompts; we also developed *Gecko*—a SOTA VQA auto-eval metric. We conclude the following takeaways: (1) Fine-grained annotation templates (e.g., WL, DSG(H)) are more consistent with each other than coarse-grained (e.g., Likert and SxS) and vice versa. (2) If using reliable prompts and templates with high inter-annotator agreement, we get a consistent model ordering (irrespective of the choice of prompts—e.g., Gecko(R)/(S)). (3) When comparing auto-eval metrics, it is better to use ‘reliable’ prompts, as they better measure alignment, give better signal from raters and a more consistent ordering of metrics and models.

Our work highlights the importance of standardising the evaluation of models with respect to both the benchmark used and also the annotation template. While our proposed metric can be reliably used for model comparisons, it is still limited by the quality of the pre-trained LLMs and VLMs in the pipeline. Consequently, an interesting future direction is to provide a confidence threshold in addition to the metric scores.

Acknowledgements We thank Nelly Papalampidi, Zi Wang, Miloš Stanojević, and Jason Baldridge for their feedback throughout the project. We are grateful to Nelly Papalampidi and Andrew Zisserman for their feedback on the manuscript. We thank Aayush Upadhyay and the rest of the Podium team for their help in running models.

A. Appendices

A Appendices	14
B Overview	15
C Gecko2K: More details	15
C.1 Automatic tagging for Gecko(R)	15
C.2 Prompt distribution in Gecko(R)	16
C.3 Templates to few-shot an LLM for Gecko(S)	16
C.4 Breakdown by skill/sub-skill in Gecko(S)	17
D Gecko Metric: More details	22
D.1 LLM prompting for generating coverage	22
D.2 LLM prompting for generating QAs	22
E Human annotation: more details and experiments	24
E.1 Annotation templates	24
E.2 Additional experimental results	26
E.3 Reliable prompts: Examples of image-prompt pairs with high human (dis)agreement	27
E.4 Human evaluation templates: Challenging cases	33
F T2I Models: Additional comparisons	35
F.1 Analysing Model Ratings Per Skill	35
F.2 Model comparisons with TIFA160	38
G Auto-eval metrics: Additional experiments	38
G.1 Pearson correlation Gecko2K.	38
G.2 Model-ordering Evaluation	39
G.3 Results for Additional Metrics on the Gecko Benchmark	39
G.4 Analysing Auto-Eval Metric results per skill.	40
G.5 Additional visualisations.	40
G.6 Results per Word for WL	40

B. Overview

In the Appendix, we give additional information on the benchmark, human annotation and corresponding results for T2I models, and experimental results for the auto-eval metrics.

Gecko Benchmark: For the benchmark, we give further information on how we automatically tag Gecko(R) and semi-automatically generate prompts in Gecko(S) in Appendix C.1 and C.3 respectively. We then give more detail about the skill breakdown in Gecko(R) in App. C.2. We define and give examples for the sub-skills in Gecko(S) in App. C.4.

Human Annotation: For the human annotation, we give additional details of our setup including screenshots of the annotation templates used and qualitative limitations of each setup in App. E.1. We further discuss more experimental results comparing inter-annotator agreement and the raw predictions under each template in App. E.2. Finally, we visualise the most and least reliable prompts in App. E.3, giving an intuition for the properties of the prompt that lead to more or less agreement across templates.

Additional results on T2I models: We give further results on using the annotated data to (1) compare T2I models by skill in App. F.1. We also compare how well prompts in TIFA160 are able to discriminate models under our human annotation setup in App. F.2.

Additional results for auto-eval metrics: Finally, we give additional results for the auto-eval metrics, including more correlation results in App. G.1, the raw results for the model-ordering evaluation in App. G.2 and results for different CLIP variants in App. G.3. We then explore results per skill for different auto-eval metrics in App. G.4, give additional visualisations in App. G.5 and demonstrate that we can use Gecko to evaluate the per-word accuracy of the metric (this is not possible with other auto-eval metrics) in App. G.6.

C. Gecko2K: More details

As described in Sec. 3.1, we use automatic tagging in order to tag prompts with different skills in Gecko(R). However, this has a few issues: (1) it can be error prone; (2) we are limited by the tagging mechanism in the skills that we tag; (3) we do not tag sub-skills. As a result, we devise a semi-automatic approach to build Gecko(S) by few-shot prompting an LLM, as discussed in Sec. 3.2 and curate a dataset with a number of skills and sub-skills for each skill.

C.1. Automatic tagging for Gecko(R)

As mentioned in Sec. 3.1, to obtain a better control for the skill coverage and prompt length, we resampled from the 10 datasets used in DSG1k[10]. To identify the categories covered in the prompts, we adopted an automatic tagging method similar to that used in DSG1K. This method utilizes a Language Model (LLM) to tag words in the text prompt, as shown in Listing 1. The only difference is that we also included named entities and landmarks to be the original categories, such as *whole*, *part*, *state*, *color* etc.

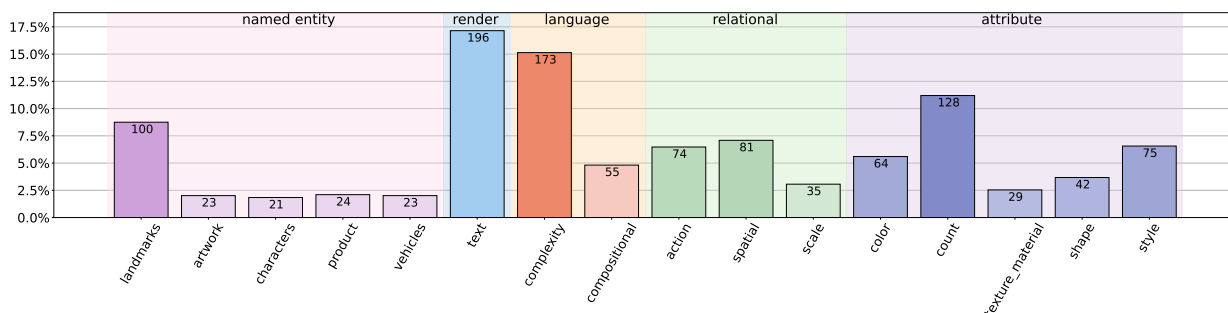


Figure 7 | **Overview of Gecko(S)**. The set of skills (coloured by the corresponding category) covered by the synthetic prompts. Note that we gather prompts by breaking each skill into sub-skills.

```

1  """
2
3  id: synthetic_v1_1
4  input: a man is holding an iPhone.
5  output: 1 | entity - whole (man)
6           2 | entity - named entity (iPhone)
7           3 | action - hold (man, iPhone)
8
9  id: diffusiondb_79
10 input: an hd painting by Vincent van Gogh. a bunch of zombified mallgoths hanging out at a hot topic store in the mall.
11 output: 1 | global - style (hd painting)
12          2 | global - style (Vincent van Gogh)
13          3 | entity - whole (mallgoths)
14          4 | attribute - state (mallgoths, zombified)
15          5 | other - count (mallgoths, ==bunch)
16          6 | entity - whole (hot topic store)
17          7 | entity - whole (mall)
18          8 | relation - spatial (mallgoths, hot topic store, at)
19          9 | relation - spatial (hot topic store, mall, in)
20 ...
21
22 id: {image_id}
23 input: {text_input}
24 output: {LLM_output}
25 """

```

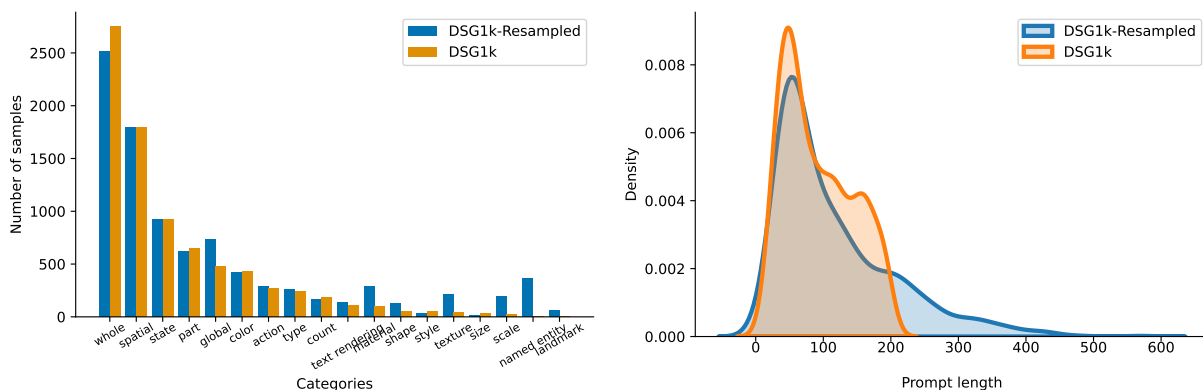
Listing 1 | The prompt used to automatically tag skills given text prompts from the base datasets in DSG1K in order to generate a more balanced Gecko(R).

C.2. Prompt distribution in Gecko(R)

We resample 1000 prompts from the base datasets used in DSG1K and ensure a more uniform distribution over skills and prompt length. To sample more prompts featuring words from under-represented skills (e.g. TEXT RENDERING, SHAPE, NAMED IDENTITY and LANDMARKS), we use automatic tagging in Appendix C.1 to categorize the words in all prompts as pertaining to a given skill. We then resample, assigning higher weights to the under-represented skills. The resulting skill distribution is shown in Fig. 8. Although the resampling increases the proportion of under-represented skills, the overall distribution remains unbalanced. This underscores the necessity of acquiring a synthetic subset with a more controlled and balanced prompt distribution. To sample long prompts, we eliminate the constraint set in DSG1K[10], which mandates that the sampled prompts should be shorter than 200 characters. This adjustment results in a more diverse prompt length distribution as shown in Fig. 8.

C.3. Templates to few-shot an LLM for Gecko(S)

As discussed in Sec. 3.2, we semi-automatically create prompts for Gecko(S) by few-shot prompting an LLM. We give an example for the TEXT RENDERING skill in Listing 2. In short, we define a set of properties based on the sub-skills we want included in our dataset. In this case, we define *text*



(a) The skill distribution (which is tagged at the word level).

(b) The prompt length distribution.

Figure 8 | Prompt distribution in DSG1K-Resampled(Gecko-R) and DSG1K.

length and language (we use *English* and *Gibberish* but we note this could be easily extended to more languages). We then create examples that have those properties to create our few-shot prompt. We can query the LLM as many times as we like to create a distribution of prompts across different text lengths and languages. We do a similar setup for each of the skills and sub-skills we define below.

```

1 """
2 Generate captions for the given text of varying length. Be creative and imagine new settings.
3
4 Text length: 20
5 Language: English
6 Text: "look at that shadow!"
7 Caption: shadow of a stone, taken from the point of view of an ant, with the caption "look at that shadow!"
8
9 ...
10
11 Text length: {text_length}
12 Language: {language}
13 Text: {LLM_output}
14 Caption: {LLM_output}
15 """

```

Listing 2 | LLM sampling template.

C.4. Breakdown by skill/sub-skill in Gecko(S)

An overview of Gecko(S) is given in Fig. 7. In this section we give more information on the skills and sub-skills within Gecko(S). We provide a detailed breakdown of each prompt sub-skill, including examples and justifications. Skills and sub-skills are listed in Table 7. We aim to cover semantic skills, some of which have already been covered in previous work (*e.g. shapes, colors or counts*), while further subdividing each skill to capture its different aspects and difficulty levels. By varying the difficulty of the prompts within a challenge we ensure we are testing the models and metrics at different difficulty levels and can find where models and metrics begin to break.

Each skill (such as *SHAPE*, *COLOR*, or *NUMERICAL*) is divided into sub-skills, so that prompts within that sub-skill can be distinguished based on difficulty or, if applicable, some other criteria that is unique to that sub-skill (*i.e. prompts inspired by literature in psychology*). We create a larger number of examples and subsample to create our final 1K set of prompts to be labelled.

C.4.1. Spatial Relationships

This skill captures a variety of spatial relationships (such as *above, on, under, far from, etc.*) between two to three objects. In the most simple case, we measure a model’s ability to understand common

Table 7 | Breakdown by skill and sub-skill including examples of prompts.

