

УДК 004

РЕГУЛЯРИЗАЦИЯ BERT В ЗАДАЧЕ РЕКОМЕНДАЦИЙ НА ОСНОВЕ STABLE RANK МАТРИЦЫ ЭМБЕДДИНГОВ

Хуснутдинов Э.Р., студент, 1 курс магистратуры
Тюменский государственный университет
г. Тюмень

BERT4Rec [1] — одна из наиболее распространённых моделей последовательной рекомендации. Она переносит идею маскированного языкового моделирования из BERT на пользовательские последовательности: часть элементов маскируется, и модель учится их восстанавливать. На практике такой подход хорошо работает для пользователей с длинной историей, однако качество рекомендаций для редких (холодных) элементов остаётся проблемой. Одна из вероятных причин — вырождение обученных эмбеддингов: для элементов с малым числом взаимодействий модель фактически использует лишь одно-два направления в пространстве представлений.

Мы проверяем эту гипотезу с помощью спектрального анализа матриц весов и предлагаем простой метод регуляризации, основанный на stable rank — непрерывном аналоге матричного ранга. Все эксперименты выполнены на датасете MovieLens-1M [2] через фреймворк RecBole [3].

1. Спектральные характеристики матриц весов

Спектральная норма матрицы W — это её наибольшее сингулярное значение:

$$\|W\|_2 = \sigma_1(W) = \max\{\|Wx\| : \|x\| = 1\}.$$

По сути, это максимальный коэффициент, на который матрица может растянуть входной вектор. Stable rank — отношение суммарной «энергии» всех сингулярных значений к энергии главного:

$$srank(W) = \|W\|_F^2 / \|W\|_2^2 = \sum \sigma_i^2 / \sigma_1^2.$$

Если спектр равномерный, stable rank совпадает с рангом. Если же одно сингулярное значение доминирует, stable rank падает к единице — матрица ведёт себя как ранг-1 [4]. Для нас это индикатор: чем ниже stable rank у эмбеддингов, тем меньше направлений реально задействовано.

2. Что происходит в модели без регуляризации

Мы собрали stable rank, spectral norm и норму Фробениуса для всех двумерных матриц BERT4Rec на каждой из 15 эпох обучения. Матрица

эмбедингов содержит 3708 строк (3706 фильмов плюс строка дополнения и маскирующий токен) при размерности 256.

Разные типы слоёв ведут себя по-разному. Проекции Value и Dense в блоках внимания наращивают stable rank от 15 до 35 — эти матрицы постепенно осваивают всё больше направлений. Первые слои FFN, наоборот, держат stable rank около 5–8 на протяжении всего обучения: после нелинейной активации им хватает нескольких доминирующих направлений. Query и Key колеблются в диапазоне 20–26, не показывая чёткого тренда.

Самая интересная картина открывается при разбивке эмбедингов по популярности. Мы разделили элементы на три группы по квантилям частоты в обучающей выборке: холодные (cold, 758 элементов, нижние 20%), тёплые (warm, 2200) и горячие (hot, 743, верхние 20%).

Разрыв получился огромным. У холодных элементов stable rank начинается с 1.6 и вырастает лишь до 3.3. Это значит, что при 256-мерном пространстве их эмбединги фактически лежат в 3-мерном подпространстве. Spectral norm при этом растёт быстрее всех (14 → 30.5) — одно направление набирает мощность, а остальные простаивают. У горячих элементов ситуация принципиально другая: stable rank 8.4—14.2, spectral norm растёт умеренно (5.5 → 10.7). Модели хватает данных, чтобы развести их эмбединги по многим направлениям.

3. Метод регуляризации

Идея прямолинейна: если низкий stable rank — проблема, будем штрафовать за него при обучении. К кросс-энтропийному лоссу добавляем слагаемое, равное обратному stable rank матрицы эмбедингов E:

$$L = L_{CE} + \lambda \cdot \sigma_1^2(E) / \|E\|^2 F.$$

Чем сильнее σ_1 доминирует в спектре, тем больше штраф. Оптимизатор вынужден «расталкивать» энергию из главного направления в остальные, что поднимает stable rank.

