

UNIwersYTET WARSZAWSKI
WYDZIAŁ FIZYKI

ROZPRAWA DOKTORSKA

Sieci tensorowe
w metrologii kwantowej

Autor:
Krzysztof
CHABUDA

Promotor:
dr hab. Rafał DEMKOWICZ-
DOBRZAŃSKI, prof. UW



Czerwiec 2020

Streszczenie

Wyznaczenie optymalnych kwantowych strategii metrologicznych w realistycznych wielocząstkowych kwantowych modelach jest w ogólności wymagającym zadaniem, któremu efektywnie nie potrafią podolać nawet najlepsze metody analityczne i numeryczne metrologii kwantowej. W tej dysertacji przedstawiamy wszechstronne podejście bazujące na metodach sieci tensorowych do rozwiązywania problemów metrologii kwantowej. Sieci tensorowe stanowią naturalny język do opisu problemów o lokalnym charakterze korelacji/splątania. Ta cecha sprawiła, że udało nam się zaproponować efektywny opis modeli metrologicznych z krótkozasięgowymi przestrzennymi oraz czasowymi korelacjami szumu przy użyciu sieci tensorowych w postaci macierzowych stanów produktowych. Wykorzystując je byliśmy w stanie znaleźć fundamentalne ograniczenia na precyzję dla takich modeli oraz zidentyfikować optymalne stany początkowe w eksperymencie. Ponadto zastosowanie nieskończonych macierzowych stanów produktowych umożliwiło nam bezpośrednie i efektywne wyznaczenie asymptotycznych wartości precyzji w optymalnych protokołach metrologicznych w granicy nieskończenie wielu cząstek. Zastosowanie naszego podejścia zostanie zilustrowane na przykładzie problemu stabilizacji zegara atomowego (czasowe korelacje szumu) oraz estymacji indukcji pola magnetycznego (przestrzenne korelacje szumu). Produktem ubocznym zaproponowanego formalizmu, opartego o sieci tensorowe, do obliczania kwantowej informacji Fishera, jest możliwość wyznaczenia wierności – parametru który jest szeroko stosowany w fizyce wielu ciał do badania kwantowych przejść fazowych.

Podziękowania

Chciałbym podziękować wszystkim moim nauczycielom. To ich pasja, chęć przekazywania wiedzy oraz zaangażowanie sprawiły, że sam też poświęciłem się zgłębianiu wiedzy o otaczającym nas świecie, czego ukoronowaniem jest niniejsza dysertacja. W szczególności pragnę podziękować:

- Rafałowi Demkowiczowi-Dobrzańskiemu – za wprowadzenie mnie w fascynujący świat kwantowej teorii informacji, wiarę we mnie i podjęcie się bycia moim promotorem oraz bycie zawsze gotowym do pomocy, dyskusji i nieustanne inspirowanie do poszukiwania nowych dróg do rozwiązania nurtujących problemów;
- Jackowi Dziarmadze – za rozjaśnienie tajników sieci tensorowych, pasjonujące dyskusje naukowe i mimo kontaktu na odległość nieustrudzone udzielanie wyczerpujących odpowiedzi na wszystkie moje pytania.

Chciałbym również podziękować Uniwersytetowi Warszawskiemu i Narodowemu Centrum Nauki za finansowanie moich badań.

Lista publikacji autora

Publikacje bezpośrednio powiązane z tą dysertacją

- K. Chabuda, I. D. Leroux i R. Demkowicz-Dobrzański, „*The quantum Allan variance*”, *New Journal of Physics* **18**, 083035 (2016).
- K. Chabuda, J. Dziarmaga, T. J. Osborne i R. Demkowicz-Dobrzański, „*Tensor-network approach for quantum metrology in many-body quantum systems*”, *Nature Communications* **11**, 250 (2020).

Publikacje wykraczające poza temat tej dysertacji

- Z. Puchała, Ł. Rudnicki, K. Chabuda, M. Paraniak i K. Życzkowski, „*Certainty relations, mutual entanglement and non-displacable manifolds*”, *Phys. Rev. A* **92**, 032109 (2015).
- K. M. Bąk, K. Chabuda, H. Montes, R. Quesada i M. J. Chmielewski, „*1,8-Diamidocarbazoles: an easily tuneable family of fluorescent anion sensors and transporters*”, *Org. Biomol. Chem.* **16**, 5188-5196 (2018).

Spis treści

Wprowadzenie	1
1 Metrologia kwantowa	7
1.1 Klasyczna teoria estymacji	7
1.1.1 Podejście oparte na informacji Fishera	8
1.1.2 Podejście Bayesowskie	10
1.2 Kwantowa teoria estymacji	11
1.2.1 Podejście oparte na kwantowej informacji Fishera	12
1.2.2 Kwantowe podejście Bayesowskie	16
1.3 Fundamentalne ograniczenia w metrologii kwantowej	17
1.3.1 Przypadek bez dekoherencji	18
1.3.2 Przypadek z dekoherencją	20
2 Sieci tensorowe	26
2.1 Diagramy sieci tensorowych	26
2.2 Rozkład według wartości osobliwych	29
2.3 Prawo powierzchniowe dla entropii splątania	31
2.4 Główne rodzaje sieci tensorowych	32
2.5 Jednowymiarowe sieci tensorowe	34
2.5.1 Macierzowe stany produktowe	34
2.5.2 Macierzowe operatory produktowe	41
2.5.3 Granica termodynamiczna	43
3 Metrologia kwantowa w języku sieci tensorowych	46
3.1 Sformułowanie badanego modelu metrologicznego	46
3.2 ρ_φ i $\dot{\rho}_\varphi$ w reprezentacji MPO	49
3.3 SLD w reprezentacji MPO	53
3.4 Optymalizacja QFI dla układów skończonych	54
3.5 Optymalizacja QFI w przypadku asymptotycznym	59
4 Przykładowe zastosowania	66
4.1 Pomiar pola magnetycznego w obecności lokalnie skorelowanego szumu	66
4.2 Stabilizacja zegara atomowego	79
4.3 Wierność wielociałowego stanu termicznego	87
5 Podsumowanie	89
Bibliografia	91

Wprowadzenie

Narodziny mechaniki kwantowej w XX wieku i jej dynamiczny rozwój całkowicie zmieniły nasze postrzeganie mikroświata. Jej idee szybko wykroczyły poza rozważania fizyków oraz dyskusje filozofów i miały ogromny wpływ na rozwój techniki zmieniając nasze codzienne życie. Współcześnie obserwujemy prężny rozwój technologii kwantowych pozwalających manipulować stanami pojedynczych cząstek, za co została przyznana w 2012 roku Nagroda Nobla w dziedzinie fizyki [Haroche 2013; Wineland 2013]. Bardzo precyzyjne pomiary pozwalają przekraczać kolejne bariery, jak np. coraz dokładniejszy pomiar czasu w zegarach atomowych. Jest on możliwy w wyniku bardzo precyzyjnych pomiarów częstotliwości za pomocą optycznego grzebienia częstotliwości, który pozwala zmierzyć częstotliwość w ponad 100 THz paśmie przenoszenia z niepewnością względną dochodzącą do 10^{-19} [L.-S. Ma i in. 2004]. Za wkład w rozwój precyzyjnej spektroskopii laserowej, w tym technikę optycznego grzebienia częstotliwości, została przyznana Nagroda Nobla w dziedzinie fizyki w 2005 roku [Glauber 2006; J. L. Hall 2006; Hänsch 2006]. Innym przykładem bardzo precyzyjnych pomiarów są interferometryczne pomiary odległości, które pozwoliły dokonać detekcji fal grawitacyjnych. Detektor fal grawitacyjnych LIGO jest w stanie wykryć zmianę odległości (powiązaną ze zmianą amplitudy sygnału) pomiędzy swoimi lustrami oddalonymi od siebie o 4 km o odległość która jest 10 000 razy mniejsza niż promień protonu. Zostało to uhonorowane Nagrodą Nobla w dziedzinie fizyki w 2017 roku [Weiss 2018; Barish 2018; Thorne 2018]. Osiągnięcia te powodują, że coraz większe znaczenie przy dokonywaniu pomiarów mają fundamentalne ograniczenia nakładane przez teorię kwantową, co zaowocowało w ostatnich latach znaczącym wzrostem zainteresowania metrologią kwantową.

Metrologia kwantowa zajmuje się nie tylko ograniczeniami, jakie na możliwą do uzyskania precyzję pomiarów narzuca mechanika kwantowa, ale też możliwościami jakie ona stwarza na poprawę tej precyzji poprzez wykorzystanie stanów splątanych czy pomiarów kolektywnych. Z matematycznego punktu widzenia metrologia kwantowa stanowi przykład zastosowania teorii estymacji (działu statystyki) do rozkładu wyników pomiarów podyktowanych prawami mechaniki kwantowej. W tym duchu została ona zapoczątkowana przez Helstroma [1976] i Holevo [1982], a później rozwinięta między innymi przez Braunsteina i Cavesa [1994] oraz Hayashiego [2005]. Takie podejście cechuje to, że pozwala nam określić fundamentalne, wynikające z praw mechaniki kwantowej, ograniczenie na precyzję, ale często nie mówi jak można je osiągnąć.

Alternatywne podejście, nie odwołujące się do ogólnej kwantowej teorii estymacji, polega na poszukiwaniu konkretnej strategii estymacyjnej (stanu początkowego w jakim należy przygotować układ, stosowanych pomiarów oraz estymatorów użytych do analizy wyników), która będzie odpowiednia dla danego

układu eksperymentalnego. Przewidywania teoretyczne uzyskane w tym podejściu znacząco przyczyniły się do wzrostu zainteresowania metrologią kwantową, ponieważ pokazały, że dla wielu układów optycznych [Caves 1981; Holland i Burnett 1993; Yurke, McCall i Klauder 1986] i atomowych [Wineland, J. J. Bollinger, Itano, F. L. Moore i in. 1992; Wineland, J. J. Bollinger, Itano i Heinzen 1994; J. J. . Bollinger i in. 1996] wystarczy zastosowanie stanów ściśniętych [J. Ma i in. 2011; Schnabel 2017] oraz standardowych metod pomiarowych, aby uzyskać istotną poprawę precyzji ponad to co można osiągnąć klasycznymi metodami.

Kolejny znaczący krok w rozwoju metrologii kwantowej stanowiły prace Giovannettiego, Lloyda i Maccone [2004; 2006], które czerpiąc inspiracje z równolegle rozwijającej się dziedziny informacji kwantowej, ustanowiły ogólny schemat do badania problemów metrologicznych, ze szczególnym uwzględnieniem roli splątania. Współcześnie metrologia kwantowa przeżywa rozkwit zarówno na polu teoretycznym, jak i doświadczalnym [Tóth i Apellaniz 2014; Demkowicz-Dobrzański, Jarzyna i Kołodyński 2015; Degen, Reinhard i Cappellaro 2017; Pezzé, Smerzi i in. 2018; Pirandola i in. 2018]. Ogromny postęp technik eksperymentalnych sprawia, że ciągle przybywa nowych przykładów poprawy precyzji pomiarów poprzez wykorzystanie własności kwantowych układów zarówno w dziedzinie eksperymentów optycznych [Mitchell, Lundeen i Steinberg 2004; Nagata i in. 2007; Kacprowicz i in. 2010; Abadie i in. 2011; Aasi i in. 2013], jak i atomowych [Leibfried i in. 2004; Schmidt i in. 2005; Roos i in. 2006; Appel i in. 2009; Leroux, Schleier-Smith i Vuletić 2010; Wasilewski i in. 2010]. Ukoronowaniem tych wysiłków i przykładem zastosowania metodologii metrologii kwantowej, nie tylko do udowodnienia słuszności teorii, ale do poprawy precyzji wykonywanych na co dzień pomiarów, jest detekcja fal grawitacyjnych, gdzie aktualnie wykorzystanie stanów ściśniętych odpowiada za wzrost liczby obserwacji aż do 50% [Acernese i in. 2019; Tse i in. 2019].

Jakkolwiek narzędzia metrologii kwantowej są potężne i udowodniły swoją użyteczność to trapi je ta sama przypadłość, która dotyka inne działy mechaniki kwantowej, które starają się badać złożone układy wielociałowe, mianowicie wymiar potrzebnej do opisu takich układów przestrzeni Hilberta skaluje się eksponencjalnie z rozmiarem układu [Feynman 1982]. Podczas gdy małe problemy można łatwo zbadać stosując bezpośrednio narzędzia numeryczne to nawet niewielki wzrost liczby elementarnych obiektów szybko sprawia, że problem staje się zbyt złożony do numerycznego zasymulowania. Stanowi to poważną koncepcyjną przeszkodę na drodze do rozwoju technologii kwantowych – dla przykładu samo przechowanie w pamięci komputera tomograficznej rekonstrukcji wielokubitowego stanu kwantowego szybko staje się niemożliwe, gdy liczba cząstek przekroczy 40 [Cramer i in. 2010].

Pomimo tych problemów, metrologia kwantowa odniosła wiele sukcesów. Jest to związane z tym, że w przypadkach bez szumu [Giovannetti, Lloyd i Maccone 2006], jak i pewnych modelach z szumem, np. dla globalnego defazowania [Dorner 2012; Macieszczak, Fraas i Demkowicz-Dobrzański 2014], opis układu może być bez straty ogólności, ograniczony do podprzestrzeni całkowicie symetrycznej. Jako że wymiar tej podprzestrzeni skaluje się liniowo z liczbą cząstek, to jest możliwe przeprowadzenie bezpośrednich efektywnych obliczeń

nawet w granicy dużych rozmiarów układu.

Postęp dokonał się również dla nieskorelowanych modeli szumu (które prowadzą z podprzestrzeni całkowicie symetrycznej) poprzez stworzenie nowych narzędzi, które pozwalają wyznaczyć ograniczenia (zazwyczaj asymptotycznie wysycalne) na możliwą do osiągnięcia precyzję w eksperymencie [Escher, Matos Filho i Davidovich 2011; Demkowicz-Dobrzański, Kołodyński i Guță 2012; Kołodyński i Demkowicz-Dobrzański 2013; Knysh, Chen i Durkin 2014; Demkowicz-Dobrzański i Maccone 2014; Sekatski i in. 2017; Demkowicz-Dobrzański, Czajkowski i Sekatski 2017; Zhou i in. 2018]. Jednakże, w przypadku gdy mamy do czynienia z częściowo skorelowanym szumem lub gdy chcemy badać stany spoza przestrzeni całkowicie symetrycznej, nie istnieją efektywne metody radzenia sobie z takimi problemami i jesteśmy zmuszeni ograniczyć się wtedy do badania układów bardzo niewielu cząstek.

Istnieje kilka sposobów, na które skorelowany szum manifestuje swoją obecność w problemach metrologicznych. Pierwszym jest oddziaływanie układu z otoczeniem, które posiada pamięć, co efektywnie prowadzi do występowania skorelowanego w czasie szumu. Najprostszym przykładem takiego problemu jest stabilizacja działania zegara atomowego, gdzie efektywny proces defazowania atomów jest skorelowany w czasie w wyniku korelacji fluktuacji częstości lokalnego oscylatora [Ludlow i in. 2015]. Z tego powodu, znalezienie optymalnej strategii stabilizacji zegara atomowego, która wzięłaby pod uwagę wszystkie możliwe stany, w których można przygotować atomy, jak również możliwe do użycia pomiary kwantowe i strategie poprawy częstości lokalnego oscylatora w kolejnych cyklach działania zegara, jest wysoce nie trywialnym zadaniem [André, Sørensen i Lukin 2004; Borregaard i Sørensen 2013; Kessler i in. 2014]. W celu sprostania temu wyzwaniu zaproponowano kwantową wariancję Allana [Chabuda, Leroux i Demkowicz-Dobrzański 2016], jednakże okazało się, że jej obliczenie nawet dla bardzo małych układów staje się szybko niewykonalne, ponieważ złożoność obliczeniowa rośnie eksponencjalnie z liczbą cykli działania zegara atomowego. Jeszcze trudniejszym przypadkiem są modele, w których czasowe korelacje nie mogą być efektywnie opisane za pomocą klasycznego procesu stochastycznego [Clerk i in. 2010; Szańkowski i in. 2017]. W związku z tym ich dynamika kwantowa jest nie-Markowowska [Chin, Huelga i Plenio 2012], jak np. w detektorach bazujących na centrach NV [Rondin i in. 2014; Paz-Silva, Norris i Viola 2017; Kwiatkowski i Cywiński 2018]. Drugi naturalny scenariusz, w którym pojawiają się nietrywialne korelacje szumu to układ wielu ciał, np. łańcuch spinów. Tutaj typowo pojawiają się przestrzenne korelacje szumu. Mają one znaczenie dla każdego modelu, w którym efektywny sygnał jest otrzymywany w wyniku oddziaływania pola z przestrzennie rozdzielonymi cząstkami [Jeske, Cole i Huelga 2014], np. w estymacji gradientu pola [Altenburg i in. 2017; Apellaniz i in. 2018].

Chociaż jest mało prawdopodobne, aby realistyczny szum był ściśle nieskorelowany, to z drugiej strony nie należy oczekiwać, że będzie on wykazywał bardzo złożony charakter korelacji. Wręcz przeciwnie, zazwyczaj czasowe korelacje szumu zanikają bardzo szybko. Podobnie w przypadku przestrzennych korelacji, między innymi z powodów energetycznych można spodziewać się, że korelacje będą jedynie krótkozasięgowe. Dla typowego krótkozasięgowego szumu

można oczekiwać, że optymalne strategie metrologiczne będą bazować na stanach początkowych, które wykazują splątanie jedynie wewnątrz małych grup cząstek, podobnie jak to ma miejsce dla nieskorelowanych modeli szumu [Jarzyna i Demkowicz-Dobrzański 2013]. To sugeruje, że aby wyjść poza reżim działania dotychczasowych metod, czyli modeli Markowskich i z nieskorelowanym szumem, należy znaleźć efektywną metodę reprezentowania wielociałowych stanów kwantowych o krótkozasięgowych korelacjach.