Prompt Category	Subcategory	Examples
Spatial Relationships rel/spatial	Simple	A cat above a dog. The lemon is in the middle of the apples. A bus is behind a truck going down the highway.
	Composed	The cat is near the banana. The banana is below the horse. The horse is on the truck.
Action rel/action	1, 2, or 3 entities	A bear is running through a field. A basketball is passed to a team member.
	Reverse actions	A ladybug is riding on the back of a flying unicorn. A koala climbs a tree, an eagle stands on the branch and a penguin flies overhead.
Scale rel/scale	Single object	A small ship in a bottle. A giant couch in a field.
	Comparative	A garlic is next to an onion and a tomato on a cutting board. The onion is larger than the garlic. The tomato is smaller than the garlic.
	Same size	The table is the same size as the cake. The mouse is the same size as the dragon.
Counting att/count	Simple modifier	2 cats. Four lemurs.
	Additive	5 burgers and one bonsai. 1 baobab, 2 cats and 3 dogs.
	Quantifiers and negations	Some shirts and some pizzas. There are more shirts than pizzas. An image with fewer dogs than cats. An image with no flowers in the vase.
Shape att/shape	Basic shapes	A line. An octagon.
	Composed shapes	A star-shaped cookie. Strawberries arranged in a heart shape.
	Hierarchical shapes	A square made of smaller letters g. A smiley face made of strawberries.
Text Rendering render/text	Rendering	A creature with a clock shaped head with the words "flimflam, bishbash, gorp" written on it. The creature has a small body and two legs, and is pointing at the ground. There are two clocks on the ground, one showing the time of 7:24 and the other showing 3:30 p.m.
	Numerical	equation of "3+4 = 7" etched into a rock. "Lorem ipsum dolor sit amet, \$%&*(), 12345 + adipiscing elit" written on a chalkboard with a piece of chalk next to it.
	Font	"congratulations" written in fancy decorative cursive font on an antique 1920's typewriter. "happiness" written in decorative font with a happy face next to it.
Color att/color	Simple colors	A pink salad. A grey vase.
	Composed expressions	A pink unicorn, a white airplane, and a green potato. A yellow couch and a green cookie.
	Colors and abstract shapes	A red rectangle on a green background. A pink circle on a green background.
	Descriptive color terms	A pastel coloured train passing through the station. A rainbow-colored bicycle leaning against a wall.
	Stroop	Text saying "yellow" in blue letters. Text saying "black" in green letters.

Prompt Category	Subcategory	Examples
Surface Characteristics att/texture+material	Texture only	A fluffy floor in the bathroom. There is a silky fabric on a bumpy couch in the room.
	Material only	There is a metal lime in the bowl. A paper snake on the table.
	Combined	There is a soft floor made of wax and and a shiny silver table in the room. A glossy diamond road.
Style att/style	style	A brightly colored canal in Venice, by Canaletto A cartoon of a cat by Goya.
	Visual medium	A sketch of a drawing of a flower in a pot. A glass vase in the style of el greco and the impressionists.
Language Complexity lang/complexity	Negation	A pencil sharpener without any pencils in it. A belt buckle with no belt. A leaflet with no text on it.
	Long prompt	An outdoors top-down view of purple sidewalk chalk on a concrete sidewalk reading, "Fear/ of/ Chores". The left side of the concrete slab has green algae growth that fades to the right. Small bits of smashed acorns are scattered across the slab.
	True paraphrase	The giraffe feeds the cat. The bystanders watch the dog.
	False paraphrase	The snake observes the kangaroo. The horse looks at the deer.
Compositional Language lang/compositional	Vary number of entities & attributes	An orange metal train. A plastic couch, a cyan blueberry, and 2 plates. 3 wooden pencils and 1 plastic fly. A brown plastic bus, a yellow plastic vase, and a yellow wooden salad.
Named Entities	Landmarks ne/landmarks	Ulvetanna Peak, Queen Maud Land, Antarctica during sunrise. Ashikaga Flower Park with beautiful wisteria in shades of violet and white. Burj Al Arab Jumeirah hotel with fireworks in the night and glowing water around.
	Animal Characters ne/characters	Grumpy Cat is sitting on a couch with a catnip toy. A cartoon of Laika playing in the snow. Yoshi is riding a skateboard down a hill.
	Vehicles ne/vehicles	A BMW M3 is on the road. A Ferrari is driving through roads in an Italian landscape. A Opel Ampera is upside down.
	Products ne/product	A bottle of Irn-Bru is sitting on a shelf. A Gucci bag with a red and white striped pattern. A Samsung Galaxy S III is on a car seat.
	Artwork ne/artwork	Charles IV of Spain and His Family is being painted on an easel. A painting of The Milkmaid hangs on the wall of a living room. A painting of Bacchus and Ariadne hanging in a stone building.

relationships between two objects. The difficulty is increased by combining simpler entities and requiring the ability to reason about implicit relationships. We use an LLM as described in Sec. C.3 to create these prompts, subsample and manually verify prompts are reasonable.

C.4.2. Action

This skill examines whether the model can bind the right action to the right object, including unusual cases where we flip the subject and the object (*i.e. Reverse actions*). An example of a reverse setup is that we swap the entities in ‘A penguin is diving while a dolphin swims’ and create ‘A penguin is swimming while a dolphin is diving’. We vary the difficulty by increasing the number of entities. We use an LLM as described in Sec. C.3 to create these prompts, subsample and manually verify prompts are reasonable.

C.4.3. Scale

We measure whether the model can reason about scale cues referring to commonly used descriptors such as *small*, *big* or *massive*. To reduce ambiguity, we typically refer to two objects, so that they can be compared in size. We test the ability to implicitly reason about scales by having *Comparative* prompts that contain several statements about objects, their relations and sizes. We use an LLM as described in Sec. C.3 to create these prompts, subsample and manually verify prompts are reasonable.

C.4.4. Counting

The simplest sub-skill *Simple modifier* contains a number (digits such as “2”, “3” or numerals such as “two”, “three”) and an entity. When selecting a vocabulary of words, we aimed to include words that occur less frequently in ordinary language (for example, “lemur” and “seahorse” occur less frequently than “dog” and “cat”) [53]. We focus on numbers 1—10, or 1—5 in more complex cases. Complexity is introduced by combining simple prompts containing just one attribute into compositional prompts containing several attributes. For example, simple prompts “1 cat” and “2 dogs” are combined into a single prompt “1 cat and 2 dogs” in the sub-skill *Additive*. We also test approximate understanding of quantities based on linguistic concepts of *many* and *few* in the *Quantifiers and negations* sub-skill.

C.4.5. Shape

We test for basic and composed shapes, where composed shapes include objects arranged in a certain shape. *Hierarchical shapes* are of the following type: “The letter H made up of smaller letters S”. This kind of challenge is used to study spatial cognition and the trade-off between global and local attention in literature on higher-order cognition [13]. We use an LLM as described in Sec. C.3 to create these prompts, subsample and manually verify prompts are reasonable.

C.4.6. Text Rendering

We investigate a model’s ability to generate text (both semantically meaningful and meaningless), including text of different lengths. We further test for the ability to generate symbols and numbers, as well as different types of fonts. We use an LLM as described in Sec. C.3 to create these prompts, subsample and manually verify prompts are reasonable.

C.4.7. Colour

The simplest prompts in this skill include basic colours bound to objects. As before, we introduce complexity by combining several simpler prompts (either two or three objects bound with a colour attribute). To include diversity of possible colour attributes, we also test descriptive colour terms such as “pastel” or “rainbow-coloured”. Finally, the sub-skill *stroop* contains prompts of the type “Text saying “blue” in green letters” similar to the incongruent condition in the Stroop task [55] used to study interference between different cognitive processes.

C.4.8. Surface Characteristics

Surface characteristics include texture and material. We first test for each sub-skill individually, and then combined. Generally, some prompts in this skill can be difficult to visualise as they might include descriptions that are typically of tactile nature (“abrasive” or “soft”).

C.4.9. Style

We divide prompts into two sub-skills: one depicting a style of an artist, and another to capture different visual mediums (such as *photo*, *stained glass* or *ceramics*). We use an LLM as described in Sec. C.3 to create these prompts, subsample and manually verify prompts are reasonable.

C.4.10. Named Entity

This skill evaluates a model’s knowledge of the world through named entities, focusing on specific entity types such as *landmarks*, *artwork*, *animal characters*, *products*, and *vehicles*, which are free of personally identifiable information (PII).

For the landmark class, we choose landmarks from the Google Landmarks V2 dataset and ensure we cover different continents and choose landmarks with high popularity [60]. Given this set of landmarks, we use an LLM as described in Sec. C.3 to create these prompts, subsample and manually verify prompts are reasonable.

For the other classes, to curate diverse named entities, we first gather candidates from Wikidata using SPARQL queries. A simple query (e.g., `instance of (P31) is painting (Q3305213)`) might yield an excessively large number of candidates. Therefore, we impose conditions to narrow down the query responses. See our criteria below.

Artwork : Created before the 20th century; any media; any movement

Animal Characters : Anthropomorphic/fictional animals; real animals with names

Products : Electric devices; food/beverage; beauty/health

Vehicle : Automobiles; aircraft

Once we have a candidate set for each entity class, we focus on selecting reasonably popular entities that are widely recognised and appropriate to present to models. We assess popularity using the number of incoming links to and contributors on their English Wikipedia pages as proxies [18]. Finally, we manually curate the final set of named entities, selecting them based on their ranked popularity scores.

C.4.11. Language complexity

We evaluate models on prompts with “tricky” language structure / wording. For this skill, we include 4 sub-skills: *negation*, *long prompt*, and *true / false phrases*. We sampled 19 prompts for *negation* from LVIS [21], COCO stuff [6], and MIT Places [68]; 58 *true paraphrases* and 58 *false paraphrases* from BLA [9]; and finally crowdsourced 38 for *long prompt* with the help from English major raters. It should be noted that while we do not cover the entire spectrum of complex language (e.g. passives, coordination, complex relative clausals, etc.), the subcategories included cover the most prominent pain points of image generation models per our experimentation.

We also include a language complexity metric which can be run over all prompts. Here we treat language complexity from two perspectives – semantic and syntactic.

- *Semantic complexity*. The quantity of semantic elements included in a prompt.
- *Syntactic complexity*. The level of complexity of the syntactic structure of a prompt.

Concretely, we define *semantic complexity* as the number of entities extracted from a prompt. Taking the visual relevance of the task into account, we apply *Stanford Scene Graph Parser* [52] for entity extraction and count the number of unique entities as the proxy for semantic complexity. For *syntactic complexity*, we implement a modified variant of [44] to look for the deepest central branch in the dependency tree of a prompt (pseudo-code below) [§] to gauge the complexity of its syntactic structure.

```
1 def central_depth(node) -> Tuple[int, int]:
2     return (max(central_depth(child)[1]+1 in node.children if child.position < node.position),
3             max(central_depth(child)[0]+1 in node.children if child.position > node.position))
```

D. Gecko Metric: More details

D.1. LLM prompting for generating coverage

```
1 """
2 Given a image description, label the visually groundable words in the description.
3 Classify each word into a type (entity, activity, attribute, counting, color, material, spatial, location, shape, style, other).
4
5 Description:
6 Portrait of a gecko wearing a train conductor's hat and holding a flag that has a yin-yang symbol on it. Woodcut.
7 The visual-groundable words and their scores are labelled below:
8 {1}[Portrait, style] of {2}[a, count] {3}[gecko, entity] {4}[wearing, activity] {5}[a, count] {6}[train conductor's hat, entity]
   and {7}[holding, entity] {8}[a, count] {9}[flag, entity] that has {10}[a yin-yang symbol, entity] on it. {11}[Woodcut,
   material].
9
10 Description:
11 square blue apples on a tree with circular yellow leaves.
12 The visual-groundable words and their scores are labelled below:
13 {1}[square, shape] {2}[blue, color] {3}[apples, entities] {4}[on, spatial] {5}[a, count] {6}[tree, entity] with {7}[circular, shape]
   {8}[yellow, color] {9}[leaves, entity]
14
15 Description:
16 ...
17 """
```

Listing 3 | LLM sampling template for labelling visually groundable words/phrases.

