

On the Limitations of Embedding Based Methods for Measuring Functional Correctness for Code Generation

Atharva Naik

Carnegie Mellon University
arnaik@andrew.cmu.edu

Abstract

The task of code generation from natural language (NL2Code) has become extremely popular, especially with the advent of Large Language Models (LLMs). However, efforts to quantify and track this progress have suffered due to a lack of reliable metrics for functional correctness. While popular benchmarks like HumanEval have test cases to enable reliable evaluation of correctness, it is time-consuming and requires human effort to collect test cases. As an alternative several reference-based evaluation metrics have been proposed, with embedding-based metrics like CodeBERTScore being touted as having a high correlation with human preferences and functional correctness. In our work, we analyze the ability of embedding-based metrics like CodeBERTScore to measure functional correctness and other helpful constructs like editing effort by analyzing outputs of ten models over two popular code generation benchmarks. Our results show that while they have a weak correlation with functional correctness ($0.16 r_{bp}$), they are strongly correlated (0.72τ) with editing effort.

1 Introduction

Code generation sometimes also called program synthesis is the task of generating a program or code from a natural language (NL) intent or specification (NL2Code) (Sun et al., 2024). Code generation has a lot of applications for both industry and academia due to the potential to increase the productivity of programmers (Sobania et al., 2022; Barke et al., 2023), connections with parsing (Shin et al., 2019; Sun et al., 2019; Ben-Nun et al., 2018), machine reasoning (Gao et al., 2023), planning (Singh et al., 2023), mathematical capabilities (Shao et al., 2024), and creating structured representations for NLP tasks (Li et al., 2023b).

Modeling approaches for code generation can be grouped into three categories (Sun et al.,

2024): neural models of code, pre-trained models (CodePTMs), and Large Language Models for code (CodeLLMs). Neural methods included recurrent and convolutional networks possibly incorporating abstract syntax tree (AST) structure (Mou et al., 2016; Zhang et al., 2019a). CodePTMs like CodeBERT (Feng et al., 2020), CodeT5 (Wang et al., 2021b), and PLBART (Ahmad et al., 2021) were pre-trained with code-oriented self-supervised training objectives followed by task-specific fine-tuning. Some CodePTMs also incorporate structures like AST, dataflow, and program dependency graphs (PDG) (Guo et al., 2020, 2022; Wang et al., 2021a). Among CodeLLMs, the earliest models were the Codex (Chen et al., 2021) and CodeGen (Nijkamp et al., 2022) series. They were followed by BigCode project models like StarCoder (Li et al., 2023c) and SantaCoder (Allal et al.) that adopted infilling ((Fried et al., 2022)) in their training. Subsequently, CodeT5+ (Wang et al., 2023a) and CodeLLaMA (Roziere et al., 2023) became the first CodeLLMs to have an encoder-decoder architecture and to be trained from an existing model (LLaMA-2 (Touvron et al., 2023)) instead of from scratch respectively. The DeepSeekCoder (Guo et al., 2024) model series added a focus on repository-level and cross-file code completion and the Lemur (Xu et al., 2023) series also based on LLaMA-2 added a focus on agents with robust coding abilities and grounding in their environment.

Evaluation metrics in this space usually target the construct of functional correctness and can be categorized broadly into reference-based, reference-free, hybrid, and human evaluation. Reference-based methods like BLEU (Papineni et al., 2002), ROUGE (Lin, 2004), METEOR (Banerjee and Lavie, 2005) and chrF (Popović, 2015) utilize n-gram matching techniques. Metrics like CodeBLEU (Ren et al., 2020), RUBY (Tran et al., 2019), and CrystalBLEU (Eghbali and Pradel, 2023) capture code-specific properties by

adding dataflow, PDG, AST information or filtering out frequent n-grams. However, prior work has shown that metrics like CodeBLEU and BLEU scores are not reliable indicators of functional correctness or human preferences (Chen et al., 2021; Evtikhiev et al., 2023). Additionally, functionally non-equivalent programs can have higher BLEU scores than functionally equivalent ones (Chen et al., 2021). Embedding-based methods like BERTScore (Zhang et al., 2019b), COMET (Rei et al., 2020), and CodeBERTScore (Zhou et al.) try to capture more semantics. Reference-free methods include executing the generated code on test cases (Chen et al., 2021), using LLMs with evaluation instructions like ICE-Score (Zhou et al., 2023) and round trip evaluation (Allamanis et al., 2024). An example of a hybrid method is CodeScore (Dong et al., 2023) which trains a language model to capture execution semantics and can utilize both the NL context and reference code. Finally despite higher costs human evaluation approaches like RealHumanEval (Mozannar et al., 2024) can capture more human-centered notions of productivity like time to complete coding tasks apart from functional correctness.

Despite the vast array of available approaches we choose to study embedding-based methods, like CodeBERTScore since among reference-based metrics they are the best at capturing semantic information, human preferences, and functional correctness (Zhou et al.), while being faster than execution-based evaluation and not requiring test-cases or test-case training data like CodeScore. They also require fewer computation resources like GPUs compared to LLM-based evaluation approaches like ICE-Score. However, claims about the ability of embedding-based metrics to capture semantics and functional correctness warrant further investigation, since the underlying models (like CodeBERT) are known to fail in capturing semantic equivalence of code (Troshin and Chirkova, 2022; Naik et al., 2022). In this work, we audit the ability of embedding-based metrics like CodeBERTScore to capture functional correctness for ten models over two popular code generation benchmarks. We also analyze their effectiveness in measuring other more human-aligned constructs like editing effort (Dibia et al., 2023). Our results show that embedding-based metrics indeed **fail to capture functional correctness** as measured by execution success, but can potentially function as metrics for editing effort.

Our results motivate the importance of exploring more reference-free evaluation methods like CodeScore for code generation.

2 Task Definition

Before we begin our analysis it is important to formally define the task of code generation.

Context Space: The context space $\mathcal{X}$ for code generation is the natural language instruction (NL) which the set of all strings can represent. $\mathcal{X} = \Sigma^*$

Decision Space: The decision space $\mathcal{Y}$ for code generation is the space of all programs (PL) which can also be represented by the set of all strings, i.e. $\mathcal{Y} = \Sigma^*$. A more specific or constrained space can be defined for specific programming languages guided by the syntax and grammar (Ben-Nun et al., 2018; Sun et al., 2019; Shin et al., 2019). Let’s say the grammar G , with non-terminals N , a start symbol S , the set of terminals Σ , a finite set of production rules P or $G = (N, \Sigma, P, S)$. We can define the decision space $\mathcal{Y}$ as the set of all possible programs that can be produced with the grammar or $\mathcal{Y} = \{w | S \Rightarrow_P^* w\}$. However, most state-of-the-art approaches like LLMs use the unconstrained decision space $\mathcal{Y} = \Sigma^*$.

Construct: The main construct of interest is *functional correctness*, i.e. whether the generated code captures the NL intent accurately. It’s important to measure because it can be challenging for humans to manually validate the correctness of generated code by eyeballing it which can lead to subtle bugs when using AI programming assistants (Bar-Zik, 2023). Since NL is an ambiguous and incomplete specification, practical efforts to measure functional correctness involve creating perfect specifications in the form of test cases on which generated code can be executed to verify the correctness, or reference programs that can be used as proxies to verify correctness. Test case evaluation is more likely to be comprehensive (even though test suites can be limited (Liu et al., 2024)) but is more expensive due to the human effort required to collect test cases, the time taken to execute code, and the need for sandboxing and environment setup to execute code (Yang et al., 2024). Execution-based functional correctness is often measured with test cases using the pass@k (Chen et al., 2021) metric which is defined as:

$$\text{pass@k} := \mathbb{E}_{\text{problems}} \left[1 - \frac{\binom{n-c}{k}}{\binom{n}{k}} \right]$$

Here k is the number of programs sampled from the model, while c is the number of correct programs that pass all test cases. A high pass rate means a high probability that the generated code is functionally correct.

Another interesting construct is the *editing effort* (or syntactic similarity) highlighted by (Dibia et al., 2023). They show that functional correctness can underestimate productivity gains as outputs/system decisions that fail test cases can still be useful if the overall effort for coding can be reduced. They propose EDIT-SIM or the normalized edit similarity defined below as a proxy for effort:

$$\text{EDIT-SIM} = 1 - \frac{\text{lev}(\text{gen}, \text{ref})}{\max(\text{len}(\text{gen}), \text{len}(\text{ref}))}$$

where gen is the code generated or system decision, ref is a reference solution to the problem, and lev is the character Levenshtein edit distance (Levenshtein, 1966). A high *EDIT-SIM* is more desirable and means low editing effort.

Similar to functional correctness, *distinguishability* (d) (Eghbali and Pradel, 2023) measures the ability of a metric to detect the similarity of semantically similar code over semantically different code. It can be mathematically formulated as:

$$d = \frac{m(\text{Pairs}_{\text{intra}})}{m(\text{Pairs}_{\text{inter}})}$$

Where m is the metric, where $\text{Pairs}_{\text{intra}}$ represents the intra-class similarity between pairs of semantically similar code (belonging to an equiv-

alence class or partition based on functional properties) while $\text{Pairs}_{\text{inter}}$ represents the similarity between pairs of code belonging to different equivalence classes. A high distinguishability is ideal as it means the metric can cluster semantically similar codes together.