Okazuje się, że taka reprezentacja jest już dobrze znana – są nią sieci tensorowe. Stanowią one naturalny opis układów, w których splątanie/korelacje ma strukturę. Opis ten jest również bardzo efektywny, gdy korelacje są krótkozasięgowe. Za pierwszy kamień milowy rozwoju sieci tensorowych uważa się wprowadzenie grupy renormalizacji macierzy gęstości (DMRG, ang. *Density Matrix Renormalization Group*) przez White’a [1992; 1993]. Jest to technika znajdowania niskoenergetycznych stanów własnych jednowymiarowych Hamiltonianów układów wielociałowych w ramach procedury renormalizacji, w której przy zwiększaniu układu pozostawiamy tylko istotne stopnie swobody. White zrozumiał, że istotne są te stopnie swobody, które są powiązane ze splątaniem wielociałowej funkcji falowej. Później zostało udowodnione, że stany które otrzymuje się w procedurze DMRG mają skończenie zasięgowe korelacje [Fannes, Nachtergaele i Werner 1992] i są to obecnie najprostsze znane nam sieci tensorowe nazywane macierzowymi stanami produktowymi (MPS, ang. *Matrix Product State*). Ponadto pokazano, że DMRG jest po prostu metodą wariacyjnej optymalizacji po stanach w reprezentacji MPS [Östlund i Rommer 1995]. DMRG osiągnęło ogromny sukces i stało się podstawową metodą do badania niskoenergetycznych, jednowymiarowych sieci kwantowych. Podobnie jak w metrologii kwantowej renesans MPS przeżyły na początku XXI wieku, na skutek inspiracji koncepcjami znanymi z kwantowej teorii informacji. Zaczęto wtedy badać splątanie w MPS oraz jego zachowanie w kwantowych przejściach fazowych. Wyznaczenie entropii splątania dla podukładu jednowymiarowego systemu w stanie krytycznym pomogło zrozumieć, że splątanie ma strukturę [Vidal i in. 2003; Calabrese i Cardy 2004]. Ponadto odkryto, że MPS idealnie nadają się do opisu jednowymiarowych sieci z lokalnymi oddziaływaniami i przerwą energetyczną właśnie ze względu na lokalną strukturę splątania, którą opisują [Hastings 2006; Hastings 2007]. Późniejsze prace uogólniły koncepcję MPS wprowadzając sieci tensorowe idealne do opisu układów dwuwymiarowych – zrzucone pary stanów splątanych (PEPS, ang. *Projected Entangled Pair States*) [Verstraete i J. I. Cirac 2004], czy układów w stanie krytycznym – wieloskalowy splątany renormalizacyjny ansatz (MERA, ang. *Multiscale Entanglement Renormalization Ansatz*) [Vidal 2007b; Vidal 2010].

Zastosowanie metod sieci tensorowych do problemów metrologii kwantowej jest ciągle niewielkie. Na chwilę obecną dokonano jedynie obliczenia – jednej z podstawowych wielkości w częstotliwościowym podejściu do metrologii kwantowej – kwantowej informacji Fishera (QFI, ang. *Quantum Fisher Information*) na stanie czystym w reprezentacji MPS [Jarzyna i Demkowicz-Dobrzański 2013]¹.

¹Rozważany problem wymagał dużo trudniejszego obliczenia QFI na stanie mieszanym, ale wprowadzając ograniczenie górne na QFI, przybliżono go sumą ważoną QFI na stanach czystych.

Dodatkowym problemem jest, że wielkości takie jak QFI czy koszt Bayesowski nie są rozważane w literaturze dotyczącej sieci tensorowych i do tej pory nie było wiadome czy mogą być efektywnie obliczone dla stanów mieszanych reprezentowanych przez sieci tensorowe.

Celem tej dysertacji jest kompleksowa adaptacja metod sieci tensorowych do zagadnień metrologicznych w obecności dekoherencji, w szczególności tych, w których występuje krótkozasięgowy szum. Pokażemy jak bezpośrednio wyznaczyć istotne z punktu widzenia metrologii kwantowej wielkości (a nie jedynie ograniczenia na nie), takie jak kwantowa informacja Fishera czy koszt Bayesowski, używając od początku do końca formalizmu sieci tensorowych. Ponadto pokażemy jak dokonać optymalizacji i wyznaczyć optymalne stany początkowe w problemach metrologicznych, używając tego formalizmu. W rezultacie przedstawimy kompletną procedurę iteracyjną pozwalającą w ramach sieci tensorowych wyznaczyć fundamentalne ograniczenie na precyzję estymacji i dokonamy tego zarówno dla skończonych układów, jak i w granicy asymptotycznej, gdy liczba cząstek w układzie dąży do nieskończoności.

Struktura pracy jest następująca. W rozdziale 1 przedstawimy sformułowanie matematyczne teorii estymacji punktowej zarówno klasycznej, jak i kwantowej. Następnie przejdziemy do omówienia fundamentalnych ograniczeń na precyzję estymacji. Skupimy się przy tym na jednym z podstawowych problemów metrologii kwantowej, czyli zagadnieniu estymacji parametru unitarnego. Pozwoli nam to przybliżyć sukcesy metrologii kwantowej, jak i problemy z którymi się zmagają oraz omówić wszystkie potrzebne pojęcia na prostych przykładach. W rozdziale 2 omówimy sieci tensorowe, starając się przedstawić jak najwięcej koncepcji za pomocą prostego języka diagramów. Skupimy się na powiązaniu sieci tensorowych ze strukturą splątania, którą one efektywnie opisują. Następnie przejdziemy do szczegółowego omówienia jednowymiarowych sieci tensorowych, które będą kluczowe dla tej pracy. Rozdział 3 jest centralny dla tej dysertacji. W nim wprowadzimy formalnie klasę modeli metrologicznych, które będziemy chcieli analizować przy pomocy sieci tensorowych. Następnie pokażemy jak ściśle zapisać ewolucję kwantową w ramach tych modeli za pomocą sieci tensorowych i przejdziemy do opisu algorytmu, który pozwoli nam wyznaczyć fundamentalne ograniczenie na precyzję (oraz optymalne stany początkowe) zarówno dla układów o skończonym rozmiarze, jak i w granicy asymptotycznej. W rozdziale 4 zaprezentujemy trzy przykłady zastosowania naszego formalizmu. Pierwszy z nich dotyczy estymacji indukcji pola magnetycznego w obecności przestrzennie skorelowanego szumu. W drugim analizujemy zagadnienie stabilizacji zegara atomowego ze skorelowanymi w czasie fluktuacjami częstości lokalnego oscylatora. W trzecim pokażemy, że zastosowania zaproponowanego przez nas formalizmu wykraczają poza dziedzinę metrologii kwantowej i wyznaczymy wierność dla wielociałowego stanu termicznego (wielkość która jest kluczowa przy badaniu kwantowych przejść fazowych). W rozdziale 5 dokonamy podsumowania i przedstawimy możliwe kolejne kierunki badań.

Główne rezultaty tej dysertacji zostały opublikowane w pracy K. Chabuda,

J. Dziarmaga, T. J. Osborne i R. Demkowicz-Dobrzański, „*Tensor-network approach for quantum metrology in many-body quantum systems*”, Nature Communications **11**, 250 [2020]. Wyjątek stanowią wyniki zaprezentowane w rozdziale 4.2, które zostały rozszerzone w efekcie przeprowadzenia nowych obliczeń numerycznych.

Rozdział 1

Metrologia kwantowa

Fundamentem, na którym opiera się metrologia kwantowa jest teoria estymacji. W rozdziale tym dokonamy przeglądu zagadnień klasycznej i kwantowej teorii estymacji skupiając się na zagadnieniu estymacji pojedynczego parametru. Rozważymy przy tym zarówno podejście oparte na informacji Fishera, jak i podejście Bayesowskie. Zakończymy ten rozdział omówieniem fundamentalnych ograniczeń na precyzję estymacji parametru unitarnego, zarówno w przypadku bez dekoherencji, jak i z.

1.1 Klasyczna teoria estymacji

Celem teorii estymacji jest jak najlepsze oszacowanie wartości parametru opisującego model statystyczny na podstawie zbioru danych opisanych tym modelem [Kay 1993; Lehmann i Casella 1998]. Możemy to zobrazować następująco, w wyniku N powtórzeń eksperymentu otrzymujemy zbiór wyników $\mathbf{x} = \{x_1, x_2, \dots, x_N\}$, który jest realizacją N niezależnych zmiennych losowych, opisanych tym samym rozkładem gęstości prawdopodobieństwa $p_\varphi(x)$, który zależy od parametru φ jaki chcemy poznać. Naszym celem jest skonstruować funkcję nazywaną estymatorem $\hat{\varphi}(\mathbf{x})$, która zwróci najlepsze możliwe oszacowanie parametru φ na podstawie dostępnych danych $\mathbf{x}$ (warto zwrócić uwagę, że estymator sam jest zmienną losową).

Wyróżnia się dwa możliwe podejścia do problemu poszukiwania optymalnego estymatora:

- podejście częstotliwościowe (będziemy je też nazywać podejściem opartym na informacji Fishera) – w którym uznaje się, że nieznany parametr jest deterministyczny i ma ściśle określoną wartość;
- podejście Bayesowskie – które traktuje nieznany parametr jako zmienną losową, opisaną rozkładem prawdopodobieństwa *a priori* $p(\varphi)$.

Wprowadzimy teraz dwie pożądane cechy estymatorów: zgodność i nieobciążoność. Estymator nazywamy zgodnym, jeżeli przy zwiększaniu liczby obserwacji ciąg odpowiednich estymatorów jest zbieżny do prawdziwej wartości parametru:

$$\lim_{N \rightarrow \infty} \hat{\varphi}(\mathbf{x}) = \varphi, \quad (1.1)$$

gdzie zbieżność rozumiemy w sensie prawdopodobieństwa¹. Obciążenie b estymatora definiujemy jako odchylenie jego wartości średniej od prawdziwej wartości:

$$b(\hat{\varphi}) = \langle \hat{\varphi} \rangle - \varphi, \quad (1.2)$$

a estymator nazywamy nieobciążonym, jeśli średnio zwraca on prawdziwą wartość parametru, dla każdej wartości tego parametru, czyli

$$b(\hat{\varphi}) = 0, \quad (1.3)$$

lub równoważnie

$$\langle \hat{\varphi} \rangle = \varphi. \quad (1.4)$$

Możemy też zdefiniować słabszą wersję nieobciążoności, mianowicie lokalną nieobciążoność estymatora w $\varphi = \varphi_0$:

$$\langle \hat{\varphi} \rangle_{\varphi=\varphi_0} = \varphi_0, \quad \left. \frac{\partial \langle \hat{\varphi} \rangle}{\partial \varphi} \right|_{\varphi=\varphi_0} = 1. \quad (1.5)$$

Jak już wspomnieliśmy estymator jest zmienną losową, zatem ma on rozkład prawdopodobieństwa. Dla estymatora zgodnego rozkład ten przy wystarczająco dużej liczbie pomiarów jest zlokalizowany dowolnie blisko prawdziwej wartości parametru, a rozkład dla estymatora zgodnego i nieobciążonego jest zawsze zlokalizowany wokół tej wartości. Intuicyjnie, precyzja estymatora jest związana z szerokością tego rozkładu, która jest odwrotnie proporcjonalna do jego wariancji.

1.1.1 Podejście oparte na informacji Fishera

Do oceny jakości estymatora będziemy używać średniego błędu kwadratowego (MSE, ang. *Mean Squared Error*)

$$\Delta^2 \hat{\varphi} = \langle (\hat{\varphi} - \varphi)^2 \rangle = \int d^N \mathbf{x} p_{\varphi}(\mathbf{x}) (\hat{\varphi}(\mathbf{x}) - \varphi)^2, \quad (1.6)$$

który informuje nas o średnim kwadracie odległości estymatora od prawdziwej wartości parametru. Warto zwrócić uwagę, że dla nieobciążonych estymatorów MSE pokrywa się z wariancją estymatora $\text{Var}(\hat{\varphi}) = \langle (\hat{\varphi} - \langle \hat{\varphi} \rangle)^2 \rangle$. Nieobciążony estymator, który minimalizuje MSE dla każdej wartości parametru, będziemy nazywać optymalnym (znany też pod akronimem MUVE, ang. *Minimum-Variance Unbiased Estimator*).

Niestety wykonanie bezpośredniej minimalizacji MSE często jest zadaniem bardzo trudnym. Może okazać się nawet, że nie istnieje pojedynczy nieobciążony estymator minimalizujący MSE dla wszystkich wartości parametru. Na szczęście z pomocą przychodzi nierówność Craméra-Rao (CRB, ang. *Cramér-Rao*

¹ Ciąg zmiennych losowych X_n zbiega w sensie prawdopodobieństwa do zmiennej losowej X jeżeli $\forall \varepsilon > 0 \forall \eta > 0 \exists N \forall n \geq N P(|X_n - X| > \varepsilon) < \eta$.

Bound), która dla regularnych rozkładów $p_\varphi(\mathbf{x})$ ² jest dolnym ograniczeniem na MSE (lokalnie) nieobciążonego estymatora [Kay 1993]:

$$\Delta^2 \hat{\varphi} \geq \frac{1}{F_{\text{cl}}[p_\varphi(\mathbf{x})]}, \quad (1.7)$$

gdzie F_{cl} nazywamy (klasyczną) informacją Fishera (FI, ang. *Fisher Information*)³

$$F_{\text{cl}}[p_\varphi(\mathbf{x})] = \left\langle \frac{\dot{p}_\varphi^2}{p_\varphi} \right\rangle = \int d^N \mathbf{x} \frac{1}{p_\varphi(\mathbf{x})} \left(\frac{\partial p_\varphi(\mathbf{x})}{\partial \varphi} \right)^2. \quad (1.8)$$

FI możemy również przedstawić w jednej z dwóch następujących postaci:

$$F_{\text{cl}}[p_\varphi(\mathbf{x})] = \left\langle \left(\frac{\partial}{\partial \varphi} \log p_\varphi \right)^2 \right\rangle = \left\langle \frac{\partial^2}{\partial \varphi^2} \log p_\varphi \right\rangle. \quad (1.9)$$

Dla pojedynczego pomiaru ($\mathbf{x} = \{x\}$) informacja Fishera $F_{\text{cl}}[p_\varphi(x)]$ mówi nam ile informacji o parametrze φ jest zawarte w rozkładzie $p_\varphi(x)$, innymi słowy, informuje nas jak dużo informacji o parametrze φ wnosi pomiar zmiennej losowej X .

Jak można by oczekiwać od miary informacji FI jest nieujemna i addytywna, czyli dla produktowych rozkładów prawdopodobieństwa jest sumą FI obliczonych na pojedynczych rozkładach. W szczególności jeśli wykonujemy N powtórzeń eksperymentu i $p_\varphi(\mathbf{x}) = \prod_{j=1}^N p_\varphi(x_j)$ to $F_{\text{cl}}[p_\varphi(\mathbf{x})] = N F_{\text{cl}}[p_\varphi(x)]$, co pokazuje że wariancja estymatora maleje jak $\frac{1}{N}$, gdy zwiększamy liczbę powtórzeń eksperymentu.

FI można łatwo wyznaczyć, jednak aby nierówność Craméra-Rao była naprawdę użyteczna to musi być ona wysycalna. Nieobciążony estymator, który wysyci CRB nazywamy efektywnym. W ogólności jedynie lokalnie możemy zawsze skonstruować estymator, który wysyci CRB. Dla wszystkich wartości parametru efektywne estymatory istnieją tylko dla tak zwanej eksponencjalnej rodziny rozkładów gęstości prawdopodobieństwa (należą do niej, np. rozkład Gaussa, Poissona, dwumianowy, chi kwadrat) [Kay 1993], które można zapisać w postaci

$$p_\varphi^{\text{exp}}(x) = h(x) \exp[\eta(\varphi)T(x) - A(\varphi)], \quad (1.10)$$

gdzie $h(x)$, $T(x)$, $\eta(\varphi)$ i $A(\varphi)$ są znanymi funkcjami. Dla takich rozkładów FI wynosi

$$F_{\text{cl}}[p_\varphi^{\text{exp}}(x)] = \dot{\eta}(\varphi) \frac{\partial}{\partial \varphi} \left(\frac{\dot{A}(\varphi)}{\dot{\eta}(\varphi)} \right), \quad (1.11)$$

a efektywny estymator to $\hat{\varphi}^{\text{exp}}(x) = T(x)$. Będzie on również nieobciążony, jeśli spełniony zostanie warunek $\dot{A}(\varphi)/\dot{\eta}(\varphi) = \varphi$, wtedy FI uprości się do $F_{\text{cl}}[p_\varphi^{\text{exp}}(x)] = \dot{\eta}(\varphi)$. Na szczęście problem z brakiem globalnego, nieobciążonego,

²Warunek regularności oznacza, że (i) $\partial p_\varphi(\mathbf{x})/\partial \varphi$ istnieje i jest skończona oraz (ii) możemy zamienić kolejność $\int d^N \mathbf{x}$ oraz $\partial/\partial \varphi$ w każdym wyrażeniu, które jest uśredniane po wynikach pomiarów $\mathbf{x}$.

³Pochodną po estymowanym parametrze będziemy często oznaczać za pomocą „kropki” nad symbolem wielkości różniczkowanej.

efektywnego estymatora występuje tylko, gdy rozważamy skończoną liczbę powtórzeń eksperymentu N . Asymptotycznie, gdy liczba powtórzeń eksperymentu $N \rightarrow \infty$, istnieje nieobciążony estymator, który dla każdej wartości parametru wysyca CRB. Nazywa się on estymatorem największej wiarygodności (ML, ang. *Maximum Likelihood*) [Kay 1993; Lehmann i Casella 1998; Vaart 1998]. Estymator ML zdefiniowany jako

$$\hat{\varphi}^{\text{ML}}(\mathbf{x}) = \arg \max_{\varphi} p_{\varphi}(\mathbf{x}) = \arg \max_{\varphi} \log p_{\varphi}(\mathbf{x}), \quad (1.12)$$

jest funkcją, która zbiorowi wyników $\mathbf{x}$ przyporządkowuje wartość parametru φ , dla którego te wyniki są najbardziej prawdopodobne. Dla skończonych N estymator ML jest obciążony, ale asymptotycznie staje się on nieobciążony $\lim_{N \rightarrow \infty} \langle \hat{\varphi}^{\text{ML}} \rangle = \varphi$ i wysyca CRB.

