

Personalized Generation In Large Model Era: A Survey

Anonymous ACL submission

Abstract

In the era of large models, content generation is gradually shifting to Personalized Generation (PGen), tailoring content to individual preferences and needs. This paper presents the first comprehensive survey on PGen, investigating existing research in this rapidly growing field. We conceptualize PGen from a unified perspective, systematically formalizing its key components, core objectives, and abstract workflows. Based on this unified perspective, we propose a multi-level taxonomy, offering an in-depth review of technical advancements, commonly used datasets, and evaluation metrics across multiple modalities, personalized contexts, and tasks. Moreover, we envision the potential applications of PGen and highlight open challenges and promising directions for future exploration. By bridging PGen research across multiple modalities, this survey serves as a valuable resource for fostering knowledge sharing and interdisciplinary collaboration, ultimately contributing to a more personalized digital landscape.

1 Introduction

Recent advancements in large generative models have catalyzed a paradigm shift in content generation, moving from generic, one-size-fits-all generation to Personalized Generation (PGen) (Wang et al., 2023c; Xu et al., 2024c; Nguyen et al., 2024b). By crafting personalized content that resonates more deeply with individual preferences, PGen holds great potential to enhance user-centric services and foster more engaging, immersive user experiences across various domains, such as customized product images in e-commerce (Yang et al., 2024a), personalized advertisements in marketing campaigns (Tang et al., 2024a), and personalized AI assistants (Zhang et al., 2024a). Given its significant potential, PGen has attracted significant attention from both academia and industry.

Despite significant progress (Alaluf et al., 2025; Salemi et al., 2024b), research efforts in PGen have largely evolved independently within different communities, such as Natural Language Processing (NLP), Computer Vision (CV), and Information Retrieval (IR). There is no survey that specifically provides a cross-community overview of PGen research. Existing surveys related to PGen separately follow either a model-centric or task-centric perspective, offering only partial summaries of relevant studies. 1) Model-centric surveys focus on specific generative models for personalization, such as Multimodal Large Language Models (MLLMs) (Wu et al., 2024b), Large Language Models (LLMs) (Zhang et al., 2024j; Chen et al., 2024e; Li et al., 2024i), and Diffusion Models (DMs) (Zhang et al., 2024g); 2) Task-centric surveys summarize personalization techniques in specific applications, such as dialogue generation (Chen et al., 2024f), role-playing (Chen et al., 2024d; Tseng et al., 2024), and generative recommendation (Ayemowa et al., 2024). None of these offers a unified framework that comprehensively summarizes PGen research across communities.

A unified framework is crucial for systematically reviewing recent advances and emerging trends in PGen, providing a comprehensive, panoramic view of this field. Moreover, it can foster communication, knowledge sharing, and collaboration between various research communities, ultimately driving the development of a more advanced and personalized digital ecosystem. However, conducting such a unified survey is challenging, as different communities prioritize distinct modalities. For instance, the NLP and IR communities primarily focus on the text modality, while CV specializes in image, video, and 3D. Since each modality presents distinct data structures and challenges, these modality-specific differences introduce inherent technical divergences, making it difficult to unify PGen research into a cohesive framework.

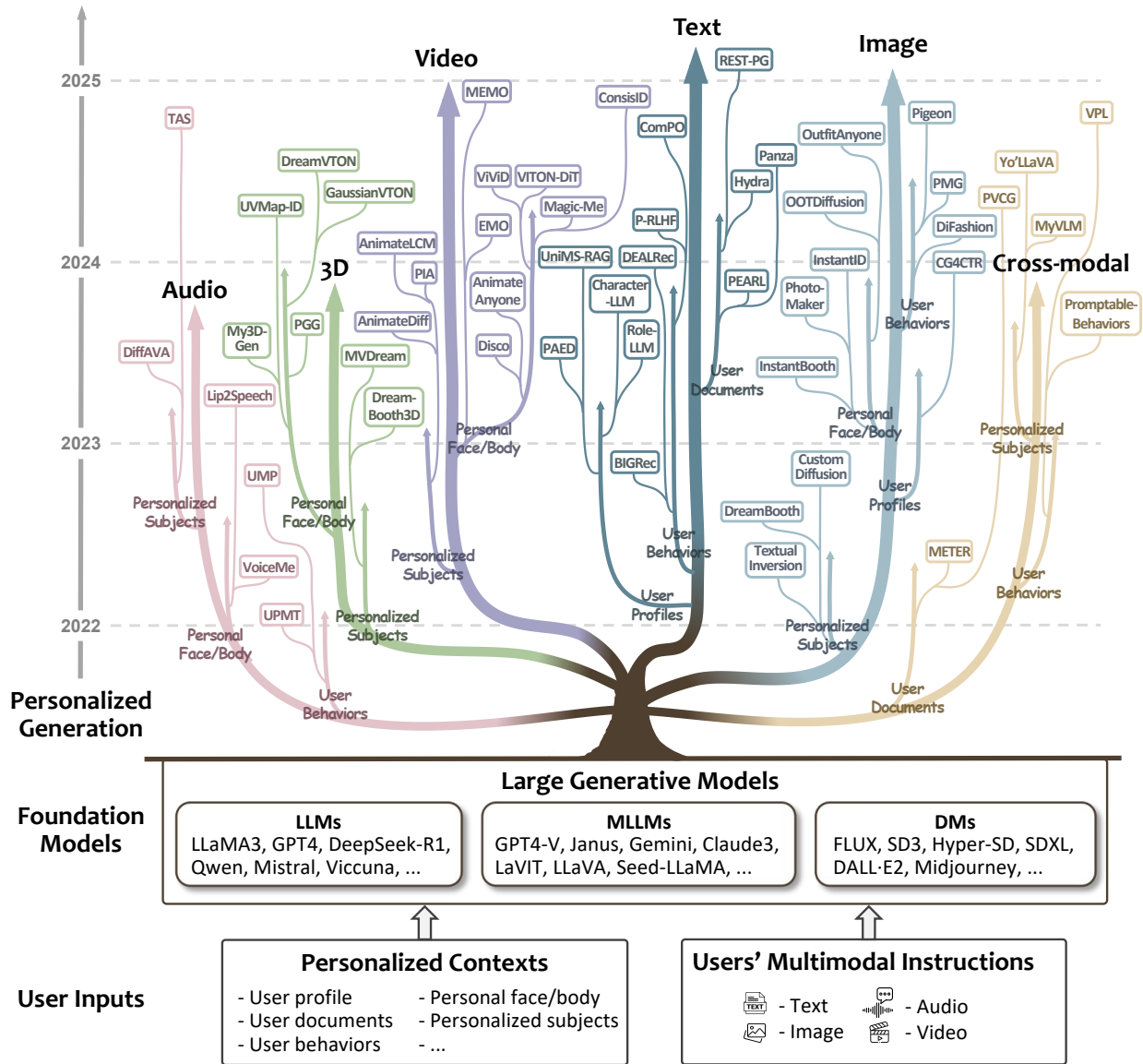


Figure 1: Overview of personalized generation across modalities, inspired by the figure in Yang et al. (2024b).

To address these challenges, it is essential to re-examine PGen from a high-level, modality-agnostic perspective. Fundamentally, PGen entails user modeling based on various personalized contexts and multimodal instructions, extracting personalized signals to guide the content generation process. As illustrated in Figure 1, existing PGen research essentially models various user inputs and leverages generative models for personalized content generation in multiple modalities.

To this end, we present the first survey dedicated to PGen. The structure of this survey and our key contributions are summarized as follows:

- **A unified user-centric perspective for PGen (Section 2).** We conceptualize PGen by formalizing the key components, core objectives, and general workflow, integrating studies across different modalities into a holistic framework.

- **A multi-level taxonomy for existing work in PGen (Section 3).** Building on the unified perspective, we introduce a novel taxonomy that systematically reviews PGen’s technical advancements, commonly used datasets, and evaluation metrics across various modalities, personalized contexts, and tasks.
- **An outlook for potential applications of PGen in enhancing user-centric services (Section 4).** We categorize potential applications of PGen by content personalization stages, with a focus on the content creation and delivery processes.
- **An overview of key open problems in PGen for future research (Section 5).** We outline the critical open problems that need to be addressed to drive innovation in future research and advance the user-centric content ecosystem.

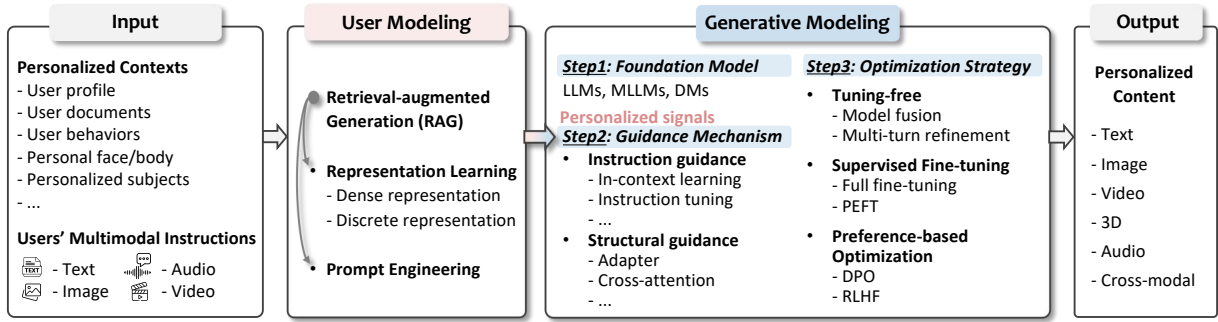


Figure 2: Personalized generation workflow.

2 A Unified User-centric Perspective for Personalized Generation

2.1 Task Formulation

PGen leverages generative models to synthesize content tailored to individual preferences and specific needs. As illustrated in Figure 1, it relies on two essential user inputs: 1) *Personalized contexts* that encapsulate user preferences; 2) *Users' multimodal instructions*, which include textual prompts, voice commands, and other modality-specific inputs that explicitly convey their content needs. Generative models learn user preferences and personal characteristics from diverse personalized contexts and follow users' multimodal instructions to generate customized content across different modalities. The personalized contexts encompass the following dimensions:

- *User profiles*: A collection of demographic and personal attributes associated with a specific user, such as age, gender, occupation, and location.
- *User documents*: User-created textual content, such as comments, emails, and social media posts, that reflects personal creative preferences.
- *User behaviors*: User interactions captured during user engagement, such as searches, clicks, likes, comments, views, shares, and purchases.
- *Personal face/body*: Individual facial and bodily traits, including both static features (*e.g.*, facial structure and body shape) and dynamic features (*e.g.*, expressions, gestures, and motions). These are widely used in tasks like portrait generation, fashion virtual try-ons, and 3D modeling.
- *Personalized subjects*: User-specific concepts or entities, such as pets, personal items, and favorite objects that reflect unique tastes.

By integrating personalized contexts with users' multimodal instructions, generative models can create highly tailored content, aligning closely with individual preferences and addressing specific needs.

2.2 Objectives

Although PGen in each modality is shaped by unique data structures, specific challenges, and distinct tasks, three essential objectives and evaluation dimensions remain consistent across modalities:

- **High quality**: Ensuring the generated content meets high standards of coherence, relevance, and aesthetics.
- **Instruction alignment**: Requiring the generated content to accurately adhere to users' multimodal instructions and effectively address their needs.
- **Personalization**: Guaranteeing that the generated content aligns with personalized contexts and caters to specific user preferences.

While text generation has consistently achieved high-quality outputs, challenges persist in other modalities, such as image, video, audio, and 3D generation. In these domains, generated content can sometimes appear chaotic or disjointed. Maintaining high-quality standards across all modalities is fundamental to achieving successful personalized generation. Furthermore, in certain domains where *factual accuracy* is particularly important, such as news, laws, policies, and expert knowledge, generative models must prioritize authenticity to ensure the reliability and trustworthiness of the content provided to users.

2.3 Workflow

As shown in Figure 2, the PGen workflow involves two key processes: 1) user modeling based on diverse user-specific data and 2) generative modeling across multiple modalities, ensuring high-quality, instruction-aligned, and personalized content.

User Modeling To effectively capture user preferences and specific content needs based on personalized contexts and users' multimodal instructions, three key techniques are commonly employed: 1) Representation learning, which encodes these inputs into dense embeddings (Ruiz et al., 2023; Tang

et al., 2024b) or summarizes them into discrete representations (e.g., texts) (Shen et al., 2024b); 2) Prompt engineering, which involves designing task-specific prompts to structure user-specific information for generative models (Chen et al., 2024g; Li et al., 2025); and 3) Retrieval-augmented generation (RAG), which enriches user-specific information by filtering out irrelevant information and integrating external relevant data (Salemi and Zamani, 2024; Mysore et al., 2024). By combining these techniques, user modeling establishes a robust foundation for PGen, extracting personalized signals to guide content personalization within the generative modeling process.

Generative Modeling To generate personalized content effectively, generative modeling follows a structured three-step process:

- **Step1: Foundation model.** In the era of large models, LLMs, MLLMs, and DMs serve as the backbone for content generation. Selecting an appropriate foundation model based on the target modality, task requirements, and user-specific data is crucial for achieving accurate and personalized content.
- **Step2: Guidance mechanism.** To effectively integrate personalized signals, two primary guidance mechanisms are employed: instruction guidance and structural guidance. Specifically, instruction guidance ensures models follow explicit user prompts and instructions using techniques such as in-context learning (Xu et al., 2023b; Chen et al., 2024g) and instruction tuning (Pi et al., 2024; Xu et al., 2024c). In contrast, structural guidance modifies the model architecture by incorporating additional modules, such as adapters (Ye et al., 2023) and cross-attention mechanisms (Wei et al., 2023), to embed personalized information.
- **Step3: Optimization Strategy.** Empowering large generative models with the capability of personalized generation involves three primary optimization strategies: 1) Tuning-free methods, which utilize pre-trained models for personalized generation without modifying model parameters. These methods often rely on model fusion to assemble multiple pre-trained models (Ding et al., 2024) or employ interactive generation processes that collect real-time user feedback for multi-turn refinement (Von Rütte et al., 2023), ensuring alignment with individual preferences. 2) Supervised fine-tuning which optimizes model pa-

rameters using explicit supervision signals, either through full fine-tuning (Xu et al., 2024b; Ruiz et al., 2023) or Parameter-Efficient Fine-Tuning (PEFT) techniques (Wu et al., 2024f; Tan et al., 2024; Zhang et al., 2024b). 3) Preference-based optimization, which incorporates user preference data to update model parameters. A key approach is Reinforcement Learning with Human Feedback (RLHF) (Li et al., 2024g; Zhang, 2024), which employs an explicit reward model to guide optimization. Alternatively, Direct Preference Optimization (DPO) offers a streamlined solution by directly aligning model parameters with pairwise user preferences (Zhang et al., 2024c; Huang et al., 2024b).

By integrating these advanced techniques and strategies, this workflow not only ensures adaptability to diverse personalized contexts and user instructions but also highlights the evolving landscape of large generative models, offering a scalable solution for PGen.

3 Personalized Generation Across Modalities

Based on the unified perspective, we present a multi-level taxonomy for PGen, systematically reviewing PGen research across various modalities, personalized contexts, and specific tasks, as outlined in Table 1. Studies on PGen are first categorized by modality, including text, image, video, audio, 3D, and cross-modal generation. Within each modality, we further classify research based on personalized contexts and examine corresponding tasks and techniques. A detailed review of image, video, and 3D modalities is provided in Appendix A. Additionally, we provide a comprehensive overview of commonly used evaluation metrics and datasets for each modality in Appendix B and summarize them in Table 2, Table 3, and Table 4.

3.1 Personalized Text Generation

Personalized text generation aims to provide textual content tailored to user preferences and needs, involving tasks ranging from information seeking to user simulation.

3.1.1 User Behaviors

User interactions with the system typically occur over time (Kelly and Belkin, 2002), allowing it to learn implicit preferences and behavioral patterns to enhance personalization and encourage

long-term engagement (Shi et al., 2013). This personalized context is valuable for the following personalized text-based tasks.

Information Seeking A primary use case of personalized text generation is crafting responses to align with user preferences, enabling more engaging interactions. The system can leverage user feedback (e.g., thumbs up/down and selected best responses) to tailor its responses to user preferences. While personalization has been extensively studied in the context of information access and search (Kasela et al., 2024; Zeng et al., 2023; Eugene et al., 2013; Guo et al., 2021), which aims to select a response from predefined options, it remains relatively underexplored in generative scenarios. This is largely due to the lack of standardized metrics and benchmarks. However, recently, Li et al. (2024g) explored the use of personalized feedback to train LLMs to generate more tailored summaries for users, as a form of information seeking. Kumar et al. (2024b) extends this approach by collecting preference feedback from a group of users as a community to optimize the model’s ability to respond to their collective information needs.

Recommendation While recommendation is not directly involved in content generation, it plays a crucial role in delivering personalized content, which has been explored across various scenarios (Hou et al., 2024; Harper and Konstan, 2015; Wan and McAuley, 2018; Wu et al., 2020a). The use of generative models in Recommendation Systems (RecSys) has been widely studied (Bao et al., 2023; Wu et al., 2024c; Yang et al., 2023a). Specifically, LLMs have been utilized either through prompting (Lyu et al., 2024) or by training them directly to perform recommendation tasks as a form of text generation (Lin et al., 2024a). Beyond transformer-based generative models, newer approaches like diffusion models have also been explored for recommendation, highlighting the versatility of generative methods in this domain (Wang et al., 2023e; Yang et al., 2024e).

Other work has also leveraged realistic user interaction to explore personalization for review generation (Ni et al., 2019; Li et al., 2020; Sun et al., 2020; Li et al., 2021) and news headline generation (Ao et al., 2021; Cai et al., 2023; Song et al., 2023). For example, Ao et al. (2021) presents a personalized headline generation benchmark by collecting user’s click history on Microsoft News.

3.1.2 User Documents

In some cases, users may not interact with the system frequently but can provide valuable information for personalization, such as user-created documents (Salemi et al., 2024b).

Writing Assistant Personalization plays a critical role in enhancing text-based writing assistants, enabling tailored text generation across diverse formats and styles. For this purpose, the LaMP benchmark (Salemi et al., 2024b,a; Salemi and Zamani, 2024; Zhuang et al., 2024) focuses on short-form text generation, such as creating headlines for news articles or subject lines for emails. In contrast, the LongLaMP benchmark (Kumar et al., 2024a) targets longer-form tasks, such as writing product reviews from a user’s perspective based on a rating, generating a post from its summary, or completing an email for a user (Salemi et al., 2025b). Additionally, the Personalized Linguistic Attributes Benchmark (Alhafni et al., 2024) is designed for text completion tasks, such as extending a post or review, leveraging user-written documents to extract stylistic features that guide LLMs in mimicking a user’s unique writing style. In this domain, RAG is the dominant approach, retrieving relevant information from users’ historically written documents to extract their writing preferences (Mysore et al., 2024; Nicolicioiu et al., 2024; Li et al., 2023a).

3.1.3 User Profiles

Generative models can infer user preferences from their profiles to guide personalized text generation.

Dialogue System In recent years, chatbots and conversational systems have been a central focus of personalized text generation (Kottur et al., 2017; Thompson et al., 2004; Kaiss et al., 2023; Kocaballi et al., 2019). The advent of advanced chat models like GPT-4 and Gemini have enabled more sophisticated and personalized interactions. By defining a persona or personality for these models based on user profiles, the system can tailor responses to individual preferences and needs (Pradhan and Lazar, 2021; Song et al., 2019). To support research in this domain, various dialogue datasets have been developed. For example, LiveChat (Gao et al., 2023) is a large-scale dataset created from live streaming interactions, FoCUS (Jang et al., 2021) focuses on conversational information-seeking scenarios, and Pchatbot (Qian et al., 2021) compiles data from Weibo and judicial forums. Enhancing LLMs’ ability to generate personalized responses involves ap-

proaches such as zero-shot prompting (Zhu et al., 2023b), in-context learning (Xu et al., 2023c), and training on limited persona datasets (Song et al., 2021; Wang et al., 2024d) or large-scale datasets (Zheng et al., 2020; Chen et al., 2023b). Additionally, chain-of-thought reasoning has proven effective in improving alignment with user preferences in dialogue systems (Li et al., 2025).

User Simulation Previous research demonstrates that LLMs excel at performing roles or personas assigned to them (Shanahan et al., 2023; Chen et al., 2024c), enabling user simulation based on their profiles to extract preferences and further personalize the system for them (Magee et al., 2024; Santurkar et al., 2023). In this domain, Shao et al. (2023) develops a test playground to interview trained agents and assess whether they effectively memorize their assigned characters and experiences. Additionally, Wang et al. (2024e) introduced a large-scale dataset for evaluating the role-playing abilities of LLMs, while Ng et al. (2024) presented a dataset specifically designed for assessing these capabilities within video game contexts.

3.2 Personalized Audio Generation

Personalized audio generation extracts users’ auditory preferences to create tailored audio content, such as music and speech.

3.2.1 User behaviors

User behaviors on music, such as listening history and ratings, are important clues for inferring user preferences for personalization.

Music Generation Methods like UMP (Ma et al., 2022) and UP-Transformer (Hu et al., 2022) infer user preferences by analyzing listening histories and ratings. UIGAN (Wang et al., 2024j) adopts an interactive approach to collect user feedback, progressively refining the generated music. Ma et al. (2024b) construct a music community graph, where nodes represent users and songs, and edges capture relationships such as likes, subscriptions, and user interactions.

3.2.2 Personalized Subjects

In some cases, users provide audio clips and aim to manipulate them using textual prompts.

Text-to-Audio Generation Personalized text-to-audio generation explores methods for synthesizing tailored audio by aligning user preferences, textual descriptions, and contextual inputs. Mo

et al. (2023) utilize Contrastive Language-Audio Pretraining (CLAP) to align features across audio and visual-text modalities in the spatiotemporal domain. Furthermore, Plitsis et al. (2024) adapt image-domain personalization techniques like Textual Inversion (Gal et al., 2023) and Dream-Booth (Ruiz et al., 2023) to audio tasks, enhancing user-centric generation. Li et al. (2024j) introduce a Semantic-Aware Fusion (SAF) module to capture text-aware audio features, establishing nuanced perceptual relationships between text and audio inputs for more contextually aligned audio outputs.

3.2.3 Personal Face/Body

Users may provide their facial images or videos, enabling generative models to extract speaker-specific attributes for customized speech generation.

Face-to-Speech Generation VioceMe (van Rijn et al., 2022) employs SpeakerNet to derive speaker embeddings and incorporates Global Style Tokens (GST) for modeling speech styles. FR-PSS (Wang et al., 2022) enhances the extraction of speech features from facial characteristics using a residual guidance strategy, which integrates prior speech information for improved efficiency. Lip2Speech (Sheng et al., 2023) leverages a Variational Autoencoder (VAE) framework to disentangle speaker identity from linguistic content in videos, achieving fine-grained control over personalized speech synthesis via speaker embeddings.

3.3 Personalized Cross-modal Generation

Personalized cross-modal generation primarily aims to produce personalized textual responses (e.g., captions, answers, or robot actions) based on multimodal personalized contexts (e.g., images, videos, historical robot trajectories).

3.3.1 User behaviors

Based on user interactions, generative models can infer user preferences to tailor robotic behaviors.

Robotics Several studies have investigated inferring user preferences from historical trajectories and associated human feedback, enabling personalized robotic decision-making. Specifically, Poddar et al. (2024) proposed Variational Preference Learning, integrating variational inference into RLHF to model diverse user preferences and enable active learning. Hwang et al. (2024) introduced “Promptable Behaviors”, using Multi-Objective Reinforcement Learning to dynamically customize policies

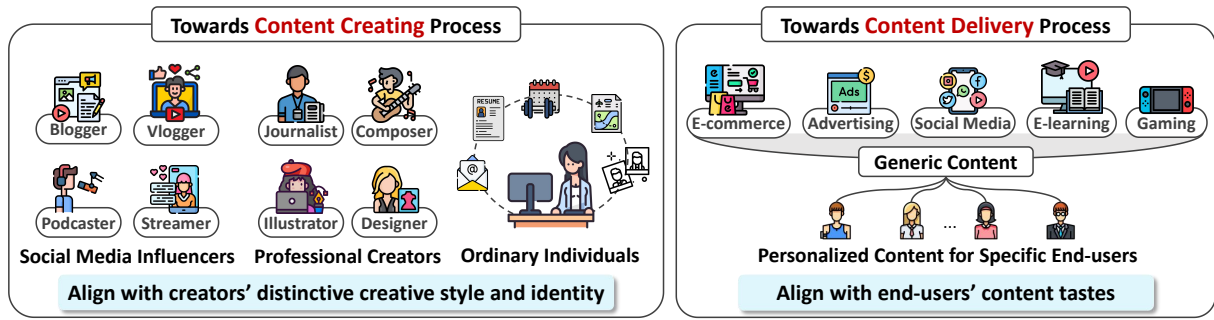


Figure 3: Applications of personalized generation.

via human demonstrations, trajectory feedback, and language instructions. Cui et al. (2024b) presented a framework with a RAG memory module to match individual users’ preferences and driving styles.

3.3.2 User Documents

User-created documents, such as comments, reviews, and captions can be used to infer their writing style and preferences for personalization.

Caption/Comment Generation Several studies utilize user-created captions and comments to develop personalized captioning and comment systems (Long et al., 2020; Zhang et al., 2020c; Geng et al., 2022). To effectively model user preferences, Shin et al. (2018) engages users by soliciting answers to targeted questions during generation. Lin et al. (2024b); Wu et al. (2024e) utilize users’ historical comments for personalized video comment generation. In addition, PV-LLM (Lin et al., 2024b) fine-tunes MLLMs using user-written comments, while PVCG (Wu et al., 2024e) learns a unique identifier for each user to enhance personalization.

3.3.3 Personalized subjects

In some cases, users may provide specific subject images, such as photos of their friends, for personalized visual question answering (e.g., “Give me a birthday gift list for my friend Peter.”).

Cross-modal Dialogue Systems Given user-specific subject images and queries, the systems are expected to identify these subjects and infer user intents for personalized responses. Existing studies in this domain mainly fall into two groups:

- Memory-based methods store user-specific subjects and activate or retrieve them as needed. For example, CSMN (Chunseong Park et al., 2017; Park et al., 2018) stores image memory, user context memory, and word output memory, and the model generates personalized captions based on memory features. Hao et al. (2024) leverages RAG techniques to personalize LMMs, allowing

LLMs to update their supported concepts without requiring additional training. PLVM (Pham et al., 2024) proposes a pre-trained Aligner to align referential concepts with the queried images.

- Optimization-based methods inject personalized information into the generation process via specific modules. For example, PerVL (Cohen et al., 2022), Yo’LLaVA (Nguyen et al., 2024b), and MC-LLaVA (An et al., 2024) learn to encode a personalized subject into a set of latent tokens based on several provided subject images. MyVLM (Alaluf et al., 2025) introduces learnable heads, each dedicated to recognizing a single user-specific subject. In addition, several studies have explored learning unique user embeddings Long et al. (2020); Zhang et al. (2020c); Xiong et al. (2020) or performing prefix tuning (Wang et al., 2023f) for personalization.

4 Applications

The prior section has highlighted the success of PGen across various modalities, demonstrating its potential to enhance user engagement and enrich experiences across diverse domains. As illustrated in Figure 3, applications of PGen can be categorized based on the stages of content personalization: 1) towards content creation process, which provides personalized tools and services for content creators of all levels to maintain their unique creative style while streamlining the creative workflows; and 2) towards content delivery process, which delivers multimodal content in a personalized manner tailored to individual preferences of end users. A more detailed discussion of PGen applications can be found in Appendix C.