D.2. LLM prompting for generating QAs

```
1 """
2 Given a image description, generate one or two multiple-choice questions that verifies if the image description is correct.
3 Classify each concept into a type (object, human, animal, food, activity, attribute, counting, color, material, spatial, location,
   shape, other), and then generate a question for each type.
4
5 Description:
6 A man posing for a selfie in a jacket and bow tie.
7 The visual-groundable words and their scores are labelled below:
8 A {1}[Man, human] {2}[posing, activity] for a {3}[selfie, object] in a {4}[jacket, object] and a {5}[bow tie, object].
```

[§]As an implementation note, we implemented [52, 44] with SpaCy 2 [26] as the workhorse parsing backend.

```
9  Generated questions and answers are below:
10 About {1}:
11 Q: is there a man in the image?
12 Choices: yes, no
13 A: yes
14 About {2}:
15 Q: is the man posing for the selfie?
16 Choices: yes, no
17 A: yes
18 About {3}:
19 Q: is the man taking a selfie?
20 Choices: yes, no
21 A: yes
22 About {4}:
23 Q: is the man wearing a jacket?
24 Choices: yes, no
25 A: yes
26 About {5}:
27 Q: is the man wearing a bow tie?
28 Choices: yes, no
29 A: yes
30
31
32 Description:
33 A horse and several cows feed on hay.
34 The visual-groundable words and their scores are labelled below:
35 A {1}[horse, animal] and {2}[several, count] {3}[cows, animal] {4}[feed, activity] on a {5}[hay, object].
36 Generated questions and answers are below:
37 About {1}:
38 Q: is there a horse?
39 Choices: yes, no
40 A: yes
41 About {2}:
42 Q: are there several cows?
43 Choices: yes, no
44 A: yes
45 About {3}:
46 Q: are there cows?
47 Choices: yes, no
48 A: yes
49 About {4}:
50 Q: are the horse and cows feeding on hay?
51 Choices: yes, no
52 A: yes
53 About {5}:
54 Q: is there hay?
55 Choices: yes, no
56 A: yes
57
58 Description:
59 ...
60 ""
```

Listing 4 | LLM sampling template for generating QAs.

E. Human annotation: more details and experiments

E.1. Annotation templates

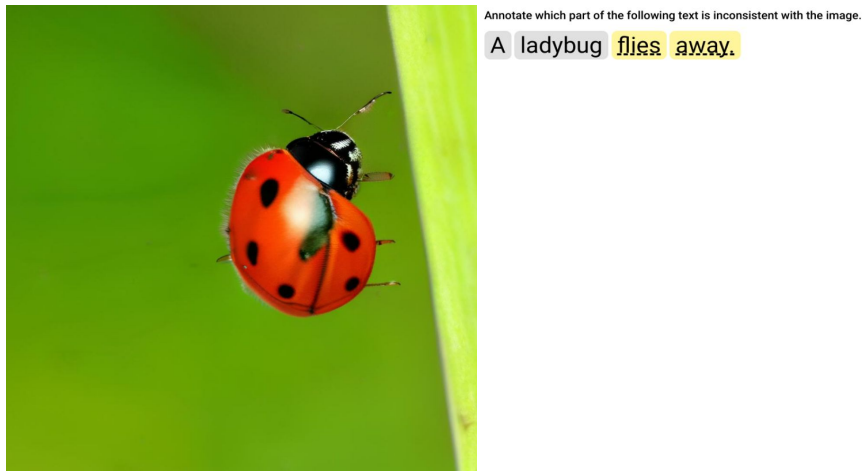


Figure 9 | **Word-level annotation template user interface.** Example depicting the interface shown to the annotators when performing evaluation tasks with the WL template. Raters are asked to click on words they find are not aligned with image, and double click on the words where they are unsure.

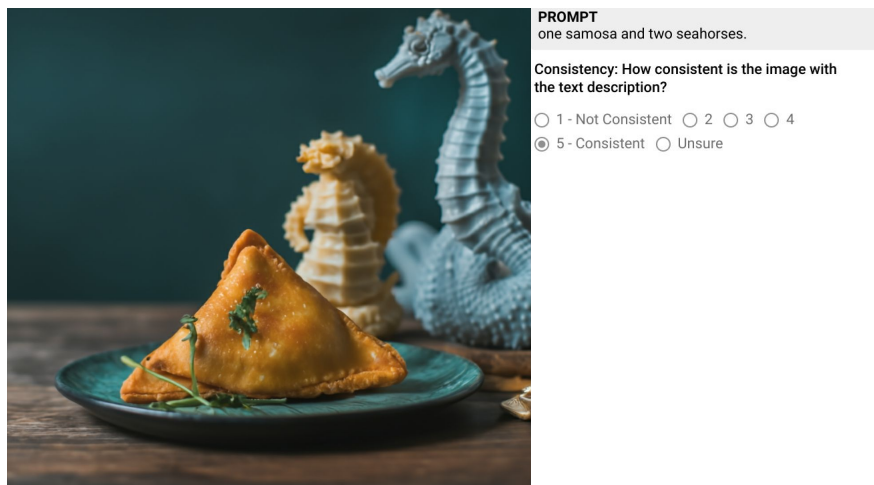


Figure 10 | **Likert annotation template user interface.** Example depicting the interface shown to the annotators when performing evaluation tasks with the Likert template. Raters are given the prompt and image and asked to rate on a 5-point scale how consistent the image is with respect to the prompt. An *Unsure* option is also given the annotators.

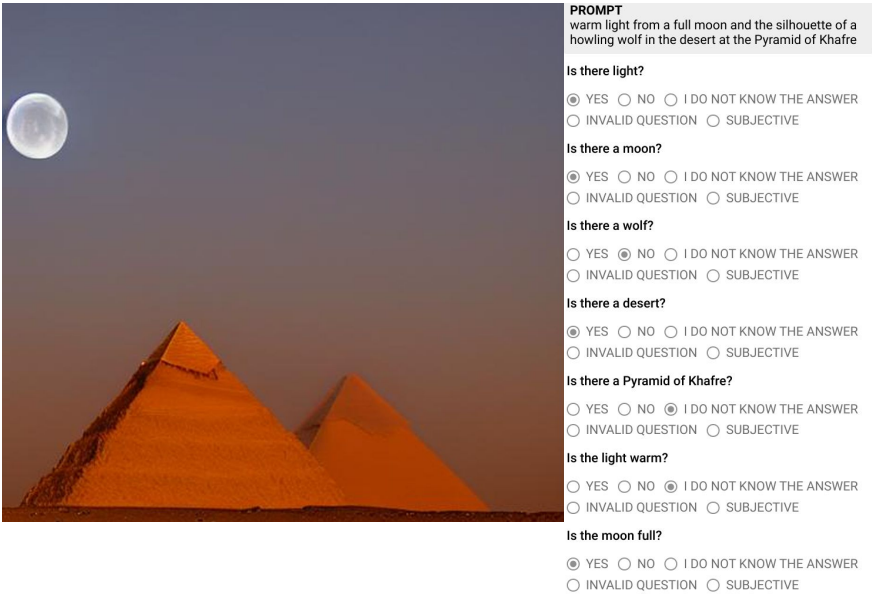


Figure 11 | **DSG(H) annotation template user interface.** Example depicting the interface shown to annotators when performing evaluation tasks with the DSG(H) template. Raters are given the image, prompt, and respective automatically generated questions. There are 5 options for answering each question. In our analysis, both *I do not know the answer* and *Subjective* answers are considered as *Unsure*.



Figure 12 | **Side-by-side comparison annotation template user interface.** Example depicting the interface shown to the annotators when performing evaluation tasks with the side-by-side template. Raters are given a pair of images from different models, the prompt used to generate them and asked to pick which one is more consistent with the prompt. An *Unsure* option is also given.

E.1.1. Percentage of unsure ratings for each annotation template/model

One of the innovations of our human evaluation setup is to allow for annotators to reflect uncertainty in their ratings. In Table 8 we show the percentage of *Unsure* ratings for each absolute comparison annotation template. Overall, we find that evaluations with Gecko(R) yield a higher percentage of *Unsure* ratings in comparison to Gecko(S).

Gen. model	WL		Likert		DSG(H)	
	Gecko(R)	Gecko(S)	Gecko(R)	Gecko(S)	Gecko(R)	Gecko(S)
Imagen	43.52	18.46	20.25	2.07	30.90	29.55
Muse	41.09	23.91	18.10	2.42	33.47	32.65
SDXL	20.09	20.94	4.24	2.05	31.77	29.64
SD1.5	13.35	26.48	10.04	4.09	37.02	33.08

Table 8 | **Percentage of *Unsure* ratings.** Overall, evaluation with Gecko(S) yields fewer *Unsure* ratings across all models and templates.

E.2. Additional experimental results

E.2.1. Correlation across templates and models.

We show the correlation between templates and models for both Gecko(R) and Gecko(S) in Table 9.

Gen. models	Gecko(R)						Gecko(S)					
	Likert vs WL		Likert vs DSG(H)		WL vs DSG(H)		Likert vs DSG(H)		Likert vs WL		WL vs DSG(H)	
	Pearson	Spearman	Pearson	Spearman	Pearson	Spearman	Pearson	Spearman	Pearson	Spearman	Pearson	Spearman
SD1.5	0.56	0.64	0.56	0.60	0.57	0.65	0.60	0.61	0.60	0.62	0.74	0.76
SDXL	0.60	0.63	0.52	0.57	0.56	0.61	0.60	0.50	0.56	0.62	0.78	0.79
Muse	0.67	0.62	0.61	0.63	0.61	0.66	0.51	0.52	0.51	0.53	0.77	0.75
Imagen	0.65	0.67	0.59	0.62	0.68	0.71	0.63	0.57	0.59	0.62	0.81	0.80

Table 9 | **Correlation between all absolute comparison templates.** We compute Pearson and Spearman correlation coefficients for all pairs of templates for both Gecko(R) and Gecko(S). We find significant results with $p < 0.001$ for all cases and that scores of all metrics are at least moderately correlated, with the finer-grained templates, WL and DSG(H), being more correlated with each other in comparison to Likert.

E.2.2. Distribution of scores per prompt-image pairs across annotation templates.

We plot the distribution of scores per each evaluated prompt-image pair for all the absolute comparison templates. The violin plots in Fig. 13-14 show the distributions for Gecko(R) and Gecko(S), respectively. It is possible to notice that scores obtained for Muse with WL and DSG(H) are more concentrated in values closer to 1 for both templates, corroborating findings from Sec. 4.3 where results showed Muse was the overall best model.

E.2.3. Side-by-side template.

In Table 10 we show inter-annotator agreement results for the side-by-side annotation template for Gecko(R) and Gecko(S), along with the respective difference in agreement when using only the reliable prompts for both Gecko2K subsets. In Table 11 we present the results of the comparison between the side-by-side template and the absolute comparison ones.



Figure 13 | **Distribution of scores for Gecko(R).** We show violin plots for scores obtained with all absolute comparison templates.

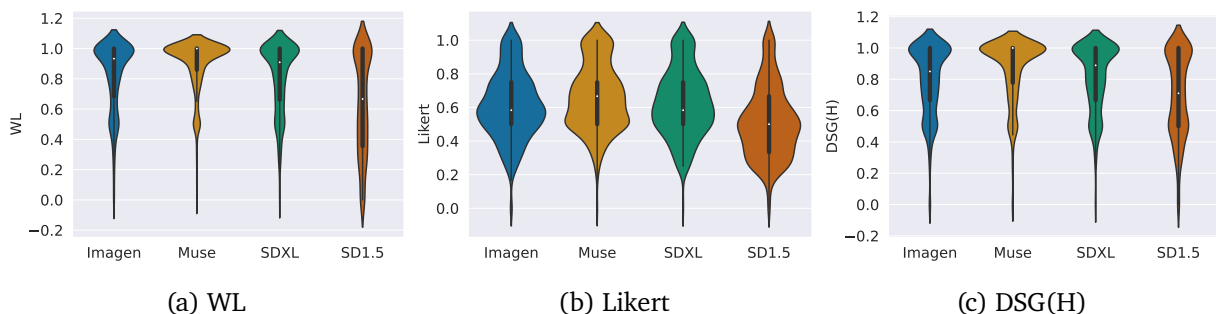


Figure 14 | **Distribution of scores for Gecko(S).** We show violin plots for scores obtained with all absolute comparison templates.

E.3. Reliable prompts: Examples of image-prompt pairs with high human (dis)agreement

In this section we show a representative list of prompts and corresponding images where human annotators were most likely to either agree or disagree in their ratings. The annotators agreed in ratings if they gave similar scores across for an image-text pair, meaning that the resulting mean variance was zero or close to zero. We refer to such prompts as “high agreement” prompts. In contrast, if annotators gave different ratings for a text-image pair, this would result in higher mean variance and we call such prompts “high disagreement” prompts.

