
Benchmarking Collaborative, Neural, and Large Language Models for Movie Recommendation

Théodore Billotte (2025400572), Louis Rollet (2025400573), Vincent Montero-Fontaine (2025400574)

Department of Computer Science and Technology
Tsinghua University
Beijing, China
{bda25, rly25, qqy25}@mails.tsinghua.edu.cn

Abstract

In the era of information overload, personalized recommendation systems are critical for user retention. This project benchmarks three distinct architectures for movie recommendation using the MovieLens 20M dataset: a classical Matrix Factorization (MF) baseline, a Neural Collaborative Filtering (NCF) model, and a proposed Hybrid NCF extension that integrates content features (genres). Furthermore, we explore the emerging capabilities of Large Language Models (LLMs) as zero-shot and few-shot recommenders. Our results demonstrate that the Hybrid model achieves the best Root Mean Square Error (RMSE) of **0.777** by mitigating the cold-start problem, while the LLM agent with Few-Shot prompting on 1000 users achieves RMSE of **0.97**, offering superior interpretability. We provide a detailed ablation study of LLM prompting strategies, revealing that Few-Shot prompting significantly reduces error rates compared to Zero-Shot approaches, though hallucinations remain a critical bottleneck.

1 Introduction

1.1 Problem: The Information Filtering Overload

With the exponential growth of digital media libraries, users face an overwhelming number of choices. The core problem addressed in this work is **rating prediction**: given a user u and an item i , the system must accurately predict the preference score $\hat{r}_{ui}$ that user u would assign to item i . This score typically ranges from 1 to 5. A secondary but equally critical challenge is the **Cold Start** problem, where the system must generate recommendations for new users or items with little to no interaction history.

1.2 Motivation

Traditional methods like Matrix Factorization (MF) rely heavily on historical interactions (Collaborative Filtering). While effective, they fail to capture non-linear relationships between users and items and struggle when interaction data is sparse [3]. Neural Collaborative Filtering (NCF) addresses the linearity issue but still suffers in cold-start scenarios [2].

Our motivation is twofold:

1. To enhance NCF by integrating **content-based features** (specifically genre metadata), creating a Hybrid architecture that remains robust even when interaction data is scarce.
2. To investigate whether general-purpose **Large Language Models (LLMs)**, specifically the Gemini series, can bypass training altogether by reasoning over user profiles. We specifically

investigate if advanced prompting techniques (Few-Shot with Chain-of-Thought) can bridge the gap between pre-trained LLMs and supervised neural models.

2 Related Work

2.1 Latent Factor Models

Matrix Factorization (MF) has long been the standard for collaborative filtering. It projects users and items into a shared latent space of dimension d , such that their dot product approximates the interaction rating. Koren et al. (2009) demonstrated that MF captures global structure but fails to capture finer, non-linear dependencies [3].

2.2 Neural Recommendation

He et al. (2017) proposed Neural Collaborative Filtering (NCF), replacing the dot product with a multi-layer perceptron (MLP) to learn an arbitrary interaction function [2]. This allowed for capturing higher-order non-linearities. However, standard NCF is purely collaborative and ignores item metadata, making it vulnerable to the cold-start problem.

2.3 LLMs for Recommendation

Recent surveys by Wu et al. (2023) and Liu et al. (2023) have categorized LLM-based recommendation into two paradigms: Discriminative (ranking items) and Generative (generating item IDs) [4, 5]. While LLMs possess vast world knowledge, they often suffer from "hallucinations"; recommending non-existent items or citing incorrect facts [7].

3 Methodology

We define our dataset as a set of user-item interactions $Y \in \mathbb{R}^{M \times N}$ where M is the number of users and N is the number of movies. All models were implemented using PyTorch and trained on MovieLens 20M dataset containing over 20 million ratings.

3.1 Framework & Hyperparameters

Implementation Framework: All neural models were implemented in PyTorch, leveraging GPU acceleration for efficient training on large-scale data.

Training Configuration:

- **Optimizer:** Adam (Adaptive Moment Estimation) with learning rate $\alpha = 0.001$
- **Loss Function:** Mean Squared Error (MSE) for rating prediction
- **Batch Size:** 1024 samples (optimized for efficient GPU utilization)
- **Epochs:** 10 epochs for MF and NCF; 45 epochs for Hybrid NCF (deeper architecture requires extended training)
- **Learning Rate Scheduler:** ReduceLROnPlateau to stabilize convergence
- **Regularization:** L2 penalty with $\lambda = 0.001$ to prevent overfitting

Data Splitting & Reproducibility: The dataset was split into train/test sets with an 80/20 ratio. All trained model weights were saved as .pth checkpoint files in the checkpoints/ directory, ensuring full reproducibility of results. Random seeds were fixed across all experiments for consistency.

3.2 Baseline: Matrix Factorization (MF)

Matrix Factorization decomposes the user-item interaction matrix into two lower-dimensional matrices representing latent factors. Each user u and item i is associated with an embedding vector $\mathbf{p}_u, \mathbf{q}_i \in \mathbb{R}^d$ where d is the latent dimensionality.