5 Open Problems

Despite the significant progress made by PGen, several key challenges remain to be addressed.

5.1 Technical Challenges

- **Scalability and Efficiency.** PGen relies on large generative models for content personalization, which often require extensive resource costs, limiting their deployment to real-time, large-scale user scenarios. Developing scalable and efficient algorithms for PGen remains a critical direction for future research (Yang et al., 2024c).
- **Deliberative Reasoning for PGen.** In certain personalized scenarios that prioritize content quality over temporal efficiency – such as digital advertising, where advertisers typically serve only several ads per user each day – inference scaling presents significant opportunities for enhancing user satisfaction. Prior work mainly focuses on multi-turn refinement to progressively enhance content relevance and personalization (Nabati et al., 2024). Inspired by the great success of LLM reasoning (Guan et al., 2024; Guo et al., 2025), deliberative reasoning may emphasize extensive logical and contextual reasoning beforehand, enabling a thorough analysis of user preferences to drive more effective content personalization (Fang et al., 2025).
- **Evolving User Preference.** As explored in traditional RecSys (Wang et al., 2023d), recognizing dynamic preferences from user behaviors is crucial to enhancing personalization. Adapting PGen to track and respond to user preference shifts remains a key research direction.
- **Multi-modal Personalization.** Existing PGen research mainly focuses on single-modality generation, while multi-modal personalization remains underexplored, such as personalized social media posts that integrate both image and text. This challenge requires high-quality, instruction-aligned, and personalized output while ensuring consistency across multiple modalities.
- **Synergy Between Generation and Retrieval.** Traditional personalized content delivery systems primarily focus on retrieval-based methods like RecSys. However, existing content may not fully meet users’ content needs. Integrating PGen with retrieval-based approaches holds great promise for building more powerful personalized content delivery systems (Wang et al., 2023c).

5.2 Benchmarks and Metrics

A fundamental challenge in PGen is establishing robust metrics and benchmark datasets. Current evaluation methods primarily rely on traditional

generation metrics (e.g., BLEU for text, CLIP-I for images), which do not fully capture how well the generated content aligns with user preferences. Future research should focus on developing more effective metrics for evaluating personalization.

5.3 Trustworthiness

Ensuring the trustworthiness of PGen is critical for fostering user confidence and promoting responsible deployment. Key considerations include:

- **Privacy.** PGen relies on user-specific data for content personalization, raising concerns about privacy issues. Striking a balance between effective personalization and robust privacy protection is crucial for advancing this field.
- **Fairness and Bias.** PGen may unintentionally reinforce biases and stereotypes present in training data, leading to skewed or discriminatory outcomes. Effective detection and mitigation strategies are essential to protect diverse user groups.
- **Safety.** Establishing transparent governance protocols, reliable moderation mechanisms, and explainable generation processes is crucial to maintaining user trust and upholding safety standards.

6 Conclusion

In this work, we present the first comprehensive survey on PGen across multiple modalities, offering an in-depth review of recent advancements and emerging trends in the field. To unify existing research, we introduce a holistic framework that formalizes diverse user-specific data, core objectives, and general workflows for PGen, providing a structured foundation for future developments. We then propose a multi-level taxonomy to categorize PGen methods based on modality, user inputs, and specific tasks. Additionally, we summarize the commonly used datasets and evaluation metrics for each modality. Beyond technical aspects, we explore PGen’s potential applications in both content creation and delivery, underscoring its significant economic and research value. Finally, we identify key research challenges that remain to be addressed. As a rapidly evolving field, PGen holds great potential to revolutionize the online content ecosystem, enabling more tailored and engaging user experiences. By unifying PGen research across multiple modalities, this survey serves as a valuable resource for fostering cross-modal knowledge sharing and collaboration in this field, contributing to a more personalized digital landscape.

Limitations

In this paper, we provide a comprehensive survey of personalized generation. However, the rapid evolution of this field makes it challenging to encompass all research efforts, as new methods, datasets, and evaluation metrics continue to emerge, requiring continuous updates to our taxonomy. Furthermore, the development of more effective and universally accepted benchmarks within different modalities remains an ongoing challenge.

References

- Rameen Abdal, Hsin-Ying Lee, Peihao Zhu, Menglei Chai, Aliaksandr Siarohin, Peter Wonka, and Sergey Tulyakov. 2023. 3davatar: Bridging domains for personalized editable avatars. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 4552–4562.
- Panos Achlioptas, Alexandros Benetatos, Iordanis Fostiropoulos, and Dimitris Skourtis. 2023. Stellar: Systematic evaluation of human-centric personalized text-to-image methods. *arXiv preprint arXiv:2312.06116*.
- Yuval Alaluf, Elad Richardson, Gal Metzer, and Daniel Cohen-Or. 2023. A neural space-time representation for text-to-image personalization. *ACM Transactions on Graphics (TOG)*, 42(6):1–10.
- Yuval Alaluf, Elad Richardson, Sergey Tulyakov, Kfir Aberman, and Daniel Cohen-Or. 2025. Myvlm: Personalizing vlms for user-specific queries. In *European Conference on Computer Vision*, pages 73–91. Springer.
- Bashar Alhafni, Vivek Kulkarni, Dhruv Kumar, and Vipul Raheja. 2024. [Personalized text generation with fine-grained linguistic control](#). In *Proceedings of the 1st Workshop on Personalization of Generative AI Systems (PERSONALIZE 2024)*, pages 88–101, St. Julians, Malta. Association for Computational Linguistics.
- Ruichuan An, Sihan Yang, Ming Lu, Kai Zeng, Yulin Luo, Ying Chen, Jiajun Cao, Hao Liang, Qi She, Shanghang Zhang, et al. 2024. Mc-llava: Multi-concept personalized vision-language model. *arXiv preprint arXiv:2411.11706*.
- Xiang Ao, Xiting Wang, Ling Luo, Ying Qiao, Qing He, and Xing Xie. 2021. Pens: A dataset and generic framework for personalized news headline generation. In *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)*, pages 82–92.

- Omri Avrahami, Kfir Aberman, Ohad Fried, Daniel Cohen-Or, and Dani Lischinski. 2023. Break-a-scene: Extracting multiple concepts from a single image. In *SIGGRAPH Asia 2023 Conference Papers*, pages 1–12.
- Matthew O Ayemowa, Roliana Ibrahim, and Muhammad Murad Khan. 2024. Analysis of recommender system using generative artificial intelligence: A systematic literature review. *IEEE Access*.
- Max Bain, Arsha Nagraani, Gül Varol, and Andrew Zisserman. 2021. Frozen in time: A joint video and image encoder for end-to-end retrieval. In *Proceedings of the IEEE/CVF international conference on computer vision*, pages 1728–1738.
- Yogesh Balaji, Martin Renqiang Min, Bing Bai, Rama Chellappa, and Hans Peter Graf. 2019. Conditional gan with discriminative filter generation for text-to-video synthesis. In *IJCAI*, volume 1, page 2.
- Alberto Baldrati, Davide Morelli, Giuseppe Cartella, Marcella Cornia, Marco Bertini, and Rita Cucchiara. 2023. Multimodal garment designer: Human-centric latent diffusion models for fashion image editing. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pages 23393–23402.
- 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.
- Keqin Bao, Jizhi Zhang, Wenjie Wang, Yang Zhang, Zhengyi Yang, Yancheng Luo, Chong Chen, Fuli Feng, and Qi Tian. 2023. [A bi-step grounding paradigm for large language models in recommendation systems](#). *Preprint*, arXiv:2308.08434.
- Ahmet Canberk Baykal, Abdul Basit Anees, Duygu Ceylan, Erkut Erdem, Aykut Erdem, and Deniz Yuret. 2023. Clip-guided stylegan inversion for text-driven real image editing. *ACM Transactions on Graphics*, 42(5):1–18.
- David Beniaguev. 2022. [Synthetic faces high quality \(sfhq\) dataset](#).
- Thierry Bertin-Mahieux, Daniel P.W. Ellis, Brian Whitman, and Paul Lamere. 2011. The million song dataset. In *Proceedings of the 12th International Conference on Music Information Retrieval (ISMIR 2011)*.
- Mikołaj Bińkowski, Danica J Sutherland, Michael Arbel, and Arthur Gretton. 2018. Demystifying mmd gans. *arXiv preprint arXiv:1801.01401*.
- M Akmal Butt and Petros Maragos. 1998. Optimum design of chamfer distance transforms. *IEEE Transactions on Image Processing*, 7(10):1477–1484.

771	Pengshan Cai, Kaiqiang Song, Sangwoo Cho, Hongwei	Hong Chen, Xin Wang, Guanning Zeng, Yipeng Zhang,	827
772	Wang, Xiaoyang Wang, Hong Yu, Fei Liu, and Dong	Yuwei Zhou, Feilin Han, and Wenwu Zhu. 2023a.	828
773	Yu. 2023. Generating user-engaging news headlines.	Videodreamer: Customized multi-subject text-to-	829
774	In <i>Proceedings of the 61st Annual Meeting of the</i>	video generation with disen-mix finetuning. <i>arXiv</i>	830
775	<i>Association for Computational Linguistics (Volume</i>	<i>preprint arXiv:2311.00990</i> .	831
776	<i>1: Long Papers)</i> , pages 3265–3280.		
777	Yufei Cai, Yuxiang Wei, Zhilong Ji, Jinfeng Bai,	Hong Chen, Xin Wang, Yipeng Zhang, Yuwei Zhou,	832
778	Hu Han, and Wangmeng Zuo. 2024. Decoupled tex-	Zeyang Zhang, Siao Tang, and Wenwu Zhu. 2024b.	833
779	tual embeddings for customized image generation.	Disenstudio: Customized multi-subject text-to-video	834
780	In <i>Proceedings of the AAAI Conference on Artificial</i>	generation with disentangled spatial control. In <i>Pro-</i>	835
781	<i>Intelligence</i> , volume 38, pages 909–917.	<i>ceedings of the 32nd ACM International Conference</i>	836
		<i>on Multimedia</i> , pages 3637–3646.	837
782	Qiong Cao, Li Shen, Weidi Xie, Omkar M Parkhi, and	Jiangjie Chen, Xintao Wang, Rui Xu, Siyu Yuan, Yikai	838
783	Andrew Zisserman. 2018. Vggface2: A dataset for	Zhang, Wei Shi, Jian Xie, Shuang Li, Ruihan Yang,	839
784	recognising faces across pose and age. In <i>2018 13th</i>	Tinghui Zhu, Aili Chen, Nianqi Li, Lida Chen, Caiyu	840
785	<i>IEEE international conference on automatic face &</i>	Hu, Siye Wu, Scott Ren, Ziquan Fu, and Yanghua	841
786	<i>gesture recognition (FG 2018)</i> , pages 67–74. IEEE.	Xiao. 2024c. From persona to personalization: A sur-	842
		vey on role-playing language agents . <i>Transactions on</i>	843
787	Mathilde Caron, Hugo Touvron, Ishan Misra, Hervé	<i>Machine Learning Research</i> . Survey Certification.	844
788	Jégou, Julien Mairal, Piotr Bojanowski, and Armand		
789	Joulin. 2021. Emerging properties in self-supervised	Jiangjie Chen, Xintao Wang, Rui Xu, Siyu Yuan, Yikai	845
790	vision transformers. In <i>Proceedings of the IEEE/CVF</i>	Zhang, Wei Shi, Jian Xie, Shuang Li, Ruihan Yang,	846
791	<i>international conference on computer vision</i> , pages	Tinghui Zhu, et al. 2024d. From persona to person-	847
792	9650–9660.	alization: A survey on role-playing language agents.	848
		<i>arXiv preprint arXiv:2404.18231</i> .	849
793	Andrés Casado-Elvira, Marc Comino Trinidad, and Dan	Jin Chen, Zheng Liu, Xu Huang, Chenwang Wu, Qi Liu,	850
794	Casas. 2022. Pergamo: Personalized 3d garments	Gangwei Jiang, Yuanhao Pu, Yuxuan Lei, Xiaolong	851
795	from monocular video. In <i>Computer Graphics Fo-</i>	Chen, Xingmei Wang, et al. 2024e. When large	852
796	<i>rum</i> , volume 41, pages 293–304. Wiley Online Li-	language models meet personalization: Perspectives	853
797	brary.	of challenges and opportunities. <i>World Wide Web</i> ,	854
		27(4):42.	855
798	Francesco Paolo Casale, Adrian Dalca, Luca Saglietti,	Lele Chen, Zhiheng Li, Ross K Maddox, Zhiyao Duan,	856
799	Jennifer Listgarten, and Nicolo Fusi. 2018. Gaussian	and Chenliang Xu. 2018. Lip movements generation	857
800	process prior variational autoencoders. <i>Advances in</i>	at a glance. In <i>Proceedings of the European confer-</i>	858
801	<i>neural information processing systems</i> , 31.	<i>ence on computer vision (ECCV)</i> , pages 520–535.	859
802	Daewon Chae, Nokyoung Park, Jinkyu Kim, and Kimin	Liang Chen, Hongru Wang, Yang Deng, Wai Chung	860
803	Lee. 2023. Instructbooth: Instruction-following per-	Kwan, Zezhong Wang, and Kam-Fai Wong. 2023b.	861
804	sonalized text-to-image generation. <i>arXiv preprint</i>	Towards robust personalized dialogue generation via	862
805	<i>arXiv:2312.03011</i> .	order-insensitive representation regularization . In	863
		<i>Findings of the Association for Computational Lin-</i>	864
806	Caroline Chan, Shiry Ginosar, Tinghui Zhou, and	<i>guistics: ACL 2023</i> , pages 7337–7345, Toronto,	865
807	Alexei A Efros. 2019. Everybody dance now. In	Canada. Association for Computational Linguistics.	866
808	<i>Proceedings of the IEEE/CVF international confer-</i>		
809	<i>ence on computer vision</i> , pages 5933–5942.	Peng Chen, Xiaobao Wei, Ming Lu, Yitong Zhu,	867
810	Di Chang, Yichun Shi, Quankai Gao, Hongyi Xu, Jes-	Naiming Yao, Xingyu Xiao, and Hui Chen. 2023c.	868
811	sica Fu, Guoxian Song, Qing Yan, Yizhe Zhu, Xiao	Diffusionaltalker: Personalization and acceleration	869
812	Yang, and Mohammad Soleymani. 2023. Magicpose:	for speech-driven 3d face diffuser. <i>arXiv preprint</i>	870
813	Realistic human poses and facial expressions retar-	<i>arXiv:2311.16565</i> .	871
814	geting with identity-aware diffusion. In <i>Forty-first</i>		
815	<i>International Conference on Machine Learning</i> .	Wen Chen, Pipei Huang, Jiaming Xu, Xin Guo, Cheng	872
816	Hila Chefer, Shiran Zada, Roni Paiss, Ariel Ephrat,	Guo, Fei Sun, Chao Li, Andreas Pfadler, Huan Zhao,	873
817	Omer Tov, Michael Rubinstein, Lior Wolf, Tali Dekel,	and Binqiang Zhao. 2019. Pog: personalized outfit	874
818	Tomer Michaeli, and Inbar Mosseri. 2024. Still-	generation for fashion recommendation at alibaba	875
819	moving: Customized video generation without cus-	ifashion. In <i>Proceedings of the 25th ACM SIGKDD</i>	876
820	tomized video data. <i>ACM Transactions on Graphics</i>	<i>international conference on knowledge discovery &</i>	877
821	<i>(TOG)</i> , 43(6):1–11.	<i>data mining</i> , pages 2662–2670.	878
822	Chaofeng Chen, Jiadi Mo, Jingwen Hou, Haoning Wu,	Wenhu Chen, Hexiang Hu, YANDONG LI, Nataniel	879
823	Liang Liao, Wenxiu Sun, Qiong Yan, and Weisi Lin.	Ruiz, Xuhui Jia, Ming-Wei Chang, and William W.	880
824	2024a. Topiq: A top-down approach from semantics	Cohen. 2023d. Subject-driven text-to-image gener-	881
825	to distortions for image quality assessment. <i>IEEE</i>	ation via apprenticeship learning . In <i>Thirty-seventh</i>	882
826	<i>Transactions on Image Processing</i> .	<i>Conference on Neural Information Processing Sys-</i>	883
		<i>tems</i> .	884

998	Ganggui Ding, Canyu Zhao, Wen Wang, Zhen Yang,	2023. Character-based outfit generation with vision-	1052
999	Zide Liu, Hao Chen, and Chunhua Shen. 2024.	augmented style extraction via llms. In <i>2023 IEEE In-</i>	1053
1000	Freecustom: Tuning-free customized image genera-	<i>ternational Conference on Big Data (BigData)</i> , pages	1054
1001	tion for multi-concept composition. In <i>Proceedings</i>	1–7. IEEE.	1055
1002	<i>of the IEEE/CVF Conference on Computer Vision</i>		
1003	<i>and Pattern Recognition</i> , pages 9089–9098.		
1004	Haoye Dong, Xiaodan Liang, Xiaohui Shen, Bochao	Jianglin Fu, Shikai Li, Yuming Jiang, Kwan-Yee Lin,	1056
1005	Wang, Hanjiang Lai, Jia Zhu, Zhiting Hu, and Jian	Chen Qian, Chen Change Loy, Wayne Wu, and Ziwei	1057
1006	Yin. 2019a. Towards multi-pose guided virtual try-on	Liu. 2022. Stylegan-human: A data-centric odyssey	1058
1007	network. In <i>Proceedings of the IEEE/CVF interna-</i>	of human generation. In <i>European Conference on</i>	1059
1008	<i>tional conference on computer vision</i> , pages 9026–	<i>Computer Vision</i> , pages 1–19. Springer.	1060
1009	9035.		
1010	Haoye Dong, Xiaodan Liang, Xiaohui Shen, Bowen Wu,	Justin Fu, Aviral Kumar, Ofir Nachum, George Tucker,	1061
1011	Bing-Cheng Chen, and Jian Yin. 2019b. Fw-gan:	and Sergey Levine. 2020. D4rl: Datasets for deep	1062
1012	Flow-navigated warping gan for video virtual try-	data-driven reinforcement learning. <i>arXiv preprint</i>	1063
1013	on. In <i>Proceedings of the IEEE/CVF international</i>	<i>arXiv:2004.07219</i> .	1064
1014	<i>conference on computer vision</i> , pages 1161–1170.		
1015	Haoye Dong, Jun Liu, and Dong Huang. 2024. Df-vton:	Stephanie Fu, Netanel Yakir Tamir, Shobhita Sundaram,	1065
1016	Dense flow guided virtual try-on network. In <i>ICASSP</i>	Lucy Chai, Richard Zhang, Tali Dekel, and Phillip	1066
1017	<i>2024-2024 IEEE International Conference on Acous-</i>	Isola. 2023. Dreamsim: Learning new dimensions	1067
1018	<i>tics, Speech and Signal Processing (ICASSP)</i> , pages	of human visual similarity using synthetic data . In	1068
1019	3175–3179. IEEE.	<i>Thirty-seventh Conference on Neural Information</i>	1069
		<i>Processing Systems</i> .	1070
1020	Benjamin Elizalde, Soham Deshmukh, Mahmoud Al Is-	Rinon Gal, Yuval Alaluf, Yuval Atzmon, Or Patash-	1071
1021	mail, and Huaming Wang. 2023. Clap learning	nik, Amit Haim Bermano, Gal Chechik, and Daniel	1072
1022	audio concepts from natural language supervision.	Cohen-or. 2023. An image is worth one word: Per-	1073
1023	In <i>ICASSP 2023-2023 IEEE International Confer-</i>	sonalizing text-to-image generation using textual in-	1074
1024	<i>ence on Acoustics, Speech and Signal Processing</i>	version. In <i>The Eleventh International Conference</i>	1075
1025	<i>(ICASSP)</i> , pages 1–5. IEEE.	<i>on Learning Representations</i> .	1076
1026	Eugene, serdyukovpv, and Will Cukierski.	Rinon Gal, Or Lichter, Elad Richardson, Or Patashnik,	1077
1027	2013. Personalized web search chal-	Amit H Bermano, Gal Chechik, and Daniel Cohen-	1078
1028	lenge. https://kaggle.com/competitions/	Or. 2024. Lcm-lookahead for encoder-based text-to-	1079
1029	yandex-personalized-web-search-challenge .	image personalization. In <i>European Conference on</i>	1080
1030	Kaggle.	<i>Computer Vision</i> , pages 322–340. Springer.	1081
1031	Gabriele Fanelli, Matthias Dantone, Juergen Gall, An-	Hanan Gani, Shariq Farooq Bhat, Muzammal Naseer,	1082
1032	drea Fossati, and Luc Van Gool. 2013. Random	Salman Khan, and Peter Wonka. 2024. LLM	1083
1033	forests for real time 3d face analysis. <i>International</i>	blueprint: Enabling text-to-image generation with	1084
1034	<i>journal of computer vision</i> , 101:437–458.	complex and detailed prompts . In <i>The Twelfth Inter-</i>	1085
		<i>national Conference on Learning Representations</i> .	1086
1035	Yi Fang, Wenjie Wang, Yang Zhang, Fengbin Zhu, Qi-	Jingsheng Gao, Yixin Lian, Ziyi Zhou, Yuzhuo Fu, and	1087
1036	fan Wang, Fuli Feng, and Xiangnan He. 2025. Large	Baoyuan Wang. 2023. LiveChat: A large-scale per-	1088
1037	language models for recommendation with deliber-	sonalized dialogue dataset automatically constructed	1089
1038	ative user preference alignment. <i>arXiv preprint</i>	from live streaming . In <i>Proceedings of the 61st An-</i>	1090
1039	<i>arXiv:2502.02061</i> .	<i>nual Meeting of the Association for Computational</i>	1091
1040	Zhouxiang Fang, Min Yu, Zhendong Fu, Boning Zhang,	<i>Linguistics (Volume 1: Long Papers)</i> , pages 15387–	1092
1041	Xuanwen Huang, Xiaoqi Tang, and Yang Yang.	15405, Toronto, Canada. Association for Computa-	1093
1042	2024a. How to generate popular post headlines on	tional Linguistics.	1094
1043	social media? <i>AI Open</i> , 5:1–9.		
1044	Zixun Fang, Wei Zhai, Aimin Su, Hongliang Song,	Xuan Gao, Chenglai Zhong, Jun Xiang, Yang Hong,	1095
1045	Kai Zhu, Mao Wang, Yu Chen, Zhiheng Liu, Yang	Yudong Guo, and Juyong Zhang. 2022. Reconstruct-	1096
1046	Cao, and Zheng-Jun Zha. 2024b. Vivid: Video vir-	ing personalized semantic facial nerf models from	1097
1047	tual try-on using diffusion models. <i>arXiv preprint</i>	monocular video. <i>ACM Transactions on Graphics</i>	1098
1048	<i>arXiv:2405.11794</i> .	<i>(TOG)</i> , 41(6):1–12.	1099
1049	Najmeh Forouzandehmehr, Yijie Cao, Nikhil Thakur-	Yuying Ge, Ruimao Zhang, Xiaogang Wang, Xiaoou	1100
1050	desai, Ramin Giahhi, Luyi Ma, Nima Farrokhsiar,	Tang, and Ping Luo. 2019. Deepfashion2: A versatile	1101
1051	Jianpeng Xu, Evren Korpeoglu, and Kannan Achan.	benchmark for detection, pose estimation, segmenta-	1102
		tion and re-identification of clothing images. In <i>Pro-</i>	1103
		<i>ceedings of the IEEE/CVF conference on computer</i>	1104
		<i>vision and pattern recognition</i> , pages 5337–5345.	1105
		Shijie Geng, Zuohui Fu, Yingqiang Ge, Lei Li, Gerard	1106
		De Melo, and Yongfeng Zhang. 2022. Improving	1107