However (Zhou et al.) shows that since distinguishability uses absolute metric values it can be easily gamed by simply exponentiating metric values making it a bad proxy or alternative for functional correctness. Additionally, they speculate that any meta-metric or construct of correctness that compares exact scores is likely to be gameable compared to ranking-based approaches. Hence we don't use it for our analysis in this study.

Metric of study: The CodeBERTScore metric proposed by (Zhou et al.) as a code-specific variant of BERTScore (Zhang et al., 2019b) utilizes embeddings of a pre-trained CodeBERT model (Feng

et al., 2020) to compute the similarity of a reference code (y^*) with a candidate ($\hat{y}$). They compute all pairs of token similarity between the candidate and reference while masking out certain tokens (according to masks m^* and $\hat{m}$) like punctuation to compute the precision (P) (matching candidate tokens $\hat{y}_j$ to most similar reference tokens y_i^*) (Eq. 1) and recall (R) (Eq. 2) (matching reference tokens to most similar candidate tokens). The F1 (Eq. 3) and F3 (Eq. 4) scores are computed using the precision and recall scores with the F3 score weighing the recall more heavily.

$$\text{CodeBERTScore}_P = \frac{1}{|\hat{y}[\hat{m}]|} \sum_{\hat{y}_j \in \hat{y}[\hat{m}]} \max_{y_i^* \in y^*[m^*]} \text{sim}(y_i^*, \hat{y}_j) \quad (1)$$

$$\text{CodeBERTScore}_R = \frac{1}{|y^*[m]|} \sum_{y_i^* \in y^*[m]} \max_{\hat{y}_j \in \hat{y}[\hat{m}]} \text{sim}(y_i^*, \hat{y}_j) \quad (2)$$

$$\text{CodeBERTScore}_{F_1} = \frac{2 \cdot \text{CodeBERTScore}_P \cdot \text{CodeBERTScore}_R}{\text{CodeBERTScore}_P + \text{CodeBERTScore}_R} \quad (3)$$

$$\text{CodeBERTScore}_{F_3} = \frac{10 \cdot \text{CodeBERTScore}_P \cdot \text{CodeBERTScore}_R}{9 \cdot \text{CodeBERTScore}_P + \text{CodeBERTScore}_R} \quad (4)$$

3 Related Work

3.1 Datasets

Due to the popularity of code generation as a task, many datasets exist out there. Early datasets like CoNaLa (Yin et al., 2018) and its multilingual

version mCoNaLa (Wang et al., 2023b) targeted the task of generating small code snippets. This was followed by the development of some of the most popular function-generation benchmarks like APPS (Hendrycks et al., 2021), MBPP (Austin et al., 2021), and HumanEval (Chen et al., 2021)

with problems that span basic “closed-domain” algorithmic challenges in Python with test cases to support evaluation. HumanEval+ is an extended version of the HumanEval dataset which has better test case coverage (Liu et al., 2024), while HumanEval-X (Zheng et al., 2023) and MultiPLE (Cassano et al., 2022) (extends HumanEval and MBPP) contain additional programming languages other than Python. Code generation datasets for more complex contexts like Jupyter notebooks (ARCADE (Yin et al., 2023), ExeDS (Huang et al., 2022) and JuIce (Agashe et al., 2019)) and classes (CoderEval (Yu et al., 2023) and ClassEval (Du et al., 2023)) have been proposed to move beyond function level code generation. Other datasets aim for specialized domains like data-science DS-1000 (Lai et al., 2023) or open domain code generation (Wang et al., 2022). Some recent datasets have also focused on interactive and multi-turn coding datasets like InterCode (Yang et al., 2024) and CodeClarQA (Li et al., 2023a) (clarification question generation). Finally, some of the latest code generation datasets target repository level and code agent tasks like SWEbench (Jimenez et al., 2023) RepoEval (Zhang et al., 2023), CodeAgentBench (Zhang et al., 2024), RepoBench (Liu et al., 2023) and Stack-Repo (Shrivastava et al., 2023).

3.2 Metrics

For this study, we mainly focus on reference-based evaluation metrics, as the embedding-based methods belong to that category. Each metric is described below:

Exact match:

$$EM := \mathbb{E}_{\text{problems}} [\mathbb{1}_{cand=ref}]$$

Here the $\mathbb{1}_{c \in \mathcal{T}}$ is an indicator variable which is 1 when the candidate is the same as the reference. This essentially plays the role of a lower limit on the performance as in every instance where the exact match is 1 all other metrics will also attain their highest value.

BLEU score:

$$BLEU = BP \times \exp \left(\sum_{n=1}^N w_n \cdot \log \left(\frac{C_{clip,n}}{C_{total,n}} \right) \right)$$

Where:

- BP is the Brevity Penalty
- N is the maximum n-gram order considered (typically $N = 4$)

- w_n is the weight assigned to the n -gram precision
- $C_{clip,n}$ is the clipped count of n -grams in the generated code
- $C_{total,n}$ is the total count of n -grams in the generated code.

The Brevity Penalty (BP) is defined as:

$$BP = \begin{cases} 1 & \text{if } \text{len}_{gen} > \text{len}_{ref} \\ e^{(1 - \frac{\text{len}_{ref}}{\text{len}_{gen}})} & \text{if } \text{len}_{gen} \leq \text{len}_{ref} \end{cases}$$

CrystalBLEU Score:

$$\text{CrystalBLEU} = BP \times \exp \left(\sum_{n=1}^N w_n \cdot \log \left(\frac{C_{clip,S_k,n}}{C_{total,S_k,n}} \right) \right)$$

The Crystal BLEU metric is defined almost identically to BLEU but with one major difference in the clipped n-gram precision computation. The modified precision now also takes in S_k or the top-k most frequent n-grams ($k = 50$ for our study) and removes them from the n-gram counts from both the numerator and denominator in the precision.

CodeBLEU:

$$\text{CodeBLEU} = \alpha \cdot \text{BLEU} + \beta \cdot \text{BLEU}_{\text{weight}} + \gamma \cdot \text{Match}_{\text{ast}} + \delta \cdot \text{Match}_{\text{df}}$$

The CodeBELU score adds a weighted BLEU score, syntax match score, and semantic dataflow match score to the BLEU score.

Here $\text{BLEU}_{\text{weight}}$ is the weighted n-gram BELU score:

$$\text{BLEU}_{\text{weight}} = BP \cdot \exp \left(\sum_{n=1}^N w_n \log p_n \right)$$

Here p_n is the weighted n-gram precision which is defined as:

$$p_n = \frac{\sum_{C \in \text{Candidates}} \sum_{i=1}^l \mu_n^i \cdot \text{Count}_{\text{clip}}(C(i, i+n))}{\sum_{C' \in \text{Candidates}} \sum_{i=1}^l \mu_n^i \cdot \text{Count}_{\text{clip}}(C'(i, i+n))}$$

Where:

- $C(i, i+n)$ is the n-gram from the position i to the position $i+n$
- $\text{Count}_{\text{clip}}(C(i, i+n))$ is the maximum number of n-grams co-occurring in a candidate code and a set of reference codes.
- μ_n^i denotes the weights of different keywords or n-grams. The weight of the keywords is usually set to 5 times the weight of other tokens.
- BP is the brevity penalty as defined for BLEU.

The syntactic match score $\text{Match}_{\text{ast}}$ is defined below:

$$\text{Match}_{\text{ast}} = \text{Count}_{\text{clip}}(T_{\text{cand}}) / \text{Count}(T_{\text{ref}})$$

Where:

- $\text{Count}(T_{\text{ref}})$ is the total number of the reference subtrees.
- $\text{Count}_{\text{clip}}(T_{\text{cand}})$ is the number of the candidate subtrees that are matched to the reference.

Finally, the data match score Match_{df} is defined below:

$$\text{Match}_{\text{df}} = \text{Count}_{\text{clip}}(DF_{\text{cand}}) / \text{Count}(DF_{\text{ref}})$$

Where:

- $\text{Count}(DF_{\text{df}})$ is the total number of the reference data-flows.
- $\text{Count}_{\text{clip}}(DF_{\text{cand}})$ is the number of matched candidate data-flows.

chrF:

$$\text{CHRF} = \max \left(1 - \frac{2 \cdot \text{unigram_errors}}{\text{reference_length} + \text{hypothesis_length}}, 0 \right) \cdot \left(1 - \frac{\text{abs}(r - h)}{r + h} \right)^\beta$$

BERTScore: The BERTScore metric is computed similarly to CodeBERTScore and we use the F1 measure (Eq. 3) throughout the study.

4 Methods

To measure the capacity of code generation metrics to capture functional correctness we use the datasets and systems described below:

4.1 Datasets

We choose the HumanEval (Chen et al., 2021) and MBPP (Austin et al., 2021) benchmarks because of their popularity in the research community as standard benchmarks, relatively simpler coding tasks that don’t involve external libraries, and the availability of test cases for directly measuring functional correctness. The HumanEval dataset has 164 examples and an average of 7.7 test cases. The MBPP (sanitized) test set has 257 examples with 3 test case on average per instance.