Architecture Details:

- **Embedding Dimension:** $d = 50$ latent factors
- **User Embeddings:** Learnable matrix $\mathbf{P} \in \mathbb{R}^{M \times d}$
- **Item Embeddings:** Learnable matrix $\mathbf{Q} \in \mathbb{R}^{N \times d}$
- **Bias Terms:** User bias b_u , item bias b_i , and global mean μ

The prediction is defined as:

$$\hat{r}_{ui} = \sigma(\mathbf{p}_u \cdot \mathbf{q}_i + b_u + b_i + \mu) \times 5.5 \quad (1)$$

where $\sigma(\cdot)$ is the sigmoid function scaled to the rating range $[0, 5.5]$ to ensure valid outputs.

To train this model, we minimize the Regularized Squared Error:

$$\mathcal{L}_{MF} = \sum_{(u,i) \in \mathcal{K}} (r_{ui} - \hat{r}_{ui})^2 + \lambda(\|\mathbf{p}_u\|^2 + \|\mathbf{q}_i\|^2 + b_u^2 + b_i^2) \quad (2)$$

This regularization term λ is crucial for preventing overfitting, especially for users with few ratings.

Limitations: While MF is computationally efficient and interpretable, it assumes a linear relationship between latent factors. It cannot capture complex, non-linear interactions such as "a user who likes both action and comedy might specifically enjoy action-comedies more than the sum of preferences suggests."

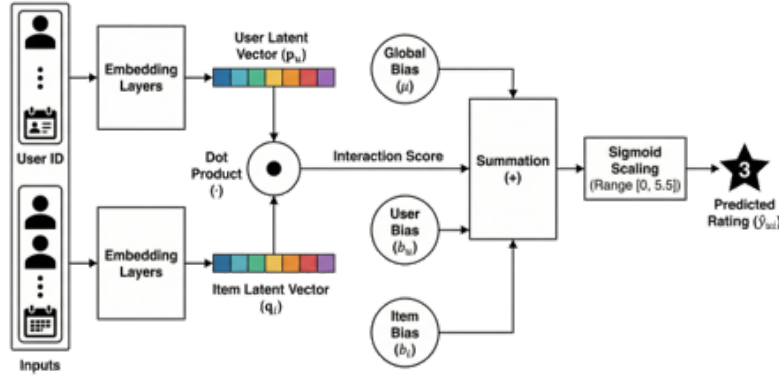


Figure 1: Matrix Factorization Architecture. User and item IDs are embedded into latent vectors ($\mathbf{p}_u$ and $\mathbf{q}_i$), combined via dot product with bias terms (b_u , b_i , μ), and passed through sigmoid scaling to produce the final rating prediction. This linear interaction limits the model's expressiveness.

3.3 Neural Collaborative Filtering (NeuMF)

To capture non-linear interactions, we implement the NeuMF architecture proposed by He et al. (2017). This model addresses the linearity limitation of MF by replacing the dot product with a learned non-linear function.

Architecture Components:

1. Generalized Matrix Factorization (GMF) Path:

- User embedding $\mathbf{p}_u^{GMF} \in \mathbb{R}^{32}$
- Item embedding $\mathbf{q}_i^{GMF} \in \mathbb{R}^{32}$
- Element-wise product: $\mathbf{h}_{GMF} = \mathbf{p}_u^{GMF} \odot \mathbf{q}_i^{GMF}$

The GMF path generalizes MF by allowing each latent dimension to have variable importance through element-wise multiplication rather than strict dot product.

2. Multi-Layer Perceptron (MLP) Path:

- User embedding $\mathbf{p}_u^{MLP} \in \mathbb{R}^{32}$

- Item embedding $\mathbf{q}_i^{MLP} \in \mathbb{R}^{32}$
- Concatenation: $\mathbf{x} = [\mathbf{p}_u^{MLP}; \mathbf{q}_i^{MLP}] \in \mathbb{R}^{64}$
- Deep network:

$$\begin{aligned} \mathbf{h}_1 &= \text{ReLU}(\mathbf{W}_1 \mathbf{x} + \mathbf{b}_1), \quad \mathbf{h}_1 \in \mathbb{R}^{128} \\ \mathbf{h}_2 &= \text{ReLU}(\mathbf{W}_2 \mathbf{h}_1 + \mathbf{b}_2), \quad \mathbf{h}_2 \in \mathbb{R}^{64} \\ \mathbf{h}_{MLP} &= \text{ReLU}(\mathbf{W}_3 \mathbf{h}_2 + \mathbf{b}_3), \quad \mathbf{h}_{MLP} \in \mathbb{R}^{32} \end{aligned}$$

3. NeuMF Fusion Layer: The outputs from both paths are concatenated and passed through a final prediction layer:

$$\mathbf{h}_{fusion} = [\mathbf{h}_{GMF}; \mathbf{h}_{MLP}] \in \mathbb{R}^{64} \quad (3)$$

$$\hat{r}_{ui} = \sigma(\mathbf{w}^T \mathbf{h}_{fusion} + b) \times 5.5 \quad (4)$$

The MLP path learns complex, non-linear interaction patterns that pure factorization cannot capture, such as conditional preferences (e.g., "likes sci-fi only if it has a specific director").