1.1.2 Podejście Bayesowskie

W podejściu Bayesowskim zakładamy, że estymowany parametr φ jest zmienną losową opisaną rozkładem *a priori* $p(\varphi)$. Rozkład ten reprezentuje naszą wiedzę o parametrze przed rozpoczęciem procedury estymacji. Żeby lepiej oddać fakt, że φ traktujemy teraz jako zmienną losową, rozkład wyników pomiarów będziemy teraz oznaczać $p(\mathbf{x}|\varphi)$ zamiast $p_{\varphi}(\mathbf{x})$. Optymalny estymator w tym podejściu musi brać pod uwagę to, jakie wartości parametru są bardziej prawdopodobne. Z tego powodu naturalnym kandydatem do oceny jakości estymatora jest uśredniony po rozkładzie *a priori* MSE (AMSE, ang. *Average Mean Squared Error*)

$$\overline{\Delta^2 \hat{\varphi}} = \int d\varphi p(\varphi) \Delta^2 \hat{\varphi} = \iint d\varphi d^N \mathbf{x} p(\varphi) p(\mathbf{x}|\varphi) (\hat{\varphi}(\mathbf{x}) - \varphi)^2. \quad (1.13)$$

W podejściu Bayesowskim optymalny estymator możemy bezpośrednio wyznaczyć korzystając z twierdzenia Bayesa $p(\mathbf{x}|\varphi)p(\varphi) = p(\varphi|\mathbf{x})p(\mathbf{x})$ i jawnie minimalizując AMSE. Nie trzeba przy tym nakładać żadnych więzów, takich jak nieobciążoność, na estymator. Używając twierdzenia Bayesa możemy zapisać AMSE w postaci

$$\overline{\Delta^2 \hat{\varphi}} = \int d^N \mathbf{x} p(\mathbf{x}) \left[\int d\varphi p(\varphi|\mathbf{x}) (\hat{\varphi}(\mathbf{x}) - \varphi)^2 \right], \quad (1.14)$$

na skutek czego widzimy, że optymalny estymator to taki, który minimalizuje wyrażenie w nawiasie kwadratowym dla każdej wartości $\mathbf{x}$. Dla ustalonej wartości $\mathbf{x}$ jest to kwadratowe wyrażenie w $\hat{\varphi}(\mathbf{x})$, więc korzystając z warunku

$$\frac{\partial}{\partial \hat{\varphi}(\mathbf{x})} \int d\varphi p(\varphi|\mathbf{x}) (\hat{\varphi}(\mathbf{x}) - \varphi)^2 = 0, \quad (1.15)$$

otrzymujemy od razu wyrażenie na estymator o minimalnym średnim błędzie kwadratowym (MMSE, ang. *Minimum Mean Squared Error*)

$$\hat{\varphi}^{\text{MMSE}}(\mathbf{x}) = \int d\varphi p(\varphi|\mathbf{x}) \varphi = \langle \varphi \rangle_{p(\varphi|\mathbf{x})}. \quad (1.16)$$

Widzimy, że optymalny estymator Bayesowski odpowiada średniej parametru po rozkładzie *a posteriori* $p(\varphi|\mathbf{x})$. Rozkład *a posteriori* reprezentuje naszą wiedzę o parametrze po uwzględnieniu zaobserwowanych danych. Możemy go zawsze wyznaczyć korzystając z twierdzenia Bayesa

$$p(\varphi|\mathbf{x}) = \frac{p(\mathbf{x}|\varphi)p(\varphi)}{p(\mathbf{x})} = \frac{p(\mathbf{x}|\varphi)p(\varphi)}{\int d\varphi p(\mathbf{x}|\varphi)p(\varphi)}. \quad (1.17)$$

Dla tak otrzymanego optymalnego estymatora AMSE wynosi

$$\begin{aligned} \overline{\Delta^2 \hat{\varphi}^{\text{MMSE}}} &= \int d^N \mathbf{x} p(\mathbf{x}) \left[\int d\varphi p(\varphi|\mathbf{x}) (\varphi - \langle \varphi \rangle_{p(\varphi|\mathbf{x})})^2 \right] = \\ &= \int d^N \mathbf{x} p(\mathbf{x}) \text{Var}(\varphi) \Big|_{p(\varphi|\mathbf{x})}, \end{aligned} \quad (1.18)$$

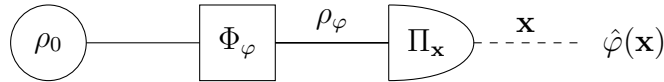
czyli jest to wariancja parametru obliczona względem rozkładu *a posteriori* i następnie uśredniona po możliwych wynikach pomiarów.

Optymalna strategia estymacji w podejściu Bayesowskim zależy silnie od rozkładu *a priori*. Jeśli rozkład *a priori* $p(\varphi)$ jest dużo bardziej czuły na zmiany parametru φ niż rozkład wyników pomiarów $p(\mathbf{x}|\varphi)$, to postać rozkładu *a posteriori* $p(\varphi|\mathbf{x})$ jest zdominowana przez rozkład *a priori* $p(\varphi|\mathbf{x}) \approx p(\mathbf{x}|\varphi)p(\varphi) \approx p(\varphi)$, co oznacza że uzyskane wyniki pomiarów mają mały wpływ na proces estymacji. To pokazuje jak ważny jest wybór odpowiedniego rozkładu *a priori*, tak aby oddawał jak najlepiej wiedzę o parametrze przed rozpoczęciem pomiarów, ale też nie zdominował uzyskiwanych wyników.

Warto zwrócić uwagę, że w podejściu Bayesowskim, w przeciwieństwie do podejścia opartego o FI, wyznaczamy bezpośrednio minimalną wartość AMSE, a więc nie musimy przejmować się zagadnieniem wysycalności otrzymanych wyników.

1.2 Kwantowa teoria estymacji

Typowy problem kwantowej teorii estymacji można przedstawić schematycznie jako [Giovannetti, Lloyd i Maccone 2004; Giovannetti, Lloyd i Maccone 2006]



Początkowy stan układu kwantowego reprezentowany jest przez macierz gęstości ρ_0 . Ewoluuje on przez zależny od parametru φ kanał kwantowy Φ_φ – czyli odwzorowanie całkowicie dodatnio zachowujące ślad (CPTP, ang. *Completely Positive Trace Preserving*) [Nielsen i Chuang 2000]. W wyniku takiej ewolucji, informacja o parametrze zostaje zapisana w stanie układu – który na wyjściu kanału reprezentowany jest przez macierz gęstości $\rho_\varphi = \Phi_\varphi[\rho_0]$. Następnie dokonujemy pomiaru. Może to być zwykły pomiar rzutowy von Neumanna, $\Pi_{\mathbf{x}}\Pi_{\mathbf{x}'} = \Pi_{\mathbf{x}}\delta_{\mathbf{x},\mathbf{x}'}$, lub pomiar uogólniony reprezentowany przez dodatnio określoną miarę operatorową (POVM, ang. *Positive Operator Valued Measure*), która musi spełniać jedynie warunek $\Pi_{\mathbf{x}} \geq 0$, $\int d^N \mathbf{x} \Pi_{\mathbf{x}} = \mathbb{1}$ [Nielsen i Chuang

2000; Bengtsson i Życzkowski 2006]. W wyniku pomiaru uzyskujemy zbiór wyników $\mathbf{x}$ z prawdopodobieństwem $p_\varphi(\mathbf{x}) = \text{Tr}(\rho_\varphi \Pi_{\mathbf{x}})$. Na końcu konstruujemy estymator $\hat{\varphi}(\mathbf{x})$, który pozwoli nam oszacować wartość parametru φ .

Naszym celem jest zminimalizować MSE lub AMSE po możliwych stanach początkowych ρ_0 , pomiarach $\Pi_{\mathbf{x}}$ i estymatorach $\hat{\varphi}(\mathbf{x})$. Można to zadanie wyraźnie podzielić na część kwantową i klasyczną. Część kwantowa obejmuje poszukiwanie optymalnych stanów i pomiarów, gdzie możemy wykorzystać takie cechy mechaniki kwantowej, jak splątanie czy pomiary kolektywne. Po ustaleniu stanu początkowego i pomiaru otrzymujemy model, który określa rozkład wyników $p_\varphi(\mathbf{x})$ i jego kwantowe pochodzenie nie ma już znaczenia. Otrzymujemy więc klasyczny problem estymacji i w celu znalezienia optymalnego estymatora możemy wykorzystać narzędzia omówione w rozdziale 1.1.

1.2.1 Podejście oparte na kwantowej informacji Fishera

W częstotliwościowym podejściu do estymacji nierówność Craméra-Rao daje nam dolne ograniczenie na MSE zoptymalizowane po wszystkich możliwych (lokalnie) nieobciążonych estymatorach. Jej uogólnieniem, uwzględniającym optymalizację po wszystkich możliwych pomiarach uogólnionych, jest kwantowa nierówność Craméra-Rao (QCRB, ang. *Quantum Cramér-Rao Bound*) [Helstrom 1976; Holevo 1982; Braunstein i Caves 1994]

$$\Delta^2 \hat{\varphi} \geq \frac{1}{F[\rho_\varphi]}, \quad (1.19)$$

gdzie $F[\rho_\varphi]$ nazywamy kwantową informacją Fishera (QFI, ang. *Quantum Fisher Information*), zdefiniowaną jako

$$F[\rho_\varphi] = \text{Tr}(\rho_\varphi \Lambda^2), \quad (1.20)$$

a Hermitowski operator Λ oznacza symetryczną pochodną logarytmiczną (SLD, ang. *Symmetric Logarithmic Derivative*), która zdefiniowana jest *implicite* przez równanie

$$\dot{\rho}_\varphi = \frac{1}{2}(\rho_\varphi \Lambda + \Lambda \rho_\varphi), \quad (1.21)$$

w którym $\dot{\rho}_\varphi = \partial \rho_\varphi / \partial \varphi$. Co warto podkreślić QFI jest całkowicie zdeterminowane przez zależność ρ_φ od estymowanego parametru φ .

Dla stanu czystego $\rho_\varphi = |\psi_\varphi\rangle\langle\psi_\varphi|$ wzór na QFI jest bardzo prosty⁴:

$$F[|\psi_\varphi\rangle] = 4 \left(\langle \dot{\psi}_\varphi | \dot{\psi}_\varphi \rangle - \left| \langle \dot{\psi}_\varphi | \psi_\varphi \rangle \right|^2 \right). \quad (1.22)$$

W ogólności jednak, dla stanu mieszanego, aby obliczyć QFI trzeba najpierw wyznaczyć SLD z równania (1.21). Wymaga to dokonania rozkładu własnego macierzy gęstości ρ_φ , czego zazwyczaj nie można wykonać analitycznie, a numerycznie jest to możliwe tylko dla bardzo małych układów. Formalnie SLD i QFI

⁴Dla stanów czystych będziemy pisać $F[|\psi\rangle]$ zamiast $F[|\psi\rangle\langle\psi|]$.

w bazie wektorów własnych $\rho_\varphi = \sum_n p_n(\varphi) |n(\varphi)\rangle\langle n(\varphi)|$ przybierają postać:

$$\Lambda = \sum_{\substack{n,m \\ p_n+p_m \neq 0}} \frac{2 \langle n(\varphi) | \dot{\rho}_\varphi | m(\varphi) \rangle}{p_n(\varphi) + p_m(\varphi)} |n(\varphi)\rangle\langle m(\varphi)|, \quad (1.23)$$

$$F[\rho_\varphi] = \sum_{\substack{n \\ p_n \neq 0}} \frac{\dot{p}_n(\varphi)^2}{p_n(\varphi)} + \sum_{\substack{n,m \\ p_n+p_m \neq 0}} \frac{2[p_n(\varphi) - p_m(\varphi)]^2}{p_n(\varphi) + p_m(\varphi)} |\langle n(\varphi) | m(\varphi) \rangle|^2. \quad (1.24)$$

Widzimy, że QFI można podzielić na dwie części: klasyczną i kwantową. Pierwsza suma przedstawia klasyczną FI (1.8) obliczoną dla rozkładu wartości własnych $p_n(\varphi)$. Natomiast druga suma opisuje zawsze nieujemny wkład kwantowy związany z obrotem wektorów własnych wraz ze zmianą parametru.

Podobnie jak w przypadku klasycznym, QFI jest nieujemna i addytywna – jeśli mamy do czynienia ze stanem produktowym $\rho_\varphi^{AB} = \rho_\varphi^A \otimes \rho_\varphi^B$ to $F[\rho_\varphi^{AB}] = F[\rho_\varphi^A] + F[\rho_\varphi^B]$. W szczególności dla stanu produktowego N identycznych cząstek $F[\rho_\varphi^{\otimes N}] = NF[\rho_\varphi]$. Oznacza to, że z punktu widzenia QCRB wykonanie N powtórzeń eksperymentu na pojedynczej cząstce i estymacji parametru ze stanu ρ_φ jest równoważne wykonaniu pojedynczego eksperymentu, w którym estymujemy parametr ze stanu $\rho_\varphi^{\otimes N}$. Jeśli w pierwszym przypadku mierzony stan otrzymujemy w wyniku ewolucji stanu ρ_0 przez kanał kwantowy Φ_φ , to w drugim przypadku mielibyśmy do czynienia z ewolucją stanu produktowego $\rho_0^{\otimes N}$ przez kanał kwantowy, który działa niezależnie i identycznie na każdą cząstkę, tak że na wyjściu otrzymujemy stan $\rho_\varphi^{\otimes N} = \Phi_\varphi[\rho_0]^{\otimes N}$. Warto podkreślić, że równoważność tych dwóch przypadków pokazuje, że jeśli mierzony stan jest produktowy $\rho_\varphi^{\otimes N}$, to zastosowanie pomiarów kolektywnych nie jest w stanie zwiększyć precyzji estymacji. Z drugiej strony oznacza to, że aby otrzymać skalowanie QFI lepsze niż liniowe z N , mierzony stan N cząstek musi być splątany [Pezze i Smerzi 2009]. Pozwala to wykorzystać QFI do wykrywania splątania stanów kwantowych, a więc jest ona świadkiem splątania [Bengtsson i Życzkowski 2006; Horodecki i in. 2009].

Jako że QCRB daje nam ograniczenie na MSE zoptymalizowane po pomiarach i estymatorach, kwestię wysycalności należy też rozpatrywać w tych dwóch kategoriach. Optymalny pomiar wysycający QCRB to pomiar rzutowy na wektory własne SLD [Braunstein i Caves 1994; Nagaoka 2005]. Niestety taki pomiar jest zazwyczaj kolektywny i może zależeć od estymowanego parametru, przez co jest trudny w realizacji eksperymentalnej. Dla ustalonego pomiaru problem redukuje się do klasycznego poszukiwania optymalnego estymatora, który aby wysycać QCRB musi spełnić takie same warunki, jak w przypadku klasycznej wersji CRB (patrz rozdział 1.1.1).

Estymowany parametr φ może być zapisany w zadanej nam z góry macierzy gęstości ρ_φ , lub informacja o nim może zostać zapisana w macierzy gęstości ρ_0 w wyniku ewolucji przez kanał kwantowy Φ_φ . W tym drugim przypadku mamy swobodę wyboru stanu początkowego. Oznacza to, że aby otrzymać fundamentalne ograniczenie na precyzję musimy znaleźć maksimum $F[\Phi_\varphi[\rho_0]]$ po

wszystkich możliwych macierzach gęstości ρ_0 . Korzystając z własności wypukłości QFI

$$F\left[\sum_j p_j \rho_\varphi^{(j)}\right] \leq \sum_j p_j F[\rho_\varphi^{(j)}], \quad (1.25)$$

gdzie współczynniki p_j spełniają warunek $\sum_j p_j = 1$, $p_j \geq 0$, możemy pokazać, że maksymalizując QFI wystarczy ograniczyć się do optymalizacji po stanach czystych⁵:

$$\begin{aligned} F[\Phi_\varphi[\rho_0]] &= F\left[\sum_j p_j \Phi_\varphi[|\psi_j\rangle]\right] \leq \sum_j p_j F[\Phi_\varphi[|\psi_j\rangle]] \leq \\ &\leq \max_j F[\Phi_\varphi[|\psi_j\rangle]] \leq \max_{|\psi_0\rangle} F[\Phi_\varphi[|\psi_0\rangle]]. \end{aligned} \quad (1.26)$$

QFI możemy też alternatywnie zdefiniować jako następujące zagadnienie maksymalizacji:

$$F[\rho_\varphi] = \max_{\substack{\Lambda \\ \Lambda^\dagger = \Lambda}} F[\rho_\varphi, \Lambda], \quad F[\rho_\varphi, \Lambda] = 2 \operatorname{Tr}(\dot{\rho}_\varphi \Lambda) - \operatorname{Tr}(\rho_\varphi \Lambda^2). \quad (1.27)$$

Formalnie biorąc pochodną wyrażenia po prawej (1.27) względem Λ i przyrównując ją do zera otrzymamy warunek na SLD (1.21), po spełnieniu którego dostaniemy pierwotną postać QFI (1.20). Przewaga definicji (1.27) objawia się, gdy chcemy zoptymalizować QFI po stanach początkowych. Takie zadanie możemy zapisać jako podwójne zagadnienie maksymalizacji:

$$F = \max_{\substack{|\psi_0\rangle \\ \|\psi_0\|^2=1}} F[\Phi_\varphi[|\psi_0\rangle]] = \max_{\substack{|\psi_0\rangle \\ \|\psi_0\|^2=1}} \max_{\substack{\Lambda \\ \Lambda^\dagger = \Lambda}} F[\Phi_\varphi[|\psi_0\rangle], \Lambda]. \quad (1.28)$$

To sformułowanie prowadzi do bardzo efektywnego numerycznego algorytmu iteracyjnego pozwalającego wyznaczyć optymalny stan początkowy [Macieszczak 2013; Macieszczak, Fraas i Demkowicz-Dobrzański 2014]: zaczynamy od losowej (lub takiej, co do której mamy przesłanki, że może być odpowiednia w naszym problemie) funkcji falowej $|\psi_0^{(1)}\rangle$ i dla niej szukamy operatora Hermitowskiego $\Lambda^{(1)}$, który maksymalizowałby $F[\Phi_\varphi[|\psi_0^{(1)}\rangle], \cdot]$. Następnie odwracamy problem, dla tak znalezionej operatora $\Lambda^{(1)}$ szukamy funkcji falowej $|\psi_0^{(2)}\rangle$, która maksymalizowałaby $F[\Phi_\varphi[\cdot], \Lambda^{(1)}]$ i powtarzamy procedurę. Doświadczenia numeryczne pokazują, że w interesujących problemach fizycznych algorytm zbiega bardzo szybko i zwraca optymalny stan początkowy oraz odpowiadającą mu wartość QFI.

Precyzja estymacji będzie tym większa im bardziej czuły będzie stan ρ_φ na zmianę parametru. O tym właśnie jak bardzo zmienia się stan przy zmianie φ informuje nas QFI. Manifestuje się to również w tym, że QFI można użyć do zdefiniowania odległości d_B pomiędzy stanami kwantowymi ρ i $\rho + d\rho$, zwanej

⁵Dla stanów czystych będziemy pisać $\Phi[|\psi\rangle]$ zamiast $\Phi[|\psi\rangle\langle\psi|]$.

odległością Buresa [Braunstein i Caves 1994; Bengtsson i Życzkowski 2006]:

$$d_B(\rho, \rho + d\rho)^2 = \frac{1}{4} \text{Tr}(\rho d\Lambda^2), \quad (1.29)$$

gdzie $d\Lambda$ jest zdefiniowana przez równanie

$$d\rho = \frac{1}{2}(\rho d\Lambda + d\Lambda \rho). \quad (1.30)$$

Tak otrzymaną metrykę w przestrzeni stanów kwantowych nazywamy metryką Buresa, a gdy ograniczymy się do stanów czystych to mówimy wtedy o metryce Fubiniego-Study.