To find prompts with high agreement across raters for all templates and all models, for each model-template combination we pick a subset of responses with low variance. Low variance is defined as the mean variance of a prompt-image pair for a model-template pair being below a certain threshold. The threshold is set as 10% of the maximum variance for that model-template set of ratings for both Gecko(R) and Gecko(S). Analogously, we also find a set of prompts with high disagreement; for this we find prompts that have mean variance above 1% of the maximum variance for a given template and for prompts from Gecko(R) and Gecko(S). The specific threshold value here is relevant only insofar as it captures at least 10 prompt-image pairs which we are interested in visualising. Then, we find prompts with high agreement by intersecting all model-template prompt sets where prompts have been selected based on the threshold. The procedure is analogous for low agreement prompts. For Gecko(R), both sets, namely the set of prompts with high agreement as well as the set of prompts with high disagreement have 34 prompts each. For Gecko(S), the set of prompts with high agreement has 62 prompts, while the set of prompts with high disagreement contains 85 prompts. A subset of 10 prompts for all different combinations is listed in Tables 12-15 and corresponding images are shown in the Figure 15-18.

	Gecko(R)	Gecko(R)-rel	Δ	Gecko(S)	Gecko(S)-rel	Δ
Imagen vs Muse	0.438	0.465	0.027	0.485	0.625	0.140
Imagen vs SDXL	0.440	0.425	-0.015	0.521	0.608	0.087
Imagen vs SD1.5	0.402	0.431	0.029	0.581	0.652	0.071
Muse vs SDXL	0.471	0.489	0.018	0.570	0.638	0.068
Muse vs SD1.5	0.539	0.592	0.053	0.600	0.617	0.017
SDXL vs SD1.5	0.389	0.438	0.049	0.522	0.562	0.040

Table 10 | **Side-by-side template: inter-annotator agreement.** We compute Krippendorff’s α for Gecko(R) and Gecko(S) and the difference (Δ) in α when using only reliable prompts for both subsets of Gecko2K. In both cases, using the reliable subsets increases the overall inter-annotator agreement.

	Gecko(R)-rel			Gecko(S)-rel		
	WL	Likert	DSG(H)	WL	Likert	DSG(H)
Imagen vs Muse	0.701	0.740	0.712	0.772	0.770	0.743
Imagen vs SDXL	0.583	0.606	0.678	0.797	0.804	0.785
Imagen vs SD1.5	0.506	0.444	0.628	0.838	0.821	0.789
Muse vs SDXL	0.654	0.699	0.676	0.753	0.793	0.729
Muse vs SD1.5	0.675	0.627	0.714	0.843	0.854	0.784
SDXL vs SD1.5	0.569	0.596	0.653	0.816	0.805	0.726

Table 11 | **Comparing side-by-side and absolute templates.** We compare the side-by-side template with the absolute comparison ones by computing the accuracy obtained by WL, Likert, and DSG(H) scores when using them to compare pairs of images. In this case, the ground-truth is assumed to be the results obtained with the side-by-side template.

Based on the analyses of such subsets, we observe several interesting trends. First, for Gecko(R) the prompts with higher agreement tend to be significantly shorter in length ($\mu = 54.32, \sigma = 32.18$) as measured by the number of characters, compared to the length of prompts with high disagreement ($\mu = 173.35, \sigma = 86.35$, Welch’s t-test $t(41.99) = -7.42$ ($p < 0.001$)). The same observation holds for Gecko(S), where high agreement prompts were also significantly shorter ($\mu = 20.77, \sigma = 18.03$), than high disagreement prompts ($\mu = 82.48, \sigma = 95.17$, Welch’s t-test $t(92.20) = -5.80$ ($p < 0.001$)). We further observe that prompts where raters tend to agree more are highly specific (*i.e.* they refer to one or just a few objects with few attributes), whereas prompts with high disagreement tend to describe more complex scenes with visual descriptors and often mentioning named entities or text rendering. Intuitively, this makes sense as longer prompts are more likely to require several skills.

High Agreement Prompts (Gecko(R))	
1	three men riding horses through a grassy field
2	a small bathroom with a shower and a toilet
3	a wooden table with four wooden chairs in front of two windows
4	a large colgate clock is by the water
5	a slice of chocolate cake is on a small plate
6	a black cat sitting in a field of grass
7	The Statue of Liberty made of gold
8	a man riding skis down a snow covered slope
9	a vast, grassy field with animals in the distance
10	a plastic bento box filled with rice, vegetables and fresh fruit

Table 12 | **Examples of Gecko(R) reliable prompts.** We show examples of Gecko(R) reliable prompts computed with a threshold equal to 10% of maximum disagreement.

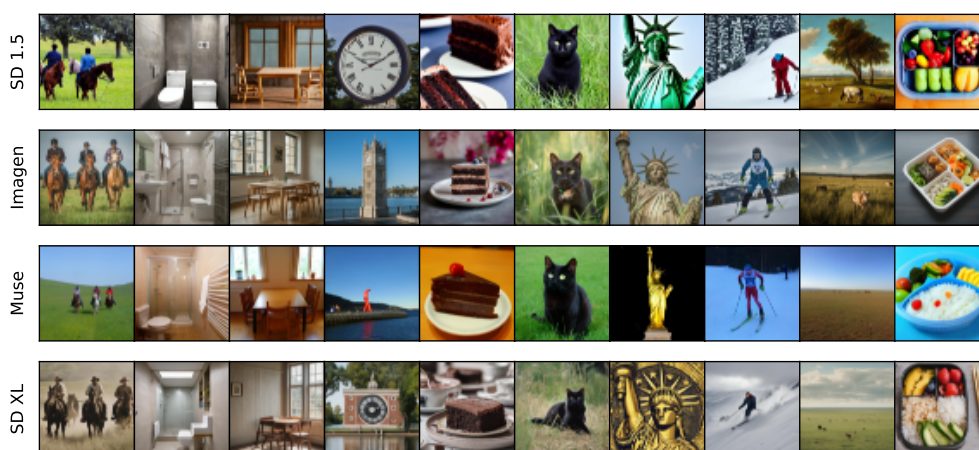


Figure 15 | **Images generated with Gecko(R) reliable prompts.** We show examples of images obtained with Gecko(S) reliable prompts computed with a threshold equal to 10% of maximum disagreement. Corresponding prompts are shown in Table 12.

High Agreement Prompts (Gecko(S))	
1	A star-shaped cookie
2	a pastel coloured train passing through the station.
3	a green boat.
4	the cat wears a gray shirt and holds a frisbee
5	a red motorcycle.
6	a black fish.
7	a pink bottle.
8	two mushrooms.
9	five cats.
10	a dog named Balto is running on a beach.

Table 13 | **Examples of Gecko(S) reliable prompts.** We show examples of Gecko(S) reliable prompts computed with a threshold equal to 10% of maximum disagreement.

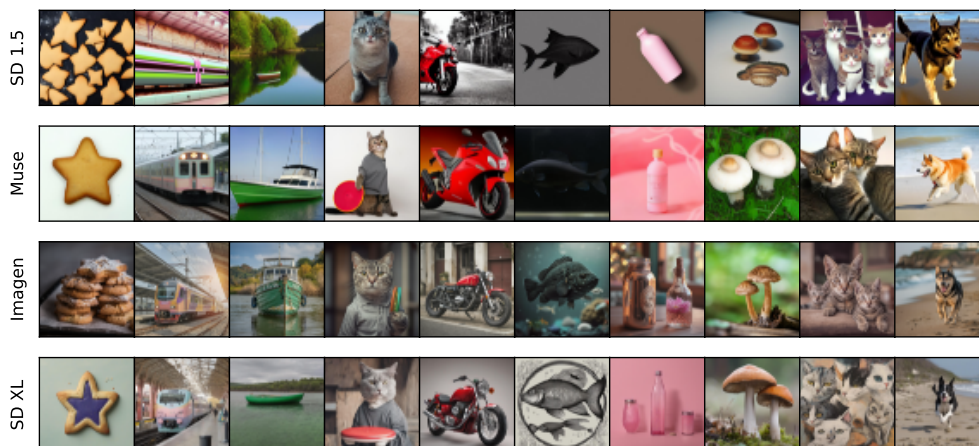


Figure 16 | **Images generated with Gecko(S) reliable prompts.** We show examples of images obtained with Gecko(S) reliable prompts computed with a threshold equal to 10% of maximum disagreement. Corresponding prompts are shown in Table 13.

High Disagreement Prompts (Gecko(R))	
1	Studio shot of sculpture of text 'cheese' made from cheese, with cheese frame.
2	vintage light monochrome six round and oval label set Illustration
3	There is a person snow boarding down a hill. There are tracks in the snow all around the snowboarder. There is a large rock in the snow next to them. There is a green pine tree in front of the snowboarder. The snowboarder is wearing blue ski pants and a blue and yellow jacket. They have a yellow snowboard on their feet.
4	pillow in the shape of words 'ready for the weekend', letterism, funny jumbled letters, [closeup]!!, breads, author unknown, flat art, swedish, diaper-shaped, 2000, white clay, surreal object photography
5	a sunflower field with a tractor about to run over a sunflower, with the caption 'after the sunflowers they will come for you'
6	a photo of a prison cell with a window and a view of the ocean, and the word 'freedom' painted on the glass
7	vehicle flying through a cyberpunk city 4 k, hyper detailed photograph
8	a scene with a city in the background, and a single cloud in the foreground, with the text 'contemplate the clouds' in rounded cursive
9	A pencil made of a tree branch with leaves
10	A wooden table that has a silver trophy in the middle of it. In front of the trophy are several bowls and dishes containing food. There is a loaf of bread on a block of wood at the front of the table.

Table 14 | **Examples of Gecko(R) non-reliable prompts.** We show examples of Gecko(R) non-reliable prompts computed with a threshold equal to 99% of maximum disagreement.

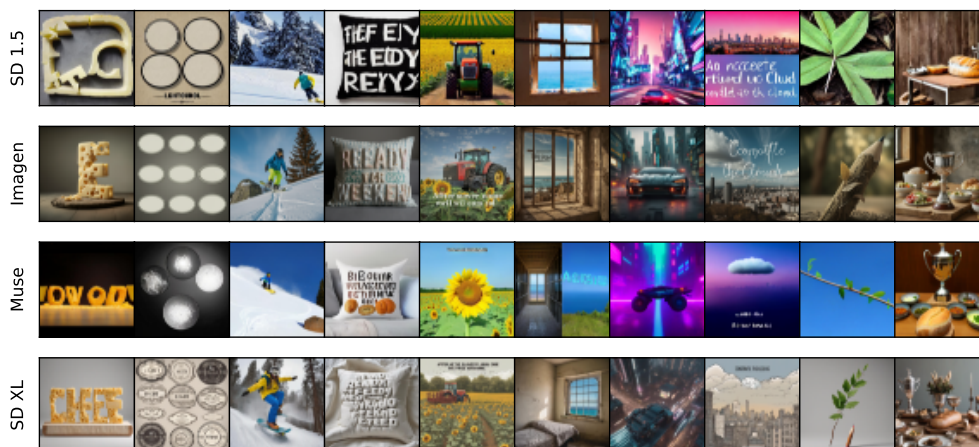


Figure 17 | **Images generated with Gecko(R) non-reliable prompts.** We show examples of images obtained with Gecko(R) non-reliable prompts computed with a threshold equal to 99% of maximum disagreement. Corresponding prompts are shown in Table 14.

High Disagreement Prompts (Gecko(S))	
1	the amazing view from the Halley Research Station in Antarctica on a clear night, the full Moon is rising and the sky is ablaze with the aurora australis, or polar lights.
2	a futuristic sculpture made of smooth metal
3	time lapse of sunrise over at the Hoover Dam
4	A huge vase in the middle of a field towering over the lawn chairs.
5	a bottle of Irn-Bru is sitting on a shelf.
6	a lord howe island palm tree with a moon rising in the distance
7	a long exposure image of the golden dunes at Playa del Ingles on the Canary Island, with a lone tourist
8	The soup is behind the cheese platter, to the left of the wine glasses, and below the crackers.
9	An alpaca and Chewbacca pose for a selfie at Machu Picchu
10	the lion cub named Simba is catching a ball.