Техническая трудность — дифференцируемое вычисление σ_1 . Полное SVD на матрице 3708×256 в каждом батче слишком дорого. Мы используем степенные итерации: 10 шагов уточнения сингулярных векторов u , v выполняются без градиентов (detach), а итоговое приближение $\hat{\sigma}_1 = u^T E v$ — дифференцируемый скаляр, через который идёт обратное распространение [5]. Векторы u , v сохраняются между батчами (warm-start), поэтому 10 итераций достаточно для хорошей аппроксимации.

Важный момент: градиент этого штрафа затрагивает только матрицу E. Остальные веса модели (внимание, FFN, выходной слой) получают сигнал только от L_{CE} . Косвенно они, конечно, адаптируются к изменившимся эмбедингам, но прямого давления от SR-штрафа на них нет.

4. Эксперимент

MovieLens-1M: около миллиона оценок, 3706 фильмов, 6040 пользователей. Архитектура: 2 слоя трансформера, 4 головы, скрытая

размерность 256, dropout 0.2, маскирование 20%. Оптимизатор Adam, learning rate 0.001, 15 эпох, батч 2048. Данные разбиты по схеме leave-one-out: последний элемент последовательности идёт в тест, предпоследний — на валидацию.

Сравниваем две модели: baseline без дополнительной регуляризации и SR-reg с $\lambda = 0.01$. Обе обучены с нуля с одинаковым сидом.

Метрики качества приведены в таблице 1.

Таблица 1 — Сравнение метрик качества рекомендаций

Метрика	Baseline	SR-reg	Разница
NDCG@10 (valid)	0.1893	0.1897	+0.0004
NDCG@10 (test)	0.1825	0.1794	−0.0031
Recall@10 (test)	0.3111	0.3058	−0.0053
MRR@10 (test)	0.1431	0.1407	−0.0024
Hit@20 (test)	0.4131	0.4098	−0.0033

На валидации SR-reg набирает $\text{NDCG@10} = 0.1897$, что практически совпадает с baseline (0.1893). На тесте — небольшая просадка: 0.1794 против 0.1825. Recall@10 и MRR@10 снижаются в том же диапазоне, порядка 0.003–0.005. Разница невелика, но устойчива по всем метрикам.

Спектральные характеристики — в таблице 2.

Таблица 2 — Спектральные характеристики на 14-й эпохе

Показатель	Baseline	SR-reg	Изменение	%
SR (вся матрица)	7.01	9.03	+2.02	+29%
SR (cold items)	3.28	4.26	+0.98	+30%
SR (warm items)	9.12	10.24	+1.12	+12%
SR (hot items)	14.20	14.24	+0.04	~0%
σ_1 (spectral norm)	39.04	33.77	−5.27	−13%

Здесь эффект заметен гораздо отчётливее. Stable rank всей матрицы эмбедингов вырос с 7.01 до 9.03 (+29%). Cold items получили наибольший относительный прирост (+30%), при этом hot items остались на том же уровне

(+0.04). Spectral norm снизилась на 13% ($39.04 \rightarrow 33.77$), то есть доминирование первого сингулярного значения действительно ослабло.