W kwantowej teorii informacji miarą podobieństwa dwóch stanów kwantowych jest wierność $\mathcal{F}$ [Nielsen i Chuang 2000]. Dla stanów czystych $|\psi_1\rangle, |\psi_2\rangle$ jest ona zdefiniowana jako

$$\mathcal{F}(|\psi_1\rangle, |\psi_2\rangle) = |\langle\psi_1|\psi_2\rangle|. \quad (1.31)$$

Możemy tę definicję uogólnić na stany mieszane ρ_1, ρ_2 , następująco:

$$\mathcal{F}(\rho_1, \rho_2) = \max_{|\Psi_1\rangle} |\langle\Psi_1|\Psi_2\rangle|, \quad (1.32)$$

gdzie $|\Psi_j\rangle \in \mathcal{H} \otimes \mathcal{H}_A$ to puryfikacje stanów ρ_j , spełniające warunek $\text{Tr}_A |\Psi_j\rangle\langle\Psi_j| = \rho_j$, a A oznacza pomocniczą przestrzeń potrzebną do stworzenia puryfikacji [Nielsen i Chuang 2000]. Korzystając z twierdzenia Uhlmann'a [Uhlmann 1976] możemy wierność dla stanów mieszanych równoważnie zapisać w postaci

$$\mathcal{F}(\rho_1, \rho_2) = \text{Tr} |\sqrt{\rho_1} \sqrt{\rho_2}| = \text{Tr} \sqrt{\sqrt{\rho_2} \rho_1 \sqrt{\rho_2}}, \quad (1.33)$$

gdzie w drugiej równości skorzystaliśmy z tego, że $|A| = \sqrt{A^\dagger A}$. Jeśli dwa stany kwantowe są nieskończenie blisko siebie to wierność pomiędzy nimi możemy powiązać z QFI [Braunstein i Caves 1994]:

$$\mathcal{F}(\rho_\varphi, \rho_{\varphi+\varepsilon}) = 1 - \frac{1}{8} F[\rho_\varphi] \varepsilon^2 + \mathcal{O}(\varepsilon^3). \quad (1.34)$$

Wzór (1.32) pozwoli nam wprowadzić nowy sposób wyznaczania QFI dla stanów mieszanych – jeśli dla dwóch stanów ρ_φ i $\rho_{\varphi+\varepsilon}$ rozwiniemy go do drugiego rzędu w ε i porównamy z równaniem (1.34) to otrzymamy, że QFI można wyrazić jako [Escher, Matos Filho i Davidovich 2011]

$$F[\rho_\varphi] = \min_{|\Psi_\varphi\rangle} F[|\Psi_\varphi\rangle] = 4 \min_{|\Psi_\varphi\rangle} \left(\langle \dot{\Psi}_\varphi | \dot{\Psi}_\varphi \rangle - \left| \langle \dot{\Psi}_\varphi | \Psi_\varphi \rangle \right|^2 \right), \quad (1.35)$$

gdzie $|\Psi_\varphi\rangle$ to puryfikacja stanu ρ_φ . Niezależnie, w pracy [Fujiwara i Imai 2008] zaproponowano inną bazującą na puryfikacjach definicję QFI, która jednakże jest równoważna powyższej:

$$F[\rho_\varphi] = 4 \min_{|\Psi_\varphi\rangle} \langle \dot{\Psi}_\varphi | \dot{\Psi}_\varphi \rangle. \quad (1.36)$$

Taki sposób wyznaczania QFI dla stanów mieszanych wcale nie upraszcza tego

zadania. Jednak jest on użyteczny, ponieważ ograniczając się do pewnej klasy puryfikacji pozwala uzyskać ograniczenie górne na QFI. Szczególnie jeżeli jesteśmy w stanie zaproponować odpowiednią dla danego problemu klasę puryfikacji to można w ten sposób otrzymać bardzo użyteczne ograniczenia [Escher, Matos Filho i Davidovich 2011; Escher, Davidovich i in. 2012; Demkowicz-Dobrzański, Kołodyński i Guţă 2012].

1.2.2 Kwantowe podejście Bayesowskie

Tak samo jak klasycznie, w kwantowym podejściu Bayesowskim chcemy zminimalizować AMSE, tylko że teraz $p_\varphi(\mathbf{x}) = \text{Tr}(\rho_\varphi \Pi_{\mathbf{x}})$:

$$\overline{\Delta^2 \hat{\varphi}} = \iint d\varphi d^N \mathbf{x} p(\varphi) \text{Tr}(\rho_\varphi \Pi_{\mathbf{x}}) (\hat{\varphi}(\mathbf{x}) - \varphi)^2. \quad (1.37)$$

W tym celu należy znaleźć optymalne pomiary i estymatory. W związku z tym wygodnie będzie połączyć te dwa elementy i indeksować operatory pomiarowe estymowanymi wartościami $\Pi_{\hat{\varphi}} = \int d^N \mathbf{x} \Pi_{\mathbf{x}} \delta(\hat{\varphi} - \hat{\varphi}(\mathbf{x}))$, co sprawi, że gdy znajdziemy optymalne POVM to od razu poznamy też optymalne estymatory. Nasze zadanie sprowadza się teraz do minimalizacji $\overline{\Delta^2 \hat{\varphi}}$ jedynie po $\Pi_{\hat{\varphi}}$ ze standardowym warunkiem na POVM $\Pi_{\hat{\varphi}} \geq 0$, $\int d\hat{\varphi} \Pi_{\hat{\varphi}} = \mathbb{1}$.

Dla uproszczenia wzorów założymy, że wartość oczekiwana φ wynosi zero $\bar{\varphi} = \int d\varphi p(\varphi) \varphi = 0$ (nie prowadzi to do utraty ogólności, ponieważ zawsze możemy wykonać podstawienie $\varphi \rightarrow \varphi - \bar{\varphi}$). Możemy teraz przepisać AMSE w następujący sposób:

$$\begin{aligned} \overline{\Delta^2 \hat{\varphi}} &= \iint d\varphi d\hat{\varphi} p(\varphi) \text{Tr}(\rho_\varphi \Pi_{\hat{\varphi}}) (\hat{\varphi} - \varphi)^2 = \\ &= \int d\varphi p(\varphi) \varphi^2 - 2 \text{Tr} \left(\int d\varphi p(\varphi) \varphi \rho_\varphi \int d\hat{\varphi} \Pi_{\hat{\varphi}} \hat{\varphi} \right) + \text{Tr} \left(\int d\varphi p(\varphi) \rho_\varphi \int d\hat{\varphi} \Pi_{\hat{\varphi}} \hat{\varphi}^2 \right) = \\ &= \text{Var}(\varphi) - 2 \text{Tr}(\bar{\rho}' \bar{\Lambda}) + \text{Tr}(\bar{\rho} \bar{\Lambda}_2), \end{aligned} \quad (1.38)$$

gdzie $\text{Var}(\varphi) = \int d\varphi p(\varphi) \varphi^2$ jest wariancją estymowanego parametru względem rozkładu *a priori*, $\bar{\rho} = \int d\varphi p(\varphi) \rho_\varphi$ oznacza stan uśredniony po rozkładzie *a priori*, $\bar{\rho}' = \int d\varphi p(\varphi) \varphi \rho_\varphi$, $\bar{\Lambda} = \int d\hat{\varphi} \Pi_{\hat{\varphi}} \hat{\varphi}$, $\bar{\Lambda}_2 = \int d\hat{\varphi} \Pi_{\hat{\varphi}} \hat{\varphi}^2$. Optymalny pomiar w tym problemie jest pomiarem rzutowym [Helstrom 1976; Macieszczak, Fraas i Demkowicz-Dobrzański 2014], dla którego $\Pi_{\hat{\varphi}}^2 = \Pi_{\hat{\varphi}}$, co oznacza, że $\bar{\Lambda}_2 = \bar{\Lambda}^2$ i ostatecznie pozwala nam zapisać

$$\overline{\Delta^2 \hat{\varphi}} = \text{Var}(\varphi) - 2 \text{Tr}(\bar{\rho}' \bar{\Lambda}) + \text{Tr}(\bar{\rho} \bar{\Lambda}^2). \quad (1.39)$$

Minimalizacja AMSE po pomiarach oznacza teraz minimalizację po operatorach Hermitowskich $\bar{\Lambda}$. Ekstremum osiągamy, gdy $\bar{\Lambda}$ spełnia warunek

$$\bar{\rho}' = \frac{1}{2} (\bar{\rho} \bar{\Lambda} + \bar{\Lambda} \bar{\rho}), \quad (1.40)$$

wtedy AMSE przyjmuje swoją minimalną wartość

$$\overline{\Delta^2 \hat{\varphi}} = \text{Var}(\varphi) - \text{Tr}(\bar{\rho} \bar{\Lambda}^2). \quad (1.41)$$

Warto zwrócić uwagę na podobieństwo powyższych wyrażeń do tych na QFI (1.20) i SLD (1.21), co sugeruje, że narzędzia matematyczne stworzone do liczenia QFI mogą być również użyte w problemach Bayesowskich.

Zauważmy, że optymalne AMSE nie jest addytywne na stanach produktowych. W ogólności dla stanu produktowego $\rho_{\varphi}^{AB} = \rho_{\varphi}^A \otimes \rho_{\varphi}^B$ optymalizując AMSE po pomiarach otrzymamy POVM, który jest kolektywny i nie może być rozdzielony na niezależne pomiary na każdym z podukładów. W związku z tym optymalne AMSE dla stanu produktowego będzie mniejsze niż optymalne AMSE w przypadku, gdy w kolejnych powtórzeniach eksperymentu wykorzystalibyśmy kolejno stan ρ_{φ}^A oraz ρ_{φ}^B (nawet jeśli uwzględnimy możliwość zaktualizowania rozkładu *a priori* pomiędzy powtórzeniami eksperymentu). W związku z tym w podejściu Bayesowskim, w przeciwieństwie do podejścia opartego o QFI, optymalne rezultaty otrzymujemy wykorzystując wszystkie dostępne zasoby (cząstki) w pojedynczym eksperymencie.

Zwróćmy uwagę na jeszcze jedną cechę podejścia Bayesowskiego, która jest szczególnie istotna w przypadku kwantowym. W podejściu Bayesowskim wyznaczamy bezpośrednio minimalną wartość AMSE, a nie jedynie ograniczenie jak to ma miejsce w podejściu opartym o QFI. W związku z tym nie musimy rozważać zagadnienia wysycalności takiego ograniczenia, które szczególnie w przypadku kwantowym potrafi być zniuansowane.

Podsumowując teraz różnicę pomiędzy dwoma omawianymi podejściami do zagadnienia estymacji możemy powiedzieć, że w podejściu Bayesowskim poszukujemy optymalnej strategii, która średnio, pozwoli w jednym eksperymencie poznać jak najlepiej nieznany parametr bazując na rozkładzie *a priori* tego parametru. W podejściu opartym o informację Fishera, chcemy wykonując wiele powtórzeń eksperymentu, poznać jak najlepiej fluktuacje parametru wokół znanej wartości.

1.3 Fundamentalne ograniczenia w metrologii kwantowej

Jednym z podstawowych celów teoretycznej metrologii kwantowej jest wyznaczanie fundamentalnych ograniczeń na precyzję estymacji oraz opracowanie praktycznych strategii pomiarowych, które dadzą najlepsze możliwe rezultaty estymacji, a w idealnym scenariuszu wysycą otrzymane ograniczenia. Należy podkreślić, że gdy mówimy o fundamentalnych ograniczeniach na precyzję estymacji to mamy na myśli ograniczenia uzyskane w wyniku optymalizacji: stanu układu na początku eksperymentu, sposobu pomiaru interesującej nas wielkości fizycznej oraz metody analizy uzyskanych wyników (estymatora). W rozważaniach przyjmujemy model fizyczny za ustalony. Innymi słowy nie zajmujemy się tym jak zmniejszyć szum w badanym układzie, ale jak zaprojektować eksperyment (stan początkowy, pomiar, estymator), tak aby wpływ szumu na precyzję estymacji był jak najmniejszy.

Podstawowym problemem rozważanym w metrologii kwantowej jest zagadnienie estymacji pojedynczego parametru unitarnego φ – czyli takiego, który

został zapisany w stanie układu w wyniku unitarnej ewolucji (przykładem parametru nieunitarnego jest np. siła dekoherencji). Sztandarowym przykładem jest tutaj estymacja fazy, z którą mamy do czynienia w różnych pomiarach interferometrycznych [Demkowicz-Dobrzański, Jarzyna i Kołodyński 2015]. Problem taki możemy opisać w ramach standardowego schematu metrologicznego opisanego w rozdziale 1.2. Będziemy modelować oddziaływanie stanu początkowego ρ_0 , złożonego z N cząstek, z układem, który chcemy zmierzyć za pomocą kanału kwantowego Φ_φ . Zakładamy, że kanał ten jest złożeniem dwóch procesów: dekoherencji, którą opisuje kanał kwantowy Φ oraz unitarnej ewolucji U_φ , w wyniku której informacja o parametrze zostaje zapisana w stanie kwantowym. W rezultacie na wyjściu kanału otrzymujemy stan

$$\rho_\varphi = \Phi_\varphi[\rho_0] = U_\varphi \Phi[\rho_0] U_\varphi^\dagger. \quad (1.42)$$

Formalnie tak zapisana ewolucja odpowiada sytuacji, w której najpierw zachodzi dekoherencja, a później unitarne nakręcanie parametru. W typowych modelach dekoherencji, które właśnie będziemy rozważać, oba te procesy komutują ze sobą i kolejność w jakiej je zapiszemy nie ma znaczenia. Należy jednak zdawać sobie sprawę, że nie zawsze takie założenie jest uprawnione [Chaves i in. 2013]. Kolejnym założeniem naszego modelu jest to, że unitarne nakręcanie parametru $U_\varphi = e^{-iH\varphi}$ jest generowane przez liniowy Hamiltonian $H = \sum_{n=1}^N h^{[n]}$, gdzie $h^{[n]}$ jest lokalnym generatorem działającym na n -tą cząstkę. Taka sytuacja jest pożądana przy projektowaniu eksperymentu, ponieważ informacja o parametrze jest zapisywana w każdej cząstce niezależnie oraz w taki sam sposób. Sprzęgnięcie zapisywania informacji o parametrze z efektami wielociałowymi jest niepożądane, szczególnie ze względu na to, że są one dużo słabiej poznane. Jednak nie zawsze możemy je pominąć, szczególnie w przypadku interferometrii atomowej (która ma zastosowanie np. w zegarach atomowych czy magnetometrii). Wpływem takich efektów na możliwą do osiągnięcia precyzję estymacji zajmuje się nieliniowa metrologia kwantowa [Boixo i in. 2007; Beau i Campo 2017; Demkowicz-Dobrzański, Czaikowski i Sekatski 2017; Czaikowski, Pawłowski i Demkowicz-Dobrzański 2019]. Omówimy teraz fundamentalne ograniczenia na precyzję estymacji parametru unitarnego w dwóch jakościowo różnych przypadkach: bez oraz z dekoherencją.

1.3.1 Przypadek bez dekoherencji

Zacniemy od przypadku bez dekoherencji, który jest dużo prostszy i dla którego wiele rezultatów można uzyskać w ścisły sposób. Skupimy się na analizie tego problemu w podejściu opartym o QFI. Jako że QCRB jest zawsze lokalnie wysycalna, to podejście takie będzie szczególnie odpowiednie, w przypadku gdy chcemy wyznaczyć małe odchylenie parametru od znanej wartości, jak to mam miejsce np. w interferometrze fal grawitacyjnych [Abadie i in. 2011; Aasi i in. 2013].

Brak dekoherencji powoduje, że cała ewolucja stanu jest unitarna, co znacząco upraszcza wyznaczanie i optymalizację QFI. Dodatkowo, jeśli stan początkowy był czysty – a jak pokazuje rozumowanie (1.26), stan czysty jest stanem

optymalnym w podejściu opartym o QFI – to takim pozostanie na wyjściu kanału. To ma duże znaczenie, ponieważ dla czystych stanów końcowych znamy jawny wzór na QFI (1.22). W związku z tym założymy, że $\rho_0 = |\psi_0\rangle\langle\psi_0|$. Otrzymamy wtedy na wyjściu kanału stan $\rho_\varphi = |\psi_\varphi\rangle\langle\psi_\varphi|$, gdzie $|\psi_\varphi\rangle = e^{-iH\varphi}|\psi_0\rangle$. Korzystając teraz ze wzoru (1.22) możemy od razu wyznaczyć QFI:

$$F[|\psi_\varphi\rangle] = 4\left(\langle\psi_0|H^2|\psi_0\rangle - \langle\psi_0|H|\psi_0\rangle^2\right) = 4\text{Var}(H), \quad (1.43)$$

która jest proporcjonalna do wariancji generatora na stanie początkowym.

Wprowadźmy dla lokalnego generatora h oznaczenia: λ_+ i $|\lambda_+\rangle$ na największą wartość własną i odpowiadający jej wektor własny oraz λ_- i $|\lambda_-\rangle$ na najmniejszą wartość własną i odpowiadający jej wektor własny. Jeśli przygotujemy początkowo układ w stanie $|\psi_0\rangle = \frac{1}{\sqrt{2}}(|\lambda_+\rangle + |\lambda_-\rangle)^{\otimes N}$ to QFI wyniesie $F = N(\lambda_+ - \lambda_-)^2$. Korzystając teraz z QCRB możemy zapisać dla takiego stanu produktowego ograniczenie na precyzję estymacji parametru unitarnego, ma ono postać szumu śrutowego

$$\Delta^2\hat{\varphi} \geq \frac{1}{N(\lambda_+ - \lambda_-)^2} \quad (1.44)$$

i zwane jest standardową granicą kwantową (SQL, ang. *Standard Quantum Limit*). Ograniczenie to można wysycić mierząc różnicę obsadzeń pomiędzy stanami $|\lambda_+\rangle$ i $|\lambda_-\rangle$. Pokazany przykład jest manifestacją ogólnej prawidłowości, że dla stanów separowalnych precyzja estymacji fazy jest ograniczona przez SQL i aby ją przekroczyć musimy użyć stanów splątanych [Pezzé i Smerzi 2009; Giovannetti, Lloyd i Maccone 2011].

Stanem optymalnym, maksymalizującym QFI w przypadku estymacji parametru unitarnego bez dekoherencji jest splątany stan $|\psi_0\rangle = \frac{1}{\sqrt{2}}(|\lambda_+\rangle^{\otimes N} + |\lambda_-\rangle^{\otimes N})$, dla którego $F = N^2(\lambda_+ - \lambda_-)^2$. Oznacza to, że fundamentalne ograniczenie na precyzję estymacji parametru unitarnego wynosi

$$\Delta^2\hat{\varphi} \geq \frac{1}{N^2(\lambda_+ - \lambda_-)^2}, \quad (1.45)$$

a zwane jest granicą Heisenberga (HL, ang. *Heisenberg Limit*). Jak wiemy, ograniczenie to można wysycić dokonując pomiaru w bazie wektorów własnych SLD, jednak taki pomiar byłby wysoce splątany. Okazuje się, że dużo prostszy (przynajmniej w teorii) pomiar parzystości również wysyci to ograniczenie.