Table 15 | **Examples of Gecko(S) non-reliable prompts.** We show examples of Gecko(R) non-reliable prompts computed with a threshold equal to 99% of maximum disagreement.

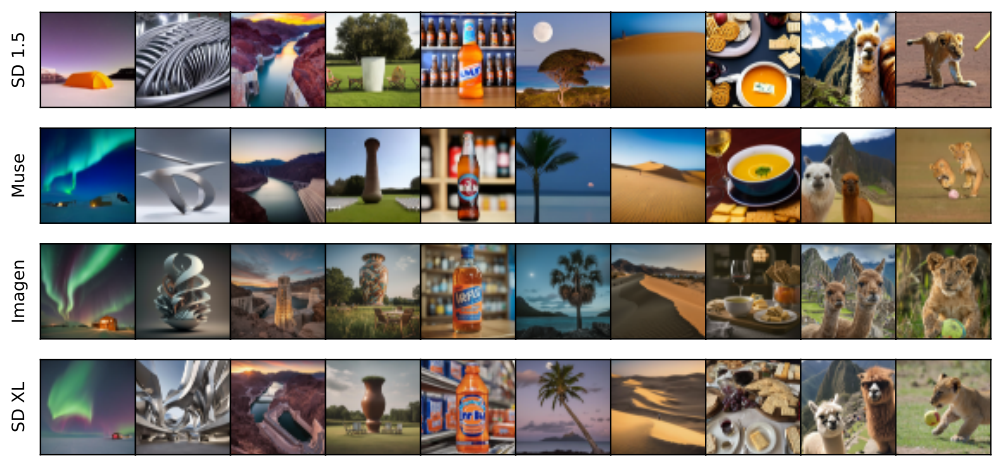


Figure 18 | **Images generated with Gecko(S) non-reliable prompts.** We show examples of images obtained with Gecko(S) non-reliable prompts computed with a threshold equal to 99% of maximum disagreement. Corresponding prompts are shown in Table 15.

E.4. Human evaluation templates: Challenging cases

In Figs. 19, 20, and 21 we show examples of challenging cases for the absolute comparison templates.

Image	Ratings
	A Nexus One is placed on a bench. A Nexus One is placed on a bench. A Nexus One is placed on a bench.
	Some shirts and some pizzas. There are more shirts than pizzas. Some shirts and some pizzas. There are more shirts than pizzas. Some shirts and some pizzas. There are more shirts than pizzas.

Figure 19 | **Examples of challenges for WL.** We show two examples of evaluated images, respective prompts and annotations from three raters. Each word is coloured according to the score given by the rater: **green** indicates *Aligned*, **red** *Not aligned*, and **yellow** *Unsure*. Both examples show that WL can be sensitive to words that are not relevant to the alignment evaluation. **Top:** All raters seem to agree it is not possible to tell whether a bench is represented in the image (hence the word is evaluated as *Unsure*). In spite of that, one of the raters disagrees on how to rate the “on a” preposition. **Bottom:** All raters seem to agree the quantity of shirts in the image does not reflect the prompt, but their ratings vary in terms of which words are rated as not aligned.

Image	Prompt	Rater 1	Rater 2	Rater 3
	A giraffe stands in the field.	4-Mostly consistent	5-Consistent	5-Consistent
	The raccoon holds the cat.	2-Mostly inconsistent	3-Somewhat consistent	4-Mostly consistent

Figure 20 | **Examples of challenges for Likert.** **Top:** Raters might take into account other aspects of the images besides alignment when evaluating a prompt-image pair. In this example, although the image is perfectly consistent with the prompt, one of the raters penalised its score. We hypothesise they took into account the fact the generated image is in grey scale. **Bottom:** “Uncalibrated” scores across raters. The scores of all three raters reflect the imperfect consistency between prompt and image, but each rater penalised the score with different intensity.

Image	Prompt	Questions
	A church without a steeple.	Is there a church? Does the church have a steeple? Is the steeple missing?
	A wood carving of an owl.	Is there an owl? Is there a wood carving? Is the wood carving made of wood?

Figure 21 | **Examples of challenges for DSG(H).** **Top:** Language complexity–Negation. As also shown in Fig. 6, the question generation is confused by the negation (asking if the church *has* a steeple as opposed to *does not have* a steeple). **Bottom:** Coverage. The question generation fails to capture that the owl should be represented as a wood carving.

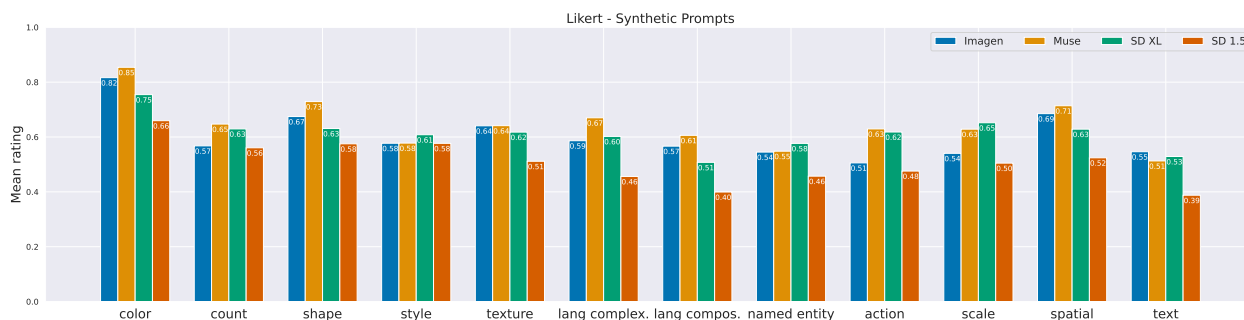


Figure 22 | **Per skill results - Likert.** Muse scores the best in nine out of the twelve categories, and SD1.5 performs the worst in all categories. Focusing on Muse, SDXL, and Imagen, the models score above 0.5 on all categories. Recalling that the Likert scale is symmetric (0.0 being inconsistent, and 1.0 being consistent), we see that these three models are more consistent than inconsistent on average (albeit only slightly for skills such as ‘lang compos.’ – , ‘named entity’ and ‘text’).

F. T2I Models: Additional comparisons

F.1. Analysing Model Ratings Per Skill

In Figures 22, 23 and 24, we plot the mean ratings in different skills for Likert, WL and DSG(H), respectively. We focus on Gecko(S) because we have a skill/sub-skill label for each prompt. Our goal is to understand how the trends in model performance on the whole prompt set relate to their performance in individual skills. We provide an overview of the results in the captions for the plots. Overall, the results broken down by skill are consistent with the averages over the whole prompt set. In other words, if a model is better or worse on the full prompt set, this is generally true for the individual categories as well. Another observation is that `COUNTING` and `COMPLEX LANGUAGE` seem to be the most difficult skills judging by WL and DSG, but this is not as clear from Likert (where many categories seem just as difficult).

Further Breaking Down Skills. We can gain more insight into the skills of the models by looking at variation within a skill. Figure 25 shows sub-skills of the `COLOUR` prompts. Two sub-skills are more challenging: the prompts require the models to combine multiple skills when generating the image (colour plus either composition or text rendering). The ‘colour:composed’ sub-skill (`COMPOSED EXPRESSIONS`) includes prompts such as ‘A brown vase, a white plate, and a red fork.’ with variations in the colors/objects. The sub-skill ‘color:strop’ (`STROOP`) contains prompts like ‘Text saying “green” in white letters.’ where the word in quotes differs from the color of the letters.

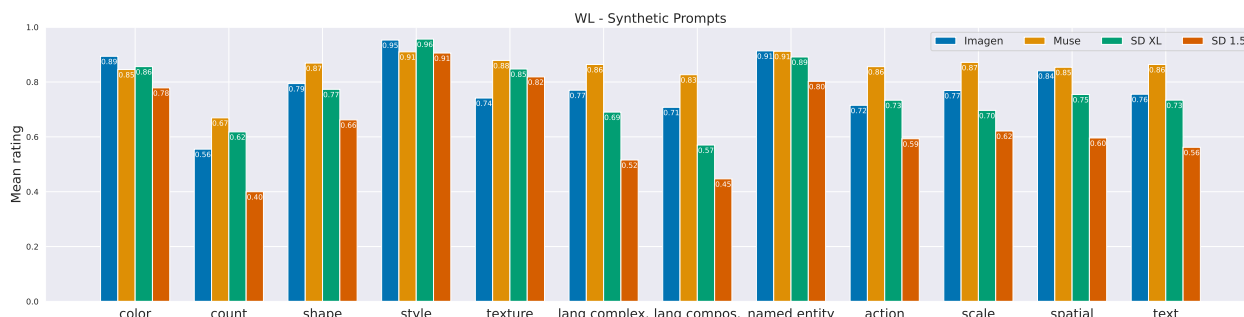


Figure 23 | **Per skill results - WL.** Muse scores the best in ten out of the twelve skills, and SD1.5 performs the worst in all skills. Moreover, Muse scores higher than the other models by a noticeable margin (≥ 0.1) for the skills ‘lang compos.’, ‘action’, ‘scale’ and ‘text’. In this case, analysing the results by skill shows that we can contribute Muse’s higher average score (over the whole prompt set) mostly to these skills.

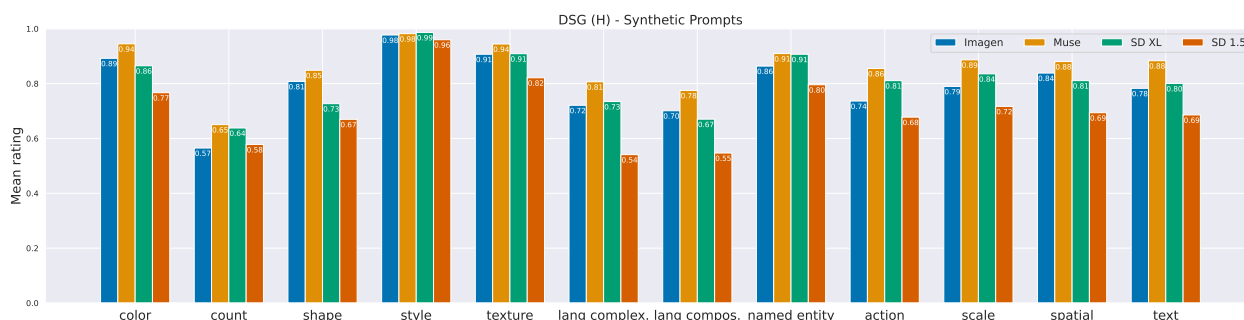


Figure 24 | **Per skill results - DSG(H).** Muse performs well across the skills, being the best eleven out of twelve times (scoring very close to the top for ‘style’). On the other hand, SD1.5 scores the worst in all the skills. This is consistent with the average scores on the overall prompt set. We see that counting is the most difficult skill for Muse, SDXL, and Imagen. Aside from counting, the hardest skills for Muse are the language ones (‘lang complex.’ and ‘lang compos.’). This relative skill deficiency is not evident from the Likert and WL ratings, and therefore, the DSG ratings are better able to capture model shortcomings for prompts with more complex linguistic structure.

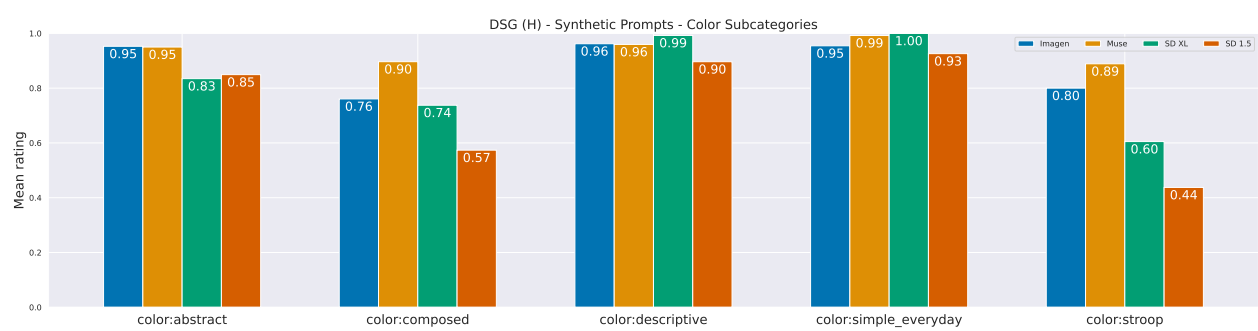


Figure 25 | **Color sub-skill results - DSG (H)**. We further break down the prompts in the color skill into five sub-skills. One observations is that two of the sub-skills (‘composed’ and ‘stroop’) are noticeably more difficult for all the models. This can be explained by the fact that they combine multiple skills: ‘composed’ includes prompts with multiple colors/objects, and ‘stroop’ includes text rendering in a certain color. Hence, while models may perform well on a skill overall, the sub-skills can illuminate where they struggle in generating images aligned with more complex prompts. On the other hand, models perform well with both abstract and everyday colors, likely because these are more commonly seen during training.