1108	personalized explanation generation through visual-	Xiaoyu Han, Shunyuan Zheng, Zonglin Li, Chenyang	1165
1109	ization. In <i>Proceedings of the 60th Annual Meeting</i>	Wang, Xin Sun, and Quanling Meng. 2024. Shape-	1166
1110	<i>of the Association for Computational Linguistics (Vol-</i>	guided clothing warping for virtual try-on. In <i>Pro-</i>	1167
1111	<i>ume 1: Long Papers)</i> , pages 244–255.	<i>ceedings of the 32nd ACM International Conference</i>	1168
		<i>on Multimedia</i> , pages 2593–2602.	1169
1112	Xue Geng, Hanwang Zhang, Jingwen Bian, and Tat-	Xintong Han, Zuxuan Wu, Zhe Wu, Ruichi Yu, and	1170
1113	Seng Chua. 2015. Learning image and user features	Larry S Davis. 2018. Viton: An image-based virtual	1171
1114	for recommendation in social networks. In <i>Proceed-</i>	try-on network. In <i>Proceedings of the IEEE con-</i>	1172
1115	<i>ings of the IEEE international conference on com-</i>	<i>ference on computer vision and pattern recognition</i> ,	1173
1116	<i>puter vision</i> , pages 4274–4282.	pages 7543–7552.	1174
1117	Junhong Gou, Siyu Sun, Jianfu Zhang, Jianlou Si, Chen	Haoran Hao, Jiaming Han, Changsheng Li, Yu-Feng	1175
1118	Qian, and Liqing Zhang. 2023. Taming the power	Li, and Xiangyu Yue. 2024. Remember, retrieve	1176
1119	of diffusion models for high-quality virtual try-on	and generate: Understanding infinite visual con-	1177
1120	with appearance flow. In <i>Proceedings of the 31st</i>	cepts as your personalized assistant. <i>arXiv preprint</i>	1178
1121	<i>ACM International Conference on Multimedia</i> , pages	<i>arXiv:2410.13360</i> .	1179
1122	7599–7607.		
1123	Shuyang Gu, Jianmin Bao, Dong Chen, and Fang Wen.	Shaozhe Hao, Kai Han, Shihao Zhao, and Kwan-Yee K	1180
1124	2020. Giga: Generated image quality assessment.	Wong. 2023. Vico: Plug-and-play visual condition	1181
1125	In <i>Computer Vision–ECCV 2020: 16th European</i>	for personalized text-to-image generation. <i>arXiv</i>	1182
1126	<i>Conference, Glasgow, UK, August 23–28, 2020, Pro-</i>	<i>preprint arXiv:2306.00971</i> .	1183
1127	<i>ceedings, Part XI 16</i> , pages 369–385. Springer.		
1128	Yuchao Gu, Xintao Wang, Jay Zhangjie Wu, Yujun Shi,	F. Maxwell Harper and Joseph A. Konstan. 2015. The	1184
1129	Yunpeng Chen, Zihan Fan, Wuyou Xiao, Rui Zhao,	movielens datasets: History and context . <i>ACM Trans.</i>	1185
1130	Shuning Chang, Weijia Wu, et al. 2024. Mix-of-	<i>Interact. Intell. Syst.</i> , 5(4).	1186
1131	show: Decentralized low-rank adaptation for multi-		
1132	concept customization of diffusion models. <i>Ad-</i>	Naomi Harte and Eoin Gillen. 2015. Tcd-timit: An	1187
1133	<i>vances in Neural Information Processing Systems</i> ,	audio-visual corpus of continuous speech. <i>IEEE</i>	1188
1134	36.	<i>Transactions on Multimedia</i> , 17(5):603–615.	1189
1135	Jiazhi Guan, Zhanwang Zhang, Hang Zhou, Tianshu Hu,	Curtis Hawthorne, Andriy Stasyuk, Adam Roberts, Ian	1190
1136	Kaisiyuan Wang, Dongliang He, Haocheng Feng,	Simon, Cheng-Zhi Anna Huang, Sander Dieleman,	1191
1137	Jingtuo Liu, Errui Ding, Ziwei Liu, et al. 2023.	Erich Elsen, Jesse Engel, and Douglas Eck. 2019.	1192
1138	Stylesync: High-fidelity generalized and personal-	Enabling factorized piano music modeling and gener-	1193
1139	ized lip sync in style-based generator. In <i>Proceedings</i>	ation with the MAESTRO dataset . In <i>International</i>	1194
1140	<i>of the IEEE/CVF Conference on Computer Vision</i>	<i>Conference on Learning Representations</i> .	1195
1141	<i>and Pattern Recognition</i> , pages 1505–1515.		
1142	Melody Y Guan, Manas Joglekar, Eric Wallace, Saachi	Junjie He, Yifeng Geng, and Liefeng Bo. 2024a.	1196
1143	Jain, Boaz Barak, Alec Heylar, Rachel Dias, Andrea	Uniportrait: A unified framework for identity-	1197
1144	Vallone, Hongyu Ren, Jason Wei, et al. 2024. Delib-	preserving single-and multi-human image personal-	1198
1145	erative alignment: Reasoning enables safer language	ization. <i>arXiv preprint arXiv:2408.05939</i> .	1199
1146	models. <i>arXiv preprint arXiv:2412.16339</i> .		
1147	Daya Guo, Dejian Yang, Haowei Zhang, Junxiao Song,	Xuanhua He, Quande Liu, Shengju Qian, Xin Wang,	1200
1148	Ruoyu Zhang, Runxin Xu, Qihao Zhu, Shirong Ma,	Tao Hu, Ke Cao, Keyu Yan, and Jie Zhang. 2024b.	1201
1149	Peiyi Wang, Xiao Bi, et al. 2025. Deepseek-r1: In-	Id-animator: Zero-shot identity-preserving human	1202
1150	centivizing reasoning capability in llms via reinforce-	video generation. <i>arXiv preprint arXiv:2404.15275</i> .	1203
1151	ment learning. <i>arXiv preprint arXiv:2501.12948</i> .		
1152	Qian Guo, Wei Chen, and Huaiyu Wan. 2021. Aol4ps:	Yingqing He, Menghan Xia, Haoxin Chen, Xiaodong	1204
1153	A large-scale data set for personalized search . <i>Data</i>	Cun, Yuan Gong, Jinbo Xing, Yong Zhang, Xintao	1205
1154	<i>Intelligence</i> , 3(4):548–567.	Wang, Chao Weng, Ying Shan, et al. 2023. Animate-	1206
1155	Yuwei Guo, Ceyuan Yang, Anyi Rao, Zhengyang Liang,	a-story: Storytelling with retrieval-augmented video	1207
1156	Yaohui Wang, Yu Qiao, Maneesh Agrawala, Dahua	generation. <i>arXiv preprint arXiv:2307.06940</i> .	1208
1157	Lin, and Bo Dai. 2024a. Animatediff: Animate your		
1158	personalized text-to-image diffusion models without	Yutong He, Alexander Robey, Naoki Murata, Yiding	1209
1159	specific tuning. In <i>The Twelfth International Confer-</i>	Jiang, Joshua Williams, George J Pappas, Hamed	1210
1160	<i>ence on Learning Representations</i> .	Hassani, Yuki Mitsufuji, Ruslan Salakhutdinov, and	1211
1161	Zinan Guo, Yanze Wu, Zhuowei Chen, Lang Chen, Peng	J Zico Kolter. 2024c. Automated black-box prompt	1212
1162	Zhang, and Qian He. 2024b. Pulid: Pure and light-	engineering for personalized text-to-image genera-	1213
1163	ning id customization via contrastive alignment. In	tion. <i>arXiv preprint arXiv:2403.19103</i> .	1214
1164	<i>Advances in Neural Information Processing Systems</i> .		
		Zecheng He, Bo Sun, Felix Juefei-Xu, Haoyu Ma,	1215
		Ankit Ramchandani, Vincent Cheung, Siddharth	1216
		Shah, Anmol Kalia, Harihar Subramanyam, Alireza	1217
		Zareian, et al. 2024d. Imagine yourself: Tuning-	1218
		free personalized image generation. <i>arXiv preprint</i>	1219
		<i>arXiv:2409.13346</i> .	1220

1221	Zijian He, Peixin Chen, Guangrun Wang, Guanbin Li,	Jiancheng Huang, Mingfu Yan, Songyan Chen,	1277
1222	Philip HS Torr, and Liang Lin. 2025. Wildvidfit:	Yi Huang, and Shifeng Chen. 2024a. Magicfight:	1278
1223	Video virtual try-on in the wild via image-based con-	Personalized martial arts combat video generation.	1279
1224	trolled diffusion models. In <i>European Conference on</i>	In <i>Proceedings of the 32nd ACM International Con-</i>	1280
1225	<i>Computer Vision</i> , pages 123–139. Springer.	<i>ference on Multimedia</i> , pages 10833–10842.	1281
1226	Martin Heusel, Hubert Ramsauer, Thomas Unterthiner,	Qihan Huang, Long Chan, Jinlong Liu, Wanggui	1282
1227	Bernhard Nessler, and Sepp Hochreiter. 2017. Gans	He, Hao Jiang, Mingli Song, and Jie Song.	1283
1228	trained by a two time-scale update rule converge to a	2024b. Patchdpo: Patch-level dpo for finetuning-	1284
1229	local nash equilibrium. <i>Advances in neural informa-</i>	free personalized image generation. <i>arXiv preprint</i>	1285
1230	<i>tion processing systems</i> , 30.	<i>arXiv:2412.03177</i> .	1286
1231	Yan Hong, Yuxuan Duan, Bo Zhang, Haoxing Chen,	Yuge Huang, Yuhan Wang, Ying Tai, Xiaoming Liu,	1287
1232	Jun Lan, Huijia Zhu, Weiqiang Wang, and Jianfu	Pengcheng Shen, Shaoxin Li, Jilin Li, and Feiyue	1288
1233	Zhang. 2025. Comfusion: Enhancing personalized	Huang. 2020. Curricularface: adaptive curriculum	1289
1234	generation by instance-scene compositing and fusion.	learning loss for deep face recognition. In <i>proceed-</i>	1290
1235	In <i>European Conference on Computer Vision</i> , pages	<i>ings of the IEEE/CVF conference on computer vision</i>	1291
1236	1–18. Springer.	<i>and pattern recognition</i> , pages 5901–5910.	1292
1237	Alain Hore and Djemel Ziou. 2010. Image quality met-	Yukun Huang, Jianan Wang, Ailing Zeng, He Cao, Xi-	1293
1238	rics: Psnr vs. ssim. In <i>2010 20th international con-</i>	anbiao Qi, Yukai Shi, Zheng-Jun Zha, and Lei Zhang.	1294
1239	<i>ference on pattern recognition</i> , pages 2366–2369.	2024c. Dreamwaltz: Make a scene with complex 3d	1295
1240	IEEE.	animatable avatars. <i>Advances in Neural Information</i>	1296
1241	Yupeng Hou, Jiacheng Li, Zhankui He, An Yan, Xiusi	<i>Processing Systems</i> , 36.	1297
1242	Chen, and Julian McAuley. 2024. Bridging lan-	Zhichao Huang, Xintong Han, Jia Xu, and Tong Zhang.	1298
1243	guage and items for retrieval and recommendation.	2021. Few-shot human motion transfer by personal-	1299
1244	<i>Preprint</i> , arXiv:2403.03952.	ized geometry and texture modeling. In <i>Proceedings</i>	1300
1245	Hexiang Hu, Kelvin CK Chan, Yu-Chuan Su, Wenhui	<i>of the IEEE/CVF Conference on Computer Vision</i>	1301
1246	Chen, Yandong Li, Kihyuk Sohn, Yang Zhao, Xue	<i>and Pattern Recognition</i> , pages 2297–2306.	1302
1247	Ben, Boqing Gong, William Cohen, et al. 2024a.	Ziqi Huang, Yinan He, Jiashuo Yu, Fan Zhang,	1303
1248	Instruct-imagen: Image generation with multi-modal	Chenyang Si, Yuming Jiang, Yuanhan Zhang, Tianx-	1304
1249	instruction. In <i>Proceedings of the IEEE/CVF Confer-</i>	ing Wu, Qingyang Jin, Nattapol Chanpaisit, et al.	1305
1250	<i>ence on Computer Vision and Pattern Recognition</i> ,	2024d. Vbench: Comprehensive benchmark suite	1306
1251	pages 4754–4763.	for video generative models. In <i>Proceedings of the</i>	1307
1252	Junxing Hu, Hongwen Zhang, Yunlong Wang, Min Ren,	<i>IEEE/CVF Conference on Computer Vision and Pat-</i>	1308
1253	and Zhenan Sun. 2023. Personalized graph genera-	<i>tern Recognition</i> , pages 21807–21818.	1309
1254	tion for monocular 3d human pose and shape estima-	Ziqi Huang, Tianxing Wu, Yuming Jiang, Kelvin C.K.	1310
1255	tion. <i>IEEE Transactions on Circuits and Systems for</i>	Chan, and Ziwei Liu. 2024e. ReVersion: Diffusion-	1311
1256	<i>Video Technology</i> .	based relation inversion from images. In <i>SIGGRAPH</i>	1312
1257	Li Hu. 2024. Animate anyone: Consistent and control-	<i>Asia 2024 Conference Papers</i> .	1313
1258	lable image-to-video synthesis for character anima-	Minyoung Hwang, Luca Weihs, Chanwoo Park, Kimin	1314
1259	tion. In <i>Proceedings of the IEEE/CVF Conference</i>	Lee, Aniruddha Kembhavi, and Kiana Ehsani. 2024.	1315
1260	<i>on Computer Vision and Pattern Recognition</i> , pages	Promptable behaviors: Personalizing multi-objective	1316
1261	8153–8163.	rewards from human preferences. In <i>Proceedings of</i>	1317
1262	Xinting Hu, Haoran Wang, Jan Eric Lenssen, and Bernt	<i>the IEEE/CVF Conference on Computer Vision and</i>	1318
1263	Schiele. 2025. PersonaHOI: Effortlessly improving	<i>Pattern Recognition</i> , pages 16216–16226.	1319
1264	personalized face with human-object interaction gen-	Catalin Ionescu, Dragos Papava, Vlad Olaru, and Cris-	1320
1265	eration. <i>arXiv preprint</i> .	tian Sminchisescu. 2013. Human3. 6m: Large scale	1321
1266	Yuke Hu, Jian Lou, Jiaqi Liu, Wangze Ni, Feng Lin,	datasets and predictive methods for 3d human sensing	1322
1267	Zhan Qin, and Kui Ren. 2024b. Eraser: Machine un-	in natural environments. <i>IEEE transactions on pat-</i>	1323
1268	learning in mlaas via an inference serving-aware ap-	<i>tern analysis and machine intelligence</i> , 36(7):1325–	1324
1269	proach. In <i>Proceedings of the 2024 on ACM SIGSAC</i>	1339.	1325
1270	<i>Conference on Computer and Communications Secu-</i>	Yasamin Jafarian and Hyun Soo Park. 2021. Learning	1326
1271	<i>rity</i> .	high fidelity depths of dressed humans by watching	1327
1272	Zhejing Hu, Yan Liu, Gong Chen, and Yongxu Liu.	social media dance videos. In <i>Proceedings of the</i>	1328
1273	2022. Can machines generate personalized music?	<i>IEEE/CVF Conference on Computer Vision and Pat-</i>	1329
1274	a hybrid favorite-aware method for user preference	<i>tern Recognition</i> , pages 12753–12762.	1330
1275	music transfer. <i>IEEE Transactions on Multimedia</i> ,		
1276	25:2296–2308.		

1444	Yuval Kirstain, Adam Polyak, Uriel Singer, Shahbuland	Cheng Li, Mingyang Zhang, Qiaozhu Mei, Yaqing	1500
1445	Matiana, Joe Penna, and Omer Levy. 2023. Pick-a-	Wang, Spurthi Amba Hombaiah, Yi Liang, and	1501
1446	pic: An open dataset of user preferences for text-to-	Michael Bendersky. 2023a. Teach llms to person-	1502
1447	image generation. <i>Advances in Neural Information</i>	alize – an approach inspired by writing education.	1503
1448	<i>Processing Systems</i> , 36:36652–36663.	<i>Preprint</i> , arXiv:2308.07968.	1504
1449	Jaehoon Ko, Kyusun Cho, Jounghbin Lee, Heeji Yoon,	Dongxu Li, Junnan Li, and Steven Hoi. 2024b. Blip-	1505
1450	Sangmin Lee, Sangjun Ahn, and Seungryong Kim.	diffusion: Pre-trained subject representation for con-	1506
1451	2024. Talk3d: High-fidelity talking portrait synthesis	trollable text-to-image generation and editing. <i>Ad-</i>	1507
1452	via personalized 3d generative prior. <i>arXiv preprint</i>	<i>advances in Neural Information Processing Systems</i> ,	1508
1453	<i>arXiv:2403.20153</i> .	36.	1509
1454	Ahmet Baki Kocaballi, Shlomo Berkovsky, Juan C	Hengjia Li, Haonan Qiu, Shiwei Zhang, Xiang Wang,	1510
1455	Quiroz, Liliana Laranjo, Huong Ly Tong, Dana	Yujie Wei, Zekun Li, Yingya Zhang, Boxi Wu, and	1511
1456	Rezazadegan, Agustina Briatore, and Enrico Coiera.	Deng Cai. 2024c. Personalvideo: High id-fidelity	1512
1457	2019. The personalization of conversational agents	video customization without dynamic and semantic	1513
1458	in health care: Systematic review. <i>J Med Internet</i>	degradation. <i>arXiv preprint arXiv:2411.17048</i> .	1514
1459	<i>Res</i> , 21(11):e15360.		
1460	Tom Kocmi and Christian Federmann. 2023. Large	Jiahe Li, Jiawei Zhang, Xiao Bai, Jun Zhou, and Lin Gu.	1515
1461	language models are state-of-the-art evaluators of	2023b. Efficient region-aware neural radiance fields	1516
1462	translation quality. <i>Preprint</i> , arXiv:2302.14520.	for high-fidelity talking portrait synthesis. In <i>Pro-</i>	1517
1463	Satwik Kottur, Xiaoyu Wang, and Vitor Carvalho. 2017.	<i>ceedings of the IEEE/CVF International Conference</i>	1518
1464	Exploring personalized neural conversational mod-	<i>on Computer Vision</i> , pages 7568–7578.	1519
1465	els. In <i>Proceedings of the Twenty-Sixth International</i>	Lei Li, Yongfeng Zhang, and Li Chen. 2020. Gener-	1520
1466	<i>Joint Conference on Artificial Intelligence, IJCAI-17</i> ,	ate neural template explanations for recommendation.	1521
1467	pages 3728–3734.	In <i>Proceedings of the 29th ACM International Con-</i>	1522
1468	Sumith Kulal, Tim Brooks, Alex Aiken, Jiajun Wu,	<i>ference on Information & Knowledge Management</i> ,	1523
1469	Jimei Yang, Jingwan Lu, Alexei A Efros, and Kr-	pages 755–764.	1524
1470	ishna Kumar Singh. 2023. Putting people in their	Lei Li, Yongfeng Zhang, and Li Chen. 2021. Person-	1525
1471	place: Affordance-aware human insertion into scenes.	alized transformer for explainable recommendation.	1526
1472	In <i>Proceedings of the IEEE/CVF Conference on Com-</i>	<i>arXiv preprint arXiv:2105.11601</i> .	1527
1473	<i>puter Vision and Pattern Recognition</i> , pages 17089–		
1474	17099.	Wei Li, Xue Xu, Jiachen Liu, and Xinyan Xiao. 2024d.	1528
1475	Ishita Kumar, Snigdha Viswanathan, Sushrita Yerra,	UNIMO-G: Unified image generation through multi-	1529
1476	Alireza Salemi, Ryan A. Rossi, Franck Dernoncourt,	modal conditional diffusion. In <i>Proceedings of the</i>	1530
1477	Hanieh Deilamsalehy, Xiang Chen, Ruiyi Zhang,	<i>62nd Annual Meeting of the Association for Compu-</i>	1531
1478	Shubham Agarwal, Nedim Lipka, Chien Van Nguyen,	<i>tational Linguistics (Volume 1: Long Papers)</i> , pages	1532
1479	Thien Huu Nguyen, and Hamed Zamani. 2024a.	6173–6188. Association for Computational Linguis-	1533
1480	Longlamp: A benchmark for personalized long-form	tics.	1534
1481	text generation. <i>Preprint</i> , arXiv:2407.11016.	Xiang Li, Lei Meng, Lei Wu, Manyi Li, and Xiangxu	1535
1482	Sachin Kumar, Chan Young Park, Yulia Tsvetkov,	Meng. 2024e. Dreamfont3d: personalized text-to-3d	1536
1483	Noah A. Smith, and Hannaneh Hajishirzi. 2024b.	artistic font generation. In <i>ACM SIGGRAPH 2024</i>	1537
1484	Compo: Community preferences for language model	<i>Conference Papers</i> , pages 1–11.	1538
1485	personalization. <i>Preprint</i> , arXiv:2410.16027.	Xiaoming Li, Xinyu Hou, and Chen Change Loy. 2024f.	1539
1486	Nupur Kumari, Bingliang Zhang, Richard Zhang, Eli	When stylegan meets stable diffusion: a w+ adapter	1540
1487	Shechtman, and Jun-Yan Zhu. 2023. Multi-concept	for personalized image generation. In <i>Proceedings of</i>	1541
1488	customization of text-to-image diffusion. In <i>Pro-</i>	<i>ceedings of the IEEE/CVF Conference on Computer Vision and</i>	1542
1489	<i>ceedings of the IEEE/CVF Conference on Computer</i>	<i>Pattern Recognition</i> , pages 2187–2196.	1543
1490	<i>Vision and Pattern Recognition</i> , pages 1931–1941.	Xiaoming Li, Shiguang Zhang, Shangchen Zhou, Lei	1544
1491	Sangyun Lee, Gyojung Gu, Sunghyun Park, Seunghwan	Zhang, and Wangmeng Zuo. 2022. Learning dual	1545
1492	Choi, and Jaegul Choo. 2022. High-resolution vir-	memory dictionaries for blind face restoration. <i>IEEE</i>	1546
1493	tual try-on with misalignment and occlusion-handled	<i>Transactions on Pattern Analysis and Machine Intel-</i>	1547
1494	conditions. In <i>European Conference on Computer</i>	<i>ligence</i> , 45(5):5904–5917.	1548
1495	<i>Vision</i> , pages 204–219. Springer.	Xinyu Li, Ruiyang Zhou, Zachary C. Lipton, and	1549
1496	Boyi Li, Jathushan Rajasegaran, Yossi Gandelsman,	Liu Leqi. 2024g. Personalized language model-	1550
1497	Alexei A Efros, and Jitendra Malik. 2024a. Synthe-	ing from personalized human feedback. <i>Preprint</i> ,	1551
1498	sizing moving people with 3d control. <i>arXiv preprint</i>	<i>arXiv:2402.05133</i> .	1552
1499	<i>arXiv:2401.10889</i> .		

1553	Yameng Li, Shi Feng, Daling Wang, Yifei Zhang, and	Gongye Liu, Menghan Xia, Yong Zhang, Haoxin Chen,	1607
1554	Xiaocui Yang. 2025. Paper: A persona-aware chain-	Jinbo Xing, Yibo Wang, Xintao Wang, Ying Shan,	1608
1555	of-thought learning framework for personalized di-	and Yujiu Yang. 2024a. Stylecrafter: Taming artistic	1609
1556	alogue response generation. In <i>Natural Language</i>	video diffusion with reference-augmented adapter	1610
1557	<i>Processing and Chinese Computing</i> , pages 215–227,	learning. <i>ACM Transactions on Graphics (TOG)</i> ,	1611
1558	Singapore. Springer Nature Singapore.	43(6):1–10.	1612
1559	Yang Li, Songlin Yang, Wei Wang, and Jing Dong.	Hanwen Liu, Zhicheng Sun, and Yadong Mu. 2024b.	1613
1560	2024h. Sefi-ide: Semantic-fidelity identity embed-	Countering personalized text-to-image generation	1614
1561	ding for personalized diffusion-based generation.	with influence watermarks. In <i>Proceedings of the</i>	1615
1562	<i>arXiv preprint arXiv:2402.00631</i> .	<i>IEEE/CVF Conference on Computer Vision and Pat-</i>	1616
1563	Yuanchun Li, Hao Wen, Weijun Wang, Xiangyu Li,	<i>tern Recognition</i> , pages 12257–12267.	1617
1564	Yizhen Yuan, Guohong Liu, Jiacheng Liu, Wenx-	Haohe Liu, Zehua Chen, Yi Yuan, Xinhao Mei, Xubo	1618
1565	ing Xu, Xiang Wang, Yi Sun, et al. 2024i. Per-	Liu, Danilo Mandic, Wenwu Wang, and Mark D	1619
1566	sonal llm agents: Insights and survey about the	Plumbley. 2023a. Audioldm: Text-to-audio gener-	1620
1567	capability, efficiency and security. <i>arXiv preprint</i>	ation with latent diffusion models. <i>arXiv preprint</i>	1621
1568	<i>arXiv:2401.05459</i> .	<i>arXiv:2301.12503</i> .	1622
1569	Zhaojian Li, Bin Zhao, and Yuan Yuan. 2024j. Tas:	Haotian Liu, Chunyuan Li, Qingyang Wu, and Yong Jae	1623
1570	Personalized text-guided audio spatialization. In <i>Pro-</i>	Lee. 2023b. Visual instruction tuning. In <i>Thirty-</i>	1624
1571	<i>ceedings of the 32nd ACM International Conference</i>	<i>seventh Conference on Neural Information Process-</i>	1625
1572	<i>on Multimedia</i> , pages 9029–9037.	<i>ing Systems</i> .	1626
1573	Zhen Li, Mingdeng Cao, Xintao Wang, Zhongang	Jia Liu, Changlin Li, Qirui Sun, Jiahui Ming, Chen	1627
1574	Qi, Ming-Ming Cheng, and Ying Shan. 2024k.	Fang, Jue Wang, Bing Zeng, and Shuaicheng Liu.	1628
1575	Photomaker: Customizing realistic human photos	2024c. Ada-adapter: Fast few-shot style personl-	1629
1576	via stacked id embedding. In <i>Proceedings of the</i>	ization of diffusion model with pre-trained image	1630
1577	<i>IEEE/CVF Conference on Computer Vision and Pat-</i>	encoder. <i>arXiv preprint arXiv:2407.05552</i> .	1631
1578	<i>tern Recognition</i> , pages 8640–8650.	Yang Liu, Dan Iter, Yichong Xu, Shuohang Wang,	1632
1579	Zhi Li, Christos Bampis, Julie Novak, Anne Aaron,	Ruochen Xu, and Chenguang Zhu. 2023c. G-eval:	1633
1580	Kyle Swanson, Anush Moorthy, and JD Cock. 2018.	NLG evaluation using gpt-4 with better human align-	1634
1581	Vmaf: The journey continues. <i>Netflix Technology</i>	ment . In <i>Proceedings of the 2023 Conference on</i>	1635
1582	<i>Blog</i> , 25(1).	<i>Empirical Methods in Natural Language Processing</i> ,	1636
1583	Jie Liang, Hui Zeng, Miaomiao Cui, Xuansong Xie,	pages 2511–2522, Singapore. Association for Com-	1637
1584	and Lei Zhang. 2021. Ppr10k: A large-scale portrait	putational Linguistics.	1638
1585	photo retouching dataset with human-region mask	Ye Liu, Siyuan Li, Yang Wu, Chang-Wen Chen, Ying	1639
1586	and group-level consistency. In <i>Proceedings of the</i>	Shan, and Xiaohu Qie. 2022. Umt: Unified multi-	1640
1587	<i>IEEE/CVF Conference on Computer Vision and Pat-</i>	modal transformers for joint video moment retrieval	1641
1588	<i>tern Recognition</i> , pages 653–661.	and highlight detection. In <i>Proceedings of the</i>	1642
1589	Chin-Yew Lin. 2004. ROUGE: A package for auto-	<i>IEEE/CVF Conference on Computer Vision and Pat-</i>	1643
1590	matic evaluation of summaries . In <i>Text Summariza-</i>	<i>tern Recognition</i> , pages 3042–3051.	1644
1591	<i>tion Branches Out</i> , pages 74–81, Barcelona, Spain.	Yexin Liu and Lin Wang. 2024. Mycloth: Towards	1645
1592	Association for Computational Linguistics.	intelligent and interactive online t-shirt customiza-	1646
1593	Xinyu Lin, Wenjie Wang, Yongqi Li, Shuo Yang, Fuli	tion based on user’s preference. <i>arXiv preprint</i>	1647
1594	Feng, Yinwei Wei, and Tat-Seng Chua. 2024a. Data-	<i>arXiv:2404.15801</i> .	1648
1595	efficient fine-tuning for llm-based recommendation .	Yilun Liu, Minggui He, Feiyu Yao, Yuhe Ji, Shimin Tao,	1649
1596	<i>Preprint</i> , arXiv:2401.17197.	Jingzhou Du, Duan Li, Jian Gao, Li Zhang, Hao Yang,	1650
1597	Xudong Lin, Ali Zare, Shiyuan Huang, Ming-Hsuan	et al. 2024d. What do you want? user-centric prompt	1651
1598	Yang, Shih-Fu Chang, and Li Zhang. 2024b. Person-	generation for text-to-image synthesis via multi-turn	1652
1599	alized video comment generation. In <i>Findings of the</i>	guidance. <i>arXiv preprint arXiv:2408.12910</i> .	1653
1600	<i>Association for Computational Linguistics: EMNLP</i>	Yixin Liu, Ruoxi Chen, Xun Chen, and Lichao Sun.	1654
1601	2024, pages 16806–16820.	2024e. Rethinking and defending protective perturba-	1655
1602	Fangfu Liu, Hanyang Wang, Weiliang Chen, Haowen	tion in personalized diffusion models. <i>arXiv preprint</i>	1656
1603	Sun, and Yueqi Duan. 2025. Make-your-3d: Fast	<i>arXiv:2406.18944</i> .	1657
1604	and consistent subject-driven 3d content generation.	Zhilei Liu, Xiaoxing Liu, Sen Chen, Jiaying Liu, Long-	1658
1605	In <i>European Conference on Computer Vision</i> , pages	biao Wang, and Chongke Bi. 2024f. Multimodal	1659
1606	389–406. Springer.	fusion for talking face generation utilizing speech-	1660
		related facial action units. <i>ACM Transactions on</i>	1661
		<i>Multimedia Computing, Communications and Appli-</i>	1662
		<i>cations</i> , 20(9):1–24.	1663