Przykładem estymacji parametru unitarnego jest estymacja fazy, z którą spotykamy się w dziedzinie optycznej np. w interferometrze Macha-Zehndera, a dla układów atomowych np. w interferometrii Ramseya. W takich problemach efektywnie generatorem ewolucji jest (bezwymiarowa) składowa z całkowitego momentu pędu $J_z = \frac{1}{2} \sum_{n=1}^N \sigma_z^{[n]}$, gdzie $\sigma_i^{[n]}$ oznacza odpowiednią macierz Pauliego działającą na n -tą cząstkę (dla takiego lokalnego generatora $\lambda_+ = \frac{1}{2}$, $\lambda_- = -\frac{1}{2}$) [Lee, Kok i Dowling 2002]. Układ w takim eksperymencie możemy efektywnie opisać jako N spinów- $\frac{1}{2}$ (N fotonów podróżujących jednym z dwóch ramion interferometru Macha-Zehndera lub N atomów dwupoziomowych w interferometrii Ramseya). Jako że przy braku dekoherencji możemy ograniczyć

się do podprzestrzeni stanów całkowicie symetrycznych to wygodnie (szczególnie dla fotonów) jest korzystać z bazy obsadzeniowej, w której $|n\rangle|N-n\rangle$ reprezentuje stan układu o n spinach „do góry” $|+\frac{1}{2}\rangle$ i $N-n$ spinach „do dołu” $|-\frac{1}{2}\rangle$. Dla początkowego stanu produktowego otrzymujemy $F = N$ i SQL $\Delta^2\hat{\varphi} \geq \frac{1}{N}$. Maksymalną wartość QFI – czyli największą wariancję J_z dla danego N – otrzymujemy dla splątanego stanu Greenbergerga-Hornego-Zeilingera (stan GHZ) [Greenberger, Horne i Zeilinger 1989], który po zapisaniu w bazie obsadzeniowej znany jest w optyce kwantowej jako stan NOON [J. J. . Bollinger i in. 1996]:

$$|\psi_0\rangle = \frac{1}{\sqrt{2}}\left(|+\frac{1}{2}\rangle^{\otimes N} + |-\frac{1}{2}\rangle^{\otimes N}\right) = \frac{1}{\sqrt{2}}(|N\rangle|0\rangle + |0\rangle|N\rangle). \quad (1.46)$$

Dla tego stanu $F = N^2$ i mamy do czynienia z HL: $\Delta^2\hat{\varphi} \geq \frac{1}{N^2}$. Należy zwrócić uwagę, że wraz ze wzrostem N dramatycznie wzrasta trudność eksperymentalnego przygotowania stanu NOON. Obecnie udało się to dla $N = 4$ [Nagata i in. 2007] i $N = 5$ [Afek, Ambar i Silberberg 2010]. Nawet jeśli uda się go przygotować to jest on bardzo wrażliwy na dekoherencję wraz ze wzrostem N [Escher, Matos Filho i Davidovich 2011; Demkowicz-Dobrzański, Kołodyński i Guță 2012; Huelga i in. 1997]. Sprawia to, że przydatność stanu NOON poza bardzo małymi wartościami N jest w praktyce znikoma, szczególnie jeśli weźmie się pod uwagę, że dużo łatwiejsze do eksperymentalnego przygotowania stany ściśnięte oferują podobne rezultaty [Pezzè i Smerzi 2008].

Przytoczmy teraz wyniki uzyskane w podejściu Bayesowskim, ograniczając się do estymacji konkretnie fazy i zakładając płaski rozkład *a priori* $p(\varphi) = \frac{1}{2\pi}$, czyli zakładając brak jakiejkolwiek informacji o rozkładzie estymowanego parametru [Hradil i in. 1996; Berry i Wiseman 2000]. Po pierwsze zauważmy, że standardowa funkcja kosztu kwadratowego $(\hat{\varphi} - \varphi)^2$ stosowana w MSE i AMSE nie uwzględnia geometrii badanego problemu, a mianowicie tego, że estymowany parametr należy do grupy $U(1)$ i w efekcie $\varphi = \varphi + 2\pi$. W podejściu opartym na QFI nie stanowiło to problemu, ponieważ jest ono lokalne, ale w przypadku analizy Bayesowskiej musimy wybrać inną funkcję kosztu. Najbardziej naturalnym wyborem jest funkcja kosztu $4\sin^2[(\hat{\varphi} - \varphi)/2]$, która w pierwszym rzędzie rozwinięcia redukuje się do standardowego kosztu kwadratowego. Minimalizacja takiej funkcji kosztu uśrednionej po wynikach pomiarów, można zostać wykonana wykorzystując własności symetrii problemu i tego, że możemy się ograniczyć do klasy pomiarów kowariantnych [Holevo 1982]. W efekcie otrzymujemy w granicy dużych N ograniczenie $\Delta^2\hat{\varphi} \geq \frac{1}{N}$ dla produktowych stanów początkowych oraz $\Delta^2\hat{\varphi} \geq \frac{\pi^2}{N^2}$ dla optymalnego (splątanego) stanu początkowego zwanego stanem SIN:

$$|\psi_0\rangle = \sum_{n=0}^N \sqrt{\frac{2}{2+N}} \sin\left(\frac{n+1}{N+2}\pi\right) |n\rangle|N-n\rangle, \quad (1.47)$$

który zapisaliśmy w bazie obsadzeniowej.

1.3.2 Przypadek z dekoherencją

Zajmiemy się teraz wyzwaniami jakie pojawiają się, gdy chcemy zbadać wpływ dekoherencji na precyzję estymacji parametru unitarnego. Dekoherencja jest

wynikiem niekontrolowanego oddziaływania układu z otoczeniem i sprawia, że stan układu staje się coraz bardziej mieszany. Z tego powodu nie można skorzystać ze wzoru (1.22) w celu wyznaczenia QFI i trzeba najpierw wyznaczyć SLD, co zazwyczaj wiąże się z koniecznością diagonalizacji ρ_φ . Poza przypadkiem globalnego defazowania, obecność dekoherencji wyprowadza z podprzestrzeni całkowicie symetrycznej, co sprawia, że obliczenia numeryczne muszą być przeprowadzone w pełnej przestrzeni Hilberta. Oznacza to z kolei, że przy wzroście rozmiaru układu bardzo szybko stają się one niemożliwe do praktycznego wykonania. Na szczęście nie oznacza to, że nie możemy nic powiedzieć o takich układach. Powstał szereg metod – takich jak np. minimalizacja po puryfikacjach kanału, klasyczna symulacja, kwantowa symulacja czy metoda rozszerzenia kanału – które umożliwiają wyznaczenie ograniczenia na QFI, które często przynajmniej asymptotycznie jest wysycalne [Escher, Matos Filho i Davidovich 2011; Demkowicz-Dobrzański, Kołodyński i Guţă 2012; Kołodyński i Demkowicz-Dobrzański 2013; Knysh, Chen i Durkin 2014; Demkowicz-Dobrzański i Maccone 2014; Sekatski i in. 2017; Demkowicz-Dobrzański, Czajkowski i Sekatski 2017; Zhou i in. 2018]. W praktyce jednak metody te są dość specyficzne i mogą być efektywnie zastosowane jedynie do wąskiego katalogu modeli metrologicznych. Spośród tych metod najbardziej wszechstronna jest metoda rozszerzenia kanału, która pozwala efektywnie badać nieskorelowane modele szumu i ją teraz omówimy.

Działanie dowolnego kanału kwantowego może być zapisane za pomocą reprezentacji Krausa [Nielsen i Chuang 2000]:

$$\rho_\varphi = \Phi_\varphi[\rho_0] = \sum_j K_\varphi^j \rho_0 K_\varphi^{j\dagger}, \quad (1.48)$$

gdzie K_φ^j to operatory Krausa spełniające warunek $\sum_j K_\varphi^{j\dagger} K_\varphi^j = \mathbb{1}$. Reprezentacja ta nie jest jednoznaczna, a równoważne zbiory operatorów Krausa są ze sobą powiązane za pomocą transformacji unitarnej

$$\tilde{K}_\varphi^i = \sum_j (\mathcal{U}_\varphi)_{i,j} K_\varphi^j, \quad (1.49)$$

gdzie $\mathcal{U}_\varphi$ to macierz unitarna zależna od parametru φ . Z każdym zestawem operatorów Krausa można powiązać operator U_φ^{SE} opisujący łączną ewolucję unitarną układu (S) oddziałującego z otoczeniem (E , które bez straty ogólności możemy przyjąć, że znajduje się początkowo w stanie czystym):

$$\rho_\varphi^{SE} = U_\varphi^{SE} (\rho_0^S \otimes |0\rangle_E \langle 0|) U_\varphi^{SE\dagger}, \quad (1.50)$$

$$\begin{aligned} \rho_\varphi^S &= \text{Tr}_E \rho_\varphi^{SE} = \sum_j \langle j| U_\varphi^{SE} (\rho_0^S \otimes |0\rangle_E \langle 0|) U_\varphi^{SE\dagger} |j\rangle_E = \\ &= \sum_j \langle j| U_\varphi^{SE} |0\rangle_E \rho_0^S \langle 0| U_\varphi^{SE\dagger} |j\rangle_E = \sum_j K_\varphi^j \rho_0^S K_\varphi^{j\dagger}, \end{aligned} \quad (1.51)$$

$$K_\varphi^j = \langle j| U_\varphi^{SE} |0\rangle_E. \quad (1.52)$$

Metoda rozszerzenia kanału pozwala wyznaczyć ograniczenie górne na QFI

zoptymalizowaną po stanach początkowych, w oparciu jedynie o dynamikę wyrażoną przez operatory Krausa [Fujiwara i Imai 2008; Demkowicz-Dobrzański, Kołodyński i Guţă 2012; Demkowicz-Dobrzański i Maccone 2014]. Opiera się ona na obserwacji, że dodatkowe trywialne działanie kanału na rozszerzonej przestrzeni może jedynie poprawić precyzję estymacji oraz wykorzystuje definicję QFI w postaci (1.36):

$$F = \max_{|\psi_0\rangle} F[\Phi_\varphi[|\psi_0\rangle]] \leq \max_{|\psi_0^{\text{ext}}\rangle} F[\Phi_\varphi \otimes \mathbb{1}[|\psi_0^{\text{ext}}\rangle]] = 4 \max_{|\psi_0^{\text{ext}}\rangle} \min_{|\Psi_\varphi\rangle} \langle \dot{\Psi}_\varphi | \dot{\Psi}_\varphi \rangle, \quad (1.53)$$

gdzie $|\psi_0^{\text{ext}}\rangle$ oznacza stan początkowy na rozszerzonej przestrzeni, a $|\Psi_\varphi\rangle$ to puryfikację stanu mieszanego $\Phi_\varphi \otimes \mathbb{1}[|\psi_0^{\text{ext}}\rangle]$. Oznaczmy przez S przestrzeń układu, E przestrzeń, na którą rozszerzyliśmy działanie kanału, A przestrzeń potrzebną do stworzenia puryfikacji. Oznacza to, że $|\psi_0\rangle \in \mathcal{H}_S$, $|\psi_0^{\text{ext}}\rangle \in \mathcal{H}_S \otimes \mathcal{H}_E$, $|\Psi_0\rangle \in \mathcal{H}_S \otimes \mathcal{H}_E \otimes \mathcal{H}_A$. Ewolucję puryfikacji możemy zapisać jako $|\Psi_\varphi\rangle = U_\varphi |\Psi_0\rangle = U_\varphi |\psi_0^{\text{ext}}\rangle |0\rangle_A$, gdzie $\langle j | U_\varphi | 0 \rangle_A = \tilde{K}_\varphi^j \otimes \mathbb{1}$, a $\{\tilde{K}_\varphi^j\}_j$ to jedna z równoważnych reprezentacji Krausa opisujących działanie kanału Φ_φ . Pozwala nam to zapisać powyższe zagadnienie optymalizacyjne jako:

$$\begin{aligned} 4 \max_{|\psi_0^{\text{ext}}\rangle} \min_{|\Psi_\varphi\rangle} \text{Tr}_{SEA} |\dot{\Psi}_\varphi\rangle\langle\dot{\Psi}_\varphi| &= 4 \max_{|\psi_0^{\text{ext}}\rangle} \min_{\tilde{K}_\varphi} \text{Tr}_{SEA} [\dot{U}_\varphi |\Psi_0\rangle\langle\Psi_0| \dot{U}_\varphi^\dagger] = \\ &= 4 \max_{|\psi_0^{\text{ext}}\rangle} \min_{\tilde{K}_\varphi} \text{Tr}_{SE} \left[\sum_j \left(\dot{\tilde{K}}_\varphi^j \otimes \mathbb{1} \right) |\psi_0^{\text{ext}}\rangle\langle\psi_0^{\text{ext}}| \left(\dot{\tilde{K}}_\varphi^{j\dagger} \otimes \mathbb{1} \right) \right] = \\ &= 4 \max_{|\psi_0^{\text{ext}}\rangle} \min_{\tilde{K}_\varphi} \text{Tr}_S \left[\sum_j \dot{\tilde{K}}_\varphi^j \rho_0^{\text{ext}} \dot{\tilde{K}}_\varphi^{j\dagger} \right] = 4 \max_{|\psi_0^{\text{ext}}\rangle} \min_{\tilde{K}_\varphi} \text{Tr}_S \left[\rho_0^{\text{ext}} \sum_j \dot{\tilde{K}}_\varphi^{j\dagger} \dot{\tilde{K}}_\varphi^j \right], \quad (1.54) \end{aligned}$$

gdzie $\rho_0^{\text{ext}} = \text{Tr}_E |\psi_0^{\text{ext}}\rangle\langle\psi_0^{\text{ext}}|$. Możemy teraz zamienić kolejność maksymalizacji oraz minimalizacji [Fujiwara i Imai 2008] i używając normy operatorowej $\|\cdot\|$ zapisać końcową postać ograniczenia:

$$F \leq \min_{\tilde{K}_\varphi} \left\| \sum_j \dot{\tilde{K}}_\varphi^{j\dagger} \dot{\tilde{K}}_\varphi^j \right\|, \quad (1.55)$$

gdzie minimalizację należy wykonać po wszystkich równoważnych reprezentacjach Krausa. Wykonanie takiej minimalizacji nie jest prostym zadaniem. Jednak jeżeli jesteśmy w stanie zaproponować odpowiednią dla danego problemu reprezentację, to można otrzymać w ten sposób użyteczne ograniczenie.

Przykładem zastosowania metody rozszerzenia kanału, w którym można łatwo skonstruować użyteczną reprezentację Krausa, są modele z nieskorelowanym szumem. W tych modelach szum działa niezależnie i identycznie na każdą cząstkę. W związku z tym wydaje się dobrym pomysłem, aby optymalizować po reprezentacjach Krausa, które mają strukturę produktową:

$$\rho_\varphi = \Phi_\varphi[\rho_0] = \sum_{j_1, \dots, j_N} k_\varphi^{j_1} \otimes \dots \otimes k_\varphi^{j_N} \rho_0 k_\varphi^{j_1\dagger} \otimes \dots \otimes k_\varphi^{j_N\dagger}, \quad (1.56)$$

gdzie $\{k_\varphi^j\}_j$ to zbiór lokalnych operatorów Krausa opisujący działanie szumu na pojedynczą cząstkę. Pozwala to zapisać ograniczenie (1.55) w postaci

$$F \leq 4 \min_{\tilde{k}_\varphi} \left(N \|\alpha_{\tilde{k}_\varphi}\| + N(N-1) \|\beta_{\tilde{k}_\varphi}\|^2 \right), \quad (1.57)$$

gdzie

$$\alpha_{\tilde{k}_\varphi} = \sum_j \dot{\tilde{k}}_\varphi^{j\dagger} \dot{\tilde{k}}_\varphi^j, \quad \beta_{\tilde{k}_\varphi} = i \sum_j \dot{\tilde{k}}_\varphi^{j\dagger} \tilde{k}_\varphi^j. \quad (1.58)$$

Optymalizację w tym wypadku wykonujemy jedynie po lokalnych równoważnych reprezentacjach Krausa $\{k_\varphi^j\}_j$, które są powiązane ze sobą za pomocą lokalnej transformacji unitarnej u_φ następująco: $\tilde{k}_\varphi^i = \sum_j (u_\varphi)_{i,j} k_\varphi^j$. Dla każdego ustalonego $\varphi = \varphi_0$, ograniczenie (1.57) zależy jedynie od operatorów $\tilde{k}_\varphi^j$ i ich pierwszych pochodnych w punkcie φ_0 . Ponadto, ograniczenie to nie jest czułe na zmianę reprezentacji Krausa za pomocą operacji unitarnej, która jest niezależna od φ . Oznacza to, że równoważne reprezentacje Krausa można sparаметryzować za pomocą Hermitowskiej macierzy h , która jest generatorem $u_\varphi = e^{-ih(\varphi-\varphi_0)}$. Redukuje to problem optymalizacyjny (1.57) do minimalizacji po h . Dodatkowo, można ten problem przeformułować korzystając z faktu, że dla większości typowych modeli szumu QFI skaluje się liniowo z N . Oznacza to, że istnieje reprezentacja, dla której $\beta_{\tilde{k}_\varphi} = 0$ lub równoważnie

$$\sum_{i,j} h_{i,j} k_\varphi^{i\dagger} k_\varphi^j = i \sum_j \dot{k}_\varphi^{j\dagger} k_\varphi^j. \quad (1.59)$$

W przypadku takich modeli ograniczenie (1.57) można zapisać w postaci

$$F \leq 4N \min_{\substack{h \\ \beta_{\tilde{k}_\varphi}=0}} \|\alpha_{\tilde{k}_\varphi}\|. \quad (1.60)$$

Występująca tu minimalizacja, to przykład problemu wypukłego, który można efektywnie rozwiązać stosując metodę SDP (ang. *Semi-Definite Programming*) [Demkowicz-Dobrzański, Kołodyński i Guţă 2012; Demkowicz-Dobrzański i Maccone 2014].