F.2. Model comparisons with TIFA160

We augment the experiment from Fig. 4 in Sec. 4.3 by performing a similar analysis with TIFA160. We generate images with this set of prompts and perform human evaluation following the same protocol as for Gecko2K and its subsets. In Fig. 26, we show the results of pairwise model comparisons with TIFA160 carried out with the Wilcoxon signed-rank test with $p < 0.001$. Results show that all three versions of Gecko prompts are able to better distinguish models by finding more significant comparisons between them.

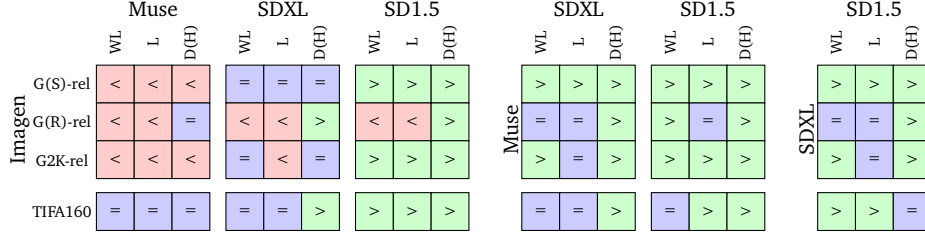


Figure 26 | **Comparing models using human annotations.** We compare model rankings on the reliable subsets of Gecko(S) (G(S)-rel), Gecko(R) (G(R)-rel), both subsets (G2K-rel), and TIFA160. We perform the Wilcoxon signed-rank test for all pairs of models ($p < 0.001$) and post-hoc comparison based on average ratings. Each grid represents a comparison between two models. Entries in the grid depict results for WL, Likert (L), and DSG(H) (D(H)) scores. The $>$ sign indicates the left-side model is better, worse ($<$), or not significantly different ($=$) than the model on the top. All three versions of Gecko prompts are able to better distinguish models by finding more significant comparisons between them.

G. Auto-eval metrics: Additional experiments

G.1. Pearson correlation Gecko2K.

Metrics	FT	Gecko(R)				Gecko(S)				Gecko(R)-Rel				Gecko(S)-Rel			
		WL	Likert	DSG(H)	SxS	WL	Likert	DSG(H)	SxS	WL	Likert	DSG(H)	SxS	WL	Likert	DSG(H)	SxS
		Pearson			Acc	Pearson			Acc	Pearson			Acc	Pearson			Acc
CLIP	✗	0.15	0.18	0.16	54.4	0.26	0.19	0.25	67.2	0.13	0.12	0.10	59.7	0.26	0.19	0.25	67.1
PyramidCLIP	✗	0.29	0.30	0.26	64.3	0.28	0.27	0.25	70.7	0.31	0.28	0.25	65.8	0.28	0.28	0.26	71.0
TIFA	✗	0.21	0.32	0.25	41.7	0.39	0.32	0.39	53.2	0.27	0.35	0.27	47.3	0.43	0.35	0.41	52.8
DSG	✗	0.28	0.41	0.38	49.6	0.43	0.42	0.44	58.1	0.39	0.43	0.44	53.9	0.45	0.43	0.46	56.9
Gecko	✗	0.38	0.51	0.42	62.1	0.46	0.48	0.46	74.6	0.50	0.55	0.51	71.3	0.52	0.52	0.50	75.2
VNLI	✓	0.34	0.48	0.39	54.4	0.41	0.55	0.42	72.7	0.25	0.41	0.22	65.6	0.35	0.49	0.34	72.7

Table 16 | **Pearson correlation between VQA-based, contrastive, and fine-tuned (FT) auto-eval metrics and human ratings across annotation templates on Gecko2K and Gecko2K-Rel.** Similar to the comparisons on Spearman correlation, Gecko again outperforms other auto-eval metrics with higher overall Pearson correlation.

We report the Spearman Rank correlation of different auto-eval metrics on Gecko2K in 4, here we report the Pearson correlation in 16 as well.

G.2. Model-ordering Evaluation

We provide detailed results for the model ordering experiment in Sec. 6.3. The goal is to examine if each auto-eval metric can predict how two models are ordered according to the human ratings. We use the results obtained by comparing T2I models with scores from all absolute comparison templates reported in Fig. 4 and consider the majority relation assigned to a model pair as the ground truth: if two of the templates find model 1 is better than model 2, then we assume model 1 > model 2. We notice that in all comparisons at least two templates agree in their results, always making it possible to define a ground-truth. We then count the number of times an auto-eval metric successfully found the same relation, in terms of both significance and direction. In addition, we count the number of times an auto-eval metric captured the ground-truth relationship following the most common practice in the literature: ordering models by solely comparing the mean values of scores regardless of statistical significance.

In Figs. 27 and 28, we show the results for all auto-eval metrics on both Gecko(R) and Gecko(S), respectively. Results highlight that the number of significant successful comparisons is lower than comparisons with mean values. All metrics except for DSG and CLIP get the same number of successful comparisons for both criteria. Note that Gecko not only is within the metrics that best capture pairwise model orderings as per the human evaluation ground-truth values, but Gecko scores are also the best correlated with the human scores themselves (as we can see in Table 4).

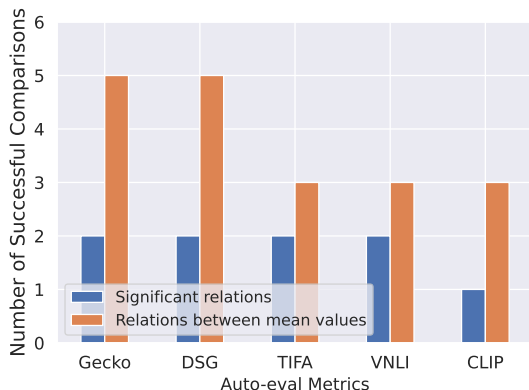


Figure 27 | **Gecko(R)-rel: Comparing correct relations between models captured by auto-eval metrics.** The ground-truth to computed the number of successful relations is given by human annotation scores.

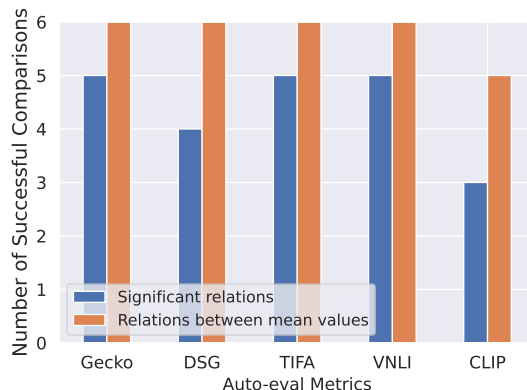


Figure 28 | **Gecko(S)-rel: Comparing correct relations between models captured by auto-eval metrics.** The ground-truth to computed the number of successful relations is given by human annotation scores.

G.3. Results for Additional Metrics on the Gecko Benchmark

We compare the Gecko metric with several score-based auto-eval metrics [36, 31, 62, 56, 67, 16, 66], as well as QA-based metrics such as TIFA [29] and DSG [10], and NLI models [61] in Tab. 17.

Generally, our Gecko metric outperforms the others and shows a higher correlation with most of the human annotation templates. DSG is the second best metric, except on SxS where it ranks third. It outperforms TIFA by a clear margin but falls behind Gecko. Finally, we note that Gecko even shows higher correlation than the supervised VNLI model.

Looking at efficient, score-based metrics, we find that PyramidCLIP achieves competitive correlations. Moreover, a larger pre-training corpus leads to better metrics; e.g., as seen by comparing CLIP-B/32 (trained on 400M images) and CLIP-B/32_{LAION-2B} (trained on 2B images). Finally, larger

Metrics	Gecko (R)						Gecko (S)					
	WL		Likert		DSG(H)		WL		Likert		DSG(H)	
	Pearson	Spearman-R	Pearson	Spearman-R	Pearson	Spearman-R	Pearson	Spearman-R	Pearson	Spearman-R	Pearson	Spearman-R
BLIP-2 _{ITM}	0.25	0.22	0.24	0.19	0.23	0.21	0.28	0.23	0.13	0.16	0.25	0.23
CLIP-B/32	0.15	0.14	0.18	0.16	0.16	0.13	0.26	0.25	0.19	0.18	0.25	0.26
CLIP-B/32 _{LAION-2B}	0.26	0.21	0.24	0.22	0.26	0.21	0.28	0.24	0.23	0.22	0.26	0.24
CLIP-B/16	0.16	0.11	0.15	0.12	0.17	0.12	0.26	0.22	0.15	0.14	0.24	0.22
CLIP-L/14	0.18	0.16	0.16	0.14	0.18	0.15	0.26	0.23	0.17	0.16	0.25	0.24
CLIP-H/14 _{LAION-2B}	0.29	0.24	0.25	0.22	0.27	0.23	0.29	0.24	0.23	0.22	0.27	0.25
CLIP-g/14 _{LAION-2B}	<u>0.30</u>	0.24	0.26	0.23	0.28	0.23	<u>0.30</u>	0.25	0.24	0.23	0.28	0.25
CLIP-G/14 _{LAION-2B}	0.29	0.23	0.25	0.21	0.27	0.23	<u>0.30</u>	0.24	0.25	0.23	0.28	0.25
CoCa-L/14	0.28	0.23	0.26	0.22	0.26	0.21	0.29	0.25	0.24	0.22	0.28	0.26
EVA-02-CLIP-L/14	0.27	0.24	0.23	0.21	0.24	0.22	<u>0.30</u>	<u>0.26</u>	0.21	0.20	0.28	<u>0.27</u>
EVA-02-CLIP-E/14	0.28	0.23	0.24	0.20	0.27	0.22	0.28	0.23	0.23	0.22	0.26	0.24
EVA-02-CLIP-E/14+	<u>0.30</u>	0.24	0.24	0.21	0.27	0.23	0.29	0.24	0.25	0.23	0.28	0.25
SigLIP-B/16	0.26	0.21	0.22	0.18	0.27	0.21	0.29	0.25	0.22	0.21	<u>0.29</u>	0.26
SigLIP-L/16	0.28	0.24	0.26	0.22	<u>0.29</u>	0.25	0.29	<u>0.26</u>	0.23	0.22	<u>0.29</u>	<u>0.27</u>
PyramidCLIP-B/16	0.29	<u>0.26</u>	<u>0.30</u>	<u>0.27</u>	<u>0.29</u>	<u>0.26</u>	0.28	0.22	<u>0.27</u>	<u>0.25</u>	0.25	0.23
X-VLM _{16M}	0.17	0.11	0.21	0.09	0.25	0.15	0.26	0.23	0.23	0.16	0.24	0.23
TIFA	0.21	0.26	0.32	0.34	0.25	0.28	0.39	0.39	0.32	0.32	0.39	0.39
DSG	0.28	0.35	0.41	0.47	0.38	0.42	0.43	0.45	0.42	0.45	0.44	0.45
Gecko	0.38	0.41	0.51	0.55	0.42	0.46	0.46	0.47	0.48	0.52	0.46	0.45
VNLI	0.34	0.37	0.48	0.49	0.39	0.42	0.41	0.45	0.55	0.55	0.42	0.45

Table 17 | **Correlation between auto-eval metrics and human ratings across three annotation templates on Gecko2K.** Best results per model type are underlined; best results are in **bold**.

models are often better (e.g., SigLIP-L/16 vs. SigLIP-B/16), although these trends are less consistent (e.g., EVA-02-CLIP-L/14+ $\approx$ EVA-02-CLIP-E/14).

G.4. Analysing Auto-Eval Metric results per skill.