1664	Cuirong Long, Xiaoshan Yang, and Changsheng Xu.	Liam Magee, Vanicka Arora, Gus Gollings, and Norma	1721
1665	2020. Cross-domain personalized image captioning.	Lam-Saw. 2024. The drama machine: Simulating	1722
1666	<i>Multimedia Tools and Applications</i> , 79(45):33333–	character development with llm agents . <i>Preprint</i> ,	1723
1667	33348.	arXiv:2408.01725.	1724
1668	Xiaoxiao Long, Yuan-Chen Guo, Cheng Lin, Yuan Liu,	Yifang Men, Yiming Mao, Yuning Jiang, Wei-Ying Ma,	1725
1669	Zhiyang Dou, Lingjie Liu, Yuexin Ma, Song-Hai	and Zhouhui Lian. 2020. Controllable person image	1726
1670	Zhang, Marc Habermann, Christian Theobalt, et al.	synthesis with attribute-decomposed gan. In <i>Pro-</i>	1727
1671	2024. Wonder3d: Single image to 3d using cross-	<i>ceedings of the IEEE/CVF conference on computer</i>	1728
1672	domain diffusion. In <i>Proceedings of the IEEE/CVF</i>	<i>vision and pattern recognition</i> , pages 5084–5093.	1729
1673	<i>Conference on Computer Vision and Pattern Recog-</i>		
1674	<i>nition</i> , pages 9970–9980.	Anish Mittal, Anush Krishna Moorthy, and Alan Con-	1730
1675	Zhi Lu, Yang Hu, Yunchao Jiang, Yan Chen, and Bing	rad Bovik. 2012a. No-reference image quality assess-	1731
1676	Zeng. 2019. Learning binary code for personal-	ment in the spatial domain. <i>IEEE Transactions on</i>	1732
1677	ized fashion recommendation. In <i>Proceedings of</i>	<i>image processing</i> , 21(12):4695–4708.	1733
1678	<i>the IEEE/CVF conference on computer vision and</i>		
1679	<i>pattern recognition</i> , pages 10562–10570.	Anish Mittal, Rajiv Soundararajan, and Alan C Bovik.	1734
1680	Hanjia Lyu, Song Jiang, Hanqing Zeng, Yinglong Xia,	2012b. Making a “completely blind” image quality	1735
1681	Qifan Wang, Si Zhang, Ren Chen, Chris Leung, Jiajie	analyzer. <i>IEEE Signal processing letters</i> , 20(3):209–	1736
1682	Tang, and Jiebo Luo. 2024. LLM-rec: Personalized	212.	1737
1683	recommendation via prompting large language mod-	Shentong Mo, Jing Shi, and Yapeng Tian. 2023. Dif-	1738
1684	els . In <i>Findings of the Association for Computational</i>	fava: Personalized text-to-audio generation with vi-	1739
1685	<i>Linguistics: NAACL 2024</i> , pages 583–612, Mexico	sual alignment. <i>arXiv preprint arXiv:2305.12903</i> .	1740
1686	City, Mexico. Association for Computational Lin-		
1687	guistics.	Davide Morelli, Alberto Baldrati, Giuseppe Cartella,	1741
1688	Yueming Lyu, Tianwei Lin, Fu Li, Dongliang He, Jing	Marcella Cornia, Marco Bertini, and Rita Cucchiara.	1742
1689	Dong, and Tieniu Tan. 2023. Notice of removal:	2023. Ladi-vton: Latent diffusion textual-inversion	1743
1690	Deltaedit: Exploring text-free training for text-driven	enhanced virtual try-on. In <i>Proceedings of the 31st</i>	1744
1691	image manipulation. In <i>2023 IEEE/CVF Conference</i>	<i>ACM International Conference on Multimedia</i> , pages	1745
1692	<i>on Computer Vision and Pattern Recognition (CVPR)</i> ,	8580–8589.	1746
1693	pages 6894–6903.	Davide Morelli, Matteo Fincato, Marcella Cornia, Fed-	1747
1694	Jian Ma, Junhao Liang, Chen Chen, and Haonan Lu.	erico Landi, Fabio Cesari, and Rita Cucchiara. 2022.	1748
1695	2024a. Subject-diffusion: Open domain personal-	Dress code: High-resolution multi-category virtual	1749
1696	ized text-to-image generation without test-time fine-	try-on. In <i>Proceedings of the IEEE/CVF conference</i>	1750
1697	tuning. In <i>ACM SIGGRAPH 2024 Conference Pa-</i>	<i>on computer vision and pattern recognition</i> , pages	1751
1698	<i>pers</i> , pages 1–12.	2231–2235.	1752
1699	Xichu Ma, Yuchen Wang, and Ye Wang. 2022. Content	Sheshera Mysore, Zhuoran Lu, Mengting Wan, Longqi	1753
1700	based user preference modeling in music generation.	Yang, Bahareh Sarrafzadeh, Steve Menezes, Tina	1754
1701	In <i>Proceedings of the 30th ACM International Con-</i>	Baghaee, Emmanuel Barajas Gonzalez, Jennifer	1755
1702	<i>ference on Multimedia</i> , pages 2473–2482.	Neville, and Tara Safavi. 2024. Pearl: Person-	1756
1703	Xichu Ma, Yuchen Wang, and Ye Wang. 2024b. Sym-	alizing large language model writing assistants	1757
1704	bolic music generation from graph-learning-based	with generation-calibrated retrievers . <i>Preprint</i> ,	1758
1705	preference modeling and textual queries. <i>IEEE Trans-</i>	arXiv:2311.09180.	1759
1706	<i>actions on Multimedia</i> .	Ofir Nabati, Guy Tennenholtz, ChihWei Hsu,	1760
1707	Yifeng Ma, Shiwei Zhang, Jiayu Wang, Xiang Wang,	Moonkyung Ryu, Deepak Ramachandran, Yinlam	1761
1708	Yingya Zhang, and Zhidong Deng. 2023. Dreamtalk:	Chow, Xiang Li, and Craig Boutilier. 2024. Personal-	1762
1709	When expressive talking head generation meets	ized and sequential text-to-image generation. <i>arXiv</i>	1763
1710	diffusion probabilistic models. <i>arXiv preprint</i>	<i>preprint arXiv:2412.10419</i> .	1764
1711	<i>arXiv:2312.09767</i> .	Arsha Nagrani, Joon Son Chung, Weidi Xie, and	1765
1712	Ze Ma, Daquan Zhou, Chun-Hsiao Yeh, Xue-She	Andrew Zisserman. 2020. Voxceleb: Large-scale	1766
1713	Wang, Xiuyu Li, Huanrui Yang, Zhen Dong, Kurt	speaker verification in the wild. <i>Computer Speech &</i>	1767
1714	Keutzer, and Jiashi Feng. 2024c. Magic-me: Identity-	<i>Language</i> , 60:101027.	1768
1715	specific video customized diffusion. <i>arXiv preprint</i>	Jisu Nam, Heesu Kim, DongJae Lee, Siyoon Jin, Seun-	1769
1716	<i>arXiv:2402.09368</i> .	gryong Kim, and Seunggyu Chang. 2024. Dream-	1770
1717	Zhiyuan Ma, Xiangyu Zhu, Guojun Qi, Chen Qian,	matcher: Appearance matching self-attention for	1771
1718	Zhaoxiang Zhang, and Zhen Lei. 2024d. Diffspeaker:	semantically-consistent text-to-image personaliza-	1772
1719	Speech-driven 3d facial animation with diffusion	tion. In <i>Proceedings of the IEEE/CVF Conference</i>	1773
1720	transformer. <i>arXiv preprint arXiv:2402.05712</i> .	<i>on Computer Vision and Pattern Recognition</i> , pages	1774
		8100–8110.	1775

1776	Niranjan D Narvekar and Lina J Karam. 2011. A no-reference image blur metric based on the cumulative probability of blur detection (cpbd). <i>IEEE Transactions on Image Processing</i> , 20(9):2678–2683.	1833
1777		1834
1778		1835
1779		1836
1780	Man Tik Ng, Hui Tung Tse, Jen-Tse Huang, Jingjing Li, Wenxuan Wang, and Michael R. Lyu. 2024. How well can llms echo us? evaluating ai chatbots’ role-play ability with echo . <i>ArXiv</i> , abs/2404.13957.	1837
1781		1838
1782		1839
1783		1840
1784	Hung Nguyen, Quang Qui-Vinh Nguyen, Khoi Nguyen, and Rang Nguyen. 2024a. Swifttry: Fast and consistent video virtual try-on with diffusion models. <i>arXiv preprint arXiv:2412.10178</i> .	1841
1785		1842
1786		
1787		
1788	Thao Nguyen, Haotian Liu, Yuheng Li, Mu Cai, Utkarsh Ojha, and Yong Jae Lee. 2024b. Yo’LLaVA: Your personalized language and vision assistant. In <i>The Thirty-eighth Annual Conference on Neural Information Processing Systems</i> .	1843
1789		1844
1790		1845
1791		1846
1792		1847
1793	Jianmo Ni, Jiacheng Li, and Julian McAuley. 2019. Justifying recommendations using distantly-labeled reviews and fine-grained aspects . In <i>Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP)</i> , pages 188–197, Hong Kong, China. Association for Computational Linguistics.	1848
1794		
1795		
1796		
1797		
1798		
1799		
1800		
1801		
1802	Armand Nicolicioiu, Eugenia Iofinova, Eldar Kurtic, Mahdi Nikdan, Andrei Panferov, Ilia Markov, Nir Shavit, and Dan Alistarh. 2024. Panza: A personalized text writing assistant via data playback and local fine-tuning . <i>Preprint</i> , arXiv:2407.10994.	1849
1803		1850
1804		1851
1805		1852
1806		
1807	Shuliang Ning, Duomin Wang, Yipeng Qin, Zirong Jin, Baoyuan Wang, and Xiaoguang Han. 2024. Picture: Photorealistic virtual try-on from unconstrained designs. In <i>Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition</i> , pages 6976–6985.	1853
1808		1854
1809		1855
1810		1856
1811		1857
1812		
1813	Yotam Nitzan, Kfir Aberman, Qiurui He, Orly Liba, Michal Yarom, Yossi Gandelsman, Inbar Mosseri, Yael Pritch, and Daniel Cohen-Or. 2022. Mystyle: A personalized generative prior. <i>arXiv preprint arXiv:2203.17272</i> .	1858
1814		1859
1815		1860
1816		1861
1817		1862
1818	Takuto Onikubo and Yusuke Matsui. 2024. High-frequency anti-dreambooth: Robust defense against personalized image synthesis. <i>arXiv preprint arXiv:2409.08167</i> .	1863
1819		1864
1820		1865
1821		1866
1822	Maxime Oquab, Timothée Darcet, Théo Moutakanni, Huy V. Vo, Marc Szafraniec, Vasil Khalidov, Pierre Fernandez, Daniel HAZIZA, Francisco Massa, Alaaeldin El-Nouby, Mido Assran, Nicolas Ballas, Wojciech Galuba, Russell Howes, Po-Yao Huang, Shang-Wen Li, Ishan Misra, Michael Rabbat, Vasu Sharma, Gabriel Synnaeve, Hu Xu, Herve Jegou, Julien Mairal, Patrick Labatut, Armand Joulin, and Piotr Bojanowski. 2024. DINOv2: Learning robust visual features without supervision. <i>Transactions on Machine Learning Research</i> .	1867
1823		1868
1824		1869
1825		1870
1826		1871
1827		1872
1828		
1829		
1830		
1831		
1832		
	Nadav Orzech, Yotam Nitzan, Ulysse Mizrahi, Dov Danon, and Amit H Bermano. 2024. Masked extended attention for zero-shot virtual try-on in the wild. <i>arXiv preprint arXiv:2406.15331</i> .	1873
		1874
		1875
		1876
		1877
	Lianyu Pang, Jian Yin, Baoquan Zhao, Feize Wu, Fu Lee Wang, Qing Li, and Xudong Mao. 2024. Atndreambooth: Towards text-aligned personalized text-to-image generation. In <i>The Thirty-eighth Annual Conference on Neural Information Processing Systems</i> .	1878
		1879
		1880
		1881
		1882
	Kishore Papineni, Salim Roukos, Todd Ward, and Wei-Jing Zhu. 2002. Bleu: a method for automatic evaluation of machine translation . In <i>Proceedings of the 40th Annual Meeting on Association for Computational Linguistics, ACL ’02</i> , page 311–318, USA. Association for Computational Linguistics.	1883
		1884
		1885
		1886
		1887
		1888
	Rishubh Parihar, Harsh Gupta, Sachidanand VS, and R Venkatesh Babu. 2024. Text2place: Affordance-aware text guided human placement. In <i>European Conference on Computer Vision</i> , pages 57–77.	
	Cesc Chunseong Park, Byeongchang Kim, and Gunhee Kim. 2018. Towards personalized image captioning via multimodal memory networks. <i>IEEE transactions on pattern analysis and machine intelligence</i> , 41(4):999–1012.	
	Junseo Park, Beomseok Ko, and Hyeryung Jang. 2024. Text-to-image synthesis for any artistic styles: Advancements in personalized artistic image generation via subdivision and dual binding. <i>arXiv preprint arXiv:2404.05256</i> .	
	Or Patashnik, Zongze Wu, Eli Shechtman, Daniel Cohen-Or, and Dani Lischinski. 2021. Styleclip: Text-driven manipulation of stylegan imagery. In <i>Proceedings of the IEEE/CVF international conference on computer vision</i> , pages 2085–2094.	
	Maitreya Patel, Tejas Gokhale, Chitta Baral, and Yezhou Yang. 2024a. Conceptbed: Evaluating concept learning abilities of text-to-image diffusion models. In <i>Proceedings of the AAAI Conference on Artificial Intelligence</i> , volume 38, pages 14554–14562.	
	Maitreya Patel, Sangmin Jung, Chitta Baral, and Yezhou Yang. 2024b. λ -ECLIPSE: Multi-concept personalized text-to-image diffusion models by leveraging CLIP latent space. <i>Transactions on Machine Learning Research</i> .	
	Yuang Peng, Yuxin Cui, Haomiao Tang, Zekun Qi, Runpei Dong, Jing Bai, Chunrui Han, Zheng Ge, Xiangyu Zhang, and Shu-Tao Xia. 2024. Dreambench++: A human-aligned benchmark for personalized image generation. <i>arXiv preprint arXiv:2406.16855</i> .	
	Chau Pham, Hoang Phan, David Doermann, and Yunjie Tian. 2024. Personalized large vision-language models. <i>arXiv preprint arXiv:2412.17610</i> .	
	Renjie Pi, Jianshu Zhang, Tianyang Han, Jipeng Zhang, Rui Pan, and Tong Zhang. 2024. Personalized visual instruction tuning. <i>arXiv preprint arXiv:2410.07113</i> .	

1889	Manos Plitsis, Theodoros Kouzelis, Georgios	2023. Dreambooth3d: Subject-driven text-to-3d gen-	1945
1890	Paraskevopoulos, Vassilis Katsouros, and Yannis	eration. In <i>Proceedings of the IEEE/CVF interna-</i>	1946
1891	Panagakakis. 2024. Investigating personalization	<i>tional conference on computer vision</i> , pages 2349–	1947
1892	methods in text to music generation. In <i>ICASSP</i>	2359.	1948
1893	<i>2024-2024 IEEE International Conference on</i>		
1894	<i>Acoustics, Speech and Signal Processing (ICASSP)</i> ,	Andreas Rossler, Davide Cozzolino, Luisa Verdoliva,	1949
1895	pages 1081–1085. IEEE.	Christian Riess, Justus Thies, and Matthias Nießner.	1950
		2019. Faceforensics++: Learning to detect manipu-	1951
1896	Sriyash Poddar, Yanming Wan, Hamish Ivison, Ab-	lated facial images. In <i>Proceedings of the IEEE/CVF</i>	1952
1897	hishek Gupta, and Natasha Jaques. 2024. Person-	<i>international conference on computer vision</i> , pages	1953
1898	alizing reinforcement learning from human feedback	1–11.	1954
1899	with variational preference learning. In <i>The Thirty-</i>		
1900	<i>eighth Annual Conference on Neural Information</i>	Nataniel Ruiz, Yuanzhen Li, Varun Jampani, Yael	1955
1901	<i>Processing Systems</i> .	Pritch, Michael Rubinstein, and Kfir Aberman. 2023.	1956
		Dreambooth: Fine tuning text-to-image diffusion	1957
1902	Alisha Pradhan and Amanda Lazar. 2021. Hey google,	models for subject-driven generation. In <i>Proceed-</i>	1958
1903	do you have a personality? designing personality and	<i>ings of the IEEE/CVF conference on computer vision</i>	1959
1904	personas for conversational agents . In <i>Proceedings</i>	<i>and pattern recognition</i> , pages 22500–22510.	1960
1905	<i>of the 3rd Conference on Conversational User Inter-</i>		
1906	<i>faces</i> , CUI ’21, New York, NY, USA. Association for	Masaki Saito, Shunta Saito, Masanori Koyama, and	1961
1907	Computing Machinery.	Sosuke Kobayashi. 2020. Train sparsely, generate	1962
		densely: Memory-efficient unsupervised training of	1963
1908	K R Prajwal, Rudrabha Mukhopadhyay, Vinay P. Nam-	high-resolution temporal gan. <i>International Journal</i>	1964
1909	boodiri, and C.V. Jawahar. 2020. A lip sync expert is	<i>of Computer Vision</i> , 128(10):2586–2606.	1965
1910	all you need for speech to lip generation in the wild .		
1911	In <i>Proceedings of the 28th ACM International Con-</i>	Alireza Salemi, Surya Kallumadi, and Hamed Zamani.	1966
1912	<i>ference on Multimedia</i> , page 484–492. Association	2024a. Optimization methods for personalizing large	1967
1913	for Computing Machinery.	language models through retrieval augmentation . In	1968
		<i>Proceedings of the 47th International ACM SIGIR</i>	1969
1914	Xavier Puig, Eric Undersander, Andrew Szot,	<i>Conference on Research and Development in Infor-</i>	1970
1915	Mikael Dallaire Cote, Tsung-Yen Yang, Ruslan Part-	<i>mation Retrieval</i> , SIGIR ’24, page 752–762, New	1971
1916	sey, Ruta Desai, Alexander William Clegg, Michal	York, NY, USA. Association for Computing Machin-	1972
1917	Hlavac, So Yeon Min, et al. 2023. Habitat 3.0: A	ery.	1973
1918	co-habitat for humans, avatars and robots. <i>arXiv</i>		
1919	<i>preprint arXiv:2310.13724</i> .	Alireza Salemi, Julian Killingback, and Hamed Zamani.	1974
		2025a. Expert: Effective and explainable evaluation	1975
1920	Luchao Qi, Jiaye Wu, Shengze Wang, and Soumyadip	of personalized long-form text generation . <i>Preprint</i> ,	1976
1921	Sengupta. 2023a. My3dgen: Building lightweight	arXiv:2501.14956.	1977
1922	personalized 3d generative model. <i>arXiv preprint</i>		
1923	<i>arXiv:2307.05468</i> .	Alireza Salemi, Cheng Li, Mingyang Zhang, Qiaozhu	1978
		Mei, Weize Kong, Tao Chen, Zhuowan Li, Michael	1979
1924	Zipeng Qi, Xulong Zhang, Ning Cheng, Jing Xiao, and	Bendersky, and Hamed Zamani. 2025b. Reasoning-	1980
1925	Jianzong Wang. 2023b. DiffTalker: Co-driven audio-	enhanced self-training for long-form personalized	1981
1926	image diffusion for talking faces via intermediate	text generation . <i>Preprint</i> , arXiv:2501.04167.	1982
1927	landmarks. <i>arXiv preprint arXiv:2309.07509</i> .		
		Alireza Salemi, Sheshera Mysore, Michael Bendersky,	1983
1928	Hongjin Qian, Xiaohe Li, Hanxun Zhong, Yu Guo,	and Hamed Zamani. 2024b. LaMP: When large lan-	1984
1929	Yueyuan Ma, Yutao Zhu, Zhanliang Liu, Zhicheng	guage models meet personalization. In <i>Proceedings</i>	1985
1930	Dou, and Ji-Rong Wen. 2021. Pchatbot: A large-	<i>of the 62nd Annual Meeting of the Association for</i>	1986
1931	scale dataset for personalized chatbot . In <i>Proceed-</i>	<i>Computational Linguistics (Volume 1: Long Papers)</i> ,	1987
1932	<i>ings of the 44th International ACM SIGIR Confer-</i>	pages 7370–7392. Association for Computational	1988
1933	<i>ence on Research and Development in Information</i>	Linguistics.	1989
1934	<i>Retrieval</i> , SIGIR ’21, page 2470–2477, New York,		
1935	NY, USA. Association for Computing Machinery.	Alireza Salemi and Hamed Zamani. 2024. Comparing	1990
		retrieval-augmentation and parameter-efficient fine-	1991
1936	Alec Radford, Jong Wook Kim, Chris Hallacy, Aditya	tuning for privacy-preserving personalization of large	1992
1937	Ramesh, Gabriel Goh, Sandhini Agarwal, Girish Sas-	language models . <i>Preprint</i> , arXiv:2409.09510.	1993
1938	try, Amanda Askell, Pamela Mishkin, Jack Clark,		
1939	et al. 2021. Learning transferable visual models from	Tim Salimans, Ian Goodfellow, Wojciech Zaremba,	1994
1940	natural language supervision. In <i>International confer-</i>	Vicki Cheung, Alec Radford, and Xi Chen. 2016.	1995
1941	<i>ence on machine learning</i> , pages 8748–8763. PMLR.	Improved techniques for training gans. <i>Advances in</i>	1996
		<i>neural information processing systems</i> , 29.	1997
1942	Amit Raj, Srinivas Kaza, Ben Poole, Michael Niemeyer,	Shibani Santurkar, Esin Durmus, Faisal Ladhak, Cino	1998
1943	Nataniel Ruiz, Ben Mildenhall, Shiran Zada, Kfir	Lee, Percy Liang, and Tatsunori Hashimoto. 2023.	1999
1944	Aberman, Michael Rubinstein, Jonathan Barron, et al.	Whose opinions do language models reflect? In	2000

2001	<i>Proceedings of the 40th International Conference on</i>	Yichun Shi, Peng Wang, Jianglong Ye, Long Mai, Kejie	2056
2002	<i>Machine Learning, ICML'23. JMLR.org.</i>	Li, and Xiao Yang. 2023b. Mvdream: Multi-view	2057
2003	Florian Schroff, Dmitry Kalenichenko, and James	diffusion for 3d generation. In <i>The Twelfth Interna-</i>	2058
2004	Philbin. 2015. Facenet: A unified embedding for	<i>tional Conference on Learning Representations.</i>	2059
2005	face recognition and clustering. In <i>Proceedings of</i>	Veronika Shilova, Ludovic Dos Santos, Flavian Vasile,	2060
2006	<i>the IEEE conference on computer vision and pattern</i>	Gaëtan Racic, and Ugo Tanielian. 2023. Adbooster:	2061
2007	<i>recognition</i> , pages 815–823.	Personalized ad creative generation using stable diffu-	2062
2008	Murray Shanahan, Kyle McDonell, and Laria Reynolds.	sion outpainting. <i>arXiv preprint arXiv:2309.11507.</i>	2063
2009	2023. Role-play with large language models.	Andrew Shin, Yoshitaka Ushiku, and Tatsuya Harada.	2064
2010	<i>Preprint</i> , arXiv:2305.16367.	2018. Customized image narrative generation via	2065
2011	Yunfan Shao, Linyang Li, Junqi Dai, and Xipeng Qiu.	interactive visual question generation and answering.	2066
2012	2023. Character-LLM: A trainable agent for role-	In <i>Proceedings of the IEEE Conference on Computer</i>	2067
2013	playing . In <i>Proceedings of the 2023 Conference on</i>	<i>Vision and Pattern Recognition</i> , pages 8925–8933.	2068
2014	<i>Empirical Methods in Natural Language Process-</i>	Aliaksandr Siarohin, Oliver J Woodford, Jian Ren, Men-	2069
2015	<i>ing</i> , pages 13153–13187, Singapore. Association for	glei Chai, and Sergey Tulyakov. 2021. Motion repre-	2070
2016	Computational Linguistics.	sentations for articulated animation. In <i>Proceedings</i>	2071
2017	Fei Shen, Xin Jiang, Xin He, Hu Ye, Cong Wang,	<i>of the IEEE/CVF Conference on Computer Vision</i>	2072
2018	Xiaoyu Du, Zechao Li, and Jinhui Tang. 2024a.	<i>and Pattern Recognition</i> , pages 13653–13662.	2073
2019	Imagdressing-v1: Customizable virtual dressing.	Kihyuk Sohn, Lu Jiang, Jarred Barber, Kimin Lee,	2074
2020	<i>arXiv preprint arXiv:2407.12705.</i>	Nataniel Ruiz, Dilip Krishnan, Huiwen Chang,	2075
2021	Shuai Shen, Wenliang Zhao, Zibin Meng, Wanhua Li,	Yuanzhen Li, Irfan Essa, Michael Rubinstein, Yuan	2076
2022	Zheng Zhu, Jie Zhou, and Jiwen Lu. 2023. Difftalk:	Hao, Glenn Entis, Irina Blok, and Daniel Castro Chin.	2077
2023	Crafting diffusion models for generalized audio-	2023. Styledrop: Text-to-image synthesis of any	2078
2024	driven portraits animation. In <i>Proceedings of the</i>	style. In <i>Thirty-seventh Conference on Neural Infor-</i>	2079
2025	<i>IEEE/CVF Conference on Computer Vision and Pat-</i>	<i>mation Processing Systems.</i>	2080
2026	<i>tern Recognition</i> , pages 1982–1991.	Joon Son Chung, Andrew Senior, Oriol Vinyals, and	2081
2027	Xiaoteng Shen, Rui Zhang, Xiaoyan Zhao, Jieming Zhu,	Andrew Zisserman. 2017. Lip reading sentences in	2082
2028	and Xi Xiao. 2024b. Pmg: Personalized multimodal	the wild. In <i>Proceedings of the IEEE conference</i>	2083
2029	generation with large language models. In <i>Proceed-</i>	<i>on computer vision and pattern recognition</i> , pages	2084
2030	<i>ings of the ACM on Web Conference 2024</i> , pages	6447–6456.	2085
2031	3833–3843.	Haoyu Song, Yan Wang, Kaiyan Zhang, Wei-Nan	2086
2032	Zheng-Yan Sheng, Yang Ai, and Zhen-Hua Ling. 2023.	Zhang, and Ting Liu. 2021. BoB: BERT over BERT	2087
2033	Zero-shot personalized lip-to-speech synthesis with	for training persona-based dialogue models from lim-	2088
2034	face image based voice control. In <i>ICASSP 2023-</i>	ited personalized data . In <i>Proceedings of the 59th</i>	2089
2035	<i>2023 IEEE International Conference on Acoustics,</i>	<i>Annual Meeting of the Association for Computational</i>	2090
2036	<i>Speech and Signal Processing (ICASSP)</i> , pages 1–5.	<i>Linguistics and the 11th International Joint Confer-</i>	2091
2037	IEEE.	<i>ence on Natural Language Processing (Volume 1:</i>	2092
2038	Jing Shi, Wei Xiong, Zhe Lin, and Hyun Joon Jung.	<i>Long Papers)</i> , pages 167–177, Online. Association	2093
2039	2024. Instantbooth: Personalized text-to-image gen-	for Computational Linguistics.	2094
2040	eration without test-time finetuning. In <i>Proceedings</i>	Haoyu Song, Wei-Nan Zhang, Yiming Cui, Dong Wang,	2095
2041	<i>of the IEEE/CVF Conference on Computer Vision</i>	and Ting Liu. 2019. Exploiting persona information	2096
2042	<i>and Pattern Recognition</i> , pages 8543–8552.	for diverse generation of conversational responses.	2097
2043	Lei Shi, Alexandra Ioana Cristea, Malik	In <i>Proceedings of the Twenty-Eighth International</i>	2098
2044	Shahzad Kaleem Awan, Craig D. Stewart, and	<i>Joint Conference on Artificial Intelligence, IJCAI-19,</i>	2099
2045	Maurice Hendrix. 2013. Towards understanding	pages 5190–5196. International Joint Conferences on	2100
2046	learning behavior patterns in social adaptive person-	Artificial Intelligence Organization.	2101
2047	alized e-learning systems . In <i>Americas Conference</i>	Kunpeng Song, Yizhe Zhu, Bingchen Liu, Qing Yan,	2102
2048	<i>on Information Systems.</i>	Ahmed Elgammal, and Xiao Yang. 2025. Moma:	2103
2049	Xiaoyu Shi, Zhaoyang Huang, Weikang Bian, Da-	Multimodal llm adapter for fast personalized image	2104
2050	song Li, Manyuan Zhang, Ka Chun Cheung, Simon	generation. In <i>European Conference on Computer</i>	2105
2051	See, Hongwei Qin, Jifeng Dai, and Hongsheng Li.	<i>Vision</i> , pages 117–132. Springer.	2106
2052	2023a. Videoflow: Exploiting temporal cues for	Wenfeng Song, Xuan Wang, Shi Zheng, Shuai Li, Aimin	2107
2053	multi-frame optical flow estimation. In <i>Proceedings</i>	Hao, and Xia Hou. 2024a. Talkingstyle: Personal-	2108
2054	<i>of the IEEE/CVF International Conference on Com-</i>	ized speech-driven 3d facial animation with style	2109
2055	<i>puter Vision</i> , pages 12469–12480.	preservation. <i>IEEE Transactions on Visualization</i>	2110
		<i>and Computer Graphics.</i>	2111