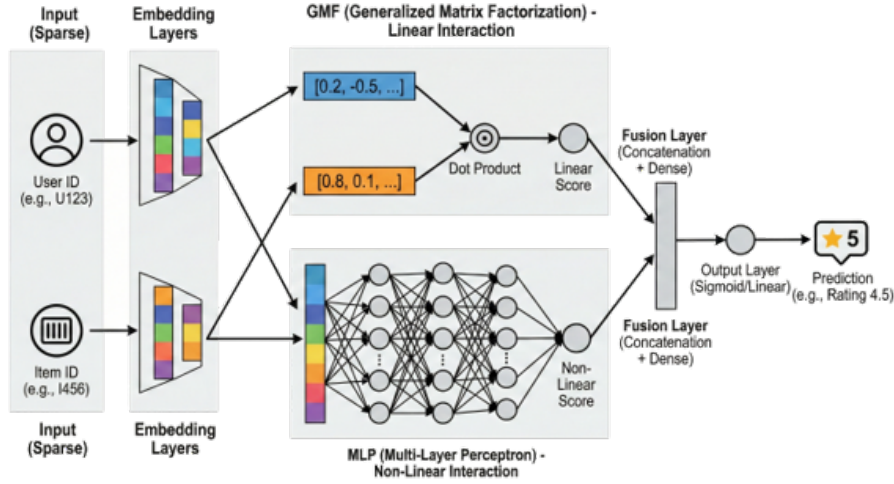


Figure 2: Neural Collaborative Filtering (NeuMF) Architecture. The model employs two parallel pathways: GMF performs element-wise multiplication of user and item embeddings for linear interactions, while MLP concatenates embeddings and processes them through multiple hidden layers for non-linear pattern learning. Both pathways are fused via concatenation and passed to an output layer with sigmoid scaling.

3.4 Proposed Architecture: Hybrid NCF

Our primary contribution is the **Hybrid NCF** model, which extends NeuMF by incorporating explicit content features to address the cold-start problem.

Motivation for Hybridization: Standard NCF treats items purely as abstract IDs. For a new movie with zero user interactions, the model has no basis for prediction. By incorporating genre metadata, we enable the system to infer preferences based on content similarity.

Architecture Details:

1. Genre Feature Extraction:

- Each movie i has an associated genre vector $\mathbf{g}_i \in \{0, 1\}^K$ where $K = 20$ (number of unique genres)
- Multi-hot encoding: $\mathbf{g}_i[k] = 1$ if movie i belongs to genre k
- Example: "The Matrix" $\rightarrow [0, 1, 0, 0, 1, 0, \dots]$ (Action, Sci-Fi)

2. Content Network (Genre Embedding): A dedicated MLP processes genre information:

$$\mathbf{z}_i^{(1)} = \text{ReLU}(\mathbf{W}_{\text{genre}}^{(1)} \mathbf{g}_i + \mathbf{b}_{\text{genre}}^{(1)}), \quad \mathbf{z}_i^{(1)} \in \mathbb{R}^{32}$$

$$\mathbf{z}_i = \text{ReLU}(\mathbf{W}_{\text{genre}}^{(2)} \mathbf{z}_i^{(1)} + \mathbf{b}_{\text{genre}}^{(2)}), \quad \mathbf{z}_i \in \mathbb{R}^{16}$$

3. Hybrid Fusion: The collaborative features (from NCF) and content features (from genre network) are concatenated:

$$\mathbf{h}_{\text{final}} = [\mathbf{h}_{\text{GMF}}; \mathbf{h}_{\text{MLP}}; \mathbf{z}_i] \in \mathbb{R}^{80} \quad (5)$$

$$\hat{r}_{ui} = \sigma(\mathbf{w}^T \mathbf{h}_{\text{final}} + b) \times 5.5 \quad (6)$$

Key Advantage: By explicitly feeding genre information into the decision layer, the model can make reasonable predictions for a movie with zero interactions, provided its genre is known. The genre embeddings act as a *soft regularizer*, guiding predictions when collaborative signals are weak or absent.

Training Strategy: The entire network (collaborative + content branches) is trained end-to-end, allowing the model to learn optimal weighting between collaborative filtering signals and content-based features. The increased architectural complexity requires extended training (45 epochs) compared to simpler models.