We present the per-skill Spearman Ranked Correlation between different auto-eval metrics and human annotation templates in 29. We observe a similar trend across the three plots, as we discussed in 6: Gecko is the best on handling prompts with “language complexity”, which can be attributed to the coverage tagging and filtering steps in its pipeline that make Gecko less prone to errors when processing long and complicated prompts. DSG is better on “compositional prompts”, as it can leverage its utilization of dependency graphs. VNLI demonstrates advantages in assessing “shape”, “color”, and “text” (e.g. `TEXT RENDERING`) prompts, as the model is trained to discriminate keywords within these categories. Meanwhile, it is also worth noting that all metrics exhibit relatively poor performance on “text” and “named identity”, highlighting the current lack of OCR and named recognition ability in existing contrastive, NLI and VQA models.

G.5. Additional visualisations.

We visualise additional examples in Fig. 30 for different categories (spatial, counting, text rendering, linguistic complexity, etc.). These examples demonstrate both differences arising from different annotation templates and also different metrics.

G.6. Results per Word for WL

Here we evaluate how well the Gecko metric can identify whether words are or are not grounded as rated by WL. This experiment can *only* be done for Gecko as *no other* metric gives word level

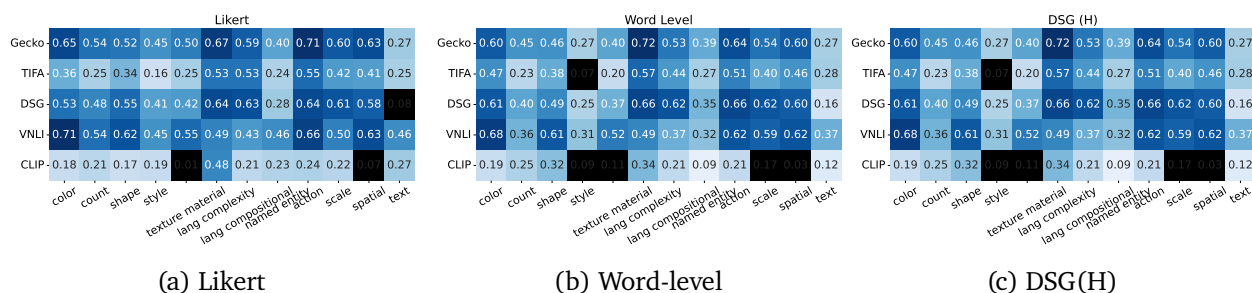


Figure 29 | **Per skill results by metric.** We visualise the correlation for each skill. Where p-values are > 0.05, we color the square black.

annotations that can be traced back to the original words in the prompt. We note that this is much more challenging than giving an overall image rating. In order to perform this experiment, we first parse the coverage prediction to ensure we can match words in the original prompt with those in the coverage prediction. For example, if we have the original prompt ‘a red-colored dog’ and ‘a {1}[red-colored] {2}[dog]’ as the generated coverage one, we can map from the word index (e.g. ‘{1}’ to the phase ‘red-colored’).

We then take all word level predictions where all annotators either annotated the word as grounded or not grounded (we removed those for which a subset of annotators annotated the word as ‘unsure’). For these words, we take the ones where the coverage model predicts that it should be covered (e.g. in the example above, even if ‘a’ was always annotated as grounded, we would ignore it as it is not considered groundable by the coverage step). Given this final set of words, we look at whether the VQA prediction was accurate and compare this to whether the annotators thought that the word was grounded or not.

We report three numbers in Table 18: (1) the number of words we are left with, (2) the accuracy and (3) the error consistency κ [17], equation (1),(3). We report error consistency as many of the words ($\sim 90\%$) are rated as grounded. Accuracy does not account for the fact that a metric which predicts ‘grounded’ $\sim 90\%$ of the time would actually get $\sim 80\%$ accuracy by chance. Error consistency takes this into account such that $\kappa = -1$ means that two sets of results never agree, $\kappa = 0$ that the overlap is explained by chance and $\kappa = 1$ means results agree perfectly. As shown by the results, Gecko is able to predict grounding at the word level with reasonable accuracy. Moreover, results are not simply explained by chance (as $\kappa > 0.$); the error consistency results indicate in general some but not substantial agreement.

Gecko(R)				Gecko(S)		
	# Words evaluated	Accuracy	Error Consistency (κ)	# Words evaluated	Accuracy	Error Consistency (κ)
SD1.5	3727	75.5	0.20	2729	77.5	0.52
SDXL	3901	80.0	0.13	3304	79.5	0.34
Muse	2792	82.5	0.24	3298	82.0	0.19
Imagen	2549	76.5	0.29	3472	80.2	0.39

Table 18 | **Word level results comparing Gecko to the WL annotation template.** We can see that Gecko achieves high accuracy but also that results are not simply explained by chance, as $\kappa > 0$ and in general indicates fair agreement.

Prompt:		A snake is on the elephant.				a bird flying over the mountains in the sunset, with the text "bla bla bla, this is the sound of a helicopter"			
									
		Imagen	Muse	SDXL	SD1.5	Imagen	Muse	SDXL	SD1.5
APIC:		1.	1.0	0.44	0.33	0.53	0.83	0.47	0.21
Likert:		1.	0.80	0.6	0.4	0.6	0.6	0.6	0.4
DSG (H):		1.	0.50	1.	0.5	0.92	0.8	0.71	0.25
Gecko:		0.98	0.84	0.73	0.26	0.85	0.58	0.86	0.21
DSG:		1.0	1.0	1.0	0.33	0.67	0.50	0.89	0.11
VNLI:		0.94	0.93	0.14	0.23	0.19	0.30	0.34	0.27

Prompt:		There are 5 apples, two of them are yellow and two are black but none are red.				the dog who wears a white shirt holds a beer			
									
		Imagen	Muse	SDXL	SD1.5	Imagen	Muse	SDXL	SD1.5
APIC:		0.88	0.73	0.57	0.61	1.0	1.0	1.	0.4
Likert:		0.73	0.60	0.53	0.60	0.8	0.93	1.	0.53
DSG (H):		0.80	0.53	0.5	0.40	0.75	1.0	1.	0.25
Gecko:		0.88	0.91	0.84	0.69	0.98	0.97	0.97	0.7
DSG:		1.0	1.0	0.9	0.8	1.	0.83	1.	0.17
VNLI:		0.23	0.20	0.19	0.16	0.9	0.95	0.94	0.17

Prompt:		A fortune cookie that has the fortune "the best way to predict the future is to create it."				A brown glass salad bowl on a grey metal table.			
									
		Imagen	Muse	SDXL	SD1.5	Imagen	Muse	SDXL	SD1.5
APIC:		0.35	0.39	0.	0.02	1.0	0.88	1.0	0.63
Likert:		0.53	0.47	0.4	0.27	0.73	0.67	0.73	0.53
DSG (H):		1.0	0.5	0.67	0.17	.83	1.0	0.83	0.83
Gecko:		0.96	0.79	0.75	0.62	0.92	0.88	0.87	0.70
DSG:		0.88	0.25	1.0	0.25	0.93	0.93	0.86	0.79
VNLI:		0.37	0.27	0.25	0.35	0.81	0.65	0.55	0.31

Figure 30 | **Additional visualisations of scores from different auto-eval metrics.** We show the image generations by the four generative models given two prompts from Gecko(S), with the alignment scores predicted by human annotators and auto-eval metrics respectively.

References

- [1] Anil, R., Dai, A.M., Firat, O., Johnson, M., Lepikhin, D., Passos, A., Shakeri, S., Taropa, E., Bailey, P., Chen, Z., et al.: Palm 2 technical report. arXiv preprint arXiv:2305.10403 (2023)
- [2] Bang, Y., Cahyawijaya, S., Lee, N., Dai, W., Su, D., Wilie, B., Lovenia, H., Ji, Z., Yu, T., Chung, W., et al.: A multitask, multilingual, multimodal evaluation of chatgpt on reasoning, hallucination, and interactivity. arXiv preprint arXiv:2302.04023 (2023)
- [3] Betker, J., Goh, G., Jing, L., Brooks, T., Wang, J., Li, L., Ouyang, L., Zhuang, J., Lee, J., Guo, Y., et al.: Improving image generation with better captions. Computer Science. <https://cdn.openai.com/papers/dall-e-3.pdf> **2**(3), 8 (2023)
- [4] Bitton-Guetta, N., Bitton, Y., Hessel, J., Schmidt, L., Elovici, Y., Stanovsky, G., Schwartz, R.: Breaking common sense: Whoops! a vision-and-language benchmark of synthetic and compositional images. In: ICCV (2023)
- [5] Bugliarello, E., Sartran, L., Agrawal, A., Hendricks, L.A., Nematzadeh, A.: Measuring progress in fine-grained vision-and-language understanding. arXiv preprint arXiv:2305.07558 (2023)
- [6] Caesar, H., Uijlings, J., Ferrari, V.: Coco-stuff: Thing and stuff classes in context. In: Computer vision and pattern recognition (CVPR), 2018 IEEE conference on. IEEE (2018)
- [7] Chang, H., Zhang, H., Barber, J., Maschinot, A., Lezama, J., Jiang, L., Yang, M.H., Murphy, K.P., Freeman, W.T., Rubinstein, M., Li, Y., Krishnan, D.: Muse: Text-to-image generation via masked generative transformers. In: Krause, A., Brunskill, E., Cho, K., Engelhardt, B., Sabato, S., Scarlett, J. (eds.) Proceedings of the 40th International Conference on Machine Learning. Proceedings of Machine Learning Research, vol. 202, pp. 4055–4075. PMLR (23–29 Jul 2023), <https://proceedings.mlr.press/v202/chang23b.html>
- [8] Chen, X., Wang, X., Changpinyo, S., Piergiovanni, A., Padlewski, P., Salz, D., Goodman, S., Grycner, A., Mustafa, B., Beyer, L., Kolesnikov, A., Puigcerver, J., Ding, N., Rong, K., Akbari, H., Mishra, G., Xue, L., Thapliyal, A., Bradbury, J., Kuo, W., Seyedhosseini, M., Jia, C., Ayan, B.K., Riquelme, C., Steiner, A., Angelova, A., Zhai, X., Houlsby, N., Soricut, R.: Pali: A jointly-scaled multilingual language-image model. arXiv preprint arXiv:2209.06794 (2022)
- [9] Chen, X., Fernández, R., Pezzelle, S.: The bla benchmark: Investigating basic language abilities of pre-trained multimodal models. In: Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing. pp. 5817–5830 (2023)
- [10] Cho, J., Hu, Y., Garg, R., Anderson, P., Krishna, R., Baldridge, J., Bansal, M., Pont-Tuset, J., Wang, S.: Davidsonian scene graph: Improving reliability in fine-grained evaluation for text-image generation. arXiv preprint arXiv:2310.18235 (2023)
- [11] Cho, J., Zala, A., Bansal, M.: Dall-eval: Probing the reasoning skills and social biases of text-to-image generation models. In: ICCV (2023)
- [12] Clark, E., August, T., Serrano, S., Haduong, N., Gururangan, S., Smith, N.A.: All that’s ‘human’ is not gold: Evaluating human evaluation of generated text. In: Zong, C., Xia, F., Li, W., Navigli, R. (eds.) Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference on Natural Language Processing (Volume 1: Long Papers) (2021)