2112	Yizhi Song, Zhifei Zhang, Zhe Lin, Scott Cohen, Brian Price, Jianming Zhang, Soo Ye Kim, He Zhang, Wei Xiong, and Daniel Aliaga. 2024b. Imprint: Generative object compositing by learning identity-preserving representation. In <i>Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition</i> , pages 8048–8058.	2168
2113		2169
2114		2170
2115		
2116		2171
2117		2172
2118		2173
2119	Yun-Zhu Song, Yi-Syuan Chen, Lu Wang, and Hong-Han Shuai. 2023. General then personal: Decoupling and pre-training for personalized headline generation. <i>Transactions of the Association for Computational Linguistics</i> , 11:1588–1607.	2174
2120		2175
2121		
2122		2176
2123		2177
2124	K Soomro. 2012. Ucf101: A dataset of 101 human actions classes from videos in the wild. <i>arXiv preprint arXiv:1212.0402</i> .	2178
2125		2179
2126		2180
2127	Ke Sun, Jian Cao, Qi Wang, Linrui Tian, Xindi Zhang, Lian Zhuo, Bang Zhang, Liefeng Bo, Wenbo Zhou, Weiming Zhang, et al. 2024. Outfitanyone: Ultra-high quality virtual try-on for any clothing and any person. <i>arXiv preprint arXiv:2407.16224</i> .	2181
2128		2182
2129		
2130		2183
2131		2184
2132	Peijie Sun, Le Wu, Kun Zhang, Yanjie Fu, Richang Hong, and Meng Wang. 2020. Dual learning for explainable recommendation: Towards unifying user preference prediction and review generation. In <i>Proceedings of The Web Conference 2020</i> , pages 837–847.	2185
2133		2186
2134		2187
2135		
2136		2188
2137		2189
2138	Kim Sung-Bin, Lee Chae-Yeon, Gihun Son, Oh Hyun-Bin, Janghoon Ju, Suekyeong Nam, and Tae-Hyun Oh. 2024. Multitalk: Enhancing 3d talking head generation across languages with multilingual video dataset. In <i>Interspeech 2024</i> , pages 1380–1384.	2190
2139		2191
2140		
2141		2192
2142		2193
2143	Shuai Tan, Bin Ji, Mengxiao Bi, and Ye Pan. 2025. Edtalk: Efficient disentanglement for emotional talking head synthesis. In <i>European Conference on Computer Vision</i> , pages 398–416. Springer.	2194
2144		2195
2145		2196
2146		2197
2147	Zhaoxuan Tan, Qingkai Zeng, Yijun Tian, Zheyuan Liu, Bing Yin, and Meng Jiang. 2024. Democratizing large language models via personalized parameter-efficient fine-tuning. <i>arXiv preprint arXiv:2402.04401</i> .	2198
2148		2199
2149		2200
2150		2201
2151		2202
2152	Brian Jay Tang, Kaiwen Sun, Noah T Curran, Florian Schaub, and Kang G Shin. 2024a. Genai advertising: Risks of personalizing ads with llms. <i>arXiv preprint arXiv:2409.15436</i> .	2203
2153		2204
2154		2205
2155		2206
2156	Yihong Tang, Bo Wang, Dongming Zhao, Jinxiao Jia, Jinxiao Jia, Zhangjijun Zhangjijun, Ruifang He, and Yuexian Hou. 2024b. MORPHEUS: Modeling role from personalized dialogue history by exploring and utilizing latent space. In <i>Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing</i> , pages 7664–7676, Miami, Florida, USA. Association for Computational Linguistics.	2207
2157		2208
2158		2209
2159		2210
2160		
2161		2211
2162		2212
2163		2213
2164	Yoad Tewel, Rinon Gal, Gal Chechik, and Yuval Atzmon. 2023. Key-locked rank one editing for text-to-image personalization. In <i>ACM SIGGRAPH 2023 Conference Proceedings</i> , pages 1–11.	2214
2165		2215
2166		2216
2167		
	Cynthia A. Thompson, Mehmet H. Göker, and Pat Langley. 2004. A personalized system for conversational recommendations. <i>J. Artif. Int. Res.</i> , 21(1):393–428.	2217
		2218
	Linrui Tian, Qi Wang, Bang Zhang, and Liefeng Bo. 2025. Emo: Emote portrait alive generating expressive portrait videos with audio2video diffusion model under weak conditions. In <i>European Conference on Computer Vision</i> , pages 244–260. Springer.	2219
		2220
	Yu-Min Tseng, Yu-Chao Huang, Teng-Yun Hsiao, Wei-Lin Chen, Chao-Wei Huang, Yu Meng, and Yun-Nung Chen. 2024. Two tales of persona in LLMs: A survey of role-playing and personalization. In <i>Findings of the Association for Computational Linguistics: EMNLP 2024</i> , pages 16612–16631. Association for Computational Linguistics.	2221
		2222
	Thomas Unterthiner, Sjoerd Van Steenkiste, Karol Kurach, Raphael Marinier, Marcin Michalski, and Sylvain Gelly. 2018. Towards accurate generative models of video: A new metric & challenges. <i>arXiv preprint arXiv:1812.01717</i> .	2223
		2224
	Dani Valevski, Danny Lumen, Yossi Matias, and Yaniv Leviathan. 2023. Face0: Instantaneously conditioning a text-to-image model on a face. In <i>SIGGRAPH Asia 2023 Conference Papers</i> , pages 1–10.	
	Thanh Van Le, Hao Phung, Thuan Hoang Nguyen, Quan Dao, Ngoc N Tran, and Anh Tran. 2023. Anti-dreambooth: Protecting users from personalized text-to-image synthesis. In <i>Proceedings of the IEEE/CVF International Conference on Computer Vision</i> , pages 2116–2127.	
	Pol van Rijn, Silvan Mertes, Dominik Schiller, Piotr Dura, Hubert Siuzdak, Peter Harrison, Elisabeth André, and Nori Jacoby. 2022. Voiceme: Personalized voice generation in tts. In <i>Interspeech 2022</i> , pages 2588–2592.	
	Shanu Vashishtha, Abhinav Prakash, Lalitesh Morishetti, Kaushiki Nag, Yokila Arora, Sushant Kumar, and Kannan Achan. 2024. Chaining text-to-image and large language model: A novel approach for generating personalized e-commerce banners. In <i>Proceedings of the 30th ACM SIGKDD Conference on Knowledge Discovery and Data Mining</i> , pages 5825–5835.	
	Timo Von Marcard, Roberto Henschel, Michael J Black, Bodo Rosenhahn, and Gerard Pons-Moll. 2018. Recovering accurate 3d human pose in the wild using imus and a moving camera. In <i>Proceedings of the European conference on computer vision (ECCV)</i> , pages 601–617.	
	Dimitri Von Rütte, Elisabetta Fedele, Jonathan Thomm, and Lukas Wolf. 2023. Fabric: Personalizing diffusion models with iterative feedback. <i>arXiv preprint arXiv:2307.10159</i> .	
	Andrey Voynov, Qinghao Chu, Daniel Cohen-Or, and Kfir Aberman. 2023. p+: Extended textual conditioning in text-to-image generation. <i>arXiv preprint arXiv:2303.09522</i> .	

2225	Cong Wan, Yuhang He, Xiang Song, and Yihong Gong.	Kam-Fai Wong. 2024d. Unims-rag: A unified multi-	2281
2226	2024. Prompt-agnostic adversarial perturbation for	source retrieval-augmented generation for personal-	2282
2227	customized diffusion models. In <i>The Thirty-eighth</i>	ized dialogue systems . <i>Preprint</i> , arXiv:2401.13256.	2283
2228	<i>Annual Conference on Neural Information Process-</i>		
2229	<i>ing Systems</i> .		
2230	Mengting Wan and Julian J. McAuley. 2018. Item rec-	Jianrong Wang, Zixuan Wang, Xiaosheng Hu, Xuewei	2284
2231	ommendation on monotonic behavior chains . In	Li, Qiang Fang, and Li Liu. 2022. Residual-guided	2285
2232	<i>Proceedings of the 12th ACM Conference on Rec-</i>	personalized speech synthesis based on face image.	2286
2233	<i>ommender Systems, RecSys 2018, Vancouver, BC,</i>	In <i>ICASSP 2022-2022 IEEE International Confer-</i>	2287
2234	<i>Canada, October 2-7, 2018</i> , pages 86–94. ACM.	<i>ence on Acoustics, Speech and Signal Processing</i>	2288
2235		(<i>ICASSP</i>), pages 4743–4747. IEEE.	2289
2236	Mengting Wan, Rishabh Misra, Ndapa Nakashole, and	Kaisiyuan Wang, Qianyi Wu, Linsen Song, Zhuoqian	2290
2237	Julian J. McAuley. 2019. Fine-grained spoiler detec-	Yang, Wayne Wu, Chen Qian, Ran He, Yu Qiao, and	2291
2238	tion from large-scale review corpora . In <i>Proceedings</i>	Chen Change Loy. 2020. Mead: A large-scale audio-	2292
2239	<i>of the 57th Conference of the Association for Compu-</i>	visual dataset for emotional talking-face generation.	2293
2240	<i>tational Linguistics, ACL 2019, Florence, Italy, July</i>	In <i>European Conference on Computer Vision</i> , pages	2294
2241	<i>28- August 2, 2019, Volume 1: Long Papers</i> , pages	700–717. Springer.	2295
2242	2605–2610. Association for Computational Linguis-		
2243	tics.	Noah Wang, Z.y. Peng, Haoran Que, Jiaheng Liu,	2296
2244	Siqi Wan, Yehao Li, Jingwen Chen, Yingwei Pan, Ting	Wangchunshu Zhou, Yuhang Wu, Hongcheng Guo,	2297
2245	Yao, Yang Cao, and Tao Mei. 2025. Improving vir-	Ruitong Gan, Zehao Ni, Jian Yang, Man Zhang,	2298
2246	tual try-on with garment-focused diffusion models.	Zhaoxiang Zhang, Wanli Ouyang, Ke Xu, Wenhao	2299
2247	In <i>European Conference on Computer Vision</i> , pages	Huang, Jie Fu, and Junran Peng. 2024e. RoleLLM:	2300
	184–199. Springer.	Benchmarking, eliciting, and enhancing role-playing	2301
2248		abilities of large language models . In <i>Findings of</i>	2302
2249	Bochao Wang, Huabin Zheng, Xiaodan Liang, Yimin	<i>the Association for Computational Linguistics: ACL</i>	2303
2250	Chen, Liang Lin, and Meng Yang. 2018a. Toward	2024, pages 14743–14777, Bangkok, Thailand. As-	2304
2251	characteristic-preserving image-based virtual try-on	sociation for Computational Linguistics.	2305
2252	network. In <i>Proceedings of the European conference</i>		
	<i>on computer vision (ECCV)</i> , pages 589–604.	Qixun Wang, Xu Bai, Haofan Wang, Zekui Qin, An-	2306
2253	Danqing Wang, Kevin Yang, Hanlin Zhu, Xiaomeng	thony Chen, Huaxia Li, Xu Tang, and Yao Hu. 2024f.	2307
2254	Yang, Andrew Cohen, Lei Li, and Yuandong Tian.	Instantid: Zero-shot identity-preserving generation	2308
2255	2024a. Learning personalized alignment for evalu-	in seconds. <i>arXiv preprint arXiv:2401.07519</i> .	2309
2256	ating open-ended text generation . In <i>Proceedings</i>		
2257	<i>of the 2024 Conference on Empirical Methods in</i>	Quan Wang, Sheng Li, Xinpeng Zhang, and Guorui	2310
2258	<i>Natural Language Processing</i> , pages 13274–13292,	Feng. 2023a. Rethinking neural style transfer: Gen-	2311
2259	Miami, Florida, USA. Association for Computational	erating personalized and watermarked stylized im-	2312
2260	Linguistics.	ages. In <i>Proceedings of the 31st ACM International</i>	2313
		<i>Conference on Multimedia</i> , pages 6928–6937.	2314
2261	Fu-Yun Wang, Zhaoyang Huang, Weikang Bian, Xi-		
2262	aoyu Shi, Keqiang Sun, Guanglu Song, Yu Liu, and	Tan Wang, Linjie Li, Kevin Lin, Chung-Ching Lin,	2315
2263	Hongsheng Li. 2024b. Animatecm: Computation-	Zhengyuan Yang, Hanwang Zhang, Zicheng Liu, and	2316
2264	efficient personalized style video generation without	Lijuan Wang. 2023b. Disco: Disentangled control	2317
2265	personalized video data. In <i>SIGGRAPH Asia 2024</i>	for referring human dance generation in real world.	2318
2266	<i>Technical Communications</i> , pages 1–5.	<i>arXiv e-prints</i> , pages arXiv–2307.	2319
2267	Hao Wang, Yitong Wang, Zheng Zhou, Xing Ji, Di-	Ting-Chun Wang, Ming-Yu Liu, Jun-Yan Zhu, Guilin	2320
2268	hong Gong, Jingchao Zhou, Zhifeng Li, and Wei Liu.	Liu, Andrew Tao, Jan Kautz, and Bryan Catan-	2321
2269	2018b. Cosface: Large margin cosine loss for deep	zaro. 2018c. Video-to-video synthesis. In <i>Proceed-</i>	2322
2270	face recognition. In <i>Proceedings of the IEEE con-</i>	<i>ings of the 32nd International Conference on Neu-</i>	2323
2271	<i>ference on computer vision and pattern recognition</i> ,	<i>ral Information Processing Systems, NIPS’18</i> , page	2324
2272	pages 5265–5274.	1152–1164. Curran Associates Inc.	2325
2273	Haotian Wang, Yuzhe Weng, Yueyan Li, Zilu Guo, Jun	Weijie Wang, Jichao Zhang, Chang Liu, Xia Li,	2326
2274	Du, Shutong Niu, Jiefeng Ma, Shan He, Xiaoyan Wu,	Xingqian Xu, Humphrey Shi, Nicu Sebe, and Bruno	2327
2275	Qiming Hu, et al. 2024c. Emotivetalk: Expressive	Lepri. 2024g. Uvmap-id: A controllable and person-	2328
2276	talking head generation through audio information	alized uv map generative model. In <i>Proceedings of</i>	2329
2277	decoupling and emotional video diffusion. <i>arXiv</i>	<i>the 32nd ACM International Conference on Multime-</i>	2330
2278	<i>preprint arXiv:2411.16726</i> .	<i>dia</i> , pages 10725–10734.	2331
2279	Hongru Wang, Wenyu Huang, Yang Deng, Rui Wang,	Wenjie Wang, Xinyu Lin, Fuli Feng, Xiangnan He, and	2332
2280	Zezhong Wang, Yufei Wang, Fei Mi, Jeff Z. Pan, and	Tat-Seng Chua. 2023c. Generative recommendation:	2333
		Towards next-generation recommender paradigm.	2334
		<i>arXiv preprint arXiv:2304.03516</i> .	2335

2336	Wenjie Wang, Xinyu Lin, Liuhui Wang, Fuli Feng, Yunshan Ma, and Tat-Seng Chua. 2023d. Causal disentangled recommendation against user preference shifts. <i>ACM Transactions on Information Systems</i> , 42(1):1–27.	2393
2337		2394
2338		2395
2339		2396
2340		2397
2341	Wenjie Wang, Yiyan Xu, Fuli Feng, Xinyu Lin, Xiangan He, and Tat-Seng Chua. 2023e. Diffusion recommender model . In <i>Proceedings of the 46th International ACM SIGIR Conference on Research and Development in Information Retrieval</i> , SIGIR '23, page 832–841, New York, NY, USA. Association for Computing Machinery.	2398
2342		2399
2343		2400
2344		
2345		2401
2346		2402
2347		2403
		2404
2348	X Wang, Siming Fu, Qihan Huang, Wanggui He, and Hao Jiang. 2024h. Ms-diffusion: Multi-subject zero-shot image personalization with layout guidance. <i>arXiv preprint arXiv:2406.07209</i> .	2405
2349		2406
2350		2407
2351		
2352	Xianquan Wang, Likang Wu, Shukang Yin, Zhi Li, Yanjiang Chen, Hufeng Hufeng, Yu Su, and Qi Liu. 2024i. I-am-g: Interest augmented multimodal generator for item personalization. In <i>Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing</i> , pages 21303–21317.	2408
2353		2409
2354		2410
2355		2411
2356		2412
2357		2413
2358	Xuan Wang, Guan hong Wang, Wenhao Chai, Jiayu Zhou, and Gaoang Wang. 2023f. User-aware prefix-tuning is a good learner for personalized image captioning. In <i>Chinese Conference on Pattern Recognition and Computer Vision (PRCV)</i> , pages 384–395. Springer.	2414
2359		2415
2360		2416
2361		2417
2362		2418
2363		2419
2364	Yanan Wang, Yan Pei, Zerui Ma, and Jianqiang Li. 2024j. A user-guided generation framework for personalized music synthesis using interactive evolutionary computation. In <i>Proceedings of the Genetic and Evolutionary Computation Conference Companion</i> , pages 1762–1769.	2420
2365		2421
2366		2422
2367		2423
2368		2424
2369		2425
2370	Yaqing Wang, Jiepu Jiang, Mingyang Zhang, Cheng Li, Yi Liang, Qiaozhu Mei, and Michael Bendersky. 2023g. Automated evaluation of personalized text generation using large language models . <i>Preprint</i> , arXiv:2310.11593.	2426
2371		2427
2372		2428
2373		2429
2374		
2375	Yuanbin Wang, Weilun Dai, Long Chan, Huanyu Zhou, Aixi Zhang, and Si Liu. 2024k. Gpd-vvto: Preserving garment details in video virtual try-on. In <i>Proceedings of the 32nd ACM International Conference on Multimedia</i> , pages 7133–7142.	2430
2376		2431
2377		2432
2378		2433
2379		2434
2380	Zhao Wang, Aoxue Li, Lingting Zhu, Yong Guo, Qi Dou, and Zhenguo Li. 2024l. Customvideo: Customizing text-to-video generation with multiple subjects. <i>arXiv preprint arXiv:2401.09962</i> .	2435
2381		2436
2382		2437
2383		
2384	Zhou Wang, Alan C Bovik, Hamid R Sheikh, and Eero P Simoncelli. 2004. Image quality assessment: from error visibility to structural similarity. <i>IEEE transactions on image processing</i> , 13(4):600–612.	2438
2385		2439
2386		2440
2387		2441
2388	Zhou Wang, Eero P Simoncelli, and Alan C Bovik. 2003. Multiscale structural similarity for image quality assessment. In <i>The Thrity-Seventh Asilomar Conference on Signals, Systems & Computers</i> , 2003, volume 2, pages 1398–1402. Ieee.	2442
2389		2443
2390		
2391		2444
2392		2445
		2446
		2447
		2448
		2449
		2450
	Fanyue Wei, Wei Zeng, Zhenyang Li, Dawei Yin, Lixin Duan, and Wen Li. 2025a. Powerful and flexible: Personalized text-to-image generation via reinforcement learning. In <i>European Conference on Computer Vision</i> , pages 394–410. Springer.	
	Huawei Wei, Zejun Yang, and Zhisheng Wang. 2024a. Aniportrait: Audio-driven synthesis of photorealistic portrait animation. <i>arXiv preprint arXiv:2403.17694</i> .	
	Yujie Wei, Shiwei Zhang, Zhiwu Qing, Hangjie Yuan, Zhiheng Liu, Yu Liu, Yingya Zhang, Jingren Zhou, and Hongming Shan. 2024b. Dreamvideo: Composing your dream videos with customized subject and motion. In <i>Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition</i> , pages 6537–6549.	
	Yujie Wei, Shiwei Zhang, Hangjie Yuan, Xiang Wang, Haonan Qiu, Rui Zhao, Yutong Feng, Feng Liu, Zhizhong Huang, Jiaxin Ye, et al. 2024c. Dreamvideo-2: Zero-shot subject-driven video customization with precise motion control. <i>arXiv preprint arXiv:2410.13830</i> .	
	Yuxiang Wei, Zhilong Ji, Jinfeng Bai, Hongzhi Zhang, Lei Zhang, and Wangmeng Zuo. 2025b. Masterweaver: Taming editability and face identity for personalized text-to-image generation. In <i>European Conference on Computer Vision</i> , pages 252–271. Springer.	
	Yuxiang Wei, Yabo Zhang, Zhilong Ji, Jinfeng Bai, Lei Zhang, and Wangmeng Zuo. 2023. Elite: Encoding visual concepts into textual embeddings for customized text-to-image generation. In <i>Proceedings of the IEEE/CVF International Conference on Computer Vision</i> , pages 15943–15953.	
	Matthias Wright and Björn Ommer. 2022. Artfid: Quantitative evaluation of neural style transfer. In <i>DAGM German Conference on Pattern Recognition</i> , pages 560–576. Springer.	
	Fangzhao Wu, Ying Qiao, Jiun-Hung Chen, Chuhan Wu, Tao Qi, Jianxun Lian, Danyang Liu, Xing Xie, Jianfeng Gao, Winnie Wu, and Ming Zhou. 2020a. MIND: A large-scale dataset for news recommendation . In <i>Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics</i> , pages 3597–3606, Online. Association for Computational Linguistics.	
	Fangzhao Wu, Ying Qiao, Jiun-Hung Chen, Chuhan Wu, Tao Qi, Jianxun Lian, Danyang Liu, Xing Xie, Jianfeng Gao, Winnie Wu, et al. 2020b. Mind: A large-scale dataset for news recommendation. In <i>Proceedings of the 58th annual meeting of the association for computational linguistics</i> , pages 3597–3606.	
	Haoning Wu, Erli Zhang, Liang Liao, Chaofeng Chen, Jingwen Hou, Annan Wang, Wenxiu Sun, Qiong Yan, and Weisi Lin. 2023a. Exploring video quality assessment on user generated contents from aesthetic and technical perspectives. In <i>Proceedings of the IEEE/CVF International Conference on Computer Vision</i> , pages 20144–20154.	

2451	Jianzong Wu, Xiangtai Li, Yanhong Zeng, Jiangning	Guangxuan Xiao, Tianwei Yin, William T. Freeman,	2507
2452	Zhang, Qianyu Zhou, Yining Li, Yunhai Tong, and	Frédo Durand, and Song Han. 2024b. Fastcomposer:	2508
2453	Kai Chen. 2024a. Motionbooth: Motion-aware cus-	Tuning-free multi-subject image generation with lo-	2509
2454	tomized text-to-video generation. <i>arXiv preprint</i>	calized attention. <i>International Journal of Computer</i>	2510
2455	<i>arXiv:2406.17758</i> .	<i>Vision</i> .	2511
2456	Junda Wu, Hanjia Lyu, Yu Xia, Zhehao Zhang, Joe Bar-	Zhujun Xiao, Jenna Cryan, Yuanshun Yao, Yi Hong	2512
2457	row, Ishita Kumar, Mehrnoosh Mirtaheri, Hongjie	Cheo, Yuanchao Shu, Stefan Saroiu, Ben Y Zhao,	2513
2458	Chen, Ryan A Rossi, Franck Dernoncourt, et al.	and Haitao Zheng. 2023. "my face, my rules": En-	2514
2459	2024b. Personalized multimodal large language mod-	abling personalized protection against unacceptable	2515
2460	els: A survey. <i>arXiv preprint arXiv:2412.02142</i> .	face editing. <i>Proceedings on Privacy Enhancing</i>	2516
2461	Likang Wu, Zhi Zheng, Zhaopeng Qiu, Hao Wang,	<i>Technologies</i> , 2023(3).	2517
2462	Hongchao Gu, Tingjia Shen, Chuan Qin, Chen Zhu,	Zhenyu Xie, Haoye Dong, Yufei Gao, Zehua Ma, and	2518
2463	Hengshu Zhu, Qi Liu, Hui Xiong, and Enhong Chen.	Xiaodan Liang. 2024. Dreamvton: Customizing 3d	2519
2464	2024c. A survey on large language models for rec-	virtual try-on with personalized diffusion models. In	2520
2465	ommendation . <i>Preprint</i> , arXiv:2305.19860.	<i>Proceedings of the 32nd ACM International Confer-</i>	2521
2466	Tao Wu, Yong Zhang, Xintao Wang, Xianpan Zhou,	<i>ence on Multimedia</i> , pages 10784–10793.	2522
2467	Guangcong Zheng, Zhongang Qi, Ying Shan, and	Kun Xiong, Liu Jiang, Xuan Dang, Guolong Wang,	2523
2468	Xi Li. 2024d. Customcrafter: Customized video	Wenwen Ye, and Zheng Qin. 2020. Towards person-	2524
2469	generation with preserving motion and concept com-	alized aesthetic image caption. In <i>2020 International</i>	2525
2470	position abilities. <i>arXiv preprint arXiv:2408.13239</i> .	<i>Joint Conference on Neural Networks (IJCNN)</i> , pages	2526
2471	Xiaoshi Wu, Keqiang Sun, Feng Zhu, Rui Zhao, and	1–8. IEEE.	2527
2472	Hongsheng Li. 2023b. Human preference score: Bet-	Yuliang Xiu, Yufei Ye, Zhen Liu, Dimitris Tzionas, and	2528
2473	ter aligning text-to-image models with human prefer-	Michael J Black. 2024. Puzzleavatar: Assembling 3d	2529
2474	ence. In <i>Proceedings of the IEEE/CVF International</i>	avatars from personal albums. <i>ACM Transactions on</i>	2530
2475	<i>Conference on Computer Vision</i> , pages 2096–2105.	<i>Graphics (TOG)</i> , 43(6):1–15.	2531
2476	Yi Wu, Ziqiang Li, Heliang Zheng, Chaoyue Wang, and	Jiazheng Xu, Xiao Liu, Yuchen Wu, Yuxuan Tong,	2532
2477	Bin Li. 2025a. Infinite-id: Identity-preserved per-	Qinkai Li, Ming Ding, Jie Tang, and Yuxiao Dong.	2533
2478	sonalization via id-semantics decoupling paradigm.	2023a. Imagereward: Learning and evaluating hu-	2534
2479	In <i>European Conference on Computer Vision</i> , pages	man preferences for text-to-image generation. <i>Ad-</i>	2535
2480	279–296. Springer.	<i>vances in Neural Information Processing Systems</i> ,	2536
2481	Yihan Wu, Ruihua Song, Xu Chen, Hao Jiang, Zhao	36:15903–15935.	2537
2482	Cao, and Jin Yu. 2024e. Understanding human pref-	Tao Xu, Pengchuan Zhang, Qiuyuan Huang, Han Zhang,	2538
2483	erences: Towards more personalized video to text	Zhe Gan, Xiaolei Huang, and Xiaodong He. 2018.	2539
2484	generation. In <i>Proceedings of the ACM on Web Con-</i>	Attngan: Fine-grained text to image generation with	2540
2485	<i>ference 2024</i> , pages 3952–3963.	attentional generative adversarial networks. In <i>Pro-</i>	2541
2486	Yongliang Wu, Shiji Zhou, Mingzhuo Yang, Lianzhe	<i>ceedings of the IEEE conference on computer vision</i>	2542
2487	Wang, Heng Chang, Wenbo Zhu, Xinting Hu, Xiao	<i>and pattern recognition</i> , pages 1316–1324.	2543
2488	Zhou, and Xu Yang. 2025b. Unlearning concepts in	Xinchao Xu, Zeyang Lei, Wenquan Wu, Zheng-Yu Niu,	2544
2489	diffusion model via concept domain correction and	Hua Wu, and Haifeng Wang. 2023b. Towards zero-	2545
2490	concept preserving gradient. In <i>Proceedings of the</i>	shot persona dialogue generation with in-context	2546
2491	<i>AAAI Conference on Artificial Intelligence (AAAI)</i> .	learning. In <i>Findings of the Association for Com-</i>	2547
2492	Yujia Wu, Yiming Shi, Jiwei Wei, Chengwei Sun,	<i>putational Linguistics: ACL 2023</i> , pages 1387–1398.	2548
2493	Yuyang Zhou, Yang Yang, and Heng Tao Shen.	Xinchao Xu, Zeyang Lei, Wenquan Wu, Zheng-Yu Niu,	2549
2494	2024f. Diffhora: Generating personalized low-rank	Hua Wu, and Haifeng Wang. 2023c. Towards zero-	2550
2495	adaptation weights with diffusion. <i>arXiv preprint</i>	shot persona dialogue generation with in-context	2551
2496	<i>arXiv:2408.06740</i> .	learning . In <i>Findings of the Association for Com-</i>	2552
2497	Weihao Xia, Yujiu Yang, Jing-Hao Xue, and Baoyuan	<i>putational Linguistics: ACL 2023</i> , pages 1387–1398,	2553
2498	Wu. 2021. Tedigan: Text-guided diverse face im-	Toronto, Canada. Association for Computational Lin-	2554
2499	age generation and manipulation. In <i>Proceedings of</i>	guistics.	2555
2500	<i>the IEEE/CVF conference on computer vision and</i>	Yifei Xu, Xiaolong Xu, Honghao Gao, and Fu Xiao.	2556
2501	<i>pattern recognition</i> , pages 2256–2265.	2024a. Sgdm: An adaptive style-guided diffusion	2557
2502	Guangxuan Xiao, Tianwei Yin, William T Freeman,	model for personalized text to image generation.	2558
2503	Frédo Durand, and Song Han. 2024a. Fastcomposer:	<i>IEEE Transactions on Multimedia</i> .	2559
2504	Tuning-free multi-subject image generation with lo-	Yiyan Xu, Wenjie Wang, Fuli Feng, Yunshan Ma, Jizhi	2560
2505	calized attention. <i>International Journal of Computer</i>	Zhang, and Xiangnan He. 2024b. Diffusion models	2561
2506	<i>Vision</i> , pages 1–20.		