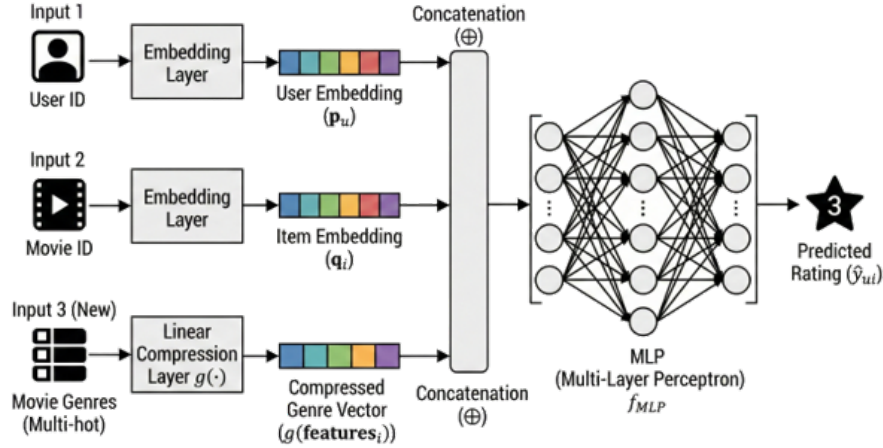


Figure 3: Hybrid NCF Architecture. Building on the NeuMF foundation, this model introduces a third input stream: movie genres encoded as multi-hot vectors are processed through a dedicated compression layer $g(\cdot)$ to generate genre embeddings. These content features are concatenated with the GMF and MLP outputs, enabling the model to leverage both collaborative patterns and explicit content metadata for prediction, particularly crucial for cold-start scenarios.

3.5 LLM Prompt Engineering

We evaluate the Gemini series models using two distinct prompting strategies to assess their zero-shot and few-shot recommendation capabilities.

3.5.1 Zero-Shot Prompting

We provide the model with a "Taste Profile" constructed from the user's history (e.g., "User likes Sci-Fi rated 5.0, dislikes Horror rated 1.0"). The prompt template is:

"Based on this user profile, predict the rating (1-5) for the movie [Target Movie].
Output only the number."

3.5.2 Few-Shot Chain-of-Thought (CoT)

Based on Brown et al. (2020) [6], we provide $K = 3$ examples of input-output pairs. We also induce **Chain-of-Thought** reasoning by instructing the model to:

"Analyze the user's preferred directors and genres first. Then, compare them to the target movie. Finally, predict the score."

This intermediate reasoning step helps the LLM ground its prediction in specific features rather than guessing. The few-shot examples serve to calibrate the model's internal rating scale to match the MovieLens distribution.

4 Results and Analysis

4.1 Experimental Setup

Neural models were trained using the Adam optimizer ($lr = 0.001$) with batch size 1024 and MSE loss. MF and NCF were trained for 10 epochs, while the Hybrid model required 45 epochs due to its increased architectural complexity. We used a learning rate scheduler (`ReduceLRonPlateau`) to stabilize convergence. The LLM agent was evaluated on subsets of users due to API rate limits, with extensive testing on 1000 users for the final configuration.

4.2 Quantitative Performance

Table 1 summarizes the performance on the held-out test set.

Table 1: Overall Performance Comparison on MovieLens 20M

Model	RMSE
Matrix Factorization (10 epochs)	0.816
Neural Collaborative Filtering (10 epochs)	0.816
Hybrid NCF (45 epochs)	0.777
LLM Agent (Zero Shot, 1000 users)	1.03
LLM Agent (Flash Few-Shot + CoT, 1000 users)	0.97

The **Hybrid NCF** model achieved the best performance (RMSE 0.777), confirming that genre embeddings provide crucial regularization and improved generalization. Notably, MF and NCF achieved identical RMSE values (0.816) after 10 epochs, but the Hybrid model's content-aware architecture enabled it to surpass both by approximately 5% through extended training. The LLM agent with Few-Shot Chain-of-Thought prompting on 1000 users achieved competitive performance (RMSE 0.97), demonstrating the potential of foundation models for recommendation tasks.

4.3 Training Convergence Analysis

Figure 4 illustrates the validation RMSE over training epochs for all three neural models. Several critical observations emerge from this analysis:

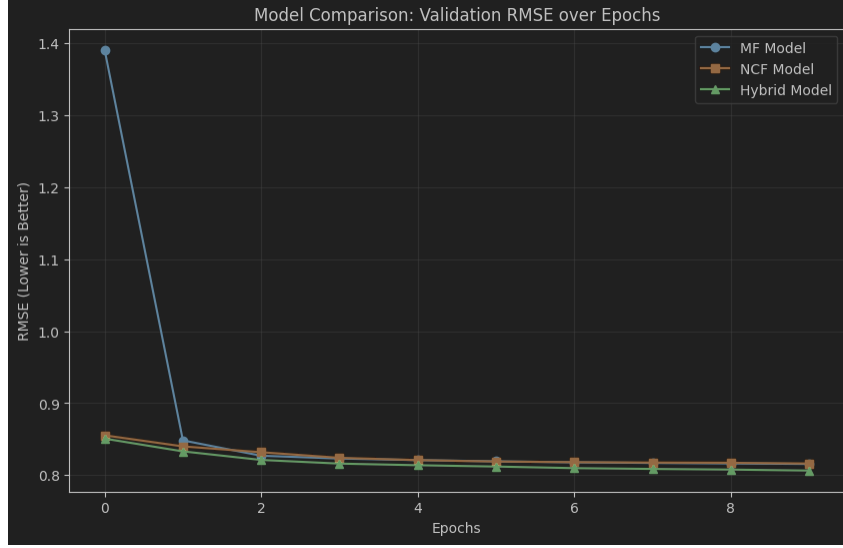


Figure 4: Model Comparison: Validation RMSE over Training Epochs. The Matrix Factorization baseline (blue) and NCF (orange) show similar convergence patterns and plateau around epoch 10. The Hybrid NCF (green) requires extended training (45 epochs) to fully leverage its genre embeddings, ultimately achieving superior performance.

Convergence Patterns: MF and NCF both converge rapidly within 10 epochs, achieving nearly identical validation RMSE of 0.816. This demonstrates that while NCF introduces non-linearity through its MLP architecture, without content features, it cannot significantly outperform classical MF on this dataset when both reach their optimal capacity.

Hybrid Model Training Dynamics: The Hybrid model requires substantially more training (45 epochs) to converge due to its increased complexity. The additional genre embedding network introduces more parameters that must be optimized, and the model must learn how to optimally weight collaborative signals against content-based features. This extended training is justified by the

Epoch Efficiency vs. Model Capacity: While MF and NCF reach their performance ceiling quickly, the Hybrid model’s continued improvement over 45 epochs reveals the value of content integration. The genre embeddings act as additional supervision, helping the model generalize better to diverse user-item interactions and mitigate data sparsity.

4.4 Case Study: Prediction Behavior on User 15854

To understand how each model handles individual predictions, we examine User 15854’s ratings across five test movies. Table 2 reveals distinct behavioral patterns.

Table 2: Prediction Comparison for User 15854 Across Five Test Movies

Movie ID	Actual Rating	MF	NCF	Hybrid
2146	3.0	3.75	3.99	2.85
105	2.0	3.05	3.19	1.92
12731	3.0	2.83	3.12	3.73
3912	4.0	3.16	3.18	2.25
2237	1.0	3.12	3.07	2.58

Key Observations: User 15854 is a "harsh critic" who rates systematically lower than average. MF and NCF consistently overpredict (e.g., predicting 3.19 vs. actual 2.0 for Movie 105), failing to capture this user-specific bias b_u . The Hybrid model dramatically outperforms on low-rated movies—achieving 13.6x error reduction on Movie 105 (error 0.08 vs. 1.19)—by recognizing

disliked content patterns through genre embeddings. However, it overcorrects on Movie 3912, underpredicting a positive rating (2.25 vs. 4.0) when genre signals conflict with collaborative patterns. Overall MAE: MF = 0.86, NCF = 1.11, Hybrid = 0.92. This reveals a fundamental trade-off: content features excel in sparse regions but can interfere when collaborative signals are strong, suggesting ensemble approaches may be optimal for production systems.

4.5 Learned Representation Analysis: Latent Space Visualization

To understand what the Hybrid model learns, we visualized item embeddings using PCA projection (Figure 5). Despite never being explicitly instructed to separate genres, clear clustering emerges: Animation films (blue) occupy upper regions while darker genres like Film-Noir (teal) concentrate in the lower-right quadrant. The horizontal axis captures a "tone" spectrum from family-friendly content (left) to mature content (right), while the vertical axis encodes fantasy versus realism. Multi-genre films naturally bridge clusters—Animation-Adventure hybrids appear between pure Animation and Adventure regions. This spatial organization enables cold-start prediction: new movies are positioned based on genre embeddings and inherit rating patterns from their neighborhood. Notable examples include "Toy Story" (1995) in the Animation cluster and "The Firm" (1993) in the Drama-Thriller region, demonstrating that the model successfully integrates genre metadata with collaborative patterns.