- [13] Delis, D.C., Robertson, L.C., Efron, R.: Hemispheric specialization of memory for visual hierarchical stimuli. *Neuropsychologia* **24**(2), 205–214 (1986)
- [14] Delmas, G., Weinzaepfel, P., Lucas, T., Moreno-Noguer, F., Rogez, G.: Posescript: 3d human poses from natural language. In: ECCV (2022)
- [15] Dosovitskiy, A., Beyer, L., Kolesnikov, A., Weissenborn, D., Zhai, X., Unterthiner, T., Dehghani, M., Minderer, M., Heigold, G., Gelly, S., et al.: An image is worth 16x16 words: Transformers for image recognition at scale. arXiv preprint arXiv:2010.11929 (2020)
- [16] Gao, Y., Liu, J., Xu, Z., Zhang, J., Li, K., Ji, R., Shen, C.: Pyramidclip: Hierarchical feature alignment for vision-language model pretraining. In: NeurIPS (2022)
- [17] Geirhos, R., Meding, K., Wichmann, F.A.: Beyond accuracy: quantifying trial-by-trial behaviour of cnns and humans by measuring error consistency. *Advances in Neural Information Processing Systems* **33**, 13890–13902 (2020)
- [18] Geva, M., Khashabi, D., Segal, E., Khot, T., Roth, D., Berant, J.: Did Aristotle Use a Laptop? A Question Answering Benchmark with Implicit Reasoning Strategies. *TACL* **9**, 346–361 (2021)
- [19] Gokhale, T., Palangi, H., Nushi, B., Vineet, V., Horvitz, E., Kamar, E., Baral, C., Yang, Y.: Benchmarking spatial relationships in text-to-image generation. arXiv preprint arXiv:2212.10015 (2022)
- [20] Guerreiro, N.M., Alves, D.M., Waldendorf, J., Haddow, B., Birch, A., Colombo, P., Martins, A.F.: Hallucinations in large multilingual translation models. *TACL* **11**, 1500–1517 (2023)
- [21] Gupta, A., Dollar, P., Girshick, R.: LVIS: A dataset for large vocabulary instance segmentation. In: *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition* (2019)
- [22] Hayes, A.F., Krippendorff, K.: Answering the call for a standard reliability measure for coding data. *Communication methods and measures* **1**(1), 77–89 (2007)
- [23] Heo, C.Y., Kim, B., Park, K., Back, R.M.: A comparison of best-worst scaling and likert scale methods on peer-to-peer accommodation attributes. *Journal of business research* **148**, 368–377 (2022)
- [24] Hessel, J., Holtzman, A., Forbes, M., Le Bras, R., Choi, Y.: Clipscore: A reference-free evaluation metric for image captioning. In: EMNLP. pp. 7514–7528 (2021)
- [25] Heusel, M., Ramsauer, H., Unterthiner, T., Nessler, B., Hochreiter, S.: GANs trained by a two time-scale update rule converge to a local nash equilibrium. *NeurIPS* **30** (2017)
- [26] Honnibal, M., Montani, I.: Spacy 2: Natural language understanding with bloom embeddings, convolutional neural networks and incremental parsing. *spacy.io* (2017)
- [27] Honovich, O., Aharoni, R., Herzig, J., Taitelbaum, H., Kukliansy, D., Cohen, V., Scialom, T., Szpektor, I., Hassidim, A., Matias, Y.: True: Re-evaluating factual consistency evaluation. arXiv preprint arXiv:2204.04991 (2022)
- [28] Honovich, O., Choshen, L., Aharoni, R., Neeman, E., Szpektor, I., Abend, O.: Q2: Evaluating factual consistency in knowledge-grounded dialogues via question generation and question answering. arXiv preprint arXiv:2104.08202 (2021)

- [29] Hu, Y., Liu, B., Kasai, J., Wang, Y., Ostendorf, M., Krishna, R., Smith, N.A.: Tifa: Accurate and interpretable text-to-image faithfulness evaluation with question answering. arXiv preprint arXiv:2303.11897 (2023)
- [30] Huang, K., Sun, K., Xie, E., Li, Z., Liu, X.: T2i-compbench: A comprehensive benchmark for open-world compositional text-to-image generation. NeurIPS **36** (2024)
- [31] Ilharco, G., Wortsman, M., Wightman, R., Gordon, C., Carlini, N., Taori, R., Dave, A., Shankar, V., Namkoong, H., Miller, J., Hajishirzi, H., Farhadi, A., Schmidt, L.: Openclip (Jul 2021). doi: 10.5281/zenodo.5143773, <https://doi.org/10.5281/zenodo.5143773>, if you use this software, please cite it as below.
- [32] Jayasumana, S., Ramalingam, S., Veit, A., Glasner, D., Chakrabarti, A., Kumar, S.: Rethinking fid: Towards a better evaluation metric for image generation. arXiv preprint arXiv:2401.09603 (2023)
- [33] Krause, J., Johnson, J., Krishna, R., Fei-Fei, L.: A hierarchical approach for generating descriptive image paragraphs. In: CVPR (2017)
- [34] Kryściński, W., McCann, B., Xiong, C., Socher, R.: Evaluating the factual consistency of abstractive text summarization. arXiv preprint arXiv:1910.12840 (2019)
- [35] Lee, T., Yasunaga, M., Meng, C., Mai, Y., Park, J.S., Gupta, A., Zhang, Y., Narayanan, D., Teufel, H., Bellagente, M., et al.: Holistic evaluation of text-to-image models. NeurIPS **36** (2024)
- [36] Li, J., Li, D., Savarese, S., Hoi, S.: Blip-2: bootstrapping language-image pre-training with frozen image encoders and large language models. In: Proceedings of the 40th International Conference on Machine Learning. ICML'23, JMLR.org (2023)
- [37] Liang, W., Zou, J., Yu, Z.: Beyond user self-reported likert scale ratings: A comparison model for automatic dialog evaluation. In: Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics. pp. 1363–1374 (2020)
- [38] Liang, Y., He, J., Li, G., Li, P., Klimovskiy, A., Carolan, N., Sun, J., Pont-Tuset, J., Young, S., Yang, F., et al.: Rich human feedback for text-to-image generation. arXiv preprint arXiv:2312.10240 (2023)
- [39] Likert, R.: A technique for the measurement of attitudes. Archives of psychology (1932)
- [40] Lin, T.Y., Maire, M., Belongie, S., Hays, J., Perona, P., Ramanan, D., Dollár, P., Zitnick, C.L.: Microsoft coco: Common objects in context. In: ECCV (2014)
- [41] Liu, R., Garrette, D., Saharia, C., Chan, W., Roberts, A., Narang, S., Blok, I., Mical, R., Norouzi, M., Constant, N.: Character-aware models improve visual text rendering. arXiv preprint arXiv:2212.10562 (2022)
- [42] Lu, C., Krishna, R., Bernstein, M., Fei-Fei, L.: Visual relationship detection with language priors. In: ECCV (2016)
- [43] Maynez, J., Narayan, S., Bohnet, B., McDonald, R.: On faithfulness and factuality in abstractive summarization. arXiv preprint arXiv:2005.00661 (2020)
- [44] Ohta, S., Fukui, N., Sakai, K.L.: Computational principles of syntax in the regions specialized for language: integrating theoretical linguistics and functional neuroimaging. Frontiers in Behavioral Neuroscience (2013)

- [45] Paiss, R., Ephrat, A., Tov, O., Zada, S., Mosseri, I., Irani, M., Dekel, T.: Teaching clip to count to ten. arXiv preprint arXiv:2302.12066 (2023)
- [46] Podell, D., English, Z., Lacey, K., Blattmann, A., Dockhorn, T., Müller, J., Penna, J., Rombach, R.: Sd-xl: Improving latent diffusion models for high-resolution image synthesis. arXiv preprint arXiv:2307.01952 (2023)
- [47] Podell, D., English, Z., Lacey, K., Blattmann, A., Dockhorn, T., Müller, J., Penna, J., Rombach, R.: SDXL: Improving latent diffusion models for high-resolution image synthesis. In: The Twelfth International Conference on Learning Representations (2024), <https://openreview.net/forum?id=di52zR8xgf>
- [48] Pont-Tuset, J., Uijlings, J., Changpinyo, S., Soricut, R., Ferrari, V.: Connecting vision and language with localized narratives. In: ECCV (2020)
- [49] Radford, A., Kim, J.W., Hallacy, C., Ramesh, A., Goh, G., Agarwal, S., Sastry, G., Askell, A., Mishkin, P., Clark, J., Krueger, G., Sutskever, I.: Learning transferable visual models from natural language supervision. In: ICML (2021)
- [50] Rombach, R., Blattmann, A., Lorenz, D., Esser, P., Ommer, B.: High-resolution image synthesis with latent diffusion models. In: CVPR (2022)
- [51] Saharia, C., Chan, W., Saxena, S., Li, L., Whang, J., Denton, E.L., Ghasemipour, K., Gontijo Lopes, R., Karagol Ayan, B., Salimans, T., et al.: Photorealistic text-to-image diffusion models with deep language understanding. NeurIPS **35**, 36479–36494 (2022)
- [52] Schuster, S., Krishna, R., Chang, A., Fei-Fei, L., Manning, C.D.: Generating semantically precise scene graphs from textual descriptions for improved image retrieval. In: Workshop on Vision and Language (VL15). Association for Computational Linguistics, Lisbon, Portugal (September 2015)
- [53] Speer, R.: Python library: rspeer/wordfreq: v3.0 (Sep 2022), <https://doi.org/10.5281/zenodo.7199437>
- [54] Stein, G., Cresswell, J., Hosseinzadeh, R., Sui, Y., Ross, B., Vilecroze, V., Liu, Z., Caterini, A.L., Taylor, E., Loaiza-Ganem, G.: Exposing flaws of generative model evaluation metrics and their unfair treatment of diffusion models. NeurIPS **36** (2024)
- [55] Stroop, J.R.: Studies of interference in serial verbal reactions. Journal of experimental psychology **18**(6), 643 (1935)
- [56] Sun, Q., Fang, Y., Wu, L., Wang, X., Cao, Y.: Eva-clip: Improved training techniques for clip at scale. arXiv preprint arXiv:2303.15389 (2023)
- [57] Tuo, Y., Xiang, W., He, J.Y., Geng, Y., Xie, X.: Anytext: Multilingual visual text generation and editing. arXiv preprint arXiv:2311.03054 (2023)
- [58] Turc, I., Nemade, G.: Midjourney user prompts & generated images (250k) (2023)
- [59] Wang, Z.J., Montoya, E., Munechika, D., Yang, H., Hoover, B., Chau, D.H.: Diffusiondb: A large-scale prompt gallery dataset for text-to-image generative models. arXiv preprint arXiv:2210.14896 (2022)
- [60] Weyand, T., Araujo, A., Cao, B., Sim, J.: Google landmarks dataset v2-a large-scale benchmark for instance-level recognition and retrieval. In: Proceedings of the IEEE/CVF conference on computer vision and pattern recognition. pp. 2575–2584 (2020)

- [61] Yarom, M., Bitton, Y., Changpinyo, S., Aharoni, R., Herzig, J., Lang, O., Ofek, E., Szpektor, I.: What you see is what you read? improving text-image alignment evaluation. *NeurIPS* **36** (2024)
- [62] Yu, J., Wang, Z., Vasudevan, V., Yeung, L., Seyedhosseini, M., Wu, Y.: Coca: Contrastive captioners are image-text foundation models. *Transactions on Machine Learning Research* (2022), <https://openreview.net/forum?id=Ee277P3AYC>
- [63] Yu, J., Xu, Y., Koh, J.Y., Luong, T., Baid, G., Wang, Z., Vasudevan, V., Ku, A., Yang, Y., Ayan, B.K., et al.: Scaling autoregressive models for content-rich text-to-image generation. *arXiv preprint arXiv:2206.10789* **2**(3), 5 (2022)
- [64] Yuksekgonul, M., Bianchi, F., Kalluri, P., Jurafsky, D., Zou, J.: When and why vision-language models behave like bags-of-words, and what to do about it? In: *ICLR* (2022)
- [65] Zapf, A., Castell, S., Morawietz, L., Karch, A.: Measuring inter-rater reliability for nominal data— which coefficients and confidence intervals are appropriate? *BMC medical research methodology* **16**, 1–10 (2016)
- [66] Zeng, Y., Zhang, X., Li, H.: Multi-grained vision language pre-training: Aligning texts with visual concepts. In: Chaudhuri, K., Jegelka, S., Song, L., Szepesvari, C., Niu, G., Sabato, S. (eds.) *Proceedings of the 39th International Conference on Machine Learning*. *Proceedings of Machine Learning Research*, vol. 162, pp. 25994–26009. PMLR (17–23 Jul 2022), <https://proceedings.mlr.press/v162/zeng22c.html>
- [67] Zhai, X., Mustafa, B., Kolesnikov, A., Beyer, L.: Sigmoid loss for language image pre-training. *arXiv preprint arXiv:2303.15343* (2023)
- [68] Zhou, B., Lapedriza, A., Khosla, A., Oliva, A., Torralba, A.: Places: A 10 million image database for scene recognition. *IEEE Transactions on Pattern Analysis and Machine Intelligence* (2017)
- [69] Zhu, X., Sun, P., Wang, C., Liu, J., Li, Z., Xiao, Y., Huang, J.: A contrastive compositional benchmark for text-to-image synthesis: A study with unified text-to-image fidelity metrics. *arXiv preprint arXiv:2312.02338* (2023)