2562	for generative outfit recommendation. In <i>Proceedings of the 47th International ACM SIGIR Conference on Research and Development in Information Retrieval</i> , pages 1350–1359.	2618
2563		2619
2564		2620
2565		2621
2566	Yiyan Xu, Wenjie Wang, Yang Zhang, Tang Biao, Peng Yan, Fuli Feng, and Xiangnan He. 2024c. Personalized image generation with large multimodal models. <i>arXiv preprint arXiv:2410.14170</i> .	2622
2567		
2568		
2569		
2570	Yuhao Xu, Tao Gu, Weifeng Chen, and Chengcai Chen. 2024d. Ootdiffusion: Outfitting fusion based latent diffusion for controllable virtual try-on. <i>arXiv preprint arXiv:2403.01779</i> .	
2571		
2572		
2573		
2574	Zhengze Xu, Mengting Chen, Zhao Wang, Linyu Xing, Zhonghua Zhai, Nong Sang, Jinsong Lan, Shuai Xiao, and Changxin Gao. 2024e. Tunnel try-on: Excavating spatial-temporal tunnels for high-quality virtual try-on in videos. In <i>Proceedings of the 32nd ACM International Conference on Multimedia</i> , pages 3199–3208.	
2575		
2576		
2577		
2578		
2579		
2580		
2581	Zhongcong Xu, Jianfeng Zhang, Jun Hao Liew, Hanshu Yan, Jia-Wei Liu, Chenxu Zhang, Jiashi Feng, and Mike Zheng Shou. 2024f. Magicanimate: Temporally consistent human image animation using diffusion model. In <i>Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition</i> , pages 1481–1490.	
2582		
2583		
2584		
2585		
2586		
2587		
2588	Yuxuan Yan, Chi Zhang, Rui Wang, Yichao Zhou, Gege Zhang, Pei Cheng, Gang Yu, and Bin Fu. 2023. Facestudio: Put your face everywhere in seconds. <i>arXiv preprint arXiv:2312.02663</i> .	
2589		
2590		
2591		
2592	Fan Yang, Zheng Chen, Ziyang Jiang, Eunah Cho, Xiaojian Huang, and Yanbin Lu. 2023a. Palr: Personalization aware llms for recommendation . <i>Preprint</i> , arXiv:2305.07622.	
2593		
2594		
2595		
2596	Hao Yang, Jianxin Yuan, Shuai Yang, Linhe Xu, Shuo Yuan, and Yifan Zeng. 2024a. A new creative generation pipeline for click-through rate with stable diffusion model. In <i>Companion Proceedings of the ACM on Web Conference 2024</i> , pages 180–189.	
2597		
2598		
2599		
2600		
2601	Jianan Yang, Haobo Wang, Yanming Zhang, Ruixuan Xiao, Sai Wu, Gang Chen, and Junbo Zhao. 2023b. Controllable textual inversion for personalized text-to-image generation. <i>arXiv preprint arXiv:2304.05265</i> .	
2602		
2603		
2604		
2605		
2606	Jingfeng Yang, Hongye Jin, Ruixiang Tang, Xiaotian Han, Qizhang Feng, Haoming Jiang, Shaochen Zhong, Bing Yin, and Xia Hu. 2024b. Harnessing the power of llms in practice: A survey on chatgpt and beyond. <i>ACM Transactions on Knowledge Discovery from Data</i> , 18(6):1–32.	
2607		
2608		
2609		
2610		
2611		
2612	Sidi Yang, Tianhe Wu, Shuwei Shi, Shanshan Lao, Yuan Gong, Mingdeng Cao, Jiahao Wang, and Yujiu Yang. 2022. Maniqa: Multi-dimension attention network for no-reference image quality assessment. In <i>Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition</i> , pages 1191–1200.	
2613		
2614		
2615		
2616		
2617		
	Wanting Yang, Zehui Xiong, Song Guo, Shiwen Mao, Dong In Kim, and Merouane Debbah. 2024c. Efficient multi-user offloading of personalized diffusion models: A drl-convex hybrid solution. <i>arXiv preprint arXiv:2411.15781</i> .	2623
		2624
	Yang Yang, Wen Wang, Liang Peng, Chaotian Song, Yao Chen, Hengjia Li, Xiaolong Yang, Qinglin Lu, Deng Cai, Boxi Wu, et al. 2024d. Lora-composer: Leveraging low-rank adaptation for multi-concept customization in training-free diffusion models. <i>arXiv preprint arXiv:2403.11627</i> .	2625
		2626
		2627
		2628
	Zhengyi Yang, Jiancan Wu, Zhicai Wang, Yancheng Yuan, Xiang Wang, and Xiangnan He. 2024e. Generate what you prefer: reshaping sequential recommendation via guided diffusion. In <i>Proceedings of the 37th International Conference on Neural Information Processing Systems, NIPS '23</i> , Red Hook, NY, USA. Curran Associates Inc.	2629
		2630
		2631
		2632
		2633
		2634
		2635
	Zhijing Yang, Mingliang Yang, Siyuan Peng, Jieming Xie, and Chao Long. 2024f. Animated clothing fitting: Semantic-visual fusion for video customization virtual try-on with diffusion models. In <i>2024 5th International Conference on Electronic Communication and Artificial Intelligence (ICECAI)</i> , pages 675–680. IEEE.	2636
		2637
		2638
		2639
		2640
		2641
		2642
	Shunyu Yao, RuiZhe Zhong, Yichao Yan, Guangtao Zhai, and Xiaokang Yang. 2022. Dfa-nerf: Personalized talking head generation via disentangled face attributes neural rendering. <i>arXiv preprint arXiv:2201.00791</i> .	2643
		2644
		2645
		2646
		2647
	Zebin Yao, Fangxiang Feng, Ruifan Li, and Xiaojie Wang. 2024. Concept conductor: Orchestrating multiple personalized concepts in text-to-image synthesis. <i>arXiv preprint arXiv:2408.03632</i> .	2648
		2649
		2650
		2651
	Hu Ye, Jun Zhang, Sibao Liu, Xiao Han, and Wei Yang. 2023. Ip-adapter: Text compatible image prompt adapter for text-to-image diffusion models. <i>arXiv preprint arXiv:2308.06721</i> .	2652
		2653
		2654
		2655
	Zixuan Ye, Huijuan Huang, Xintao Wang, Pengfei Wan, Di Zhang, and Wenhan Luo. 2024. Stylemaster: Stylize your video with artistic generation and translation. <i>arXiv preprint arXiv:2412.07744</i> .	2656
		2657
		2658
		2659
	Yu-Ying Yeh, Jia-Bin Huang, Changil Kim, Lei Xiao, Thu Nguyen-Phuoc, Numair Khan, Cheng Zhang, Manmohan Chandraker, Carl S Marshall, Zhao Dong, et al. 2024. Texturedreamer: Image-guided texture synthesis through geometry-aware diffusion. In <i>Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition</i> , pages 4304–4314.	2660
		2661
		2662
		2663
		2664
		2665
		2666
	Ran Yi, Zipeng Ye, Juyong Zhang, Hujun Bao, and Yong-Jin Liu. 2020. Audio-driven talking face video generation with learning-based personalized head pose. <i>arXiv preprint arXiv:2002.10137</i> .	2667
		2668
		2669
		2670
	Cong Yu, Yang Hu, Yan Chen, and Bing Zeng. 2019. Personalized fashion design. In <i>Proceedings of the IEEE/CVF international conference on computer vision</i> , pages 9046–9055.	2671
		2672
		2673
		2674

2675	Jianhui Yu, Hao Zhu, Liming Jiang, Chen Change Loy,	shape estimation from clothed 3d scan sequences. In	2731
2676	Weidong Cai, and Wayne Wu. 2023. Celebv-text: A	<i>Proceedings of the IEEE Conference on Computer</i>	2732
2677	large-scale facial text-video dataset. In <i>Proceedings</i>	<i>Vision and Pattern Recognition</i> , pages 4191–4200.	2733
2678	of the <i>IEEE/CVF Conference on Computer Vision</i>		
2679	and <i>Pattern Recognition</i> , pages 14805–14814.		
2680	Zhengyang Yu, Zhaoyuan Yang, and Jing Zhang. 2024.	Chenxu Zhang, Saifeng Ni, Zhipeng Fan, Hongbo Li,	2734
2681	Dreamsteerer: Enhancing source image conditioned	Ming Zeng, Madhukar Budagavi, and Xiaohu Guo.	2735
2682	editability using personalized diffusion models . In	2021a. 3d talking face with personalized pose dy-	2736
2683	<i>The Thirty-eighth Annual Conference on Neural In-</i>	namics. <i>IEEE Transactions on Visualization and</i>	2737
2684	<i>formation Processing Systems</i> .	<i>Computer Graphics</i> , 29(2):1438–1449.	2738
2685	Ge Yuan, Xiaodong Cun, Yong Zhang, Maomao	Chenxu Zhang, Chao Wang, Jianfeng Zhang, Hongyi	2739
2686	Li, Chenyang Qi, Xintao Wang, Ying Shan, and	Xu, Guoxian Song, You Xie, Linjie Luo, Yapeng	2740
2687	Huicheng Zheng. 2023. Inserting anybody in dif-	Tian, Xiaohu Guo, and Jiashi Feng. 2023b. Dream-	2741
2688	fusion models via celeb basis . In <i>Thirty-seventh Con-</i>	talk: diffusion-based realistic emotional audio-driven	2742
2689	<i>ference on Neural Information Processing Systems</i> .	method for single image talking face generation.	2743
2690		<i>arXiv preprint arXiv:2312.13578</i> .	2744
2691	Shenghai Yuan, Jinfa Huang, Xianyi He, Yunyuan Ge,	Jiarui Zhang. 2024. Guided profile generation improves	2745
2692	Yujun Shi, Liuhan Chen, Jiebo Luo, and Li Yuan.	personalization with large language models. In <i>Find-</i>	2746
2693	2024. Identity-preserving text-to-video genera-	<i>ings of the Association for Computational Linguistics:</i>	2747
2694	tion by frequency decomposition. <i>arXiv preprint</i>	<i>EMNLP 2024</i> , pages 4005–4016.	2748
2695	<i>arXiv:2411.17440</i> .		
2696	Dongxu Yue, Maomao Li, Yunfei Liu, Qin Guo, Ailing	Kai Zhang, Yangyang Kang, Fubang Zhao, and Xi-	2749
2697	Zeng, Tianyu Yang, and Yu Li. 2025. Addme: Zero-	aozhong Liu. 2024a. Llm-based medical assistant	2750
2698	shot group-photo synthesis by inserting people into	personalization with short-and long-term memory co-	2751
2699	scenes. In <i>European Conference on Computer Vision</i> ,	ordination. In <i>Proceedings of the 2024 Conference</i>	2752
2700	pages 222–239. Springer.	<i>of the North American Chapter of the Association for</i>	2753
2701	Polina Zablotskaia, Aliaksandr Siarohin, Bo Zhao, and	<i>Computational Linguistics: Human Language Tech-</i>	2754
2702	Leonid Sigal. 2019. Dwnet: Dense warp-based net-	<i>nologies (Volume 1: Long Papers)</i> , pages 2386–2398.	2755
2703	work for pose-guided human video generation. <i>arXiv</i>		
2704	<i>preprint arXiv:1910.09139</i> .	Kai Zhang, Lizhi Qing, Yangyang Kang, and Xiaozhong	2756
2705	Heiga Zen, Viet Dang, Rob Clark, Yu Zhang, Ron J	Liu. 2024b. Personalized llm response generation	2757
2706	Weiss, Ye Jia, Zhifeng Chen, and Yonghui Wu. 2019.	with parameterized memory injection. <i>arXiv preprint</i>	2758
2707	Libritts: A corpus derived from librispeech for text-	<i>arXiv:2404.03565</i> .	2759
2708	to-speech. <i>arXiv preprint arXiv:1904.02882</i> .		
2709	Andy Zeng, Pete Florence, Jonathan Tompson, Stefan	Kaipeng Zhang, Zhanpeng Zhang, Zhifeng Li, and	2760
2710	Welker, Jonathan Chien, Maria Attarian, Travis Arm-	Yu Qiao. 2016. Joint face detection and alignment	2761
2711	strong, Ivan Krasin, Dan Duong, Vikas Sindhwani,	using multitask cascaded convolutional networks.	2762
2712	et al. 2021. Transporter networks: Rearranging the	<i>IEEE signal processing letters</i> , 23(10):1499–1503.	2763
2713	visual world for robotic manipulation. In <i>Conference</i>		
2714	<i>on Robot Learning</i> , pages 726–747. PMLR.	Mozhi Zhang, Pengyu Wang, Chenkun Tan, Mianqiu	2764
2715	Hansi Zeng, Surya Kallumadi, Zaid Alibadi, Rodrigo	Huang, Dong Zhang, Yaqian Zhou, and Xipeng Qiu.	2765
2716	Nogueira, and Hamed Zamani. 2023. A personalized	2024c. Metaalign: Align large language models with	2766
2717	dense retrieval framework for unified information ac-	diverse preferences during inference time. <i>arXiv</i>	2767
2718	cess . In <i>Proceedings of the 46th International ACM</i>	<i>preprint arXiv:2410.14184</i> .	2768
2719	<i>SIGIR Conference on Research and Development</i>		
2720	<i>in Information Retrieval</i> , SIGIR '23, page 121–130,	Richard Zhang, Phillip Isola, Alexei A Efros, Eli Shecht-	2769
2721	New York, NY, USA. Association for Computing	man, and Oliver Wang. 2018. The unreasonable ef-	2770
2722	Machinery.	fectiveness of deep features as a perceptual metric.	2771
2723	Bowen Zhang, Chenyang Qi, Pan Zhang, Bo Zhang,	In <i>Proceedings of the IEEE conference on computer</i>	2772
2724	HsiangTao Wu, Dong Chen, Qifeng Chen, Yong	<i>vision and pattern recognition</i> , pages 586–595.	2773
2725	Wang, and Fang Wen. 2023a. Metaportrait: Identity-		
2726	preserving talking head generation with fast person-	Shufang Zhang, Minxue Ni, Shuai Chen, Lei Wang,	2774
2727	alized adaptation. In <i>Proceedings of the IEEE/CVF</i>	Wenxin Ding, and Yuhong Liu. 2024d. A two-stage	2775
2728	<i>Conference on Computer Vision and Pattern Recog-</i>	personalized virtual try-on framework with shape	2776
2729	<i>nition</i> , pages 22096–22105.	control and texture guidance. <i>IEEE Transactions on</i>	2777
2730	Chao Zhang, Sergi Pujades, Michael J Black, and Ger-	<i>Multimedia</i> .	2778
	ard Pons-Moll. 2017. Detailed, accurate, human	Tianshu Zhang, Buzhen Huang, and Yangang Wang.	2779
		2020a. Object-occluded human shape and pose es-	2780
		timation from a single color image. In <i>Proceedings</i>	2781
		<i>of the IEEE/CVF conference on computer vision and</i>	2782
		<i>pattern recognition</i> , pages 7376–7385.	2783

2784	Tianyi Zhang, Varsha Kishore, Felix Wu, Kilian Q.	Jiang. 2025. Magdiff: Multi-alignment diffusion	2840
2785	Weinberger, and Yoav Artzi. 2020b. Bertscore: Eval-	for high-fidelity video generation and editing. In	2841
2786	uating text generation with bert . In <i>International</i>	<i>European Conference on Computer Vision</i> , pages	2842
2787	<i>Conference on Learning Representations</i> .	205–221. Springer.	2843
2788	Wei Zhang, Yue Ying, Pan Lu, and Hongyuan Zha.	Jian Zhao, Jianshu Li, Yu Cheng, Terence Sim,	2844
2789	2020c. Learning long-and short-term user literal-	Shuicheng Yan, and Jiashi Feng. 2018. Understand-	2845
2790	preference with multimodal hierarchical transformer	ing humans in crowded scenes: Deep nested adversar-	2846
2791	network for personalized image caption. In <i>Proceed-</i>	ial learning and a new benchmark for multi-human	2847
2792	<i>ings of the AAAI Conference on Artificial Intelligence</i> ,	parsing. In <i>Proceedings of the 26th ACM interna-</i>	2848
2793	volume 34, pages 9571–9578.	<i>tional conference on Multimedia</i> , pages 792–800.	2849
2794	Weixia Zhang, Guangtao Zhai, Ying Wei, Xiaokang	Jun Zheng, Fuwei Zhao, Youjiang Xu, Xin Dong, and	2850
2795	Yang, and Kede Ma. 2023c. Blind image quality	Xiaodan Liang. 2024a. Viton-dit: Learning in-the-	2851
2796	assessment via vision-language correspondence: A	wild video try-on from human dance videos via diffu-	2852
2797	multitask learning perspective. In <i>Proceedings of</i>	sion transformers. <i>arXiv preprint arXiv:2405.18326</i> .	2853
2798	<i>the IEEE/CVF conference on computer vision and</i>	Longtao Zheng, Yifan Zhang, Hanzhong Guo, Jiachun	2854
2799	<i>pattern recognition</i> , pages 14071–14081.	Pan, Zhenxiong Tan, Jiahao Lu, Chuanxin Tang,	2855
2800	Xujie Zhang, Ente Lin, Xiu Li, Yuxuan Luo, Michael	Bo An, and Shuicheng Yan. 2024b. Memo: Memory-	2856
2801	Kampffmeyer, Xin Dong, and Xiaodan Liang. 2024e.	guided diffusion for expressive talking video genera-	2857
2802	Mmtryon: Multi-modal multi-reference control for	tion. <i>arXiv preprint arXiv:2412.04448</i> .	2858
2803	high-quality fashion generation. <i>arXiv preprint</i>	Yinglin Zheng, Hao Yang, Ting Zhang, Jianmin Bao,	2859
2804	<i>arXiv:2405.00448</i> .	Dongdong Chen, Yangyu Huang, Lu Yuan, Dong	2860
2805	Xulu Zhang, Xiao-Yong Wei, Jinlin Wu, Tianyi Zhang,	Chen, Ming Zeng, and Fang Wen. 2022. General	2861
2806	Zhaoxiang Zhang, Zhen Lei, and Qing Li. 2024f.	facial representation learning in a visual-linguistic	2862
2807	Compositional inversion for stable diffusion models.	manner. In <i>Proceedings of the IEEE/CVF conference</i>	2863
2808	In <i>Proceedings of the AAAI Conference on Artificial</i>	<i>on computer vision and pattern recognition</i> , pages	2864
2809	<i>Intelligence</i> , volume 38, pages 7350–7358.	18697–18709.	2865
2810	Xulu Zhang, Xiao-Yong Wei, Wengyu Zhang, Jinlin Wu,	Yinhe Zheng, Rongsheng Zhang, Minlie Huang, and	2866
2811	Zhaoxiang Zhang, Zhen Lei, and Qing Li. 2024g. A	Xiaoxi Mao. 2020. A pre-training based personalized	2867
2812	survey on personalized content synthesis with diffu-	dialogue generation model with persona-sparse data .	2868
2813	sion models. <i>arXiv preprint arXiv:2405.05538</i> .	<i>Proceedings of the AAAI Conference on Artificial</i>	2869
2814	Yiming Zhang, Zhening Xing, Yanhong Zeng, Youqing	<i>Intelligence</i> , 34(05):9693–9700.	2870
2815	Fang, and Kai Chen. 2024h. Pia: Your personalized	Xiaoqing Zhong, Zhonghua Wu, Taizhe Tan, Guosheng	2871
2816	image animator via plug-and-play modules in text-	Lin, and Qingyao Wu. 2021. Mv-ton: Memory-based	2872
2817	to-image models. In <i>Proceedings of the IEEE/CVF</i>	video virtual try-on network. In <i>Proceedings of the</i>	2873
2818	<i>Conference on Computer Vision and Pattern Recog-</i>	<i>29th ACM International Conference on Multimedia</i> ,	2874
2819	<i>nition</i> , pages 7747–7756.	pages 908–916.	2875
2820	Yuxuan Zhang, Yiren Song, Jiaming Liu, Rui Wang,	Yong Zhong, Min Zhao, Zebin You, Xiaofeng Yu,	2876
2821	Jinpeng Yu, Hao Tang, Huaxia Li, Xu Tang, Yao Hu,	Changwang Zhang, and Chongxuan Li. 2025. Pose-	2877
2822	Han Pan, et al. 2024i. Ssr-encoder: Encoding selec-	crafter: One-shot personalized video synthesis fol-	2878
2823	tive subject representation for subject-driven gener-	lowing flexible pose control. In <i>European Confer-</i>	2879
2824	ation. In <i>Proceedings of the IEEE/CVF Conference</i>	<i>ence on Computer Vision</i> , pages 243–260. Springer.	2880
2825	<i>on Computer Vision and Pattern Recognition</i> , pages	Dingfu Zhou, Jin Fang, Xibin Song, Chenye Guan,	2881
2826	8069–8078.	Junbo Yin, Yuchao Dai, and Ruigang Yang. 2019.	2882
2827	Zhehao Zhang, Ryan A Rossi, Branislav Kveton, Yijia	Iou loss for 2d/3d object detection. In <i>2019 interna-</i>	2883
2828	Shao, Diyi Yang, Hamed Zamani, Franck Dernon-	<i>tional conference on 3D vision (3DV)</i> , pages 85–94.	2884
2829	court, Joe Barrow, Tong Yu, Sungchul Kim, et al.	IEEE.	2885
2830	2024j. Personalization of large language models: A	Yufan Zhou, Ruiyi Zhang, Jiuxiang Gu, and Tong Sun.	2886
2831	survey. <i>arXiv preprint arXiv:2411.00027</i> .	2024. Customization assistant for text-to-image gen-	2887
2832	Zhimeng Zhang, Lincheng Li, Yu Ding, and Changjie	eration. In <i>Proceedings of the IEEE/CVF Conference</i>	2888
2833	Fan. 2021b. Flow-guided one-shot talking face gen-	<i>on Computer Vision and Pattern Recognition</i> , pages	2889
2834	eration with a high-resolution audio-visual dataset.	9182–9191.	2890
2835	In <i>Proceedings of the IEEE/CVF Conference on Com-</i>	Yufan Zhou, Ruiyi Zhang, Tong Sun, and Jinhui Xu.	2891
2836	<i>puter Vision and Pattern Recognition</i> , pages 3661–	2023. Enhancing detail preservation for customized	2892
2837	3670.	text-to-image generation: A regularization-free ap-	2893
2838	Haoyu Zhao, Tianyi Lu, Jiayi Gu, Xing Zhang, Qing-	proach. <i>arXiv preprint arXiv:2305.13579</i> .	2894
2839	ping Zheng, Zuxuan Wu, Hang Xu, and Yu-Gang		

Chenyang Zhu, Kai Li, Yue Ma, Chunming He, and Li Xiu. 2024. Multiboost: Towards generating all your concepts in an image from text. *arXiv preprint arXiv:2404.14239*.

Luyang Zhu, Dawei Yang, Tyler Zhu, Fitsum Reda, William Chan, Chitwan Saharia, Mohammad Norouzi, and Ira Kemelmacher-Shlizerman. 2023a. Tryondiffusion: A tale of two unets. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 4606–4615.

Luyao Zhu, Wei Li, Rui Mao, Vlad Pandealea, and Erik Cambria. 2023b. **PAED: Zero-shot persona attribute extraction in dialogues**. In *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 9771–9787, Toronto, Canada. Association for Computational Linguistics.

Yizhe Zhua, Chunhui Zhanga, Qiong Liub, and Xi Zhou. 2023. Audio-driven talking head video generation with diffusion model. In *ICASSP 2023-2023 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pages 1–5. IEEE.

Yuchen Zhuang, Haotian Sun, Yue Yu, Rushi Qiang, Qifan Wang, Chao Zhang, and Bo Dai. 2024. **HYDRA: Model factorization framework for black-box LLM personalization**. In *The Thirty-eighth Annual Conference on Neural Information Processing Systems*.

A Personalized Generation Across Modalities

A.1 Personalized Image Generation

Personalized image generation aims to synthesize images that reflect individual preferences and requirements. By incorporating various personalized contexts, existing studies have made significant strides in enhancing the capability of generative models to produce images tailored to specific needs, ranging from general-purpose generation to more specialized tasks.