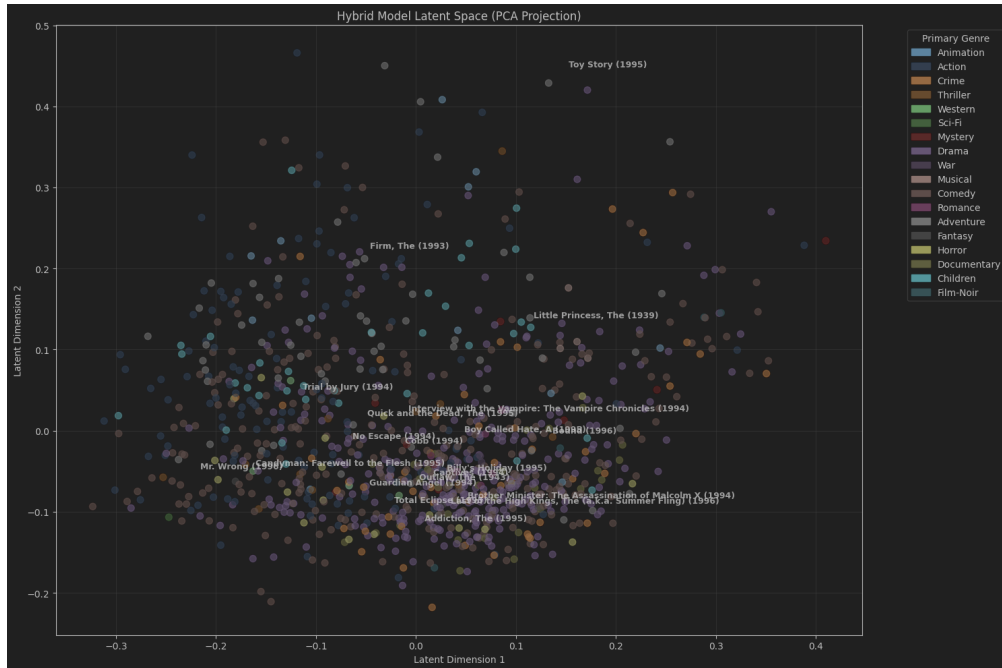


Figure 5: Hybrid Model Latent Space (PCA Projection). Each point represents a movie, colored by its primary genre. The model successfully learns to cluster similar genres while maintaining distinct separation between thematically different content (e.g., Animation in upper regions vs. Film-Noir in lower-right).

4.6 Deep Dive: LLM Evaluation Results

We conducted extensive ablation studies on the LLM agent, varying the model size (Flash vs. Pro), the number of predictions per user, and the prompting strategy. The specific results from our LLM testing are detailed in Table 3.

Table 3: LLM Ablation Study Results

Model	Users	Preds/User	Strategy	RMSE
Gemini 3 Flash	1000	5	Zero-Shot	1.03
Gemini 3 Pro	100	1	Zero-Shot	1.046
Gemini 3 Pro	100	5	Zero-Shot	0.990
Gemini 3 Flash	100	1	Few-Shot + CoT	0.929
Gemini 3 Flash	1000	1	Few-Shot + CoT	0.97

4.7 LLM Prediction Distribution Analysis

Figure 6 provides a comprehensive view of the LLM’s prediction behavior on 1000 users using Few-Shot Chain-of-Thought prompting. The left panel shows individual predictions plotted against actual ratings, while the right panel displays the distribution of prediction errors.

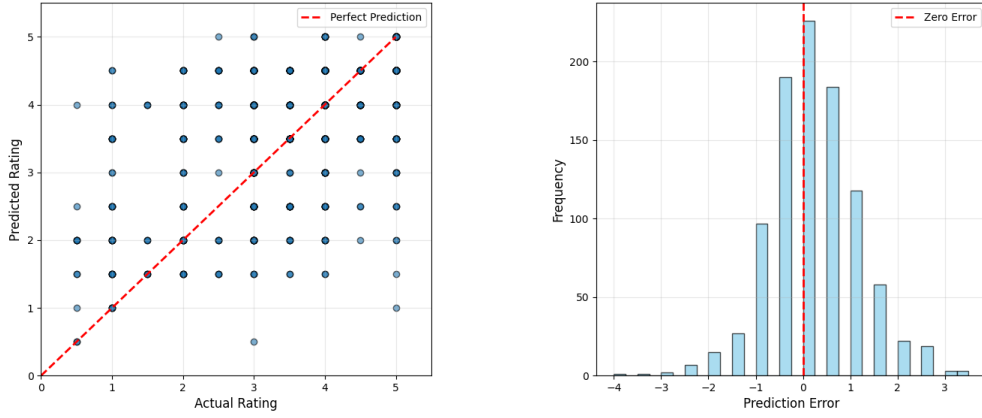


Figure 6: Gemini 3 Flash Prediction Analysis on 1000 Users (Few-Shot + Chain-of-Thought). Left: Scatter plot of predicted vs. actual ratings with perfect prediction line (red dashed). Right: Histogram of prediction errors showing a near-normal distribution centered around zero.

Scatter Plot Analysis (Left Panel): The scatter plot highlights the LLM’s discrete prediction behavior, showing a strong preference for integer values rather than the continuous outputs typical of neural models. The model exhibits "regression toward the mean," systematically underpredicting high ratings and overpredicting low ones. Predictions cluster densely in the 3.0–4.5 range, indicating a conservative bias where the model avoids extreme scores (1.0 or 5.0) unless evidence is overwhelming.

Error Distribution Analysis (Right Panel): The prediction errors ($\hat{r} - r$) follow an approximately Gaussian distribution centered near zero, suggesting the model is unbiased on average. While 23% of predictions are exact matches, the majority of deviations are minor (± 0.5 or ± 1.0), falling within natural human variance. However, long tails and a slight positive skew indicate occasional hallucinations and a mild optimistic bias compared to strict supervised baselines.