4.2 Systems

To analyze the effectiveness of the evaluation metrics we pick some of the most popular open-source language models (StarCoder2 (Lozhkov et al., 2024), CodeLLaMA (Roziere et al., 2023), WizardCoder (Luo et al., 2023), DeepSeekCoder (Guo et al., 2024), Magicoder (Wei et al., 2023), CodeQwen (Bai et al., 2023) CodeGemma (Team et al., 2024), LLaMA-3 (AI@Meta, 2024)) in the 6.7B-15B parameter range as well as some popular closed source models (GPT-3.5-Turbo (Ouyang et al., 2022) and GPT-4 (Achiam et al., 2023)). The

results obtained by the models on all the metrics and the constructs of interest (functional correctness and editing effort) for HumanEval and MBPP are shown in Table 1 and Table 2 respectively. The CodeBERTScore metrics like CodeBERTScore-P etc. are abbreviated as CB-P and so on.

5 Metric Properties

We analyze and compare various properties of CodeBERTScore with respect to other metrics.

5.1 Boundedness

The CodeBERTScore metric theoretically spans values from -1 to 1 similar to BERTScore (Zhang et al.), but empirically it only spans from 0 to 1 empirically. The bounds for all metrics are shown in Table 3.

5.2 Score Distribution

We also analyze the score distribution of all the metrics in Figure 1 and summarize the centrality and shape measures in Table 4. The score distributions show that the embedding-based metrics BERTScore and CodeBERTScore overly reward models while the other metrics, especially exact match overly penalize the models. The centrality measures like mean, median, and mid-hinge also point towards that. The Kurtosis (excess) measures for all metrics except chrF are Leptokurtic, meaning they have fatter tails than a normal distribution. The chrF metric exhibits the least skewness and excess Kurtosis while the exact match metric is most skewed with a fat tail, as most of the values

Model	Metrics										Constructs	
	EM	chrF	BLEU	CodeBLEU	CrystalBLEU	BERTScore	CB-P	CB-R	CB-F1	CB-F3	Functional Correctness	Editing Effort
StarChat2-15B	0.000	0.238	0.111	0.188	0.030	0.570	0.479	0.525	0.500	0.519	0.366	0.191
CodeLLaMA-13B-Hf	0.000	0.156	0.032	0.063	0.008	0.835	0.743	0.673	0.705	0.679	0.018	0.221
WizardCoder-13B	0.000	0.191	0.058	0.105	0.029	0.858	0.783	0.700	0.737	0.707	0.024	0.255
Magicoder-S-DS-6.7B	0.030	0.398	0.251	0.248	0.070	0.897	0.816	0.803	0.809	0.804	0.543	0.440
CodeGemma-7B-It	0.018	0.381	0.213	0.287	0.054	0.893	0.774	0.832	0.799	0.824	0.500	0.369
DeepSeekCoder-6.7B-Instruct	0.030	0.427	0.260	0.267	0.084	0.906	0.838	0.829	0.832	0.829	0.598	0.459
LLaMA-3-8B-Instruct	0.012	0.368	0.199	0.269	0.053	0.890	0.785	0.821	0.800	0.816	0.488	0.377
CodeQwen1.5-7B-Chat	0.043	0.433	0.267	0.311	0.101	0.904	0.820	0.846	0.830	0.843	0.652	0.459
GPT-3.5-Turbo	0.024	0.433	0.266	0.284	0.080	0.904	0.840	0.836	0.838	0.837	0.567	0.456
GPT-4	0.055	0.482	0.325	0.323	0.107	0.914	0.861	0.855	0.857	0.855	0.665	0.499

Table 1: Results for all models over all metrics and constructs for the HumanEval dataset. The CodeQwen and DeepSeekCoder seem to have the best performance over the embedding-based metrics and n-gram-based metrics respectively and the functional correctness and editing effort constructs respectively. Among the close source models, GPT-4 is better on all metrics and constructs and better overall among all models.

Model	Metrics										Constructs	
	EM	chrF	BLEU	CodeBLEU	CrystalBLEU	BERTScore	CB-P	CB-R	CB-F1	CB-F3	Functional Correctness	Editing Effort
StarChat2-15B	0.000	0.353	0.151	0.250	0.016	0.893	0.798	0.820	0.807	0.817	0.533	0.353
CodeLLaMA-13B-Hf	0.000	0.201	0.032	0.044	0.000	0.868	0.806	0.712	0.755	0.720	0.027	0.239
WizardCoder-13B	0.000	0.207	0.036	0.043	0.000	0.861	0.802	0.709	0.752	0.717	0.012	0.237
Magicoder-S-DS-6.7B	0.000	0.360	0.160	0.224	0.018	0.898	0.838	0.815	0.825	0.817	0.529	0.383
CodeGemma-7B-It	0.000	0.299	0.085	0.241	0.009	0.861	0.718	0.785	0.749	0.778	0.494	0.229
DeepSeekCoder-6.7B-Instruct	0.000	0.317	0.128	0.180	0.008	0.889	0.814	0.792	0.802	0.794	0.366	0.352
LLaMA-3-8B-Instruct	0.000	0.349	0.140	0.221	0.015	0.887	0.810	0.810	0.809	0.809	0.451	0.357
CodeQwen1.5-7B-Chat	0.004	0.487	0.317	0.306	0.066	0.910	0.854	0.863	0.857	0.862	0.619	0.474
GPT-3.5-Turbo	0.000	0.445	0.246	0.249	0.025	0.904	0.845	0.845	0.844	0.845	0.786	0.426
GPT-4	0.004	0.469	0.301	0.267	0.051	0.911	0.873	0.854	0.862	0.856	0.802	0.458

Table 2: Results for all models over all metrics and constructs for the MBPP dataset. The CodeQwen model has the best performance over all metrics among the open-source models and the constructs. Among the close source models, GPT-4 is better on all metrics and best for most metrics overall (except the BLEU metrics, chrF and CodeBERTScore-F3).

Metric	Theoretical Bounds		Empirical Bounds	
	Max	Min	Max	Min
Exact Match	0	1	0	1
chrF	0	1	0	1
BLEU	0	1	0	1
CodeBLEU	0	1	0	1
CrystalBLEU	0	1	0	1
BERTScore-F1	-1	1	0	1
CodeBERTScore-P	-1	1	0	1
CodeBERTScore-R	-1	1	0	1
CodeBERTScore-F1	-1	1	0	1
CodeBERTScore-F3	-1	1	0	1

Table 3: The theoretical bounds and empirically attained bounds for all metrics

are zero (high peak) and then there is a fat tail of all the non-zero values.

5.3 Number of ties

We compute the percentage of ties across a given context for all pairs of model comparisons and show the results in Table 5. For continuous metrics that can attain any real values between zero to one on an instance level, we count it as a tie if the metric values differ by a small value $\epsilon = 10^{-6}$. The exact match metric has the most ties because of sparsity and being 0 for most cases while, BERTScore and CodeBERTScore metrics have fewest ties.

5.4 Convergent and Discriminant Validity

We analyze the Kendall Tau (τ) correlation between all the metrics (EM, chrF, BLEU, CodeBLEU, CrystalBLEU, and BERTScore-F1) and all the CodeBERTScore metrics (P, R, F1, F3) and visualize the results via a heatmap in Figure 2. All the correlations are statistically significant with very low p-values. We observe that except for exact match, CodeBERTScore exhibits moderate correlation with all metrics.

5.5 Discriminative Power

We plot the discriminative power of all the metrics compared using the student t-test with the hypothesis that the mean of the underlying distribution of the first model is less than the mean of the distribution of the second model with Bonferroni correction (Armstrong, 2014) applied to the acceptance threshold $\alpha = \frac{0.05}{n}$ (n is the number of hypotheses) to account for multiple hypotheses. Since there are $n = 2 \times \binom{10}{2} = 90$ comparisons, the acceptance threshold $\alpha = \frac{0.05}{90} = 0.00055$. We plot the achieved significance level (Sakai, 2014) for each metric (sorted p-values on the y-axis and run comparison pairs on the x-axis) in Figure 3. However, since the acceptance threshold is very low, to make the plots easier to visualize we also show a clipped version where all the p-values are clipped if they are above 0.008 in Figure 4. We note that CodeBERTScore-F1 has the most discriminative power being able to achieve statistical significance for 58/90 pairs, and the exact match being the worst, not being able to distinguish any pairs. Additionally, the code-specific BLEU scores (CodeBLEU, CrystalBLEU) are worse than regular BLEU which is comparable to BERTScore and chrF.

6 Empirical Analysis

6.1 Hypotheses

Since the primary focus of the work is to test the reliability of CodeBERTScore as a measure of the primary construct of functional correctness and the secondary measure of editing effort, we can formulate the following hypotheses:

H1. CodeBERTScore is moderate to strongly correlated with execution results (pass@1).