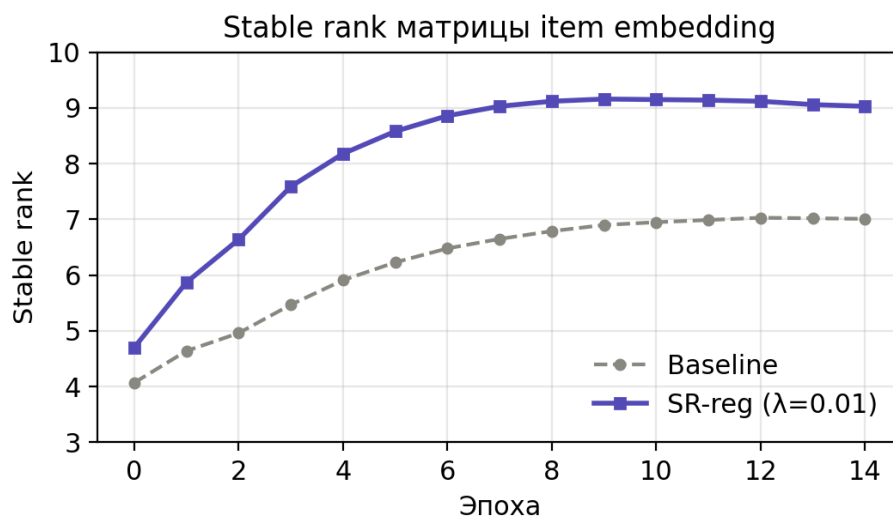


Рисунок 1 — Динамика stable rank матрицы эмбедингов по эпохам

На рис. 1 видно, как формируется разрыв: SR-reg быстро уходит от baseline в первые 5 эпох и выходит на плато около 9.0, тогда как baseline стабилизируется на 7.0.

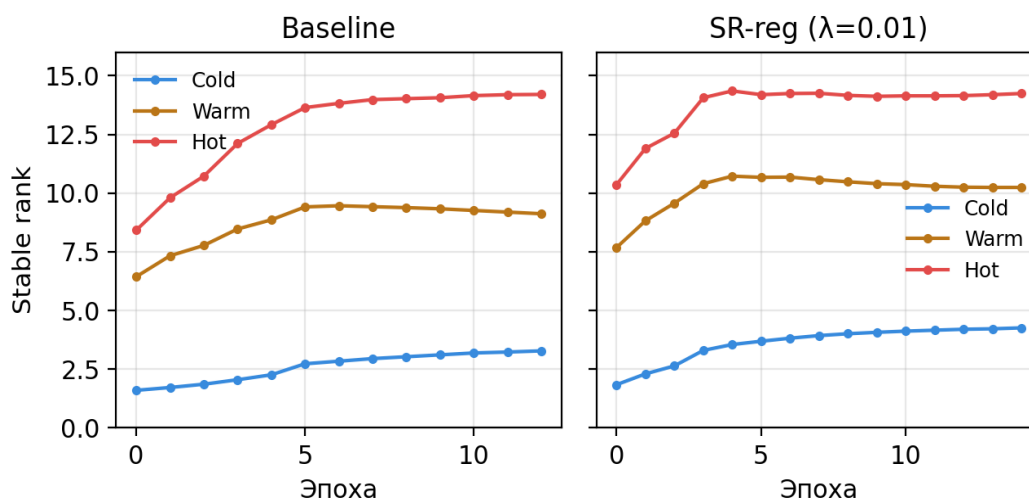


Рисунок 2 — Stable rank эмбедингов по группам популярности

Побакетная разбивка (рис. 2) подтверждает, что регуляризация помогает прежде всего холодным элементам: их stable rank поднимается с 3.3 до 4.3. У тёплых элементов прирост есть, но после 5-й эпохи начинается лёгкий откат — spectral norm продолжает расти и «съедает» часть набранного запаса.