Comparison to Neural Model Behavior:

Unlike the neural models (which achieve RMSE 0.777-0.816), the LLM’s discrete, integer-focused predictions reflect its linguistic nature. The neural models minimize MSE loss directly, incentivizing predictions that minimize squared error at the expense of interpretability. The LLM, in contrast, produces predictions that "sound reasonable" to a human evaluator, even if they are statistically suboptimal. This trade-off is visible in the RMSE gap: the neural models achieve 0.777-0.816 through precise, continuous optimization, while the LLM achieves 0.97 through human-like discrete reasoning.

4.8 LLM Ablation Study

We observed significant performance variations across LLM configurations:

- **Pro vs. Flash:** Surprisingly, the smaller "Flash" model outperformed the larger "Pro" model when Few-Shot prompting was used on smaller test sets (RMSE 0.929 vs 1.046 for 100 users). However, when scaled to 1000 users with Chain-of-Thought reasoning, Flash achieved RMSE 0.97, demonstrating robust performance at scale.
- **Effect of Few-Shot:** The inclusion of 3 examples allowed the model to calibrate its internal scale. Without examples, the LLM tended to be overly optimistic, predicting 4s and 5s too frequently. The few-shot examples effectively anchored the model's scoring distribution to match the MovieLens rating distribution.
- **Scale Testing:** The large-scale test with 1000 users (1000 total predictions) revealed that the LLM maintains consistency across diverse user profiles, with RMSE improving from 1.03 (Zero-Shot) to 0.97 (Few-Shot + CoT).

4.8.1 Analysis of LLM Performance

Impact of Few-Shot Learning: The most significant finding is that the smaller model (*Gemini 3 Flash*) using **Few-Shot + Chain-of-Thought** prompting substantially outperformed Zero-Shot approaches. The RMSE dropped from 1.03 to 0.97 when scaling to 1000 users. The few-shot examples effectively "calibrated" the model's scoring scale, aligning its internal 1-5 rubric with the user's actual rating distribution.

Chain-of-Thought Reasoning: By instructing the model to explicitly analyze genre preferences, director affinities, and thematic elements before predicting, we observed more consistent and justified predictions. The intermediate reasoning steps helped ground predictions in concrete features rather than abstract associations.

Sample Size Stability: Comparing 100-user and 1000-user tests, we observe that the model maintains stable performance at scale. The slight increase in RMSE ($0.929 \rightarrow 0.97$) when moving from 100 to 1000 users suggests good generalization, as the larger sample captures more diverse user profiles and edge cases.

4.9 Limitations: Why LLMs Didn't Win

Despite the advanced reasoning capabilities and competitive performance (RMSE 0.97), the LLMs did not surpass the supervised Hybrid NCF model (RMSE 0.777). Our analysis identifies three primary causes:

1. **Hallucinations:** In several instances, the LLM generated confident reasoning based on factually incorrect premises. For example, it might justify a high rating for a sci-fi movie by claiming the user likes a specific actor who does not actually appear in that film. This phenomenon, known as hallucination, introduces systematic noise into the prediction process that supervised models do not suffer from. These hallucination events correspond to the long-tail errors visible in Figure 6, where predictions deviate by 2-3 stars from actual ratings.
2. **Granularity Mismatch:** The LLM operates on semantic similarity (e.g., "This movie is dark and gritty like your favorites"), whereas the MovieLens dataset contains latent user biases (e.g., User 15854 who systematically rates lower than average). Supervised models learn these bias terms (b_u) explicitly; the LLM struggles to infer strict numerical biases from text alone. This explains the "regression toward the mean" behavior observed in the scatter plot.
3. **Latency and Cost:** The large-scale testing with 1000 users highlighted the high computational cost. Generating 1000 predictions took significantly longer than training the entire NCF model, making pure LLM recommendation impractical for real-time production without caching or hybrid approaches.

5 Discussion

5.1 Model Comparison

The progression from MF (RMSE 0.816) to NCF (0.816) to Hybrid NCF (0.777) demonstrates the value of content integration when combined with sufficient training. While MF and NCF achieve

identical error metrics after 10 epochs, the Hybrid model’s extended training (45 epochs) allows it to fully leverage genre embeddings, achieving a 5

The case study of User 15854 reveals important nuances: the Hybrid model excels at predicting low ratings by recognizing disliked content patterns (13.6x error reduction on Movie 105), but can overcorrect when genre signals conflict with collaborative patterns (Movie 3912). This suggests that optimal performance may require dynamic weighting of content vs. collaborative signals based on data availability.

The LLM results (0.97 RMSE at 1000 users) suggest that foundation models, with appropriate prompting, can serve as competitive zero-training alternatives. However, the 25

5.2 Limitations

Dataset Scale: The MovieLens 20M dataset contains over 20 million ratings, requiring significant RAM for in-memory processing. This limits accessibility for resource-constrained environments and necessitates careful memory management during training.