H2. CodeBERTScore is robust to surface-level perturbation, i.e. it exhibits strong auto-correlation under function-name and variable input perturbations (test of criterion validity with invariant transformations)

H3. CodeBERTScore is moderate to strongly correlated with editing effort (EDIT-SIM)

6.2 Hypothesis Testing

To test **H1** we measure the correlation between the pass@1 metric and CodeBERTScore-F1 using point-biserial correlation r_{bp} which is suitable for measuring the correlation between a binary or dichotomous variable like whether an instance passes all test cases and a continuous variable like CodeBERTScore. We use F1 for all subsequent hypotheses as it is the most sensitive metric according to the discriminative power and combines both precision and recall without any specific bias to either. In addition to this correlation, we also compute the point-biserial correlation between editing effort (EDIT-SIM) and pass@1 and BERTScore and pass@1 to see if the two constructs are correlated and compare the reliability of BERTScore and CodeBERTScore respectively.

We find an r_{bp} of 0.162 with p-value of 3.28×10^{-26} between CodeBERTScore and pass@1 and an r_{bp} of 0.186 between EDIT-SIM and pass@1 with a p-value of 5.83×10^{-34} and an r_{bp} of 0.103 with a p-value of 2.11×10^{-11} for BERTScore.

To test **H2** we measure the auto-correlation between the results obtained with CodeBERTScore (and additionally BERTScore) before and after semantics preserving input perturbations, using Kendall Tau (τ) and Spearman (ρ) correlations. The perturbations are described below:

1. **Var: Cand only** The variable names in the candidate are replaced with generic variable names (e.g. if sum is the third variable in the code, then it will be replaced with var2). We expect this transformation to **lower the metric**

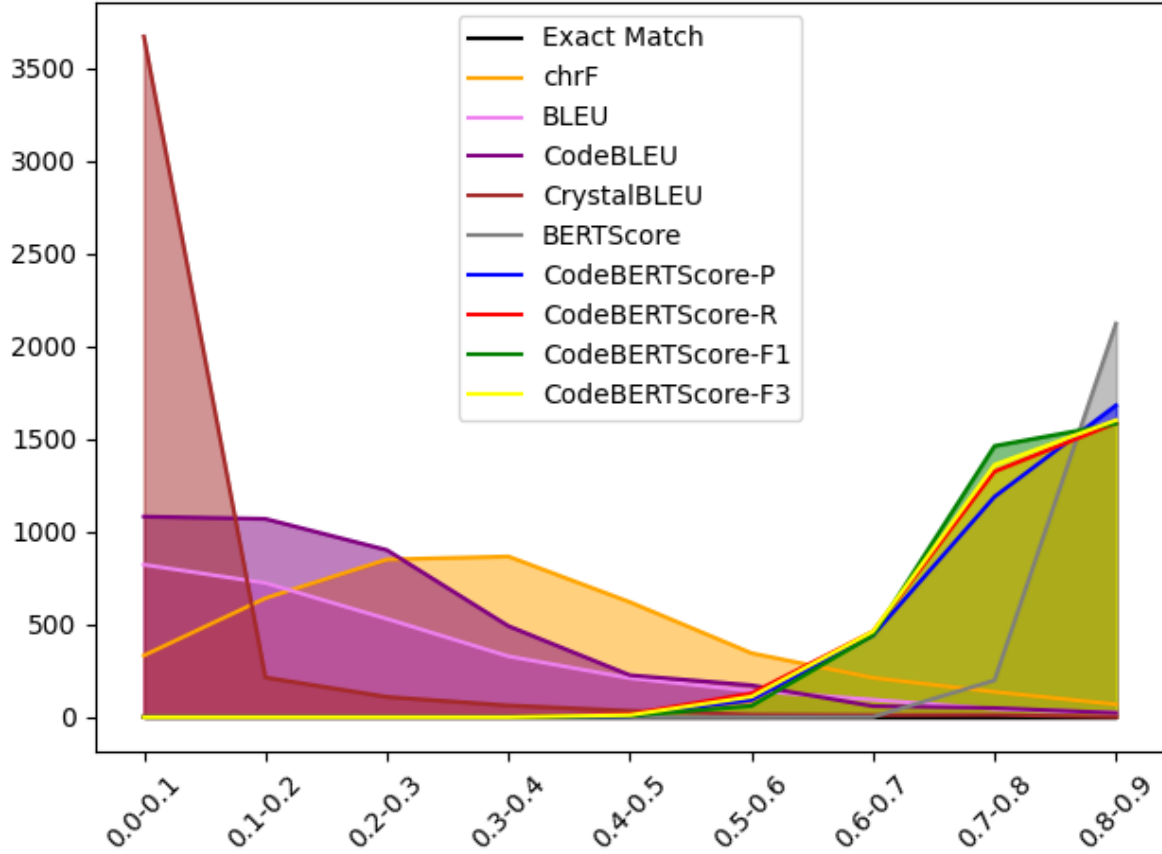


Figure 1: The distribution of metric values. The **x-axis** shows the range of values while the **y-axis** shows the number of model decisions with values in that range. The distributions clearly show that the embedding metrics tend to over-reward decisions while the other metrics tend to over-penalize them. The **chrF** metric has the least skew

value

2. **Var: Ref only** The variable names in the reference are replaced with generic variable names. We expect this transformation to **lower the metric value**
3. **Var: Cand and Ref** The variable names in both reference and candidate are replaced with generic variable names. We expect this transformation to **increase the metric value**
4. **Func: same** The function names of both candidate and reference functions are replaced with the same name (candidate_function). We expect this transformation to **increase the metric value**
5. **Func: different** The function names of both candidate and reference functions are replaced with different names (candidate_function and reference_function respectively). We

expect this transformation to **decrease the metric value**

We show the correlations and p-values in Table 6. Note that all the p-values are 0 and all autocorrelation values are in the strong range, even though the metric values themselves are affected slightly by the perturbations (all the trends of increasing/decreasing values are as expected).

To test **H3** we compute the Kendall Tau (τ) and Spearman (ρ) correlations between EDIT-SIM and CodeBERTScore getting $\tau = 0.72$ and $\rho = 0.89$ with p-values of 0 in both cases.

7 Discussion

The results for **H1** show that while there is a statistically significant correlation between functional correctness pass@1 and CodeBERTScore, it is rather weak and the same is the case for BERTScore. It is less correlated with it than the construct of editing effort EDIT-SIM, so we can discard **H1**. The results for **H2** show that while the input pertur-

Metric	Median	Midhinge	Mean	Std. Dev.	Skewness	Kurtosis
Exact Match	0	0	0.009	0.093	10.526	108.793
chrF	0.323	0.330	0.350	0.205	0.813	0.645
BLEU	0.108	0.132	0.175	0.212	1.611	2.642
CodeBLEU	0.186	0.189	0.215	0.177	1.314	2.302
CrystalBLEU	0	0	0.037	0.117	4.772	27.645
BERTScore-F1	0.892	0.889	0.876	0.118	-6.035	42.250
CodeBERTScore-P	0.817	0.815	0.800	0.135	-3.072	15.600
CodeBERTScore-R	0.806	0.804	0.789	0.137	-2.829	13.819
CodeBERTScore-F1	0.804	0.804	0.792	0.131	-3.187	17.127
CodeBERTScore-F3	0.805	0.803	0.790	0.135	-2.937	14.796

Table 4: The centrality and shape measures for all the metrics. The metrics highlighted in red are heavily skewed towards the left (i.e. they over-penalize the models) while the ones highlighted in red are heavily skewed towards the right (i.e. they over-reward the models). Additionally almost all of the metrics are Leptokurtic, i.e. they have fatter tails than a normal distribution, especially Exact Match and BERT-Score.

Metric	Rate of Ties
Exact Match	98.8
chrF	3.73
BLEU	16.84
CodeBLEU	7.7
CrystalBLEU	47.77
BERTScore-F1	3.17
CodeBERTScore-P	3.17
CodeBERTScore-R	3.17
CodeBERTScore-F1	3.18
CodeBERTScore-F3	3.17

Table 5: Percentage of ties for each metric. For real-valued metrics, we count it as a tie if the metric values differ by a very small value $\epsilon = 10^{-6}$. The exact match metric has the most ties because of sparsity and being 0 for most cases while BERTScore and CodeBERTScore metrics have the fewest ties.

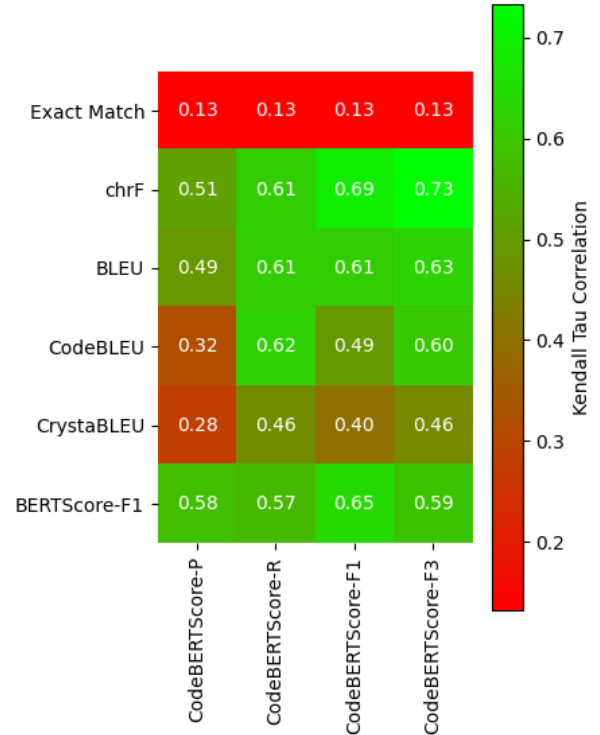


Figure 2: Correlations of the CodeBERTScore metrics with the other metrics. Except for exact match, most of the metrics have a moderate correlation with CodeBERTScore

bations can affect the metric values in expected ways, the metrics remain strongly and statistically significantly auto-correlated so we can accept **H2**. The results for **H3** show a strong and statistically significant correlation so we can accept it.