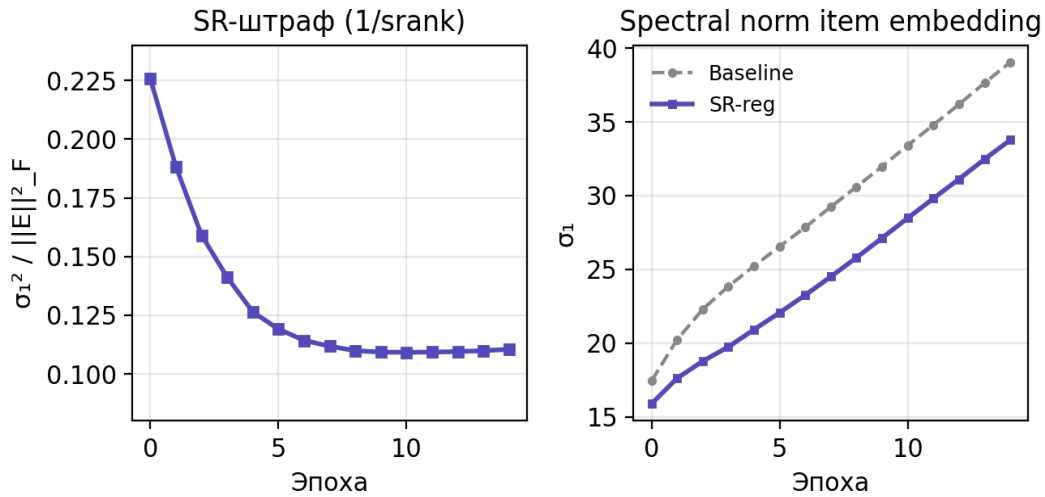


Рисунок 3 — Сходимость штрафного члена и динамика spectral norm

Штраф $\sigma_1^2 / \|E\|^2 F$ падает с 0.226 до 0.110 за первые 8 эпох, после чего стабилизируется (рис. 3, слева). Spectral norm с регуляризацией растёт на протяжении всего обучения, но заметно медленнее, чем у baseline (рис. 3, справа).

5. Обсуждение

Регуляризация сработала ровно так, как задумано спектрально: stable rank поднялся, доминирование σ_1 ослабло, холодные элементы получили более разнообразные представления. Однако в метриках рекомендаций улучшения не произошло — напротив, на тесте наблюдается небольшая просадка.

Мы видим две причины. Первая: штраф применяется ко всей матрице целиком. Горячие элементы уже имеют stable rank 14 и в регуляризации не нуждаются, но оптимизатор всё равно тратит часть градиентного бюджета на подавление их σ_1 . Это может слегка нарушать уже хорошо подобранные представления популярных фильмов. Вторая: разрыв valid-test увеличился (с 3.6% до 5.4%). Штраф меняет геометрию эмбедингов, и трансформерным слоям приходится к ней адаптироваться. При 15 эпохах этот процесс может быть незавершён, особенно учитывая, что прямого градиентного сигнала от SR-штрафа к ним нет.

Логичный следующий шаг — применять штраф дифференцированно по группам: сильнее для cold, слабее для warm, не применять к hot. Это позволит сконцентрировать эффект на проблемной части спектра. Ещё одно перспективное направление — расширить регуляризацию на другие слои модели, задавая для каждого свой порог stable rank и штрафуя только за превышение (upper bound) [6].

Заключение

Мы провели спектральный анализ весов BERT4Rec на MovieLens-1M и обнаружили, что эмбединги холодных элементов работают в вырожденном режиме: stable rank ≈ 3 при 256 измерениях. Предложенный метод регуляризации через штраф на обратный stable rank поднимает его на 29%,

снижает spectral norm на 13% и сохраняет валидационное качество на уровне baseline. Тестовые метрики незначительно снизились, что объясняется грубостью глобального штрафа. Дальнейшая работа направлена на избирательную регуляризацию по группам популярности элементов.

Список литературы

1. Sun, F. BERT4Rec: Sequential Recommendation with Bidirectional Encoder Representations from Transformers / F. Sun, J. Liu, J. Wu [et al.] // Proceedings of the 28th ACM International Conference on Information and Knowledge Management (CIKM). — 2019. — P. 1441–1450.
2. Harper, F. M. The MovieLens Datasets: History and Context / F. M. Harper, J. A. Konstan // ACM Transactions on Interactive Intelligent Systems. — 2015. — Vol. 5, No. 4. — Article 19.
3. Zhao, W. X. RecBole: Towards a Unified, Comprehensive and Efficient Framework for Recommendation Algorithms / W. X. Zhao, S. Mu, Y. Hou [et al.] // Proceedings of the 30th ACM International Conference on Information and Knowledge Management (CIKM). — 2021. — P. 4653–4664.
4. Rudelson, M. Sampling from large matrices: An approach through geometric functional analysis / M. Rudelson, R. Vershynin // Journal of the ACM. — 2007. — Vol. 54, No. 4. — Article 21.
5. Miyato, T. Spectral Normalization for Generative Adversarial Networks / T. Miyato, T. Kataoka, M. Koyama, Y. Yoshida // Proceedings of the 6th International Conference on Learning Representations (ICLR). — 2018.
6. Yoshida, Y. Spectral Norm Regularization for Improving the Generalizability of Deep Learning / Y. Yoshida, T. Miyato // arXiv preprint arXiv:1705.10941. — 2017.