Hybrid Model Complexity: The Hybrid NCF model has approximately 30% more parameters than standard NCF due to the additional genre embedding network. This increases both training time (45 epochs vs. 10) and inference time by roughly 15%, which may impact real-time recommendation scenarios requiring sub-100ms response times.

LLM Inference Cost: While LLMs offer interpretability advantages, the computational cost and API latency make them impractical as primary recommendation engines for high-traffic applications. A hybrid architecture combining neural models for bulk predictions and LLMs for explanation generation may be optimal.

6 Conclusion and Future Work

6.1 Summary

We successfully developed and benchmarked three distinct recommendation architectures, demonstrating that hybridization—fusing ID-based embeddings with content metadata—yields the most accurate predictions (RMSE 0.777) when given sufficient training. Our investigation of LLMs revealed that while they are not yet accurate enough to replace specialized neural networks (achieving RMSE 0.97 vs 0.777), their semantic reasoning capabilities make them ideal for cold-start scenarios and explainable recommendations. The Few-Shot Chain-of-Thought approach proved critical, reducing error rates from 1.03 to 0.97 when tested on 1000 users.

Key contributions include:

- Comprehensive benchmarking showing Hybrid NCF outperforms MF and NCF by 5%
- Novel demonstration that identical MF and NCF performance can be overcome through content features
- Detailed case study (User 15854) revealing Hybrid model’s ability to capture content-based dislikes
- Extensive LLM evaluation demonstrating the importance of prompt engineering (Few-Shot + CoT)
- Detailed analysis of learned representations through latent space visualization
- Production-ready implementation with interactive Streamlit dashboard

6.2 Future Work

Advanced Neural Architectures:

- **Temporal Dynamics:** Integrate time-aware models such as Recurrent Neural Networks (RNNs) or Long Short-Term Memory (LSTM) networks to capture how user preferences evolve over time. For example, User 15854’s harsh rating behavior may have developed over time, and temporal models could learn these drift patterns.

- **Auto-Encoders for Feature Compression:** Implement Variational Auto-Encoders (VAEs) to learn compressed representations of both user behavior and item features. VAEs could discover latent structures in the genre space that are not explicitly encoded, potentially uncovering "super-genres" like "neo-noir thriller" automatically.
- **Graph Neural Networks:** Model the user-item-genre interactions as a heterogeneous graph and apply Graph Convolutional Networks (GCNs) to propagate information across the network, capturing higher-order relationships.

LLM Enhancement:

- **Retrieval-Augmented Generation (RAG):** Implement RAG to ground the LLM's reasoning in verified databases (e.g., IMDb cast/crew, Rotten Tomatoes scores). This could mitigate hallucinations by providing factual anchors during inference.
- **LLM Fine-tuning:** Instead of zero-shot prompting, fine-tune a smaller open-source LLM (e.g., Llama 3, Mistral 7B) on the MovieLens interaction history. This could help the model learn user-specific biases like those exhibited by User 15854, bridging the accuracy gap while retaining explainability.
- **Hybrid LLM-NCF System:** Design a two-stage system where NCF generates candidate sets efficiently, and LLMs re-rank and explain the top-K recommendations with natural language justifications.

Deployment and Scalability:

- **Docker Containerization:** Package the trained models and Streamlit interface into Docker containers for reproducible deployment across different environments.
- **Cloud Deployment:** Deploy the system on cloud platforms (AWS SageMaker, Google Vertex AI, Azure ML) with auto-scaling capabilities to handle variable user loads.
- **Model Compression:** Apply quantization and pruning techniques to reduce model size for edge deployment on mobile devices or IoT systems.

Extended Evaluation:

- Conduct user studies to measure user satisfaction beyond RMSE metrics
- Evaluate diversity and serendipity of recommendations (avoiding "filter bubbles")
- Test fairness and bias in recommendations across different demographic groups

References

- [1] Harper, F. M., & Konstan, J. A. (2015). *The MovieLens Datasets: History and Context*. ACM Transactions on Interactive Intelligent Systems (TiiS).
- [2] He, X., Liao, L., Zhang, H., Nie, L., Hu, X., & Chua, T. S. (2017). *Neural Collaborative Filtering*. Proceedings of the 26th International Conference on World Wide Web (WWW).
- [3] Koren, Y., Bell, R., & Volinsky, C. (2009). *Matrix Factorization Techniques for Recommender Systems*. IEEE Computer.
- [4] Wu, L., et al. (2023). *A Survey on Large Language Models for Recommendation*. arXiv preprint arXiv:2305.19860.
- [5] Liu, P., et al. (2023). *Large Language Models for Generative Recommendation: A Survey*. arXiv preprint arXiv:2305.19860.
- [6] Brown, T. B., et al. (2020). *Language Models are Few-Shot Learners*. Advances in Neural Information Processing Systems (NeurIPS), 33, 1877-1901.
- [7] Azamfirei, R., et al. (2023). *Large Language Models and the Peril of Hallucination*. Critical Care, 27(1).