From these results, we can conclude that despite

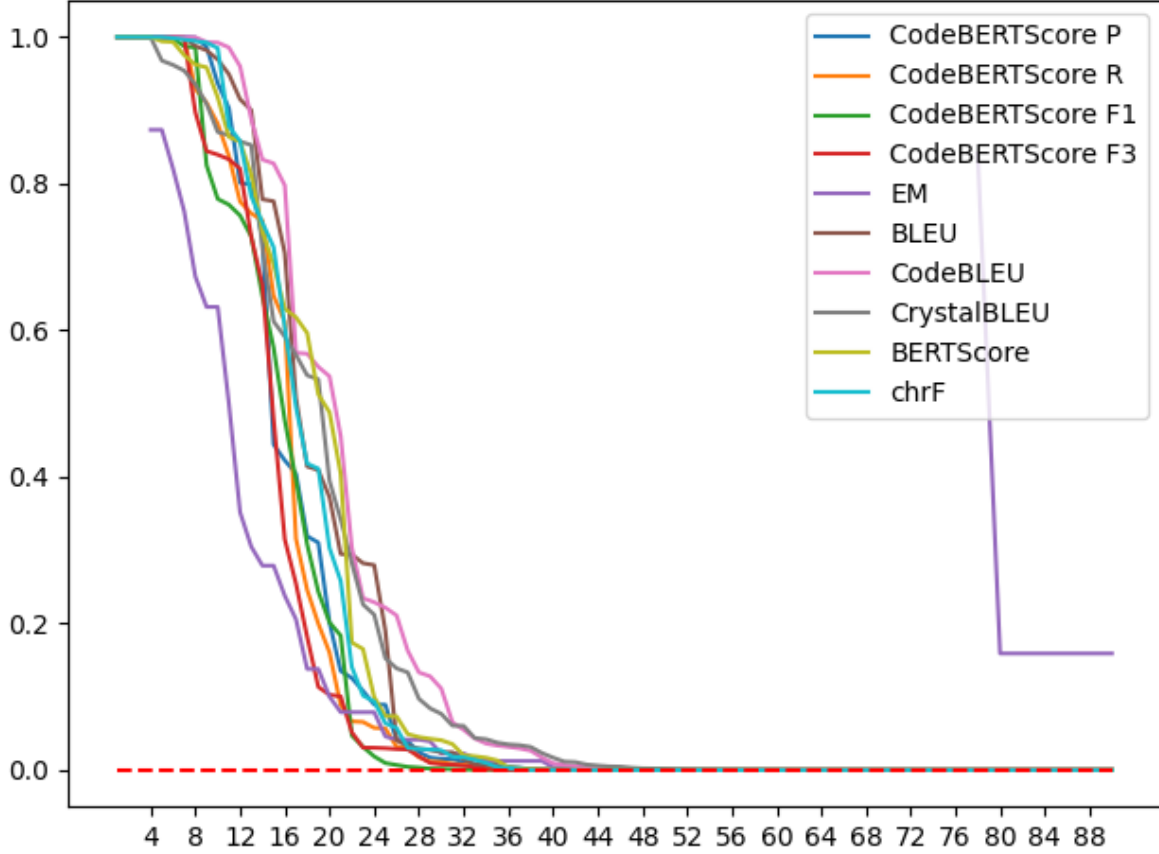


Figure 3: The achieved significance level of the metric values. **x-axis** shows pairs of runs while **y-axis** shows the p-values. The dotted line shows the acceptance threshold

claims made by (Zhou et al.), **CodeBERTScore (or BERTScore) is not a reliable proxy for functional correctness.** However, it is a relatively robust metric with the highest sensitivity (discriminative power) among reference-based metrics and could be used as a measure of syntactic similarity or editing effort for a reference/goal output. In fact, the results indicate that embedding-based metrics might be able to capture the syntax level structure of code (suggested by the robustness to input perturbations) but fail to penetrate to the level of code semantic equivalence as supported by previous studies that probe code models (Troshin and Chirkova, 2022; Naik et al., 2022).

Future Work: While this study extensively explores many state-of-the-art LLMs over two of the most popular benchmarks, it would be interesting to explore the results over more datasets like ODEX (Wang et al., 2022) that cover more challenging and open-domain code generation but still support test case evaluation to further validate the findings. It would also be interesting to explore

these results for models trained on the evaluation benchmarks for the MBPP dataset and benchmark the performance of iterative/multi-step code generators like code agents (Huang et al., 2024; Zhong et al., 2024; Hong et al., 2023; Zhou et al., 2023). Finally, it would be interesting to explore the effectiveness of metrics that learn execution behavior like CodeScore (Dong et al., 2023) and explore ways to improve embedding-based metrics by learning execution behavior.

Online Experiments: While static benchmarks like HumanEval (Chen et al., 2021) were useful at the time of their introduction for evaluating the functional correctness of LLM-generated code, they fail to completely align with the primary use case of LLMs as programming assistants and their effect on constructs like programmer productivity. As a result, benchmarking efforts like RealHumanEval (Mozannar et al., 2024) use an **online web interface, with an editor and chat window** to measure the ability of LLMs to

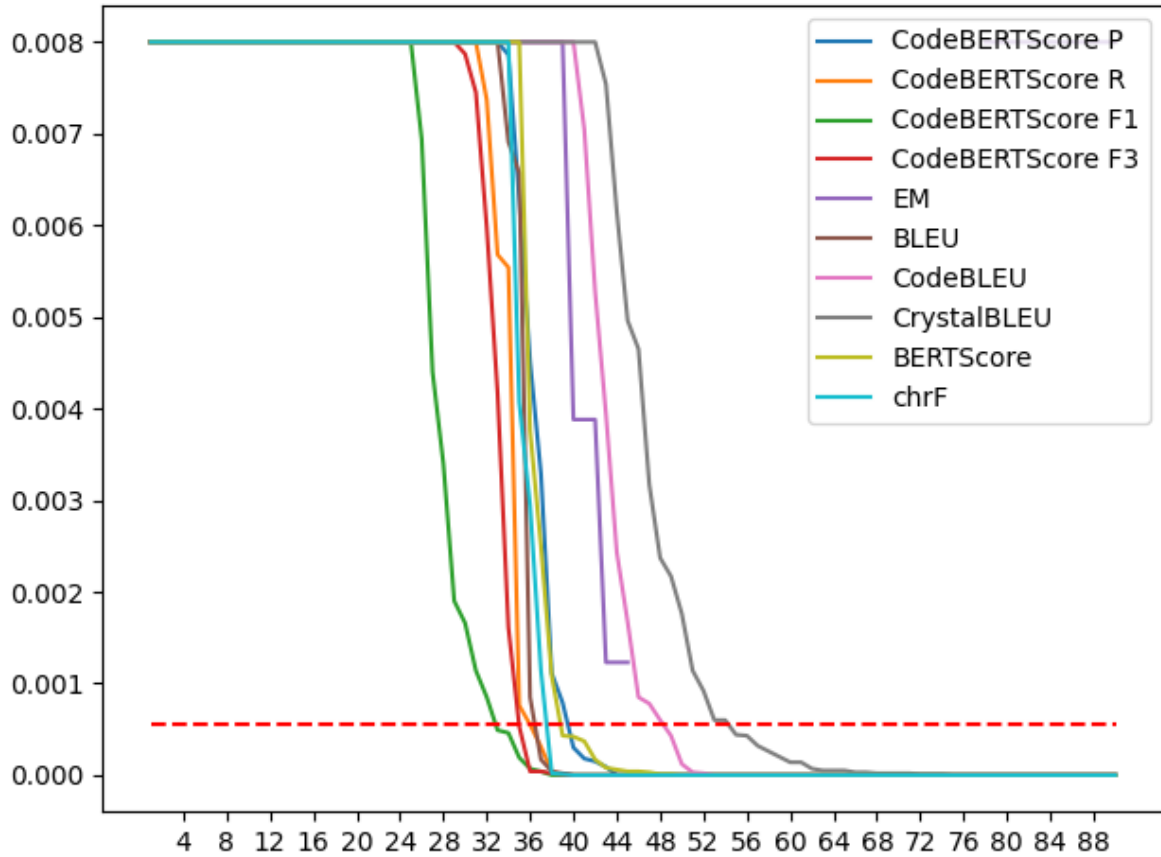


Figure 4: The achieved significance level of the metric values. **x-axis** shows pairs of runs while **y-axis** shows the p-values. The dotted line shows the acceptance threshold. All p-values above 0.008 are clipped to zoom into the portion close to the acceptance threshold.

assist programmers through auto-complete or chat support, while using proxy behavioral **signals like code acceptance and copy rates**. Experiments done by (Mozannar et al., 2024) show that while increased performance on static benchmarks leads to increased productivity the gaps in benchmark and human performance are not always proportional. They also found that programmer preferences do not correlate with actual LLM performance motivating more human-centered evaluation. Important **sub-populations** for such experiments include **novice programmers and experts**. Experts are likely to reject more code or take longer till they except a code as they might be stricter with their requirements and other aspects of code quality like clarity and efficiency. The copy rates could be higher with any population as while experts might be more strict with accepting output they might be more likely to just copy code and prefer to edit it themselves than novice users. Figuring out these differences for certain would

require further experimentation.