A.1.1 User Behaviors

User interactions serve as a key source for inferring visual preferences, guiding the personalized image generation process. Based on user behaviors such as historical engagements and real-time feedback, existing methods have explored various approaches to enhance personalization.

General-purpose Image Generation This task involves generating tailored images across various scenarios. For instance, PMG (Shen et al., 2024b), I-AM-G (Wang et al., 2024i), and Pigeon (Xu et al., 2024c) utilize historically interacted images

of users to infer their visual tastes, enabling personalized generation across various scenarios, including stickers, movie posters, fashion designs, and news posters. In specific, PMG converts historical interacted images into textual descriptions, allowing LLMs to distill user preferences effectively; I-AM-G introduces an interest rewrite strategy to address preference ambiguity and leverages retrieval-augmented generation (RAG) to enrich semantic information; and Pigeon leverages MLLMs with specialized modules to manage noisy historical data and multimodal instructions, ensuring precise user modeling. Similarly, SGDM (Xu et al., 2024a) enhances personalization by introducing a style extraction module that captures user-specific style preferences to guide image generation. Personalized PR (Chen et al., 2024g) proposes a personalized prompt-rewriting method, leveraging historical user query-image pairs to enhance text-to-image (T2I) personalization.

Beyond inferring user preferences from interaction history, some studies (Von Rütte et al., 2023; Liu et al., 2024d; Nabati et al., 2024) have explored interactive, personalized image generation by incorporating real-time user feedback through multi-turn interactions, thereby progressively refining the generated outputs.

Fashion Design Generation This task involves personalized fashion design with inferring personal style preferences from user behaviors. Yu et al. (2019) employs GANs to learn user preference vectors from interaction history, generating fashion images compatible with provided images. DiFashion (Xu et al., 2024b) adopts diffusion models to extract user preferences from interaction history for personalized outfit generation and recommendation.

E-commerce Product Image Generation This task aims to create customized, eye-catching visuals for e-commerce products to attract target consumers. Based on user behaviors, Ad-Booster (Shilova et al., 2023) utilizes the Stable Diffusion outpainting model, conditioning on individual user interest to generate appealing product images. Vashishtha et al. (2024) leverages LLMs to generate text prompts for diffusion models to craft engaging banners. Czapp et al. (2024) employs a contextual bandit algorithm to select prompts from a pool, generating personalized product backgrounds.

A.1.2 User Profiles

Some studies utilize users’ demographic attributes to infer preferences or categorize them into groups for personalized image generation.

Fashion Design Generation Base on user attributes (*e.g.*, age, gender, interests in characters), LVA-COG (Forouzandehmehr et al., 2023) utilizes LLMs to extract user preferences to guide fashion design generation for recommendation.

E-commerce Product Image Generation By categorizing users into distinct groups based on their attributes, CG4CTR (Yang et al., 2024a) proposes a self-cyclic generation pipeline to produce tailored product images for each user group.

A.1.3 Personalized Subjects

This is a primary focus of the computer vision community, which aims to capture the subject representation from a limited set of subject images and follow users’ instructions for subject-driven text-to-image (T2I) generation. For instance, given a few images of a user’s pet, the model generates new images featuring the pet in different contexts or environments while preserving its identity.

Subject-driven T2I Generation Existing research in this area can be broadly categorized into two branches:

- Optimization-based methods introduce a learnable unique identifier in the embedding space to encapsulate the semantics and visual details of each subject. Specifically, Textual Inversion (Gal et al., 2023) optimizes a pseudo-word identifier, denoted as S^* , to encode a personalized representation that guides T2I generation in diffusion models. DreamBooth (Ruiz et al., 2023) combines a unique identifier with a subject class name (*e.g.*, “A S^* cat”) to leverage the class-specific prior knowledge embedded in the model, leveraging the model’s prior knowledge of the class. Building upon these pioneering works, recent efforts have aimed to improve efficiency, subject fidelity, and instruction alignment, addressing both single-subject generation (Voynov et al., 2023; Alaluf et al., 2023; Tewel et al., 2023; Pang et al., 2024; Cai et al., 2024; Hong et al., 2025) and multi-subject generation (Avrahami et al., 2023; Kumari et al., 2023; Gu et al., 2024; Yao et al., 2024; Zhang et al., 2024f).
- Encoder-based methods utilize a pre-trained image encoder to extract subject-specific features,

which are then incorporated into text prompts or directly injected into the generator through dedicated cross-attention mechanisms or adapters. For instance, ELITE (Wei et al., 2023), IP-Adapter (Ye et al., 2023), Subject-Diffusion (Ma et al., 2024a) and SSR-Encoder (Zhang et al., 2024i) employ CLIP (Radford et al., 2021) as the feature extractor, each adopting unique encoding and injection strategies to seamlessly incorporate subject features into the image generation process. MoMA (Song et al., 2025) takes advantage of a pre-trained MLLM, LLaVA (Liu et al., 2023b), to extract subject features and refine them based on the target prompt, producing contextualized image features that are injected into the generator via cross-attention layers. In addition, encoder-based methods have been extended to support multi-subject generation (Patel et al., 2024b; Li et al., 2024d; Zhu et al., 2024; Wang et al., 2024h).

Furthermore, other research has explored diverse techniques for personalized subject-driven generation. These include prompt engineering (He et al., 2024c), which formulates structured prompts to guide generation; instruction tuning (Zhou et al., 2024; Hu et al., 2024a), which fine-tunes models on personalized instructions for improved alignment; and reinforcement learning (Chen et al., 2023d; Chae et al., 2023; Huang et al., 2024b; Wei et al., 2025a), which optimizes models through user feedback to refine subject representation.

Moreover, beyond capturing user preferences for concrete subjects like objects, some studies have focused on more abstract concepts, such as specified relations or styles, to guide personalized T2I generation (Huang et al., 2024e; Sohn et al., 2023; Wang et al., 2023a; Liu et al., 2024c; Park et al., 2024).

A.1.4 Personal Face/Body

Personal face and body images have become popular for personalized image generation, due to their high relevance to individual identity. By leveraging these images, generative models can create highly tailored and realistic images that reflect users’ distinct identities (IDs) while adhering to user-specific requirements. Existing studies primarily focus on face generation and virtual try-on applications.

Face Generation Generative models utilize personal face images to create high-fidelity portraits or avatars that preserve individual face IDs while

adhering to users’ multimodal instructions, such as modifying expressions, actions, backgrounds, and styles. Early GAN-based work mainly encodes face images into the latent space of StyleGAN (Karras et al., 2019) for face manipulation (Xia et al., 2021; Patashnik et al., 2021; Lyu et al., 2023; Baykal et al., 2023). To achieve more precise identity preservation and more flexible control, DM-based methods usually integrate a separate image encoder to convert face images into ID representations, which are then combined with user instructions to guide the face generation process. For instance, FastComposer (Xiao et al., 2024b), Face0 (Valevski et al., 2023), PhotoMaker (Li et al., 2024k), Infinite-ID (Wu et al., 2025a), MasterWeaver (Wei et al., 2025b), and AddMe (Yue et al., 2025) utilize a pre-trained CLIP image encoder or face recognition model to extract ID embeddings for identity preservation. In contrast, SeFi-IDE (Li et al., 2024h) directly optimizes one ID representation as multiple per-stage tokens to enhance semantic control. Except for ID representations, some studies have explored additional features or conditions to enhance style control (Yan et al., 2023), spatial control (Wang et al., 2024f; He et al., 2024d,a; Jiang et al., 2025), especially specific human-object interaction (Guo et al., 2024b; Hu et al., 2025), and scene affordance (Kulal et al., 2023; Parihar et al., 2024).

However, the rapid development of personalized face generation techniques has raised concerns about potential misuse and privacy risks. Recent research has investigated unlearning-related methods (Wu et al., 2025b; Hu et al., 2024b), adversarial attack-based methods (Van Le et al., 2023; Xiao et al., 2023; Wan et al., 2024; Onikubo and Matsui, 2024; Liu et al., 2024e) and watermark-based methods (Liu et al., 2024b) to protect user privacy.

Virtual Try-on This task aims to synthesize a photorealistic image of a dress person by combining their body and face images with specified garments. Early GAN-based works (Wang et al., 2018a; Dong et al., 2019a; Men et al., 2020; Choi et al., 2021; Lee et al., 2022) mainly follow a two-stage process. Initially, a dedicated warping module is employed to align garment images with the person’s body shape. Subsequently, the reshaped garment is seamlessly blended with the person’s image to generate the final try-on result. With the great success of DMs in various tasks, recent research has deployed their applications in virtual

try-ons (Zhang et al., 2024e; Ning et al., 2024; Kim et al., 2024a; Wan et al., 2025). For example, LaDITON (Morelli et al., 2023) and DCI-VTON (Gou et al., 2023) explicitly warp clothes to match the person’s body and then utilize DMs for blending. In contrast, TryOnDiffusion (Zhu et al., 2023a) proposes a Parallel-UNet architecture, which performs implicit warping and blending in a unified process. Similarly, several subsequent studies (Choi et al., 2025; Xu et al., 2024d; Sun et al., 2024; Shen et al., 2024a) utilize parallel UNets for garment feature extraction and enhance blending via self-attention and cross-attention mechanisms.

A.2 Personalized Video Generation

Personalized video generation aims to produce tailored video content that reflects individual preferences, traits, and specific needs.

A.2.1 Personalized Subjects

In some cases, users may provide one or several images of a personalized subject, such as an object or concept, along with a specified text prompt, requiring generative models to perform subject-driven text-to-video (T2V) generation. This task is conceptually similar to subject-driven T2I generation.

Subject-driven T2V Generation Given the great success of various personalized models in subject-driven T2I generation, methods such as AnimateDiff (Guo et al., 2024a), PIA (Zhang et al., 2024h), and Still-Moving (Chefer et al., 2024) adapt these models for T2V generation by incorporating motion and temporal dynamics through additional modules. MagDiff (Zhao et al., 2025) enhances subject-driven video generation with three types of alignments. VideoBooth (Jiang et al., 2024) proposes a coarse-to-fine manner to encode subject images into the generator. CustomCrafter (Wu et al., 2024d) integrates a plug-and-play module with a dynamic weighted video sampling strategy to maintain motion generation and conceptual combination abilities during subject learning. Other studies introduce additional conditions, such as motion control (Wu et al., 2024a; Wei et al., 2024b,c) and depth control (He et al., 2023), enabling more flexible subject customization. Besides, some studies have explored multi-subject T2V generation (Chen et al., 2023a, 2024b; Wang et al., 2024l). Beyond these, methods such as StyleCrafter (Liu et al., 2024a) and StyleMaster (Ye et al., 2024) integrate specified style images as subjects,

allowing for stylized T2V generation that tailors the video aesthetic to the desired style.

A.2.2 Personal Face/Body

Similarly, users may provide one or more personal face and body images, enabling generative models to synthesize personalized videos that preserve their identities while following multimodal instructions. These tasks include ID-preserving T2V generation, talking head generation, pose-guided video generation, and video virtual try-on.

ID-preserving T2V Generation This task focuses on creating personalized videos that align with personal face IDs and specified text prompts. For example, Magic-Me (Ma et al., 2024c) builds upon Textual Inversion (Gal et al., 2023) to learn ID-specific representations to guide T2V generation, requiring separate training for each ID. ID-Animator (He et al., 2024b) employs a face adapter to encode ID-related information and incorporates it into the generator via cross-attention. ConsisID (Yuan et al., 2024) decomposes facial information into frequency-aware features, which are integrated into Diffusion Transformers (DiT) for video generation. PersonalVideo (Li et al., 2024c) applies direct supervision on T2V-generated videos, aligning model tuning with the inference process.

Talking Head Generation This task aims to synthesize lip-synchronized talking videos, typically driven by personal face images and corresponding audio clips. Recent research has explored diverse approaches, such as Neural Radiance Fields (NeRFs) (Yao et al., 2022; Li et al., 2023b) and different backbone networks, including GANs (Yi et al., 2020; Ki and Min, 2023; Guan et al., 2023) and DMs (Zhang et al., 2023b; Zhua et al., 2023; Shen et al., 2023; Liu et al., 2024f; Wei et al., 2024a; Tian et al., 2025; Tan et al., 2025; Wang et al., 2024c; Zheng et al., 2024b). Beyond audio-driven generation, some studies have also investigated video-driven (Zhang et al., 2023a) and text-driven (Choi et al., 2024) methods for talking head generation.

Pose-guided Video Generation Recent studies have explored adapting personal face and body images to match specific pose sequences through various condition mechanisms for video generation (Wang et al., 2023b; Chang et al., 2023; Karras et al., 2023; Xu et al., 2024f; Hu, 2024; Zhong et al., 2025). Beyond single-person scenarios, Magic-

Fight (Huang et al., 2024a) tackles the complexities of two-person martial arts combat video generation, addressing challenges such as identity confusion and action mismatches.

Video Virtual Try-on This task seeks to seamlessly transfer a specified garment onto a person in a source video while preserving their motion and identity. Early GAN-based work (Dong et al., 2019b; Zhong et al., 2021; Jiang et al., 2022a) primarily follows a two-stage workflow, warping the specified garment and blending it with the target person by a GAN generator. Recent studies utilize DMs for video try-ons, incorporating specialized modules for garment and pose encoding, along with dedicated condition mechanisms, such as Tunnel Try-on (Xu et al., 2024e), ACF (Yang et al., 2024f), GPD-VVTO (Wang et al., 2024k), VITON-DiT (Zheng et al., 2024a), ViViD (Fang et al., 2024b), WildVidFit (He et al., 2025), and SwiftTry (Nguyen et al., 2024a).

A.3 Personalized 3D Generation

Personalized 3D generation involves transforming users’ personalized visual or textual contexts (e.g., body shapes, facial features, images, and prompts) into 3D assets.

A.3.1 Personalized Subjects

The most common paradigm for personalized 3D generation involves users providing some image-based personalized subjects, and then generating the corresponding 3D assets.

Image-to-3D Generation Personalized image-to-3D generation focuses on creating 3D assets that accurately capture the geometry and appearance of given personalized subjects. 3DAvatarGAN (Abdal et al., 2023) introduces a cross-domain adaptation framework that aligns features from 2D-GANs with those of 3D-GANs. PuzzleAvatar (Xiu et al., 2024) utilizes an enhanced Score Distillation Sampling (SDS) technique to optimize the geometry and texture of 3D portraits. TextureDreamer (Yeh et al., 2024) integrates geometric information using ControlNet, proposing a Geometry-Aware Personalized Score Distillation (PGSD) approach.

Some methods further ensure alignment with textual prompts during the 3D generation process. MVDream (Shi et al., 2023b) employs a multi-view diffusion model to generate consistent multi-view images based on text prompts. Dream-Booth3D (Raj et al., 2023) combines DreamFusion

and DreamBooth models within a three-stage optimization framework to enhance detail preservation and consistency through multi-view pseudo-data generation. Wonder3D (Long et al., 2024) introduces a cross-domain diffusion model leveraging multi-view cross-attention to produce detailed normal and color maps. DreamFont3D (Li et al., 2024e) utilizes NeRF as the 3D representation, optimizing font geometry and texture with multi-view mask constraints and progressive weight adjustments. Make-your-3D (Liu et al., 2025) implements a joint optimization framework that combines identity-aware and subject-prior optimizations, aligning a 2D personalization model with a multi-view diffusion model for accurate 3D generation.

A.3.2 Personal Face/Body

In some cases, users may provide personal face and body inputs in the form of images or monocular videos, aiming to generate identity-preserving 3D assets.

3D Face Generation For 3D face generation, (Zhang et al., 2021a) introduces PoseGAN, a module designed to generate dynamic head poses. (Gao et al., 2022) proposes a linear blend model based on multi-level voxel fields, representing expressions as neural radiance field bases. My3DGen (Qi et al., 2023a) adapts the pre-trained EG3D model using Low-Rank Adaptation (LoRA) for parameter-efficient training on a limited set of personalized images. DiffusionTalker (Chen et al., 2023c) employs contrastive learning to map speech features to personalized speaker identity. (Wang et al., 2024g) integrates a face fusion module into a fine-tuned text-to-image diffusion model for identity-driven customization. (Ko et al., 2024) utilizes VIVE3D to fine-tune the EG3D generator by inverting key frames from monocular videos. (Song et al., 2024a) applies Cross-Modal Aggregation to blend style and facial motion features, ensuring alignment between generated facial expressions, speech, and styles. DiffSpeaker (Ma et al., 2024d) proposes a diffusion model-based Transformer architecture to enhance speech-to-facial-animation mapping.

3D Human Pose generation For 3D human pose generation, (Huang et al., 2021) combines source image shape information with 2D key points to generate a personalized UV map. PGG (Hu et al., 2023) introduces a geometry-aware graph constructed from intermediate human mesh predic-

tions, enabling personalized and dynamic pose generation. 3DHM (Li et al., 2024a) and DreamWaltz (Huang et al., 2024c) enable animating people from a single image or textual prompts.

3D Virtual Try-on 3D Virtual Try-on enables the creation of high-quality, customized 3D models from minimal inputs, such as user images, clothing images, and textual prompts. (Chu et al., 2017) addresses the precision requirements of personalized facial modeling for applications like eyeglass frame design, utilizing parametric modeling techniques for 3D face creation. Pergamo (Casado-Elvira et al., 2022) addresses the challenge of reconstructing 3D clothing from 2D images, employing semantic segmentation, normal prediction, and a parameterized clothing model to optimize coarse geometry and fine details through differentiable rendering. DreamVTON (Xie et al., 2024) incorporates Multi-Concept LoRA and Normal Style LoRA into Stable Diffusion, enabling the generation of pose-consistent, detail-rich clothing models.

B Evaluation Metrics

B.1 Personalized Text Generation

Evaluating personalized text generation is challenging because only the target user can accurately determine whether the generated content aligns with their preferences and needs. One approach to evaluating personalized text generation is through human judgment, where individuals assess the quality and relevance of the generated content. Automatic evaluation of personalized text generation can be conducted using reference-based and reference-free approaches. In reference-based evaluation, it is assumed that a reference output is available for comparison. This comparison can be performed using term-matching metrics such as accuracy, ROUGE (Lin, 2004), BLEU (Papineni et al., 2002), and METEOR (Banerjee and Lavie, 2005) or semantic-matching methods using models like BERTScore (Zhang et al., 2020b), GEMBA (Kocmi and Federmann, 2023), or G-Eval (Liu et al., 2023c). Recently, ExPerT (Salemi et al., 2025a) was introduced for evaluating personalized text generation in reference-based settings. It segments both the generated text and the reference output into atomic facts and scores them based on content and writing style similarity. In reference-free evaluation, an LLM can assess the generated content directly, assigning scores based on various aspects, including how well the content aligns with

the user profile or preferences (Wang et al., 2023g, 2024a), offering a more dynamic and personalized evaluation framework.

B.2 Personalized Audio Generation

To quantify personalization and stylistic alignment in tasks like music transfer and text-to-audio generation, existing studies commonly use similarity metrics such as CLAP scores, Pattern Similarity (PS), and Embedding Distance are commonly applied to assess how well the generated content aligns with target musical styles or speaker characteristics (Sheng et al., 2023; Plitsis et al., 2024).

Audio quality assessment often involves perceptual metrics like Short-Time Objective Intelligibility (STOI) and Extended STOI (ESTOI) for speech intelligibility, Perceptual Evaluation of Speech Quality (PESQ) for clarity, and Fréchet Audio Distance (FAD) for measuring realism and diversity. For music generation tasks, metrics such as Precision, Recall, Density, and Coverage (P&R&D&C) are frequently employed to capture both creativity and output diversity (Wang et al., 2024j; Mo et al., 2023; Hu et al., 2022).

Subjective evaluations remain crucial, particularly for tasks where user experience and personalization are central. User studies are often conducted to measure factors such as perceived musicality, naturalness, emotional impact, and how closely the generated content reflects user preferences (Ma et al., 2022; Dai et al., 2022).

B.3 Personalized Cross-modal Generation

To quantify the degree of personalization in text generation tasks such as personal assistant and comment generation, existing studies commonly employ (1) *term-matching metrics*, such as ROUGE, BLEU, Meteor, CIDEr (Cohen et al., 2022; Alaluf et al., 2025; Nguyen et al., 2024b; Pi et al., 2024; An et al., 2024); (2) *semantic matching metrics*, such as CLIPScore (Cohen et al., 2022); (3) *recall, precision and F1-score* to validate whether the user-specific concept appears in the generated caption (Cohen et al., 2022; Alaluf et al., 2025; Nguyen et al., 2024b; Pi et al., 2024; An et al., 2024); (4) *human evaluation* to determine alignment with ground truths in terms of emotion, style, and relevance.

To measure whether agents align with diverse human preferences, studies employ success rate, success weighted by path length (SPL), distance to goal, and episode length (Poddar et al., 2024;

Hwang et al., 2024). Human evaluation is also an effective validation method for determining whether an agent has successfully completed a personalized task (Hwang et al., 2024; Cui et al., 2024b).

B.4 Personalized Image Generation

To assess the alignment between generated images and personalized contexts, as well as adherence to users’ multimodal instructions, most studies rely on similarity metrics like Learned Perceptual Image Patch Similarity (LPIPS) (Zhang et al., 2018) and Structural Similarity Index Measure (SSIM) (Wang et al., 2004). Additionally, pre-trained models such as CLIP (Radford et al., 2021), DINO (Oquab et al., 2024), and various face recognition models (Schroff et al., 2015; Deng et al., 2019) are often used to extract image features for computing cosine similarity, enabling a more contextual evaluation of personalization and instruction alignment. Moreover, Stellar (Achlioptas et al., 2023) introduces specialized metrics designed for subject-driven image generation and human generation.

To quantify the quality and coherence of generated images, conventional metrics such as Fréchet Inception Distance (FID) (Heusel et al., 2017) and Kernel Inception Distance (KID) (Bińkowski et al., 2018) are also commonly employed. Beyond quantitative evaluation, most studies present case studies and conduct human evaluations to assess the personalization and instruction alignment of the generated images. In e-commerce scenarios, product image generation often incorporates online tests to evaluate model performance in real-world scenarios.

B.5 Personalized Video Generation

To assess personalization and instruction alignment, similar to personalized image generation, existing studies commonly rely on similarity metrics such as LPIPS, SSIM, and Peak Signal-to-Noise Ratio (PSNR) (Hore and Ziou, 2010). Additionally, pre-trained image encoders like CLIP and DINO are frequently used to extract frame-level features and compute cosine similarity for quantitative evaluation. There are some task-specific metrics, such as SyncNet score (Casale et al., 2018), which evaluates audio-visual synchronization quality for audio-driven talking head generation, and face similarity metrics for ID-preserving human generation, which are based on backbone face recognition models

like ArcFace (Deng et al., 2019) and Curricular-Face (Huang et al., 2020).

For overall video quality assessment, standard metrics include frame-level evaluations like FID and KID, as well as video-level metrics such as VFID (Wang et al., 2018c), FID-VID (Balaji et al., 2019), FVD (Unterthiner et al., 2018), and KVD (Unterthiner et al., 2018). Additionally, some studies leverage CLIP-based cosine similarity between consecutive video frames to assess temporal consistency.

Beyond quantitative metrics, many studies complement their evaluations with qualitative case studies and human assessments to better capture personalization, instruction alignment, and overall video quality.

B.6 Personalized 3D Generation

To quantify personalization and instruction alignment in 3D generation, existing studies commonly use similarity metrics such as LPIPS, SSIM, PSNR, and CLIP scores, similar to those in image and video generation. Additionally, some task-specific scores can be evaluated through pre-trained models, such as facial attribute classifiers (Abdal et al., 2023; Qi et al., 2023a).

For 3D geometric quality, commonly used metrics include Chamfer Distance (CD) (Gao et al., 2022; Xiu et al., 2024) and Point-to-Surface Distance (P2S) (Xiu et al., 2024) for shape fidelity, as well as Normal Consistency and Volume IoU for surface detail and volumetric overlap (Xie et al., 2024). Task-specific metrics such as Mean Per Joint Position Error (MPJPE) and Mean Per Vertex Error (MPVE) are frequently used in the human body and pose estimation tasks (Hu et al., 2023).

Beyond objective metrics, qualitative assessments like user studies are often conducted to evaluate subjective aspects of 3D generation (Zhang et al., 2021a; Qi et al., 2023a; Xie et al., 2024; Huang et al., 2021), including realism, texture photorealism, and shape-texture consistency.

C Applications

C.1 Towards Content Creation Process

Generative models have transformed the realm of content creation, pushing the boundaries of productivity and creativity. By incorporating personalization techniques, PGen can further empower content creators across various domains, allowing

them to achieve higher efficiency while preserving their distinctive creative style and identity.

For social media influencers, such as bloggers, vloggers, and podcasters, PGen can analyze their past content to offer tailored suggestions for compelling headlines (Fang et al., 2024a) and introductions, or even generate new content that aligns seamlessly with their established brand. This not only streamlines the creative process but also maintains their unique style, fostering deeper connections with their audiences.

For professional creators, such as journalists, designers, illustrators, and music composers, personal style serves as their creative fingerprint, crucial for building reputation and recognition within competitive creative industries. By analyzing creators' previous work, PGen can identify and adapt to their unique stylistic traits to provide tailored ideas, drafts, or modifications, striking a perfect balance between personal style and external demands.

Ordinary individuals can also benefit from PGen for routine tasks, such as personalized email drafting, resume creation, travel planning, workout scheduling, and portrait generation.

C.2 Towards Content Delivery Process

In the era of information overload, personalized content delivery is becoming increasingly essential for helping individuals navigate through vast amounts of multimodal content on the internet. By integrating PGen into the content delivery process, generic content can be adapted into diverse, personalized formats to engage different audiences, catering to their unique tastes and content needs. Below are representative application scenarios of PGen for personalized content delivery:

Marketing and Advertising. PGen can assist organizations and brands in creating targeted marketing strategies and dynamic advertisements that resonate deeply with specific audiences, ultimately driving higher click-through and conversion rates.

Retail and E-commerce. Through personalized product descriptions and images, customized manuals, and virtual try-ons, PGen empowers retailers to attract consumers and enhance engagement, delivering unique and tailored shopping experiences.

Entertainment and Media. On digital content platforms such as Flipboard, Twitter, Netflix, and YouTube, personalized content plays a crucial role in attracting and retaining users. Examples include personalized news, posts, movie posters, video thumbnails, and other tailored media assets that

3591 can enhance user loyalty to the platform.

3592 **Education and E-learning.** Generative models
3593 have shown significant promise in education, exem-
3594 plified by platforms like Google Learn About ¹.
3595 PGen can further enhance personalized educa-
3596 tional experiences by offering customized learning
3597 roadmaps and materials, dynamically adapting to
3598 individual learning styles, goals, and progress.

3599 **Gaming.** Integrating PGen into the gaming in-
3600 dustries enables the creation of dynamic storylines,
3601 customized tasks, scalable difficulty levels, and
3602 interactive characters that adapt to players' prefer-
3603 ences and behaviors, fostering more immersive and
3604 engaging gaming experiences.

3605 **Personalized AI Assistant.** PGen can be incor-
3606 porated into AI assistants to provide specialized
3607 support, such as legal assistance, medical advice,
3608 and financial guidance, ensuring precision and user-
3609 specific customization.

¹<https://learning.google.com/experiments/learn-about>.

Table 1: Overview of personalized generation.