Rozważmy dla przykładu estymację fazy φ w obecności lokalnego defazowania o sile η . Będzie nas interesował układ N spinów- $\frac{1}{2}$, który ewoluuje przez kanał kwantowy postaci $\Phi_\varphi[\rho_0] = e^{-iH\varphi} \Phi[\rho_0] e^{iH\varphi}$, gdzie $H = \frac{1}{2} \sum_{n=1}^N \sigma_z^{[n]}$, a kanał Φ opisuje działanie lokalnego defazowania. W problemie takim szum działa niezależnie i identycznie na każdą cząstkę, a więc będziemy rozważać produktową strukturę reprezentacji Krausa kanału Φ_φ , jak w równaniu (1.56). Pozwoli to wprowadzić ograniczenie na precyzję estymacji w oparciu jedynie o dynamikę pojedynczego spinu. Fizycznie, defazowanie oznacza utratę koherencji. Pojedynczy spin, który początkowo znajduje się w stanie czystym $|\psi_0\rangle = \alpha \left|+\frac{1}{2}\right\rangle + \beta \left|-\frac{1}{2}\right\rangle$, pod wpływem lokalnego defazowania przejdzie do stanu mieszanego

$$\Phi[|\psi_0\rangle] = \begin{pmatrix} |\alpha|^2 & \eta\alpha\beta^* \\ \eta\alpha^*\beta & |\beta|^2 \end{pmatrix}. \quad (1.61)$$

Działanie takiego kanału Φ możemy zareprezentować za pomocą następujących lokalnych operatorów Krausa:

$$k^0 = \sqrt{\frac{1+\eta}{2}} \mathbb{1}, \quad k^1 = \sqrt{\frac{1-\eta}{2}} \sigma_z. \quad (1.62)$$

Natomiast lokalne operatory Krausa dla całkowitego kanału Φ_φ mają postać:

$$k_\varphi^0 = e^{-\frac{i}{2}\sigma_z\varphi} \sqrt{\frac{1+\eta}{2}} \mathbb{1}, \quad k_\varphi^1 = e^{-\frac{i}{2}\sigma_z\varphi} \sqrt{\frac{1-\eta}{2}} \sigma_z. \quad (1.63)$$

Stosując teraz metodę rozszerzenia kanału otrzymujemy $\beta_{k_\varphi} = 0$, gdy

$$h = -\frac{1}{2\sqrt{1-\eta^2}} \sigma_x. \quad (1.64)$$

Ograniczenie na precyzję estymacji przyjmuje wtedy postać [Demkowicz-Dobrzański, Kołodyński i Guţă 2012; Demkowicz-Dobrzański i Maccone 2014]

$$\Delta^2 \hat{\varphi} \geq \frac{1-\eta^2}{\eta^2 N}. \quad (1.65)$$

Przypomnijmy, że ograniczenie to jest zoptymalizowane po stanach początkowych i uwzględnia możliwość użycia stanów kwantowych, które są splątane. Jeśli zawężymy rozważania do stanów produktowych to okaże się, że najlepsza wartość precyzji estymacji to $\Delta^2 \hat{\varphi} = \frac{1}{\eta^2 N}$. Oznacza to, że użycie stanów splątanych umożliwia poprawę precyzji estymacji maksymalnie o stały czynnik $1 - \eta^2$.

Ograniczenie (1.65) ilustruje ogólną prawidłowość dla modeli z nieskorelowanym szumem. Mianowicie mimo że stany splątane umożliwiają poprawę precyzji to nie pozwalają na osiągnięcia HL. Dla takich układów fundamentalne ograniczenie na precyzję estymacji ma postać zbliżoną do SQL:

$$\Delta^2 \hat{\varphi} \geq \frac{c}{N}, \quad (1.66)$$

gdzie c jest stałą specyficzną dla konkretnego modelu szumu [Demkowicz-Dobrzański, Kołodyński i Guţă 2012].

W obecności dekoherencji analiza Bayesowska zagadnienia estymacji fazy jest zazwyczaj dużo trudniejsza. Na szczęście pokazano, że dla nieskorelowanych modeli szumu asymptotyczna wartości minimalnego AMSE otrzymana w kwantowym podejściu Bayesowskim pokrywa się z wartością QFI (dokładnie jej odwrotnością) [Jarzyna i Demkowicz-Dobrzański 2015]. Pozwala nam to wnioskować o precyzji estymacji w podejściu Bayesowskim w oparciu o wartość QFI. Oznacza to również, że dla nieskorelowanych modeli szumu asymptotyczne ograniczenia na precyzję estymacji otrzymane w oparciu o QFI są zawsze wysycalne. Jest to związane z tym, że w podejściu Bayesowskim wykonujemy bezpośrednią minimalizację, więc otrzymane rezultaty nie są obciążone koniecznością dodatkowych rozważań o wysycalności.

Jakkolwiek fundamentalne ograniczenia na precyzję estymacji mają ogromne

znaczenie teoretyczne to równie istotna jest ich praktyczna wysycalność w eksperymencie. Dla przykładu stany GHZ/NOON i SIN, optymalne w przypadku braku dekoherencji odpowiednio w podejściu opartym o QFI i o rozumowanie Bayesowskie, są bardzo trudne do przygotowania z wyjątkiem układów bardzo niewielu cząstek. Dla dużej liczby cząstek jedyne łatwo dostępne stany splątane to stany ściśnięte [J. Ma i in. 2011; Demkowicz-Dobrzański, Jarzyna i Kołodyński 2015; Pezzé, Smerzi i in. 2018]. Okazuje się jednak, że są one wystarczające, aby wysycić asymptotycznie ograniczenia na precyzję w typowych problemach metrologicznych z nieskorelowanym szumem. Obserwujemy to np. w problemie estymacji fazy w interferometrze Macha-Zehndera w obecności dekoherencji, gdzie optymalny schemat pomiarowy zakłada wpuszczenie do dwóch różnych ramion interferometru odpowiednio stanu koherentnego i stanu ściśniętej próżni, a następnie pomiar różnicy obsadzeń [Caves 1981; Demkowicz-Dobrzański, Jarzyna i Kołodyński 2015]. Podobnie w przypadku interferometrii Ramseya, gdzie optymalny stan początkowy zespołu atomów to stan jednoosiowo skręcony [Kitagawa i Ueda 1993; Ulam-Orgikh i Kitagawa 2001].

Przykładem takiego praktycznego zastosowania metrologii kwantowej jest użycie kombinacji stanu koherentnego i stanu ściśniętej próżni w detekcji fal grawitacyjnych. Pierwsze demonstracyjne zastosowania stanów ściśniętych do poprawy czułości detektorów fal grawitacyjnych GEO 600 oraz LIGO opisano w pracach [Abadie i in. 2011; Aasi i in. 2013]. Następnie pokazano między innymi, że przy występującym wtedy poziomie szumów optycznych w detektorze GEO 600 taki właśnie stan zapewnia praktycznie najlepszą możliwą czułość [Demkowicz-Dobrzański, Banaszek i Schnabel 2013]. Aktualnie stany ściśnięte są już na stałe używane w detektorach Advanced LIGO i Advanced Virgo. Były wykorzystane podczas pomiarów, w których dokonano detekcji fal grawitacyjnych i odpowiadają za wzrost liczby obserwacji aż do 50% [Acernese i in. 2019; Tse i in. 2019].

Rozdział 2

Sieci tensorowe

W tym rozdziale zostanie dokonany przegląd podstawowych wiadomości o sieciach tensorowych. Zaczniemy od wprowadzenia graficznej reprezentacji tensorów i zademonstrujemy jak w tym nowym języku wyrażają się podstawowe operacje wykonywane na tensorach. Następnie wprowadzimy rozkład na wartości osobliwe oraz rozkład Schmidta, co pozwoli nam płynnie przejść do entropii splątania i jej prawa powierzchniowego. Będziemy wtedy gotowi, aby bazując na koncepcji splątania, omówić podstawowe rodzaje sieci tensorowych. Na końcu skupimy się na szczegółowym opisie jednowymiarowych sieci tensorowych.

2.1 Diagramy sieci tensorowych

Przez tensor rozumiemy wielowymiarową tablicę liczb zespolonych (nie zakładamy żadnych własności transformacyjnych), czyli element $\mathbb{C}^{d_1 \times \dots \times d_r}$, gdzie r oznacza rząd tensora, a d_j to wymiar j -tego indeksu tensora. Skalary, wektory i macierze to najprostsze przykłady tensorów odpowiednio o rzędzie 0, 1 i 2. Sieć tensorową definiujemy jako zbiór tensorów, których część albo wszystkie indeksy są zwięzione według określonego schematu.

Graficzna reprezentacja sieci tensorowych w postaci diagramów była jednym z katalizatorów sukcesu, jaki odniosły one w modelowaniu wielu problemów fizycznych, podobnie jak miało to miejsce w przypadku diagramów Feynmana. Diagramy te umożliwiają nie tylko uproszczenie równań, ale też w przejrzysty sposób oddają geometryczną strukturę rozważanego problemu. Wiele ogólnych własności badanych obiektów, w szczególności stanów kwantowych, może być odczytanych ze struktury diagramu sieci tensorowej. Idea przedstawiania tensorów w postaci diagramów pochodzi od Penrosa [1971] i może być uważana za rozwinięcie konwencji sumacyjnej Einsteina. W tej notacji tensor rzędu r jest przedstawiony jako geometryczny kształt (np. okrąg), z którego wychodzi r linii („nóżek”)¹:

¹Użyliśmy tutaj zapisu abstrakcyjno-wskaźnikowego, w którym $T_{\alpha_1, \dots, \alpha_r}$ oznacza „cały” tensor, a nie jedynie jego konkretny element w wybranej bazie (w praktyce z kontekstu powinno być jasne czy pisząc $T_{\alpha_1, \dots, \alpha_r}$ mamy na myśli tensor czy jego składową w bazie).

$$T_{\alpha_1, \dots, \alpha_r} = \begin{array}{c} \alpha_3 \quad \alpha_2 \\ \diagdown \quad \diagup \\ \alpha_4 - (T) - \alpha_1 \\ \diagup \quad \diagdown \\ \alpha_5 \quad \dots \quad \alpha_r \end{array},$$

Jeśli wymaga tego rozważany problem to można nadać dodatkowe znaczenie użytym kształtom albo kierunkom linii, np. wprowadzić rozróżnienie na linii idące do góry/dołu i interpretować je jako indeksy kontrawariantne/kowariantne, jednak na potrzeby tej dysertacji nie będzie to potrzebne. W celu zwiększenia przejrzystości diagramów będziemy pomijać wypisywanie symboli poszczególnych indeksów. Dla przykładu przedstawimy jeszcze za pomocą diagramów podstawowe rodzaje tensorów, czyli skalar s , wektor v_α i macierz $M_{\alpha,\beta}$:

$$\begin{array}{ccc} \textcircled{s}, & \text{---} \textcircled{v}, & \text{---} \textcircled{M} \text{---}. \end{array}$$

Omówimy teraz jak za pomocą diagramów sieci tensorowych przedstawić różne operacje na tensorach. Podstawową operacją jaką wykonujemy na tensorach jest ich zwięzanie, czyli sumowanie po wspólnym indeksie, co na diagramie przedstawiamy za pomocą łączenia linii odpowiadających tym indeksom. Zobrazujemy to przykładami iloczynu skalarnego², mnożenia macierzy przez wektor, mnożenia dwóch macierzy oraz śladu macierzy:

$$\begin{aligned} \sum_{\alpha} v_{\alpha}^* v_{\alpha} &= \textcircled{v^*} \text{---} \textcircled{v}, \\ \sum_{\beta} M_{\alpha,\beta} v_{\beta} &= \text{---} \textcircled{M} \text{---} \textcircled{v}, \\ \sum_{\beta} M_{\alpha,\beta} N_{\beta,\gamma} &= \text{---} \textcircled{M} \text{---} \textcircled{N} \text{---}, \\ \sum_{\alpha} M_{\alpha,\alpha} &= \textcircled{M} \text{---} \textcircled{M}. \end{aligned}$$

Zwróćmy uwagę, że w zapisie diagramatycznym operator identycznościowy $\delta_{\alpha,\beta}$ jest reprezentowany przez linię:

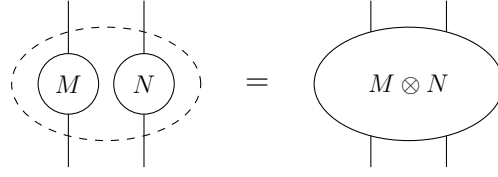
$$\delta_{\alpha,\beta} = \alpha \text{---} \beta.$$

Kolejną ważną operacją jest iloczyn tensorowy, którego wartość dla ustalonego zbioru indeksów jest iloczynem wartości poszczególnych tensorów dla tych indeksów, mianowicie

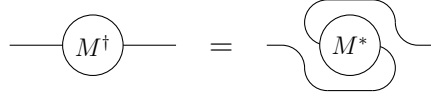
$$[M \otimes N]_{\alpha_1, \beta_1, \alpha_2, \beta_2} = M_{\alpha_1, \beta_1} \cdot N_{\alpha_2, \beta_2}. \quad (2.1)$$

Możemy to w prosty sposób przedstawić na diagramie zestawiając dwa tensory obok siebie:

²Sprzężenie zespolone oznaczamy gwiazdką.



Istotną operacją jest też sprzężenie Hermitowskie, czyli połączenie sprzężenia zespolonego z transpozycją, którą graficznie możemy wykonać zginając odpowiednio linie:



Ostatnią operacją, którą omówimy, a która jest szczególnie często wykonywana podczas pracy z sieciami tensorowymi, jest grupowanie/dzielenie indeksów. W ogólności przestrzenie liniowe $\mathbb{C}^{d_1 \times \dots \times d_r}$ i $\mathbb{C}^{d'_1 \times \dots \times d'_s}$ są izomorficzne jeśli $\prod_{j=1}^r d_j = \prod_{j=1}^s d'_j$ – oznacza to, że pojęcie rzędu tensora jest dosyć płynne, ponieważ możemy znaleźć tensor, który będzie zawierał taką samą informację, ale będzie miał inny rząd. Przykładem grupowania indeksów jest wektoryzacja macierzy $M_{\alpha,\beta}$ o wymiarach $m \times n$, podczas której łączymy kolejne kolumny macierzy w pojedynczy wektor kolumnowy $M_{\gamma,1}$ o długości $m \cdot n$, gdzie nowy indeks γ jest zdefiniowany jako $\gamma = \alpha + m(\beta - 1)$ (wektoryzacja do wektora wierszowego przebiega analogicznie). Operację tę możemy z łatwością odwrócić $\alpha = \gamma \bmod m$, $\beta = 1 + (\gamma - \alpha)/m$ i mówimy wtedy o dzieleniu indeksu. Grupowanie indeksów przedstawiamy diagramatycznie wprowadzając nową grubszą linię, która przebiega po zgrupowanym indeksie, np. w przypadku opisanej tu wektoryzacji macierzy wprowadzamy linię $\mathbb{I} = ||$ odpowiadającą podwójnemu indeksowi $\gamma = (\alpha, \beta)$ (γ jest parą uporządkowaną α oraz β), co na diagramie wygląda następująco:

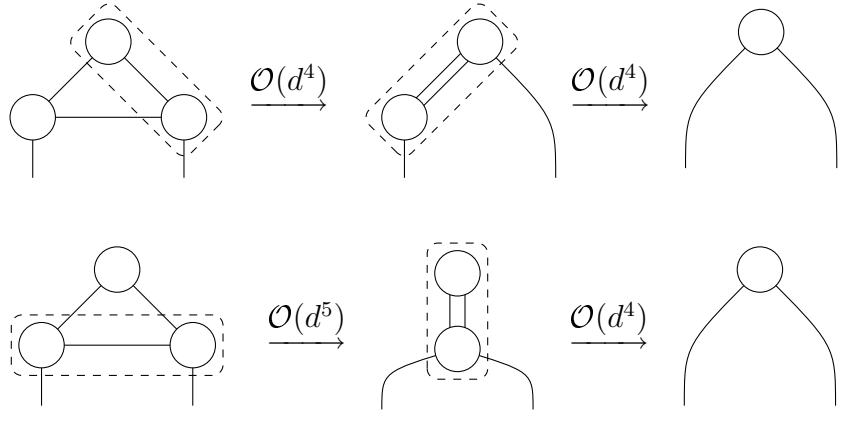


Takie grupowanie/dzielenie indeksów można z łatwością uogólnić w celu opisu obniżania/podwyższania rzędu dowolnego tensora. Biegłe opanowanie grupowania/dzielenia indeksów oraz ich permutowania jest kluczem do sukcesu przy implementacji numerycznej sieci tensorowych. Jest to związane z tym, że większość algorytmów numerycznych jest zoptymalizowana do działania na macierzach. Dla przykładu, jeśli chcemy wyznaczyć $T_{\alpha_3, \alpha_4, \alpha_5} = U_{\alpha_1, \alpha_2, \alpha_3, \alpha_4} W_{\alpha_1, \alpha_2, \alpha_5}$, czyli dokonać zwężenia tensorów U i W po wspólnych indeksach α_1 i α_2 to należy:

- dokonać permutacji indeksów tensora U , tak aby otrzymać $U_{\alpha_3, \alpha_4, \alpha_1, \alpha_2}$;
- zgrupować indeksy obu tensorów, tak aby przyjęły one postać macierzy $U_{(\alpha_3, \alpha_4), (\alpha_1, \alpha_2)}$ i $W_{(\alpha_1, \alpha_2), \alpha_5}$;
- zwęzić podwójny indeks (α_1, α_2) w wyniku czego otrzymamy macierz $T_{(\alpha_3, \alpha_4), \alpha_5}$, wykorzystując przy tym to, że kod mnożenia macierzy jest bardzo dobrze zoptymalizowaną numerycznie operacją;

- rozdzielić podwójny indeks (α_3, α_4) z powrotem na α_3 i α_4 .

Poruszymy jeszcze jeden istotny wątek związany z numeryczną implementacją sieci tensorowych, mianowicie kolejność zwięzania indeksów. Czasowa złożoność obliczeniowa mnożenia dwóch macierzy M i N o wymiarach odpowiednio $m \times n$ i $n \times p$ wynosi $\mathcal{O}(nmp)$ [Cormen i in. 1990], co można równoważnie zapisać jako $\mathcal{O}\left(\frac{|M||N|}{|M \cap N|}\right)$ gdzie $|M|$, oznacza iloczyn wymiarów macierzy M , czyli $|M| = mn$, $|N| = np$, a $|M \cap N|$ oznacza wymiar zwięzanego indeksu, czyli $|M \cap N| = n$. Jak przed chwilą pokazaliśmy, zwięzanie indeksów dowolnych tensorów można sprowadzić do mnożenia macierzy, więc wzór $\mathcal{O}\left(\frac{|T||W|}{|T \cap W|}\right)$ opisuje czasową złożoność obliczeniową zwięzania dowolnych dwóch tensorów T i W , tylko że teraz $|T|$ jest iloczynem wszystkich wymiarów tensora T , a $|T \cap W|$ jest iloczynem wymiarów wszystkich zwięzanych (wspólnych) indeksów. Rozważmy teraz dwa sposoby zwięzania tej samej sieci tensorowej, w której każdy indeks ma wymiar d (tensory zwięzane w danym kroku są otoczone dookoła przerywaną linią):



Widać że złożoność obliczeniowa zależy od kolejności w jakiej zwięzamy tensory, dlatego jest bardzo istotne, aby przed rozpoczęciem obliczeń numerycznych, zbadać jaka kolejność będzie najbardziej optymalna. W ogólności znalezienie optymalnej kolejności zwięzań jest problemem NP-zupełnym [Arad i Landau 2010; Bridgeman i Chubb 2017], jednak w przypadku sieci tensorowych, które będą dla nas istotne nie stanowi to bardzo dużego problemu.