References

- Josh Achiam, Steven Adler, Sandhini Agarwal, Lama Ahmad, Ilge Akkaya, Florencia Leoni Aleman, Diogo Almeida, Janko Altenschmidt, Sam Altman, Shyamal Anadkat, et al. 2023. Gpt-4 technical report. *arXiv preprint arXiv:2303.08774*.
- Rajas Agashe, Srinivasan Iyer, and Luke Zettlemoyer. 2019. [JuICe: A large scale distantly supervised dataset for open domain context-based code generation](#). In *Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP)*, pages 5436–5446, Hong Kong, China. Association for Computational Linguistics.
- WU Ahmad, S Chakraborty, B Ray, and KW Chang. 2021. Unified pre-training for program understanding and generation. In *Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*.

Metric	Transformation	F1	Kendall Tau Autocorrelation		Spearman Autocorrelation	
			Tau	p-value	Rho	p-value
CodeBERTScore-F1	No transform	0.879				
	Var: Cand only	0.847 (-0.032)	0.811	0	0.951	0
	Var: Ref only	0.842 (-0.037)	0.82	0	0.954	0
	Var: Cand & Ref	0.904 (+0.025)	0.806	0	0.943	0
	Func: same	0.881 (+0.002)	0.926	0	0.988	0
	Func: different	0.872 (-0.005)	0.918	0	0.989	0
BERTScore	No transform	0.937				
	Var: Cand only	0.927 (-0.01)	0.857	0	0.97	0
	Var: Ref only	0.926 (-0.011)	0.861	0	0.971	0
	Var: Cand & Ref	0.946 (+0.009)	0.864	0	0.971	0
	Func: same	0.939 (+0.002)	0.935	0	0.991	0
	Func: different	0.936 (-0.001)	0.938	0	0.993	0

Table 6: Results of auto-correlation under semantics preserving input perturbations for CodeBERTScore and BERTScore. The results show the although the metric values are affected slightly by the transformations, the metrics are still robust due to the high correlation between the old and new values after applying the transform.

AI@Meta. 2024. [Llama 3 model card](#).

Loubna Ben Allal, Hugging Face, Raymond Li, Denis Kocetkov, Chenghao Mou, Christopher Akiki, ScaDS AI Leipzig, Carlos Munoz Ferrandis, Niklas Muennighoff, Mayank Mishra, et al. Santacoder: Don’t reach for the stars!

Miltiadis Allamanis, Sheena Panthaplackel, and Pengcheng Yin. 2024. [Unsupervised evaluation of code llms with round-trip correctness](#).

Richard A Armstrong. 2014. When to use the bonferroni correction. *Ophthalmic Physiol Opt*, 34(5):502–508.

Jacob Austin, Augustus Odena, Maxwell Nye, Maarten Bosma, Henryk Michalewski, David Dohan, Ellen Jiang, Carrie Cai, Michael Terry, Quoc Le, et al. 2021. Program synthesis with large language models. *arXiv preprint arXiv:2108.07732*.

Jinze Bai, Shuai Bai, Yunfei Chu, Zeyu Cui, Kai Dang, Xiaodong Deng, Yang Fan, Wenbin Ge, Yu Han, Fei Huang, Binyuan Hui, Luo Ji, Mei Li, Junyang Lin, Runji Lin, Dayiheng Liu, Gao Liu, Chengqiang Lu, Keming Lu, Jianxin Ma, Rui Men, Xingzhang Ren, Xuancheng Ren, Chuanqi Tan, Sinan Tan, Jianhong Tu, Peng Wang, Shijie Wang, Wei Wang, Sheng-guang Wu, Benfeng Xu, Jin Xu, An Yang, Hao Yang, Jian Yang, Shusheng Yang, Yang Yao, Bowen Yu, Hongyi Yuan, Zheng Yuan, Jianwei Zhang, Xingxuan Zhang, Yichang Zhang, Zhenru Zhang, Chang Zhou, Jingren Zhou, Xiaohuan Zhou, and Tianhang Zhu. 2023. Qwen technical report. *arXiv preprint arXiv:2309.16609*.

Satanjeev Banerjee and Alon Lavie. 2005. [METEOR: An automatic metric for MT evaluation with improved correlation with human judgments](#). In *Proceedings of the ACL Workshop on Intrinsic and Extrinsic Evaluation Measures for Machine Translation and/or Summarization*, pages 65–72, Ann Arbor,

Michigan. Association for Computational Linguistics.

Ran Bar-Zik. 2023. [Guarding the gates of language: Exploring vulnerabilities in llm output](#). *Medium*.

Shraddha Barke, Michael B. James, and Nadia Polikarpova. 2023. [Grounded copilot: How programmers interact with code-generating models](#). *Proc. ACM Program. Lang.*, 7(OOPSLA1).

Tal Ben-Nun, Alice Shoshana Jakobovits, and Torsten Hoefler. 2018. Neural code comprehension: A learnable representation of code semantics. *Advances in neural information processing systems*, 31.

Federico Cassano, John Gouwar, Daniel Nguyen, Sydney Nguyen, Luna Phipps-Costin, Donald Pinckney, Ming-Ho Yee, Yangtian Zi, Carolyn Jane Anderson, Molly Q Feldman, et al. 2022. Multipl-e: A scalable and extensible approach to benchmarking neural code generation. *arXiv preprint arXiv:2208.08227*.

Mark Chen, Jerry Tworek, Heewoo Jun, Qiming Yuan, Henrique Ponde de Oliveira Pinto, Jared Kaplan, Harri Edwards, Yuri Burda, Nicholas Joseph, Greg Brockman, Alex Ray, Raul Puri, Gretchen Krueger, Michael Petrov, Heidy Khlaaf, Girish Sastry, Pamela Mishkin, Brooke Chan, Scott Gray, Nick Ryder, Mikhail Pavlov, Alethea Power, Lukasz Kaiser, Mohammad Bavarian, Clemens Winter, Philippe Tillet, Felipe Petroski Such, Dave Cummings, Matthias Plappert, Fotios Chantzis, Elizabeth Barnes, Ariel Herbert-Voss, William Hebgen Guss, Alex Nichol, Alex Paino, Nikolas Tezak, Jie Tang, Igor Babuschkin, Suchir Balaji, Shantanu Jain, William Saunders, Christopher Hesse, Andrew N. Carr, Jan Leike, Josh Achiam, Vedant Misra, Evan Morikawa, Alec Radford, Matthew Knight, Miles Brundage, Mira Murati, Katie Mayer, Peter Welinder, Bob McGrew, Dario Amodei, Sam McCandlish, Ilya Sutskever, and Wojciech Zaremba. 2021. [Evaluating large language models trained on code](#).