Modality	Personalized Contexts	Tasks	Representative Works
Text (Section 3.1)	User behaviors	Recommendation	LLM-Rec (Lyu et al., 2024), DEALRec (Lin et al., 2024a), Bi-gRec (Bao et al., 2023), DreamRec (Yang et al., 2024e)
		Information seeking	P-RLHF (Li et al., 2024g), ComPO (Kumar et al., 2024b)
	User documents	Writing Assistant	REST-PG (Salemi et al., 2025b), RSPG (Salemi et al., 2024a), Hydra (Zhuang et al., 2024), PEARL (Mysore et al., 2024), Panza (Nicolicioiu et al., 2024)
	User profiles	Dialogue System	PAED (Zhu et al., 2023b), BoB (Song et al., 2021), UniMS-RAG (Wang et al., 2024d), ORIG (Chen et al., 2023b)
		User Simulation	Drama Machine (Magee et al., 2024), Character-LLM (Shao et al., 2023), RoleLLM (Wang et al., 2024e)
Image (Section A.1)	User behaviors	General-purpose generation	PMG (Shen et al., 2024b), Pigeon (Xu et al., 2024c), PASTA (Nabati et al., 2024)
		Fashion design	DiFashion (Xu et al., 2024b), Yu et al. (2019)
		E-commerce product image	AdBooster (Shilova et al., 2023), Vashishtha et al. (2024), Czapp et al. (2024)
	User profiles	Fashion design	LVA-COG (Forouzandehmehr et al., 2023)
		E-commerce product image	CG4CTR (Yang et al., 2024a)
	Personalized subjects	Subject-driven T2I generation	Textual Inversion (Gal et al., 2023), DreamBooth (Ruiz et al., 2023), Custom Diffusion (Kumari et al., 2023)
	Personal face/body	Face generation	PhotoMaker (Li et al., 2024k), InstantBooth (Shi et al., 2024), InstantID (Wang et al., 2024f)
		Virtual try-on	IDM-VTON (Choi et al., 2025), OOTDiffusion (Xu et al., 2024d), OutfitAnyone (Sun et al., 2024)
Video (Section A.2)	Personalized subjects	Subject-driven T2V generation	AnimateDiff (Guo et al., 2024a), AnimateLCM (Wang et al., 2024b), PIA (Zhang et al., 2024h)
	Personal face/body	ID-preserving T2V generation	Magic-Me (Ma et al., 2024c), ID-Animator (He et al., 2024b), ConsisID (Yuan et al., 2024)
		Talking head generation	DreamTalk (Ma et al., 2023), EMO (Tian et al., 2025), MEMO (Zheng et al., 2024b)
		Pose-guided video generation	Disco (Wang et al., 2023b), AnimateAnyone (Hu, 2024), MagicAnimate (Xu et al., 2024f)
		Video virtual try-on	ViViD (Fang et al., 2024b), VITON-DiT (Zheng et al., 2024a), WildVidFit (He et al., 2025)
3D (Section A.3)	Personalized subjects	Image-to-3D generation	MVDream (Shi et al., 2023b), DreamBooth3D (Raj et al., 2023), Wonder3D (Long et al., 2024)
	Personal face/body	3D face generation	PoseGAN (Zhang et al., 2021a), My3DGen (Qi et al., 2023a), DiffSpeaker (Ma et al., 2024d)
		3D human pose generation	FewShotMotionTransfer (Huang et al., 2021), PGG (Hu et al., 2023), 3DHM (Li et al., 2024a), DreamWaltz (Huang et al., 2024c)
		3D virtual try-on	Pergamo (Casado-Elvira et al., 2022), DreamVTON (Xie et al., 2024)
Audio (Section 3.2)	Personal face	Face-to-speech generation	ViocMe (van Rijn et al., 2022), FR-PSS (Wang et al., 2022), Lip2Speech (Sheng et al., 2023)
	User behaviors	Music generation	UMP (Ma et al., 2022), UP-Transformer (Hu et al., 2022), UIGAN (Wang et al., 2024j)
	Personalized subjects	Text-to-audio generation	DiffAVA (Mo et al., 2023), TAS (Li et al., 2024j)
Cross-Modal (Section 3.3)	User behaviors	Robotics	VPL (Poddar et al., 2024), Promptable Behaviors (Hwang et al., 2024)
	User documents	Caption/Comment generation	PV-LLM (Lin et al., 2024b), PVCG (Wu et al., 2024e), ME-TER (Geng et al., 2022)
	Personalized subjects	Cross-modal dialogue systems	MyVLM (Alaluf et al., 2025), Yo’LLaVA (Nguyen et al., 2024b), MC-LLaVA (An et al., 2024)

Table 2: Datasets for personalized generation.

Modality	Personalized Contexts	Tasks	Datasets
Text (Section 3.1)	User behaviors	Recommendation	Amazon (Hou et al., 2024; Ni et al., 2019), MovieLens (Harper and Konstan, 2015), MIND (Wu et al., 2020a), Goodreads (Wan and McAuley, 2018; Wan et al., 2019)
		Information seeking	SE-PQA (Kasela et al., 2024), PWSC (Eugene et al., 2013), AOL4PS (Guo et al., 2021)
	User documents	Writing Assistant	LaMP (Salemi et al., 2024b), LongLaMP (Kumar et al., 2024a), PLAB (Alhafni et al., 2024)
	User profiles	Dialogue System	LiveChat (Gao et al., 2023), FoCus (Jang et al., 2021), Pchatbot (Qian et al., 2021)
		User Simulation	OpinionsQA (Santurkar et al., 2023), 3 RoleBench (Wang et al., 2024e)
Image (Section A.1)	User behaviors	General-purpose generation	Pinterest (Geng et al., 2015), MovieLens (Harper and Konstan, 2015), MIND (Wu et al., 2020b), POG (Chen et al., 2019), PASTA (Nabati et al., 2024), FABRIC (Von Rütte et al., 2023), DialPrompt (Liu et al., 2024d), PIP (Chen et al., 2024g)
		Fashion design	POG (Chen et al., 2019), Polyvore-U (Lu et al., 2019)
	User profiles	E-commerce product image	-
		Fashion design	-
	Personalized subjects	E-commerce product image	-
		Subject-driven T2I generation	Dreambench (Ruiz et al., 2023), Dreambench++ (Peng et al., 2024), CustomConcept101 (Kumari et al., 2023), ConceptBed (Patel et al., 2024a), Textual Inversion (Gal et al., 2023), ViCo (Hao et al., 2023), DreamMatcher (Nam et al., 2024), Break-A-Scene (Avrahami et al., 2023), Mix-of-Show (Gu et al., 2024), Concept Conductor (Yao et al., 2024), LoRA-Composer (Yang et al., 2024d), StyleDrop (Sohn et al., 2023)
	Personal face/body	Face generation	CelebA-HQ (Karras et al., 2018), FFHQ (Karras et al., 2021), SFHQ (Beniganev, 2022), LV-MHP-v2 (Zhao et al., 2018), Stellar (Achlioptas et al., 2023), AddMe-1.6M (Yue et al., 2025), FFHQ-FastComposer (Xiao et al., 2024a), LAION-Face (Zheng et al., 2022), PPR10K (Liang et al., 2021), LCM-Lookahead (Gal et al., 2024), CelebRef-HQ (Li et al., 2022), CelebV-T (Yu et al., 2023), FaceForensics++ (Rossler et al., 2019), VG-GFace2 (Cao et al., 2018)
		Virtual try-on	VITON (Han et al., 2018), VITON-HD (Choi et al., 2021), Dress-Code (Morelli et al., 2022), StreetTryOn (Cui et al., 2024a), DeepFashion (Ge et al., 2019), Deepfashion-Multimodal (Jiang et al., 2022b), MPV (Dong et al., 2019a), IGPair (Shen et al., 2024a), SHHQ (Fu et al., 2022)
Video (Section A.2)	Personalized subjects	Subject-driven T2V generation	WebVid-10M (Bain et al., 2021), UCF101 (Soomro, 2012), AnimateBench (Zhang et al., 2024h), VideoBooth (Jiang et al., 2024), StyleCrafter (Liu et al., 2024a), Datasets for subject-driven T2I generation...
	Personal face/body	ID-preserving T2V generation	ID-Animator (He et al., 2024b), ConsisID (Yuan et al., 2024)
		Talking head generation	LRW (Chung and Zisserman, 2017a), VoxCeleb (Nagrani et al., 2020), VoxCeleb2 (Chung et al., 2018), TCD-TIMIT (Harte and Gillen, 2015), LRS2 (Son Chung et al., 2017), HDTF (Zhang et al., 2021b), MEAD (Wang et al., 2020), GRID (Cooke et al., 2006), MultiTalk (Sung-Bin et al., 2024)
		Pose-guided video generation	FashionVideo (Zablotskaia et al., 2019), TikTok (Jafarian and Park, 2021), TED-talks (Siarohin et al., 2021), Everybody-dance-now (Chan et al., 2019)
		Video virtual try-on	VVT (Dong et al., 2019b), ViViD (Fang et al., 2024b), Fashion-Video (Zablotskaia et al., 2019), TikTok (Jafarian and Park, 2021), TikTokDress (Nguyen et al., 2024a)
3D (Section A.3)	Personalized subjects	Image-to-3D generation	Dreambench (Ruiz et al., 2023), Objaverse (Deitke et al., 2023)
	Personal face/body	3D face generation	Mystyle (Nitzan et al., 2022), BIWI (Fanelli et al., 2013), VO-CASET (Cudeiro et al., 2019)
		3D human pose generation	Human3.6M (Ionescu et al., 2013), 3DPW (Von Marcard et al., 2018), 3DOH50K (Zhang et al., 2020a)
		3D virtual try-on	BUFF (Zhang et al., 2017), DreamVTON (Xie et al., 2024)
Audio (Section 3.2)	Personal face	Face-to-speech generation	Voxceleb2 (Chung et al., 2018), LibriTTS (Zen et al., 2019), VG-GFace2 (Cao et al., 2018), GRID (Cooke et al., 2006), MultiTalk (Sung-Bin et al., 2024)
	User behaviors	Music generation	Echo (Bertin-Mahieux et al., 2011), MAESTRO (Hawthorne et al., 2019)
	Personalized subjects	Text-to-audio generation	TASBench (Li et al., 2024j), AudioCaps (Kim et al., 2019), AudioLDM (Liu et al., 2023a)
Cross-Modal (Section 3.3)	User behaviors	Robotics	D4RL (Fu et al., 2020), Ravens (Zeng et al., 2021), Habitat-Rearrange (Puig et al., 2023), RoboTHOR (Deitke et al., 2020)
	User documents	Caption/Comment generation	TripAdvisor (Geng et al., 2022), Yelp (Geng et al., 2022), PerVid-Com (Lin et al., 2024b)
	Personalized subjects	Cross-modal dialogue systems	Yo’LLaVA (Nguyen et al., 2024b), MyVLM (Alaluf et al., 2025)

Table 3: Evaluation metrics for personalized text and image generation.

Text (Section 3.1)	Metrics	1	2	3	4	5	6	Evaluation Dimensions	Representative Works
1. Recommendation 2. Information Seeking 3. Content Generation 4. Writing Assistant 5. Dialogue System 6. User Simulation	NDCG (Järvelin and Kekäläinen, 2002)	✓	✓					Overall	BIGRec (Bao et al., 2023), DEALRec (Lin et al., 2024a), AOL4PS (Guo et al., 2021)
	Hit Rate	✓	✓					Overall	BIGRec (Bao et al., 2023)
	Precision	✓	✓					Overall	LLM-Rec (Lyu et al., 2024), AOL4PS (Guo et al., 2021)
	Recall	✓	✓					Overall	LLM-Rec (Lyu et al., 2024), DEALRec (Lin et al., 2024a), AOL4PS (Guo et al., 2021)
	win-rate		✓					Overall	Personalized RLHF (Li et al., 2024g)
	ROUGE (Lin, 2004)			✓	✓	✓	✓	Overall	LaMP (Salemi et al., 2024b), RSPG (Salemi et al., 2024a), Hydra (Zhuang et al., 2024)
	BLEU (Papineni et al., 2002)			✓	✓	✓	✓	Overall	AuthorPred (Li et al., 2023a)
	BERTScore (Zhang et al., 2020b)			✓	✓	✓	✓	Overall	LongLaMP (Kumar et al., 2024a)
	GEMBA (Kocmi and Federmann, 2023)			✓	✓	✓	✓	Overall	REST-PG (Salemi et al., 2025b)
	G-Eval (Liu et al., 2023c)			✓	✓	✓	✓	Overall	REST-PG (Salemi et al., 2025b)
	ExPerT (Salemi et al., 2025a)			✓	✓			Personalization	ExPerT (Salemi et al., 2025a)
	AuPEL (Wang et al., 2023g)			✓	✓			Personalization	AuPEL (Wang et al., 2023g)
	PERSE (Wang et al., 2024a)			✓	✓			Personalization	PERSE (Wang et al., 2024a)
Image (Section A.1)	Metrics	1	2	3	4	5	6	Evaluation Dimensions	Representative Works
1. General-purpose generation 2. Fashion design 3. E-commerce product image 4. Subject-driven T2I generation 5. Face generation 6. Virtual try-on	CLIP-I (Radford et al., 2021)	✓	✓		✓	✓	✓	Personalization	Textual Inversion (Gal et al., 2023), Custom Diffuison (Kumari et al., 2023), DreamBooth (Ruiz et al., 2023)
	DINO-I (Caron et al., 2021; Quab et al., 2024)	✓	✓		✓	✓		Personalization	DreamBooth (Ruiz et al., 2023), BLIP-Diffusion (Li et al., 2024b), ELITE (Wei et al., 2023)
	LPIS (Zhang et al., 2018)	✓	✓		✓		✓	Personalization	DreamSteerer (Yu et al., 2024), DiFashion (Xu et al., 2024b), PMG (Shen et al., 2024b)
	PSNR (Hore and Ziou, 2010)					✓	✓	Personalization	GroupDiff (Jiang et al., 2025), MYCloth (Liu and Wang, 2024), SCW-VTON (Han et al., 2024)
	SSIM (Wang et al., 2004)	✓	✓		✓	✓	✓	Personalization	DreamSteerer (Yu et al., 2024), PMG (Shen et al., 2024b), OOTDiffusion (Xu et al., 2024d)
	MS-SSIM (Wang et al., 2003)	✓	✓		✓		✓	Personalization	DreamSteerer (Yu et al., 2024), Pigeon (Xu et al., 2024c), SieveNet (Jandial et al., 2020)
	DreamSim (Fu et al., 2023)				✓		✓	Personalization	IMPRINT (Song et al., 2024b), MaX4Zero (Orzech et al., 2024)
	Face similarity (Deng et al., 2019; Schroff et al., 2015; Kim et al., 2022; Wang et al., 2018b)					✓		Personalization	Infinite-ID (Wu et al., 2025a), PhotoMaker (Li et al., 2024k), ProFusion (Zhou et al., 2023)
	Face detection rate (Deng et al., 2019; Zhang et al., 2016)					✓		Personalization	SeFi-IDE (Li et al., 2024h), Celeb Basis (Yuan et al., 2023), W ₊ Adapter (Li et al., 2024f)
	CLIP-T (Radford et al., 2021)	✓	✓	✓	✓	✓		Instruction Alignment	Textual Inversion (Gal et al., 2023), Custom Diffuison (Kumari et al., 2023), DreamBooth (Ruiz et al., 2023)
	ImageReward (Xu et al., 2023a)				✓	✓	✓	Instruction Alignment	InstructBooth (Chae et al., 2023), DiffLoRA (Wu et al., 2024f), IMAGDressing-v1 (Shen et al., 2024a)
	PickScore (Kirstain et al., 2023)				✓			Instruction Alignment	InstructBooth (Chae et al., 2023), FABRIC (Von Rütte et al., 2023), Stellar (Achlioptas et al., 2023)
	HPSv1 (Wu et al., 2023b)				✓	✓		Instruction Alignment	Stellar (Achlioptas et al., 2023)
	HPSv2 (Wu et al., 2023b)				✓	✓		Instruction Alignment	Stellar (Achlioptas et al., 2023)
	R-precision (Xu et al., 2018)				✓			Instruction Alignment	COTI (Yang et al., 2023b)
	PAR score (Gani et al., 2024)			✓				Instruction Alignment	Vashishtha et al. (2024)
	FID (Heusel et al., 2017)	✓	✓		✓	✓	✓	Content Quality	COTI (Yang et al., 2023b), IMPRINT (Song et al., 2024b), DiFashion (Xu et al., 2024b)
	KID (Bińkowski et al., 2018)				✓	✓	✓	Content Quality	Custom Diffuison (Kumari et al., 2023), OOTDiffusion (Xu et al., 2024d), LaDI-VTON (Morelli et al., 2023)
	IS (Salimans et al., 2016)				✓		✓	Content Quality	PE-VTON (Zhang et al., 2024d), DF-VTON (Dong et al., 2024), Layout-and-Retouch (Kim et al., 2024b)
	LAION-Aesthetics (Christoph and Romain, 2022)				✓	✓		Content Quality	BLIP-Diffusion (Li et al., 2024b), UniPortrait (He et al., 2024a)
	TOPIQ (Chen et al., 2024a)				✓			Content Quality	DreamSteerer (Yu et al., 2024)
	MUSIQ (Ke et al., 2021)				✓		✓	Content Quality	DreamSteerer (Yu et al., 2024), PE-VTON (Zhang et al., 2024d)
	MANIQA (Yang et al., 2022)						✓	Content Quality	PE-VTON (Zhang et al., 2024d)
	LIQE (Zhang et al., 2023c)				✓			Content Quality	DreamSteerer (Yu et al., 2024)
	QS (Gu et al., 2020)				✓			Content Quality	AddMe (Yue et al., 2025)
	BRISQUE (Mittal et al., 2012a)			✓				Content Quality	Vashishtha et al. (2024)
	CTR			✓				Overall	CG4CTR (Yang et al., 2024a), Czapp et al. (2024)
	Stellar metrics (Achlioptas et al., 2023)				✓	✓		Overall	Stellar (Achlioptas et al., 2023)
	CAMI (Shen et al., 2024a)						✓	Overall	IMAGDressing-v1 (Shen et al., 2024a)

Table 4: Evaluation metrics for personalized generation across video, 3D, audio, and cross-modal domains.

Video (Section A.2)	Metrics	1	2	3	4	5	-	Evaluation Dimensions	Representative Works
1. Subject-driven T2V generation 2. ID-preserving T2V generation 3. Talking head generation 4. Pose-guided video generation 5. Video virtual try-on	CLIP-I (Radford et al., 2021)	✓	✓	✓				Personalization	PIA (Zhang et al., 2024h), PoseCrafter (Zhong et al., 2025), ID-Animator (He et al., 2024b)
	DINO-I (Caron et al., 2021; Oquab et al., 2024)	✓	✓					Personalization	DisenStudio (Chen et al., 2024b), DreamVideo (Wei et al., 2024b), Magic-Me (Ma et al., 2024c)
	SSIM (Wang et al., 2004)			✓	✓	✓		Personalization	AnimateAnyone (Hu, 2024), VITON-DiT (Zheng et al., 2024a), ViViD (Fang et al., 2024b)
	PSNR (Hore and Ziou, 2010)			✓	✓	✓		Personalization	AnimateAnyone (Hu, 2024), Yi et al. (2020), Zhua et al. (2023)
	LPIPS (Zhang et al., 2018)			✓	✓	✓		Personalization	AnimateAnyone (Hu, 2024), DiffTalk (Shen et al., 2023), DisCo (Wang et al., 2023b)
	VGG (Johnson et al., 2016)				✓			Personalization	DreamPose (Karras et al., 2023)
	L1 error				✓			Personalization	DisCo (Wang et al., 2023b), DreamPose (Karras et al., 2023), MagicAnimate (Xu et al., 2024f)
	AED				✓			Personalization	DisCo (Wang et al., 2023b), DreamPose (Karras et al., 2023)
	Face similarity (Deng et al., 2019; Huang et al., 2020; Kim et al., 2022)		✓	✓	✓			Personalization	ID-Animator (He et al., 2024b), MagicPose (Chang et al., 2023), ConsisID (Yuan et al., 2024)
	CLIP-T (Radford et al., 2021)	✓	✓		✓			Instruction Alignment	PIA (Zhang et al., 2024h), ConsisID (Yuan et al., 2024), PoseCrafter (Zhong et al., 2025)
	UMT score (Liu et al., 2022)	✓						Instruction Alignment	StyleMaster (Ye et al., 2024)
	AKD (Siarohin et al., 2021)				✓			Instruction Alignment	MagicAnimate (Xu et al., 2024f)
	MKR (Siarohin et al., 2021)				✓			Instruction Alignment	MagicAnimate (Xu et al., 2024f)
	MSE-P				✓			Instruction Alignment	PoseCrafter (Zhong et al., 2025)
	SyncNet score (Chung and Zisserman, 2017b)		✓					Instruction Alignment	DreamTalk (Ma et al., 2023), EMO (Tian et al., 2025), MEMO (Zheng et al., 2024b)
	LMD (Chen et al., 2018)		✓					Instruction Alignment	DFA-NeRF (Yao et al., 2022), DreamTalk (Ma et al., 2023), Yi et al. (2020)
	LSE-C (Prajwal et al., 2020)		✓					Instruction Alignment	StyleLipSync (Ki and Min, 2023), DiffTalker (Qi et al., 2023b), Choi et al. (2024)
	LSE-D (Prajwal et al., 2020)		✓					Instruction Alignment	StyleLipSync (Ki and Min, 2023), DiffTalker (Qi et al., 2023b), Choi et al. (2024)
	PD (Baldrati et al., 2023)					✓		Instruction Alignment	ACF (Yang et al., 2024f)
	FID (Heusel et al., 2017)	✓	✓	✓	✓	✓		Content Quality	EMO (Tian et al., 2025), ConsisID (Yuan et al., 2024), DisCo (Wang et al., 2023b)
	KID (Bińkowski et al., 2018)					✓		Content Quality	WildVidFit (He et al., 2025)
	ArtFID (Wright and Ommer, 2022)	✓						Content Quality	StyleMaster (Ye et al., 2024)
	VFID (Wang et al., 2018c)					✓		Content Quality	SwiftTry (Nguyen et al., 2024a), VITON-DiT (Zheng et al., 2024a), ViViD (Fang et al., 2024b)
	FVD (Unterthiner et al., 2018)	✓	✓	✓	✓			Content Quality	PersonalVideo (Li et al., 2024c), MotionBooth (Wu et al., 2024a), AnimateAnyone (Hu, 2024)
	FID-VID (Balaji et al., 2019)				✓			Content Quality	DisCo (Wang et al., 2023b), MagicAnimate (Xu et al., 2024f), MagicPose (Chang et al., 2023)
	KVD (Unterthiner et al., 2018)	✓						Content Quality	Animate-A-Story (He et al., 2023)
	E-FID (Tian et al., 2025)		✓					Content Quality	EMO (Tian et al., 2025), EmotiveTalk (Wang et al., 2024c)
	NIQE (Mittal et al., 2012b)				✓			Content Quality	MagicFight (Huang et al., 2024a)
	CPBD (Narvekar and Karam, 2011)		✓					Content Quality	DreamTalk (Ma et al., 2023)
3D (Section A.3) 1. Image-to-3D generation 2. 3D face generation 3. 3D human pose generation 4. 3D virtual try-on	Temporal consistency (Radford et al., 2021)	✓	✓					Content Quality	AnimateDiff (Guo et al., 2024a), Magic-Me (Ma et al., 2024c), DreamVideo (Wei et al., 2024b)
	Dynamic degree (Huang et al., 2024d)	✓	✓					Content Quality	StyleMaster (Ye et al., 2024), ID-Animator (He et al., 2024b), PersonalVideo (Li et al., 2024c)
	Video IS (Saito et al., 2020)	✓				✓		Content Quality	MagDiff (Zhao et al., 2025)
	Flow error (Shi et al., 2023a)	✓						Content Quality	MotionBooth (Wu et al., 2024a)
	Stitch score	✓						Content Quality	VideoDreamer (Chen et al., 2023a)
	Dover score (Wu et al., 2023a)		✓					Content Quality	ID-Animator (He et al., 2024b)
	Motion score (Li et al., 2018)		✓					Content Quality	ID-Animator (He et al., 2024b)
	Metrics	1	2	3	4	-	-	Evaluation Dimensions	Representative Works
	LPIPS (Zhang et al., 2018)	✓	✓	✓				Personalization	Wonder3D (Long et al., 2024), PuzzleAvatar (Xiu et al., 2024), My3DGen (Qi et al., 2023a)
	PSNR (Hore and Ziou, 2010)	✓	✓	✓				Personalization	Wonder3D (Long et al., 2024), PuzzleAvatar (Xiu et al., 2024), My3DGen (Qi et al., 2023a)
	SSIM (Wang et al., 2004)	✓	✓	✓				Personalization	Wonder3D (Long et al., 2024), PuzzleAvatar (Xiu et al., 2024), My3DGen (Qi et al., 2023a)
	Chamfer Distances (Butt and Maragos, 1998)	✓						Personalization	Wonder3D (Long et al., 2024), PuzzleAvatar (Xiu et al., 2024)
	CLIP (Radford et al., 2021)				✓			Personalization	DreamVTON (Xie et al., 2024)
	Volume IoU (Zhou et al., 2019)	✓						Personalization	Wonder3D (Long et al., 2024)
	Lip Vertex Error (LVE) (Ma et al., 2024d)		✓					Personalization	DiffSpeaker (Ma et al., 2024d)
	Facial Dynamics Deviation (FDD) (Ma et al., 2024d)		✓					Personalization	DiffSpeaker (Ma et al., 2024d), DiffusionTalker (Chen et al., 2023c)
	FReID (Huang et al., 2021)			✓				Personalization	FewShotMotionTransfer (Huang et al., 2021)
	CLIP-T (Radford et al., 2021)	✓						Instruction Alignment	MVDream (Shi et al., 2023b), DreamBooth3D (Raj et al., 2023), MakeYour3D (Liu et al., 2025)
	FID (Heusel et al., 2017)	✓			✓			Content Quality	MVDream (Shi et al., 2023b), 3DAvatarGAN (Abdal et al., 2023), TextureDreamer (Yeh et al., 2024), DreamVTON (Xie et al., 2024)
	IS (Salimans et al., 2016)	✓						Content Quality	MVDream (Shi et al., 2023b)
Audio (Section 3.2) 1. Face-to-speech generation 2. Music generation 3. Text-to-audio generation	Metrics	1	2	3	-	-	-	Evaluation Dimensions	Representative Works
	CLAP (Elizalde et al., 2023)		✓	✓				Personalization	DB&TI (Plitsis et al., 2024)
	Embedding Distance	✓	✓	✓				Personalization	UMP (Ma et al., 2022), FR-PSS (Wang et al., 2022)
	FAD (Kilgour et al., 2018)	✓	✓	✓				Personalization	UIGAN (Wang et al., 2024j), DiffAVA (Mo et al., 2023), DB&TI (Plitsis et al., 2024)
	IS (Salimans et al., 2016)			✓				Content Quality	DiffAVA (Mo et al., 2023)
Cross-modal (Section 3.3) 1. Robotics 2. Caption/Comment generation 3. Multimodal dialogue systems	STOI, ESTOI, PESQ (Sheng et al., 2023)	✓						Content Quality	Lip2Speech (Sheng et al., 2023)
	Metrics	1	2	3	-	-	-	Evaluation Dimensions	Representative Works
	BLEU, Meteor		✓	✓				Overall	PVCG (Wu et al., 2024e), METER (Geng et al., 2022), PV-LLM (Lin et al., 2024b)
	Recall, Precision, F1			✓				Overall	MyVLM (Alaluf et al., 2025), Yo'LLaVA (Nguyen et al., 2024b)
	success rate	✓						Overall	VPL (Poddar et al., 2024), Promptable Behaviors (Hwang et al., 2024)