2.2 Rozkład według wartości osobliwych

Jednym z głównych zastosowań sieci tensorowych jest reprezentacja wielocząstkowych stanów kwantowych (dla których rozmiar przestrzeni Hilberta jest bardzo duży) za pomocą macierzy o małym wymiarze. Centralnym filarem formalizmu sieci tensorowych, umożliwiającym dokonanie takiego przybliżenia, jest rozkład według wartości osobliwych (SVD, ang. *Singular Value Decomposition*) [Meyer 2000]. SVD pozwala rozłożyć dowolną macierz M , o wymiarach $m \times n$, na

$$M = USV^\dagger, \quad (2.2)$$

gdzie U i V to macierze unitarne ($UU^\dagger = \mathbb{1} = VV^\dagger$) o wymiarach odpowiednio $m \times m$ i $n \times n$, a S to macierz diagonalna o wymiarach $m \times n$ z nieujemnymi rzeczywistymi wartościami na diagonalu, zwyczajowo uporządkowanymi malejąco, zwanymi wartościami singularnymi. Liczba niezerowych wartości singularnych jest równa rzędowi macierzy M . Dla nas jedną z najważniejszych własności SVD jest to, że pozwala skonstruować najlepsze możliwe przybliżenie macierzy M o rzędzie r przez macierz M' o niższym rzędzie r' . Jeśli za błąd przybliżenia uznamy normę Frobeniusa $\|M - M'\|_F$ to spośród wszystkich macierzy o rzędzie r' minimalną wartość błędu przybliżenia uzyskamy dla macierzy $M' = US'V^\dagger$, gdzie macierz S' różni się od macierzy S tym że zawiera jedynie r' największych wartości singularnych (a pozostałe zostały przyrównane do zera) [Eckart i Young 1936]. Warto wspomnieć, że w praktyce numerycznej nie zawsze potrzebujemy SVD, często wystarczający jest mniej wymagający numerycznie rozkład QR, który pozwala zapisać dowolną macierz M o wymiarach $m \times n$ w postaci $M = QR$, gdzie Q to macierz unitarna o wymiarach $m \times m$, a R to macierz górnótrójkątna (to znaczy $R_{\alpha,\beta} = 0$ jeśli $\alpha > \beta$) o wymiarach $m \times n$.

SVD stanowi również trzon bardzo użytecznego narzędzia kwantowej teorii informacji, mianowicie rozkładu Schmidta [Nielsen i Chuang 2000], który pozwala nam scharakteryzować splątanie pomiędzy dwoma podukładami złożonego układu (zakładając, że układ fizyczny znajduje się w stanie czystym). Stan czysty układu złożonego AB można zapisać jako

$$|\psi\rangle = \sum_{\alpha,\beta} \Psi_{\alpha,\beta} |\alpha\rangle_A |\beta\rangle_B, \quad (2.3)$$

gdzie $\{|\alpha\rangle_A\}_\alpha$ i $\{|\beta\rangle_B\}_\beta$ stanowią bazę odpowiednio podukładu A i B . Współczynniki $\Psi_{\alpha,\beta}$ możemy zinterpretować jako elementy macierzy Ψ , którą możemy rozłożyć za pomocą SVD:

$$\begin{aligned} |\psi\rangle &= \sum_{\alpha,\beta} \sum_{\gamma} U_{\alpha,\gamma} S_{\gamma,\gamma} V_{\beta,\gamma}^* |\alpha\rangle_A |\beta\rangle_B = \\ &= \sum_{\gamma} \left(\sum_{\alpha} |\alpha\rangle_A \langle \alpha| U| \gamma \rangle_A \right) S_{\gamma,\gamma} \left(\sum_{\beta} |\beta\rangle_B \langle \beta| V^*| \gamma \rangle_B \right) = \\ &= \sum_{\gamma=1}^{\chi} s_{\gamma} |u_{\gamma}\rangle_A |v_{\gamma}\rangle_B, \end{aligned} \quad (2.4)$$

gdzie $|u_{\gamma}\rangle_A = U|\gamma\rangle_A$ i $|v_{\gamma}\rangle_B = V^*|\gamma\rangle_B$ to wektory ortonormalnej bazy (bazy Schmidta), a uporządkowane malejąco wartości singularne $s_{\gamma} = S_{\gamma,\gamma}$ nazywamy współczynnikami Schmidta. Liczbę niezerowych współczynników Schmidta nazywamy rzędem Schmidta χ i jest to jedna z możliwych miar splątania pomiędzy podukładami A i B – jeśli $\chi = 1$ to układ jest w stanie separowalnym, w przeciwnym wypadku mamy do czynienia ze stanem splątany (maksymalna wartość rzędu Schmidta to wymiar przestrzeni Hilberta mniejszego z podukładów). Widzimy teraz, że jeżeli chcielibyśmy przybliżyć stan $|\psi\rangle$ o rzędzie Schmidta χ stanem $|\psi'\rangle$, który wykazuje mniejsze splątanie pomiędzy podukładami (mierzone

rzędem Schmidta) – ustalmy że chcemy aby ono wynosiło χ' ($\chi' < \chi$) – to najlepsze możliwe przybliżenie otrzymamy przyjmując $|\psi'\rangle = \sum_{\gamma=1}^{\chi'} s_{\gamma} |u_{\gamma}\rangle_A |v_{\gamma}\rangle_B$ (w celu normalizacji trzeba dodatkowo odpowiednio przeskalować współczynniki Schmidta).

Znając rozkład Schmidta stanu $|\psi\rangle$ możemy z łatwością skonstruować zredukowane macierze gęstości:

$$\rho_A = \text{Tr}_B |\psi\rangle\langle\psi| = \sum_{\gamma} s_{\gamma}^2 |u_{\gamma}\rangle_A \langle u_{\gamma}|, \quad (2.5)$$

$$\rho_B = \text{Tr}_A |\psi\rangle\langle\psi| = \sum_{\gamma} s_{\gamma}^2 |v_{\gamma}\rangle_B \langle v_{\gamma}|, \quad (2.6)$$

przy czym warto zwrócić uwagę, że wartości własne zredukowanych macierzy gęstości to kwadraty współczynników Schmidta, a wektory własne to wektory bazy Schmidta. Jako miarę splątania będziemy używać entropię splątania von Neumanna $\mathcal{S}$, która wiąże się ze współczynnikami Schmidta następująco:

$$\mathcal{S} = -\text{Tr}(\rho_A \log_2 \rho_A) = -\text{Tr}(\rho_B \log_2 \rho_B) = -\sum_{\gamma} s_{\gamma}^2 \log_2 s_{\gamma}^2. \quad (2.7)$$

Jeżeli układ jest w stanie separowalnym to $\mathcal{S} = 0$. Dla ustalonego rzędu Schmidta χ maksymalna wartość entropii splątania to $\mathcal{S} = \log_2 \chi$ i realizuje się ona, gdy stan układu ma płaski rozkład wartości Schmidta $s_{\gamma} = \frac{1}{\sqrt{\chi}}$.

2.3 Prawo powierzchniowe dla entropii splątania

Losowy stan z wielocząstkowej przestrzeni Hilberta zachowuje się zgodnie z prawem objętościowym (ang. *volume law*), to znaczy entropia splątania pomiędzy dwoma podukładami jest proporcjonalna do objętości tych podukładów. W szczególności, losowy stan N d -wymiarowych cząstek wykazuje entropię splątania $\mathcal{S} \approx \frac{N}{2} \log_2 d$ przy podziale układu na dwie równe części po $\frac{N}{2}$ cząstek [Page 1993].

Okazuje się jednak, że nie jest to typowa własność dla stanów występujących w przyrodzie. Wiele Hamiltonianów opisujących układy fizyczne cechuje lokalny charakter oddziaływania pomiędzy cząsteczkami. Lokalność tych oddziaływań okazuje się mieć daleko idące konsekwencje. W szczególności stany podstawowe oraz niskoenergetyczne stany wzbudzone lokalnych Hamiltonianów wykazują zachowanie zgodne z prawem powierzchniowym (ang. *area law*), to znaczy entropia splątania pomiędzy dwoma podukładami (dla odpowiednio dużych podukładów) jest proporcjonalna do powierzchni granicznej pomiędzy nimi [Srednicki 1993; Eisert, Cramer i Plenio 2010] (dla układów bez przerwy energetycznej mogą pojawić się dodatkowe poprawki do tego skalowania, zazwyczaj logarytmiczne w funkcji rozmiaru układu [Vidal i in. 2003]).

Oznacza to, że fizycznie interesujące stany, te które spełniają prawo powierzchniowe, nie zajmują całej przestrzeni Hilberta (która dla N d -wymiarowych cząstek ma wymiar aż d^N), a jedynie jej eksponencjalnie mały kawałek. Rozważymy rozmaitość stanów, którą można otrzymać w wyniku ewolucji generowanej przez dowolny, zależny od czasu lokalny Hamiltonian przez czas, który skaluje

się wielomianowo z liczbą cząstek w układzie N . Okaże się, że otrzymamy rozmaitość, która jest eksponencjalnie małą częścią przestrzeni Hilberta [Poulin i in. 2011]. Oznacza to, że znakomita większość przestrzeni Hilberta jest osiągalna dopiero po czasie ewolucji, który jest eksponencjalny w N . Mówi się wręcz, że przestrzeń Hilberta stanowi „wygodną iluzję” [Poulin i in. 2011] – jest wygodna z perspektywy matematycznej, ale jest również iluzją, ponieważ nikt nigdy nie zaobserwuje jej w całości.

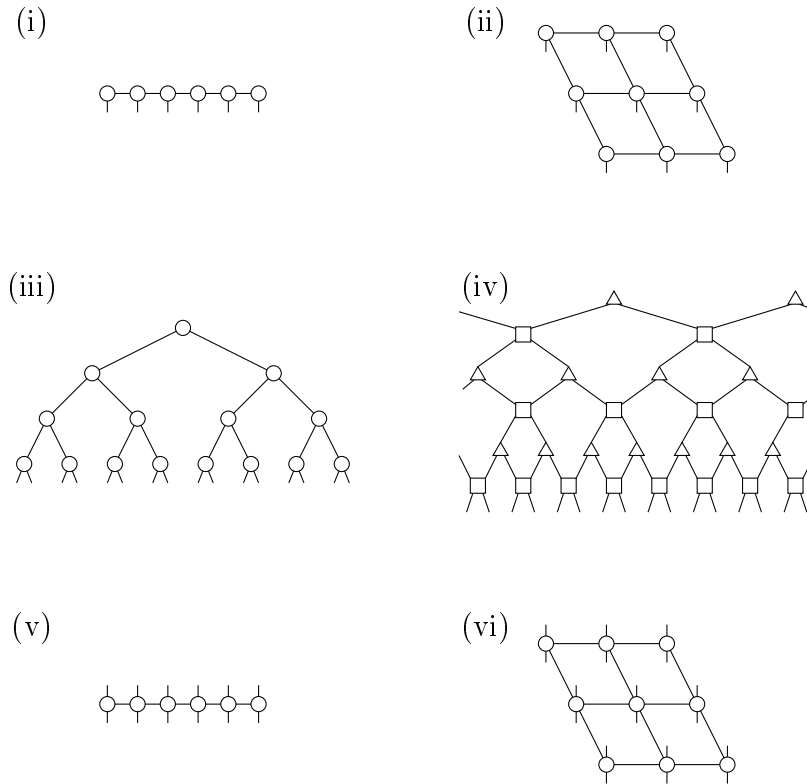
Jeśli chcemy teraz badać stany należące do fizycznie istotnego wycinka przestrzeni Hilberta to byłoby pożądanym, aby mieć do nich dostęp bezpośrednio. To właśnie umożliwiają nam sieci tensorowe. Można traktować je jako reprezentację klasy stanów kwantowych, które wykazują określony rodzaj splątania – tak w sensie geometrycznej struktury (jedno- lub dwuwymiarowe), jak i zasięgu korelacji (skończony lub nieskończony). W następnym podrozdziale powiążemy różne rodzaje sieci tensorowych z typem splątania, który naturalnie opisują.

2.4 Główne rodzaje sieci tensorowych

Dokonyamy teraz przeglądu podstawowych własności najczęściej spotykanych w literaturze sieci tensorowych [Orús 2019], które diagramatycznie przedstawiliśmy na rysunku 2.1.

Macierzowe stany produktowe (MPS, ang. *Matrix Product States*) [Fannes, Nachtergaele i Werner 1992] albo jak niektórzy matematycy je nazywają łańcuchy tensorów (ang. *Tensor Trains*) [Hackbusch 2010] – są to jednowymiarowe (1d) sieci tensorowe w kształcie łańcucha, jak na rysunku 2.1(i). Każdej cząstce w układzie jest przyporządkowany jeden tensor w łańcuchu. Charakteryzują się stałą entropią splątania, niezależnie od rozmiaru układu, czyli spełniają 1d prawo powierzchniowe dla entropii splątania. W związku z tym stanowią naturalną reprezentację dla niskoenergetycznych stanów własnych lokalnych 1d Hamiltonianów z przerwą energetyczną [Hastings 2006; Hastings 2007]. Można je efektywnie i ściśle zwęzać. Dwupunktowa funkcja korelacji na stanie w reprezentacji MPS zanika eksponencjalnie z odległością, a więc MPS opisują stany o skończonym zasięgu korelacji [Schollwöck 2011]. W związku z tym nie mogą ściśle oddać struktury splątania układów w stanie krytycznym. Jednak są one bardzo efektywne i na przestrzeni lat opracowano szereg metod umożliwiających uzyskanie przy ich pomocy wartościowych informacji o własnościach układów w stanie krytycznym [Tagliacozzo i in. 2008; Pollmann i in. 2009].

Zrzutowane pary stanów splątanych (PEPS, ang. *Projected Entangled Pair States*) [Verstraete i J. I. Cirac 2004] stanowią naturalne uogólnienie MPS na większą liczbę wymiarów przestrzennych. Wspomnimy tylko o ich wariancie dwuwymiarowym (2d), który dla sieci kwadratowej przedstawiono na rysunku 2.1(ii). Spełniają one 2d prawo powierzchniowe dla entropii splątania. Reprezentują niskoenergetyczne stany własne lokalnych 2d Hamiltonianów [Pérez-García i in. 2008]. W przeciwieństwie do MPS, zwężanie PEPS jest eksponencjalnie trudnym problemem. W związku z czym trzeba posilkować się odpowiednimi metodami przybliżonymi [Schuch i in. 2007; Orús 2014]. Kolejną różnicą w stosunku do MPS jest to, że PEPS mogą mieć nieskończony zasięg korelacji, a więc



RYSUNEK 2.1: Przykłady sieci tensorowych. (i) Macierzowy stan produktowy (MPS). (ii) Zrzutowane pary stanów splątanych (PEPS) dla sieci kwadratowej. (iii) Sieć drzew tensorowych (TTN). (iv) Wieloskalowy splątany renormalizacyjny ansatz (MERA). Kwadraty oznaczają tensory unitarne, a trójkąty tensory izometryczne. (v) Macierzowy operator produktowy (MPO). (vi) Zrzutowane pary operatorów splątanych (PEPO).

mogą oddać korelacje występujące w układach krytycznych gdzie dwupunktowa funkcja korelacji zanika potęgowo z odległością [Verstraete, Wolf i in. 2006].

Sieci drzew tensorowych (TTN, ang. *Tree Tensor Networks*) [Shi, Duan i Vidal 2006] są to sieci tensorowe o strukturze w kształcie drzewa jak na rysunku 2.1(iii). Poza fizycznym wymiarem przestrzennym mają one dodatkowy wymiar „holograficzny”. Ten dodatkowy wymiar można interpretować jako kierunek renormalizacji – odpowiada on za opis układu na różnych skalach długości. W związku z tym tego typu sieci tensorowe są naturalnym kandydatem do opisu układów niezmienniczych ze względu na skalowanie, jak np. układy w stanie krytycznym. Entropia splątania TTN zależy od sposobu podziału sieci na dwa podukłady. Istnieją takie podziały, które spełniają 1d prawo powierzchniowe oraz takie, które wprowadzają do niego dodatkowe poprawki logarytmiczne (w funkcji rozmiaru układu), jednakże uśredniając TTN spełniają 1d prawo powierzchniowe. Można je jak MPS efektywnie zwięzać. TTN dobrze nadają się do opisu 1d układów z przerwą energetyczną. Zazwyczaj stanowią trochę lepszy od MPS ansatz przy opisie 1d układów w stanie krytycznym [Silvi i in. 2010]. Są też (wraz z MPS) z powodzeniem stosowane w chemii kwantowej [Chan i Sharma 2011; Szalay i in. 2015].

Wieloskalowy splątany renormalizacyjny ansatz (MERA, ang. *Multiscale*

Entanglement Renormalization Ansatz) [Vidal 2007b; Vidal 2010] jest to typ sieci tensorowych o strukturze jak na rysunku 2.1(iv). MERA można uważać za uogólnienie TTN. W MERA na tensory są narzucone ograniczenia, mianowicie muszą one być unitarne albo izometryczne. Unitarne tensory odpowiadają za splątanie pomiędzy sąsiednimi stopniami swobody, a izometryczne tensory odpowiadają za gruboziarnistość (ang. *coarse-graining*) i zachowanie normalizacji stanu kwantowego. Sprawia to, że ich zwięźanie jest bardzo efektywne oraz powoduje, że dla MERA splątanie buduje się lokalnie, na każdej skali długości układu. To wszystko skutkuje tym, że dla 1d MERA entropia splątania wykazuje skalowanie logarytmiczne z rozmiarem układu – dokładnie takie jak konforemna teoria pola przewiduje dla 1d układów w stanie krytycznym [Holzhey, Larsen i Wilczek 1994; Calabrese i Cardy 2004]. Uważa się, że MERA jest powiązana z geometrią przestrzeni poprzez korespondencję anty de Sittera/konforemna teoria pola (AdS/CFT – dualność cechowania/grawitacji) [Swingle 2012].

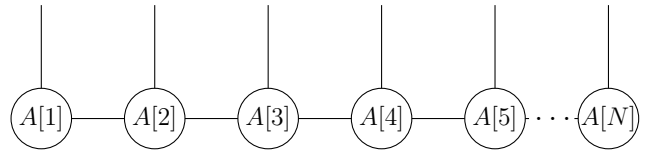
Wszystkie wymienione dotychczas sieci tensorowe służą do reprezentowania czystych stanów kwantowych. Możemy je jednak łatwo uogólnić, aby opisać operatory (w tym macierze gęstości) i tak np. w 1d mamy do czynienia z macierzowymi operatorami produktowymi (MPO, ang. *Matrix Product Operators*) [Zwolak i Vidal 2004; Pirvu i in. 2010] przedstawionymi na rysunku 2.1(v), a w 2d ze zrzutowanymi parami operatorów splątanych (PEPO, ang. *Projected Entangled Pair Operators*) [Czarnik i Dziarmaga 2015] przedstawionymi na rysunku 2.1(vi).