- Victor Dibia, Adam Fourney, Gagan Bansal, Forough Poursabzi-Sangdeh, Han Liu, and Saleema Amershi. 2023. [Aligning offline metrics and human judgments of value for code generation models](#).
- Yihong Dong, Jiazheng Ding, Xue Jiang, Zhuo Li, Ge Li, and Zhi Jin. 2023. Codescore: Evaluating code generation by learning code execution. *arXiv preprint arXiv:2301.09043*.
- Xueying Du, Mingwei Liu, Kaixin Wang, Hanlin Wang, Junwei Liu, Yixuan Chen, Jiayi Feng, Chaofeng Sha, Xin Peng, and Yiling Lou. 2023. Classeval: A manually-crafted benchmark for evaluating llms on class-level code generation. *arXiv preprint arXiv:2308.01861*.
- Aryaz Eghbali and Michael Pradel. 2023. [Crystalbleu: Precisely and efficiently measuring the similarity of code](#). In *Proceedings of the 37th IEEE/ACM International Conference on Automated Software Engineering, ASE '22*, New York, NY, USA. Association for Computing Machinery.
- Mikhail Evtikhiev, Egor Bogomolov, Yaroslav Sokolov, and Timofey Bryksin. 2023. Out of the bleu: how should we assess quality of the code generation models? *Journal of Systems and Software*, 203:111741.
- Zhangyin Feng, Daya Guo, Duyu Tang, Nan Duan, Xiaocheng Feng, Ming Gong, Linjun Shou, Bing Qin, Ting Liu, Daxin Jiang, and Ming Zhou. 2020. [CodeBERT: A pre-trained model for programming and natural languages](#). In *Findings of the Association for Computational Linguistics: EMNLP 2020*, pages 1536–1547, Online. Association for Computational Linguistics.
- Daniel Fried, Armen Aghajanyan, Jessy Lin, Sida Wang, Eric Wallace, Freda Shi, Ruiqi Zhong, Scott Yih, Luke Zettlemoyer, and Mike Lewis. 2022. Incoder: A generative model for code infilling and synthesis. In *The Eleventh International Conference on Learning Representations*.
- Luyu Gao, Aman Madaan, Shuyan Zhou, Uri Alon, Pengfei Liu, Yiming Yang, Jamie Callan, and Graham Neubig. 2023. Pal: Program-aided language models. In *International Conference on Machine Learning*, pages 10764–10799. PMLR.
- Daya Guo, Shuai Lu, Nan Duan, Yanlin Wang, Ming Zhou, and Jian Yin. 2022. Unixcoder: Unified cross-modal pre-training for code representation. In *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 7212–7225.
- Daya Guo, Shuo Ren, Shuai Lu, Zhangyin Feng, Duyu Tang, Shujie Liu, Long Zhou, Nan Duan, Alexey Svyatkovskiy, Shengyu Fu, et al. 2020. Graphcodebert: Pre-training code representations with data flow. *arXiv preprint arXiv:2009.08366*.
- Daya Guo, Qihao Zhu, Dejian Yang, Zhenda Xie, Kai Dong, Wentao Zhang, Guanting Chen, Xiao Bi, Y Wu, YK Li, et al. 2024. Deepseek-coder: When the large language model meets programming—the rise of code intelligence. *arXiv preprint arXiv:2401.14196*.
- Dan Hendrycks, Steven Basart, Saurav Kadavath, Mantas Mazeika, Akul Arora, Ethan Guo, Collin Burns, Samir Puranik, Horace He, Dawn Song, et al. 2021. Measuring coding challenge competence with apps. In *Thirty-fifth Conference on Neural Information Processing Systems Datasets and Benchmarks Track (Round 2)*.
- Sirui Hong, Mingchen Zhuge, Jonathan Chen, Xiaowu Zheng, Yuheng Cheng, Ceyao Zhang, Jinlin Wang, Zili Wang, Steven Ka Shing Yau, Zijuan Lin, Liyang Zhou, Chenyu Ran, Lingfeng Xiao, Chenglin Wu, and Jürgen Schmidhuber. 2023. [Metagpt: Meta programming for a multi-agent collaborative framework](#).
- Dong Huang, Qingwen Bu, Jie M. Zhang, Michael Luck, and Heming Cui. 2024. [Agentcoder: Multi-agent-based code generation with iterative testing and optimisation](#).
- Junjie Huang, Chenglong Wang, Jipeng Zhang, Cong Yan, Haotian Cui, Jeevana Priya Inala, Colin Clement, and Nan Duan. 2022. [Execution-based evaluation for data science code generation models](#). In *Proceedings of the Fourth Workshop on Data Science with Human-in-the-Loop (Language Advances)*, pages 28–36, Abu Dhabi, United Arab Emirates (Hybrid). Association for Computational Linguistics.
- Carlos E. Jimenez, John Yang, Alexander Wettig, Shunyu Yao, Kexin Pei, Ofir Press, and Karthik Narasimhan. 2023. [Swe-bench: Can language models resolve real-world github issues?](#) *ArXiv*, abs/2310.06770.
- Yuhang Lai, Chengxi Li, Yiming Wang, Tianyi Zhang, Ruiqi Zhong, Luke Zettlemoyer, Wen-tau Yih, Daniel Fried, Sida Wang, and Tao Yu. 2023. Ds-1000: a natural and reliable benchmark for data science code generation. In *Proceedings of the 40th International Conference on Machine Learning*, pages 18319–18345.
- Vladimir Iosifovich Levenshtein. 1966. Binary codes capable of correcting deletions, insertions and reversals. *Soviet Physics Doklady*, 10(8):707–710. Doklady Akademii Nauk SSSR, V163 No4 845-848 1965.
- Haau-Sing (Xiaocheng) Li, Mohsen Mesgar, André Martins, and Iryna Gurevych. 2023a. [Python code generation by asking clarification questions](#). In *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 14287–14306, Toronto, Canada. Association for Computational Linguistics.
- Peng Li, Tianxiang Sun, Qiong Tang, Hang Yan, Yuanbin Wu, Xuanjing Huang, and Xipeng Qiu. 2023b.

- Codeie: Large code generation models are better few-shot information extractors. *arXiv preprint arXiv:2305.05711*.
- Raymond Li, Yangtian Zi, Niklas Muennighoff, Denis Kocetkov, Chenghao Mou, Marc Marone, Christopher Akiki, Li Jia, Jenny Chim, Qian Liu, et al. 2023c. Starcoder: may the source be with you! *Transactions on Machine Learning Research*.
- Chin-Yew Lin. 2004. **ROUGE: A package for automatic evaluation of summaries**. In *Text Summarization Branches Out*, pages 74–81, Barcelona, Spain. Association for Computational Linguistics.
- Jiawei Liu, Chunqiu Steven Xia, Yuyao Wang, and Lingming Zhang. 2024. Is your code generated by chatgpt really correct? rigorous evaluation of large language models for code generation. *Advances in Neural Information Processing Systems*, 36.
- Tianyang Liu, Canwen Xu, and Julian McAuley. 2023. Repobench: Benchmarking repository-level code auto-completion systems. In *The Twelfth International Conference on Learning Representations*.
- Anton Lozhkov, Raymond Li, Loubna Ben Allal, Federico Cassano, Joel Lamy-Poirier, Nouamane Tazi, Ao Tang, Dmytro Pykhtar, Jiawei Liu, Yuxiang Wei, et al. 2024. Starcoder 2 and the stack v2: The next generation. *arXiv preprint arXiv:2402.19173*.
- Ziyang Luo, Can Xu, Pu Zhao, Qingfeng Sun, Xubo Geng, Wenxiang Hu, Chongyang Tao, Jing Ma, Qingwei Lin, and Daxin Jiang. 2023. Wizardcoder: Empowering code large language models with evol-instruct. *arXiv preprint arXiv:2306.08568*.
- Lili Mou, Ge Li, Lu Zhang, Tao Wang, and Zhi Jin. 2016. Convolutional neural networks over tree structures for programming language processing. In *Proceedings of the Thirtieth AAAI Conference on Artificial Intelligence, AAAI’16*, page 1287–1293. AAAI Press.
- Hussein Mozannar, Valerie Chen, Mohammed Alsobay, Subhro Das, Sebastian Zhao, Dennis Wei, Manish Nagireddy, Prasanna Sattigeri, Ameet Talwalkar, and David Sontag. 2024. **The realhumaneval: Evaluating large language models’ abilities to support programmers**.
- Shounak Naik, Rajaswa Patil, Swati Agarwal, and Veeky Baths. 2022. **Probing semantic grounding in language models of code with representational similarity analysis**. *ArXiv*, abs/2207.07706.
- Erik Nijkamp, Bo Pang, Hiroaki Hayashi, Lifu Tu, Huan Wang, Yingbo Zhou, Silvio Savarese, and Caiming Xiong. 2022. Codegen: An open large language model for code with multi-turn program synthesis. *arXiv preprint arXiv:2203.13474*.
- Long Ouyang, Jeffrey Wu, Xu Jiang, Diogo Almeida, Carroll Wainwright, Pamela Mishkin, Chong Zhang, Sandhini Agarwal, Katarina Slama, Alex Ray, et al. 2022. Training language models to follow instructions with human feedback. *Advances in neural information processing systems*, 35:27730–27744.
- Kishore Papineni, Salim Roukos, Todd Ward, and Wei-Jing Zhu. 2002. **Bleu: a method for automatic evaluation of machine translation**. In *Proceedings of the 40th Annual Meeting of the Association for Computational Linguistics*, pages 311–318, Philadelphia, Pennsylvania, USA. Association for Computational Linguistics.
- Maja Popović. 2015. **chrF: character n-gram F-score for automatic MT evaluation**. In *Proceedings of the Tenth Workshop on Statistical Machine Translation*, pages 392–395, Lisbon, Portugal. Association for Computational Linguistics.
- Ricardo Rei, Craig Stewart, Ana C Farinha, and Alon Lavie. 2020. **COMET: A neural framework for MT evaluation**. In *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pages 2685–2702, Online. Association for Computational Linguistics.
- Shuo Ren, Daya Guo, Shuai Lu, Long Zhou, Shujie Liu, Duyu Tang, Neel Sundaresan, Ming Zhou, Ambrosio Blanco, and Shuai Ma. 2020. Codebleu: a method for automatic evaluation of code synthesis. *arXiv preprint arXiv:2009.10297*.
- Baptiste Roziere, Jonas Gehring, Fabian Gloeckle, Sten Sootla, Itai Gat, Xiaoqing Ellen Tan, Yossi Adi, Jingyu Liu, Tal Remez, Jérémy Rapin, et al. 2023. Code llama: Open foundation models for code. *arXiv preprint arXiv:2308.12950*.
- Tetsuya Sakai. 2014. Metrics, statistics, tests. In *Bridging Between Information Retrieval and Databases - PROMISE Winter School 2013, Revised Tutorial Lectures*, Lecture Notes in Computer Science, pages 116–163. Springer Verlag.
- Zhihong Shao, Peiyi Wang, Qihao Zhu, R. X. Xu, Jun-Mei Song, Mingchuan Zhang, Y. K. Li, Yu Wu, and Daya Guo. 2024. **Deepseekmath: Pushing the limits of mathematical reasoning in open language models**. *ArXiv*, abs/2402.03300.
- Richard Shin, Miltiadis Allamanis, Marc Brockschmidt, and Oleksandr Polozov. 2019. **Program synthesis and semantic parsing with learned code idioms**. In *Neural Information Processing Systems*.
- Disha Shrivastava, Denis Kocetkov, Harm de Vries, Dzmitry Bahdanau, and Torsten Scholak. 2023. Repofusion: Training code models to understand your repository.
- Ishika Singh, Valts Blukis, Arsalan Mousavian, Ankit Goyal, Danfei Xu, Jonathan Tremblay, Dieter Fox, Jesse Thomason, and Animesh Garg. 2023. Prog-prompt: Generating situated robot task plans using large language models. In *2023 IEEE International Conference on Robotics and Automation (ICRA)*, pages 11523–11530. IEEE.

- Dominik Sobania, Martin Briesch, and Franz Rothlauf. 2022. Choose your programming copilot: a comparison of the program synthesis performance of github copilot and genetic programming. In *Proceedings of the genetic and evolutionary computation conference*, pages 1019–1027.
- Qiushi Sun, Zhirui Chen, Fangzhi Xu, Kanzhi Cheng, Chang Ma, Zhangyue Yin, Jianing Wang, Chengcheng Han, Renyu Zhu, Shuai Yuan, Qipeng Guo, Xipeng Qiu, Pengcheng Yin, Xiaoli Li, Fei Yuan, Lingpeng Kong, Xiang Li, and Zhiyong Wu. 2024. [A survey of neural code intelligence: Paradigms, advances and beyond](#).
- Zeyu Sun, Qihao Zhu, Lili Mou, Yingfei Xiong, Ge Li, and Lu Zhang. 2019. A grammar-based structural cnn decoder for code generation. In *Proceedings of the AAAI conference on artificial intelligence*, volume 33, pages 7055–7062.
- Gemma Team, Thomas Mesnard, Cassidy Hardin, Robert Dadashi, Surya Bhupatiraju, Shreya Pathak, Laurent Sifre, Morgane Rivière, Mihir Sanjay Kale, Juliette Love, et al. 2024. Gemma: Open models based on gemini research and technology. *arXiv preprint arXiv:2403.08295*.
- Hugo Touvron, Louis Martin, Kevin Stone, Peter Albert, Amjad Almahairi, Yasmine Babaei, Nikolay Bashlykov, Soumya Batra, Prajjwal Bhargava, Shruti Bhosale, et al. 2023. Llama 2: Open foundation and fine-tuned chat models. *arXiv preprint arXiv:2307.09288*.
- Ngoc Tran, Hieu Tran, Son Nguyen, Hoan Nguyen, and Tien Nguyen. 2019. Does bleu score work for code migration? In *2019 IEEE/ACM 27th International Conference on Program Comprehension (ICPC)*, pages 165–176. IEEE.
- Sergey Troshin and Nadezhda Chirkova. 2022. [Probing pretrained models of source codes](#). In *Proceedings of the Fifth BlackboxNLP Workshop on Analyzing and Interpreting Neural Networks for NLP*, pages 371–383, Abu Dhabi, United Arab Emirates (Hybrid). Association for Computational Linguistics.
- Xin Wang, Yasheng Wang, Fei Mi, Pingyi Zhou, Yao Wan, Xiao Liu, Li Li, Hao Wu, Jin Liu, and Xin Jiang. 2021a. Syncobert: Syntax-guided multi-modal contrastive pre-training for code representation. *arXiv preprint arXiv:2108.04556*.
- Yue Wang, Hung Le, Akhilesh Gotmare, Nghi Bui, Junnan Li, and Steven Hoi. 2023a. Codet5+: Open code large language models for code understanding and generation. In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pages 1069–1088.
- Yue Wang, Weishi Wang, Shafiq R. Joty, and Steven C. H. Hoi. 2021b. [Codet5: Identifier-aware unified pre-trained encoder-decoder models for code understanding and generation](#). *ArXiv*, abs/2109.00859.
- Zhiruo Wang, Grace Cuenca, Shuyan Zhou, Frank F Xu, and Graham Neubig. 2023b. Mconala: A benchmark for code generation from multiple natural languages. In *Findings of the Association for Computational Linguistics: EACL 2023*, pages 265–273.
- Zhiruo Wang, Shuyan Zhou, Daniel Fried, and Graham Neubig. 2022. Execution-based evaluation for open-domain code generation. *arXiv preprint arXiv:2212.10481*.
- Yuxiang Wei, Zhe Wang, Jiawei Liu, Yifeng Ding, and Lingming Zhang. 2023. Magicoder: Source code is all you need. *arXiv preprint arXiv:2312.02120*.
- Yiheng Xu, SU Hongjin, Chen Xing, Boyu Mi, Qian Liu, Weijia Shi, Binyuan Hui, Fan Zhou, Yitao Liu, Tianbao Xie, et al. 2023. Lemur: Harmonizing natural language and code for language agents. In *The Twelfth International Conference on Learning Representations*.
- John Yang, Akshara Prabhakar, Karthik Narasimhan, and Shunyu Yao. 2024. Intercode: Standardizing and benchmarking interactive coding with execution feedback. *Advances in Neural Information Processing Systems*, 36.
- Pengcheng Yin, Bowen Deng, Edgar Chen, Bogdan Vasilescu, and Graham Neubig. 2018. [Learning to mine aligned code and natural language pairs from stack overflow](#). In *International Conference on Mining Software Repositories, MSR*, pages 476–486. ACM.
- Pengcheng Yin, Wen-Ding Li, Kefan Xiao, Abhishek Rao, Yeming Wen, Kensen Shi, Joshua Howland, Paige Bailey, Michele Catasta, Henryk Michalewski, et al. 2023. Natural language to code generation in interactive data science notebooks. In *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 126–173.
- Hao Yu, Bo Shen, Dezhi Ran, Jiaxin Zhang, Qi Zhang, Yuchi Ma, Guangtai Liang, Ying Li, Tao Xie, and Qianxiang Wang. 2023. Codereval: A benchmark of pragmatic code generation with generative pre-trained models. *arXiv preprint arXiv:2302.00288*.
- Fengji Zhang, Bei Chen, Yue Zhang, Jacky Keung, Jin Liu, Daoguang Zan, Yi Mao, Jian-Guang Lou, and Weizhu Chen. 2023. [RepoCoder: Repository-level code completion through iterative retrieval and generation](#). In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pages 2471–2484, Singapore. Association for Computational Linguistics.
- Jian Zhang, Xu Wang, Hongyu Zhang, Hailong Sun, Kaixuan Wang, and Xudong Liu. 2019a. [A novel neural source code representation based on abstract syntax tree](#). In *2019 IEEE/ACM 41st International Conference on Software Engineering (ICSE)*, pages 783–794.

Kechi Zhang, Jia Li, Ge Li, Xianjie Shi, and Zhi Jin. 2024. Codeagent: Enhancing code generation with tool-integrated agent systems for real-world repo-level coding challenges. *arXiv preprint arXiv:2401.07339*.

Tianyi Zhang, Varsha Kishore, Felix Wu, Kilian Weinberger, and Yoav Artzi. bert_score/journal/rescale_baseline.md at master · Tiiiger/bert_score — github.com. https://github.com/Tiiiger/bert_score/blob/master/journal/rescale_baseline.md. [Accessed 26-04-2024].

Tianyi Zhang, Varsha Kishore, Felix Wu, Kilian Q Weinberger, and Yoav Artzi. 2019b. Bertscore: Evaluating text generation with bert. In *International Conference on Learning Representations*.

Qinkai Zheng, Xiao Xia, Xu Zou, Yuxiao Dong, Shan Wang, Yufei Xue, Zihan Wang, Lei Shen, Andi Wang, Yang Li, Teng Su, Zhilin Yang, and Jie Tang. 2023. Codegeex: A pre-trained model for code generation with multilingual evaluations on humaneval-x.

Lily Zhong, Zilong Wang, and Jingbo Shang. 2024. Ldb: A large language model debugger via verifying runtime execution step-by-step.

Andy Zhou, Kai Yan, Michal Shlapentokh-Rothman, Haohan Wang, and Yu-Xiong Wang. 2023. Language agent tree search unifies reasoning acting and planning in language models.

Shuyan Zhou, Uri Alon, Sumit Agarwal, and Graham Neubig. Codebertscore: Evaluating code generation with pretrained models of code.

This is a section in the appendix.