Na koniec wspomnijmy, że istnieją również ciągle sieci tensorowe służące jako ansatz dla niskoenergetycznych funkcjonałów w kwantowej teorii pola (jak i dla operatorów działających na te funkcjonały) [Verstraete i J. I. Cirac 2010; Haegeman i in. 2013].

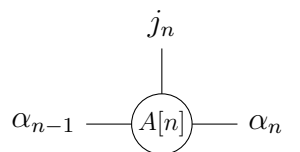
2.5 Jednowymiarowe sieci tensorowe

2.5.1 Macierzowe stany produktowe

W reprezentacji macierzowych stanów produktowych (MPS) stan kwantowy $|\psi\rangle$ opisujący N rozróżnialnych d -wymiarowych cząstek ma postać



Tensor $A[n]$ – n służy nam do oznaczenia pozycji w łańcuchu, $n \in \{1, \dots, N\}$ – ma trzy indeksy



Linie poziome odpowiadają indeksom (nazywanymi wirtualnymi indeksami) $\alpha_n \in \{1, \dots, \chi_n\}$, gdzie parametr χ_n nazywamy wymiarem wiązania ($\chi_0 =$

$\chi_N = 1$), a linia pionowa odpowiada indeksowi (nazywanemu indeksem fizycznym) $j_n \in \{0, \dots, d-1\}$, który numeruje wektory bazy n -tej cząstki. Możemy ten stan zapisać również za pomocą równania jako

$$|\psi\rangle = \sum_{j_1, j_2, \dots, j_N=0}^{d-1} A[1]^{j_1} A[2]^{j_2} \dots A[N]^{j_N} |j_1 j_2 \dots j_N\rangle, \quad (2.8)$$

z którego widzimy że amplituda prawdopodobieństwa dla danego wektora bazy (ustalonych wartości j_n) jest określona przez iloczyn (produkt) N macierzy, stąd nazwa macierzowy stan produktowy. Aby w efekcie takiego zwięzania otrzymać wielkość skalarną, na końcach łańcucha znajdują się specjalne wersje macierzy – wektory.

Największą wartość χ_n oznaczmy χ , czyli $\chi = \max_n \chi_n$, opisuje ona wymiar wiązania całej sieci tensorowej. Zauważmy, że gdy $\chi = 1$ to mamy do czynienia ze stanem produktowym (wtedy i tylko wtedy). Wymiar wiązania jest kluczowym parametrem sieci tensorowej, ściśle powiązany z maksymalną ilością splątania, którą dana sieć zawiera, co zaraz pokażemy. Wyznamy ograniczenie górne na entropię splątania stanu w reprezentacji MPS, przy podziale łańcucha N cząstek na dwa podukłady – podukład A stanowi pierwszych n cząstek łańcucha, a podukład B pozostałe $N - n$ cząstek. Zaczniemy od zapisania stanu układu w postaci

$$|\psi\rangle = \sum_{j_1, \dots, j_N=0}^{d-1} A[1]^{j_1} \dots A[N]^{j_N} |j_1 \dots j_N\rangle = \sum_{\alpha_n=1}^{\chi_n} |\alpha_n\rangle_A |\alpha_n\rangle_B, \quad (2.9)$$

gdzie wprowadziliśmy oznaczenia

$$|\alpha_n\rangle_A = \sum_{j_1, \dots, j_n=0}^{d-1} \left(A[1]^{j_1} \dots A[n]^{j_n} \right)_{1, \alpha_n} |j_1 \dots j_n\rangle, \quad (2.10)$$

$$|\alpha_n\rangle_B = \sum_{j_{n+1}, \dots, j_N=0}^{d-1} \left(A[n+1]^{j_{n+1}} \dots A[N]^{j_N} \right)_{\alpha_n, 1} |j_{n+1} \dots j_N\rangle. \quad (2.11)$$

Mimo że zapis (2.9) przypomina rozkład Schmidta to w ogólności nim nie jest, ponieważ bazy $\{|\alpha_n\rangle_A\}_{\alpha_n}$ i $\{|\alpha_n\rangle_B\}_{\alpha_n}$ odpowiednio podukładu A i B zazwyczaj nie są ortonormalne. Zredukowana do podukładu A macierz gęstości ma postać

$$\rho_A = \sum_{\alpha_n, \alpha'_n=1}^{\chi_n} X_{\alpha_n, \alpha'_n} |\alpha_n\rangle_A \langle \alpha'_n|, \quad (2.12)$$

gdzie

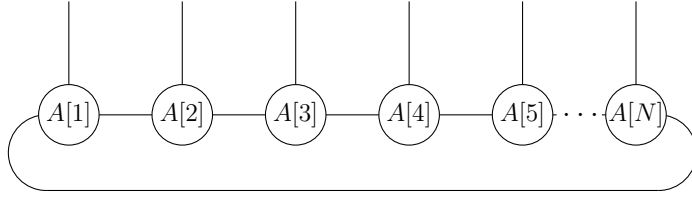
$$X_{\alpha_n, \alpha'_n} = \langle \alpha'_n | \alpha_n \rangle_B, \quad (2.13)$$

jest macierzą rzędu maksymalnie χ_n . Wynika z tego, że maksymalna wartość entropii splątania $\mathcal{S} = -\text{Tr}(\rho_A \log_2 \rho_A)$ wynosi $\log_2 \chi_n$ – gdy wszystkie wartości własne X są równe $\frac{1}{\chi_n}$. Oznacza to, że niezależnie od wyboru wielkości podukładów

$$\mathcal{S} \leq \log_2 \chi. \quad (2.14)$$

Wynik ten można uogólnić na wyżej wymiarowe sieci tensorowe – każda linia odpowiadająca wirtualnemu indeksowi, która zostanie przecięta przez powierzchnię graniczną pomiędzy dwoma podukładami daje wkład do entropii splątania równy maksymalnie $\log_2 \chi$ [Orús 2014]. Uzyskane ograniczenie pokazuje też, że entropia splątania stanu w reprezentacji MPS ma stałą wartość niezależną od długości łańcucha N , czyli spełnia 1d prawo powierzchniowe.

Prezentowana do tej pory postać MPS reprezentuje stan układu o otwartych warunkach brzegowych (OBC, ang. *Open Boundary Conditions*), jeżeli jednak chcielibyśmy uzyskać okresowe warunki brzegowe (PBC, ang. *Periodic Boundary Conditions*) to możemy je osiągnąć zastępując brzegowe wektory macierzami i łącząc ich indeksy otrzymując geometrię pierścienia:



Dodatkowa linia odpowiadająca za zwięźlenie pierwszej i ostatniej macierzy jest równoważna wprowadzeniu śladu do definicji MPS:

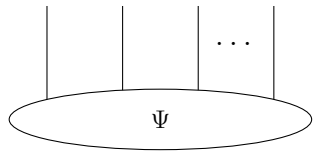
$$|\psi\rangle = \sum_{j_1, j_2, \dots, j_N=0}^{d-1} \text{Tr}(A[1]^{j_1} A[2]^{j_2} \dots A[N]^{j_N}) |j_1 j_2 \dots j_N\rangle. \quad (2.15)$$

Jeżeli w każdym oczku łańcucha mamy ten sam tensor, czyli $\forall_n A[n] = A$, to stan układu jest translacyjnie niezmienniczy (TI, ang. *Translationally Invariant*). Odwrotne stwierdzenie jest również prawdziwe – to znaczy każdy TI stan można zareprezentować za pomocą TI MPS [Pérez-García i in. 2007].

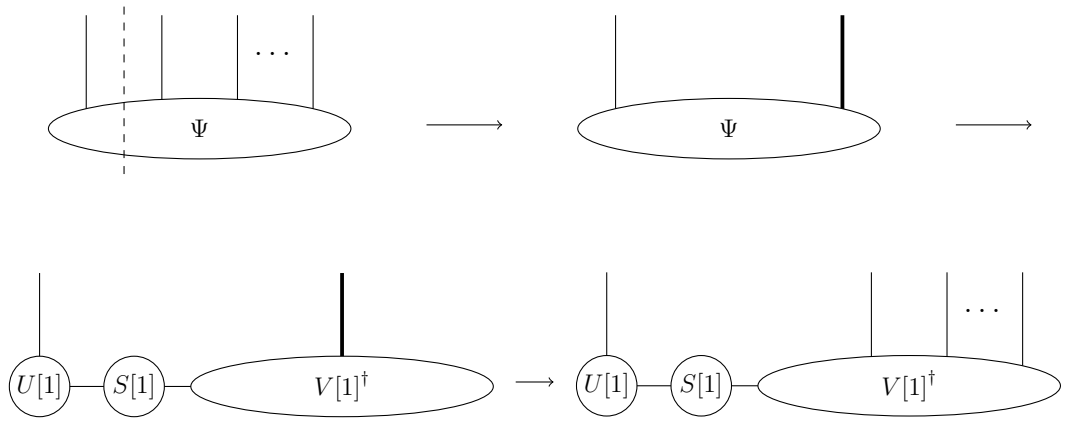
Pokażemy teraz, że dowolny stan kwantowy z N -cząstkowej przestrzeni Hilberta można zareprezentować przy użyciu MPS. Stan taki ma postać

$$|\psi\rangle = \sum_{j_1, j_2, \dots, j_N=0}^{d-1} \Psi_{j_1, j_2, \dots, j_N} |j_1 j_2 \dots j_N\rangle, \quad (2.16)$$

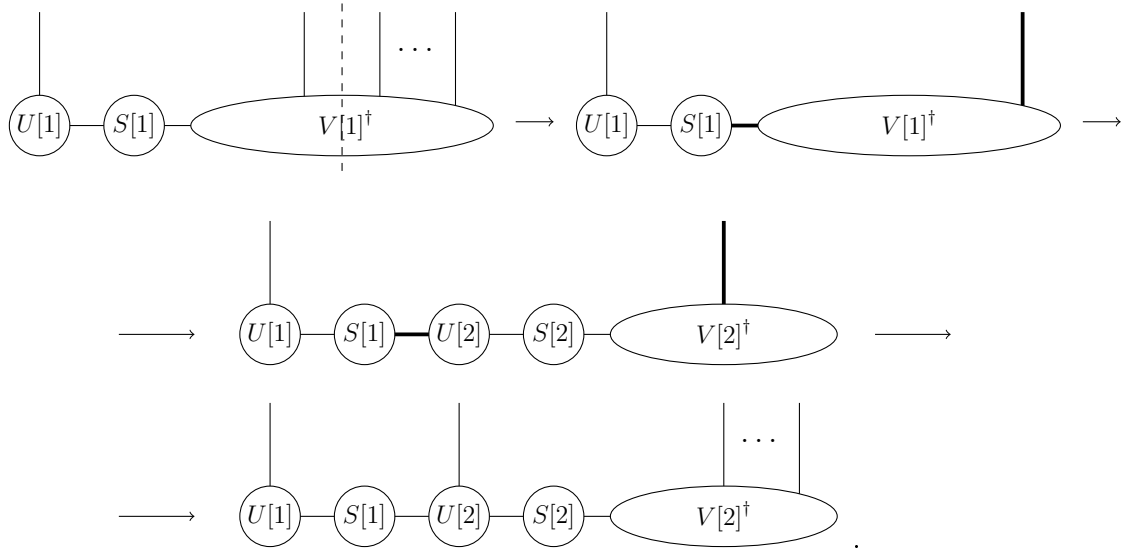
gdzie współczynniki $\Psi_{j_1, j_2, \dots, j_N}$ możemy traktować jako elementy pewnego tensora Ψ rzędu N , który diagramatycznie można przedstawić jako elipsę, z której wychodzi N linii:



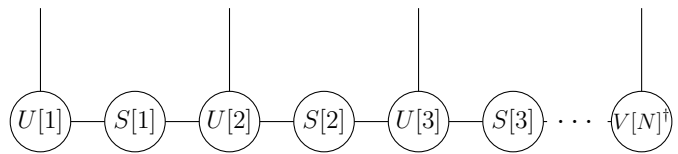
Zgrupujemy teraz wszystkie indeksy poza pierwszym, czyli od j_2 do j_N , w jeden indeks, przekształcając tym samym Ψ w macierz, do której zastosujemy SVD i następnie rozdzielimy z powrotem indeksy:



Powtarzając teraz w stosunku do $V[1]^\dagger$ operacje grupowania indeksów, SVD i rozdzielania indeksów otrzymamy:



Kontynuujemy taką procedurę dla każdego kolejnego $V[n]^\dagger$ aż otrzymamy stan w postaci



Jest to już praktycznie stan w reprezentacji MPS. Jeśli chcielibyśmy otrzymać postać identyczną jak na początku tego podrozdziału, to trzeba zwięzić ze sobą tensory $U[n]$ i $S[n]$, a wynikowy tensor oznaczyć jako $A[n]$ oraz zmienić oznaczenie tensora $V[N]^\dagger$ na $A[N]$. Rozkład ten cechuje to, że wartości singularne znajdujące się na diagonalu macierzy $S[n]$ mają ważną interpretację fizyczną – są to współczynniki Schmidta przy podziale układu pomiędzy pierwszych n cząstek i pozostałe $N - n$ cząstek, a więc jednoznacznie wyznaczają entropię

splątania pomiędzy tymi podukładami $\mathcal{S} = -\sum_{\gamma=1}^{\chi_n} (S[n])_{\gamma,\gamma}^2 \log_2 (S[n])_{\gamma,\gamma}^2$ [Vidal 2003]. Wraz ze wzrostem wymiaru wiązania χ możemy opisać coraz bardziej splątane stany kwantowe, ale aby pokryć całą przestrzeń Hilberta musimy zwiększyć wymiar wiązania aż do $d^{N/2}$. Jeżeli jednak badany stan spełnia 1d prawo powierzchniowe dla entropii splątania, jak np. stan podstawowy 1d lokalnego Hamiltonianu z przerwą energetyczną, to można go opisać przy użyciu MPS o małym wymiarze wiązania [Verstraete i J. I. Cirac 2006]. Jest to związane z tym, że dla takich stanów spektrum współczynników Schmidta zanika eksponencjalnie szybko i rzutując na podprzestrzeń powiązaną z istotnie niezerowymi wartościami singularnymi (których jest niewiele) możemy wyraźnie zredukować wymiar wiązania zachowując przy tym wiernie badany stan. Widzimy więc, że dla 1d układów kwantowych, które wykazują jedynie niewielkie lokalne korelacje, reprezentacja MPS stanowi efektywny opis, który pozwala zredukować liczbę potrzebnych parametrów z d^N (przy standardowym opisie ze pomocą pojedynczego tensora Ψ) do $d\chi^2 N$, czyli zmienić skalowanie z wykładniczego na liniowe z rozmiarem układu. Warto wspomnieć, że zaprezentowany tu rozkład stanu kwantowego jest dydaktycznie pouczający, ponieważ ilustruje typowe operacje (grupowanie/dzielenie indeksów, SVD) stosowane przy pracy z użyciem sieci tensorowych. Jednak nie jest to przykład tego jak używa się MPS w praktyce. Trzeba pamiętać, że zysk ze stosowania sieci tensorowych będzie istotny jedynie wtedy, gdy wszystkie etapy obliczeń, a nawet sam badany problem, będą wyrażone za pomocą sieci tensorowych.

Za pomocą MPS możemy nie tylko opisywać stany, które mają krótkozasięgowe korelacje, takie jak stany podstawowe lokalnych 1d Hamiltonianów z przerwą energetyczną. Możemy też opisać stany, które mają nielocalne korelacje i łamią nierówności Bella jak stan GHZ (1.46). Kluczowe do efektywnego opisu przy użyciu MPS, jest aby stan charakteryzował się niewielką entropią splątania, gdyż wtedy można go przedstawić jako MPS o wymiarze wiązania (maksymalnie) $\chi = 2^S$. Dla stanu GHZ entropia splątania wynosi jedynie $\mathcal{S}_{\text{GHZ}} = 1$ (stan maksymalnie splątany N kubitów ma entropię splątania $\mathcal{S}_{\text{max}} = N$), czyli można go wyrazić przy pomocy MPS o $\chi = 2$. Reprezentacja MPS stanu GHZ z PBC jest wyjątkowo prosta, ponieważ każdej cząstce odpowiada taki sam zestaw macierzy (nie ma zależności od n):

$$A^0 = \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix}, \quad A^1 = \begin{pmatrix} 0 & 0 \\ 0 & 1 \end{pmatrix}. \quad (2.17)$$

Reprezentacja MPS nie jest unikalna, a dany stan fizyczny może być reprezentowany przez nieskończenie wiele MPS. Transformacje MPS, które nie zmieniają fizycznego stanu, noszą nazwę transformacji cechowania. Przykładem takiej transformacji jest przekształcenie

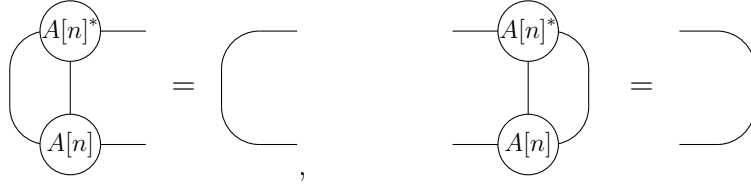
$$A[n]^{j_n} \rightarrow A[n]^{j_n} X^{-1} \quad \text{dla } n \text{ nieparzystych}, \quad (2.18)$$

$$A[n]^{j_n} \rightarrow X A[n]^{j_n} \quad \text{dla } n \text{ parzystych}, \quad (2.19)$$

gdzie od macierzy X wymagamy jedynie, aby miała lewą odwrotność, a więc może być prostokątna i zwiększać wymiar wiązania. Ta swoboda może być wykorzystana do wprowadzenia postaci kanonicznej MPS. Dla OBC można przy

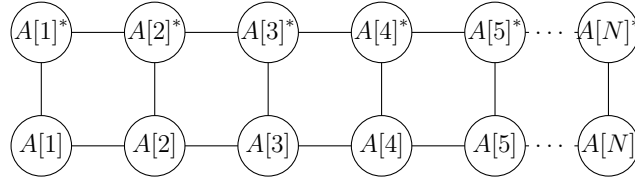
pomocy SVD łatwo sprowadzić każdy MPS do bardzo wygodnej postaci lewo- lub prawo-kanonicznej [Pérez-García i in. 2007], dla których tensory tworzące MPS spełniają warunek (odpowiednio po lewej/prawej dla postaci lewo-/prawo-kanonicznej):

$$\sum_{j_n=0}^{d-1} A[n]^{j_n \dagger} A[n]^{j_n} = \mathbb{1}, \quad \sum_{j_n=0}^{d-1} A[n]^{j_n} A[n]^{j_n \dagger} = \mathbb{1}, \quad (2.20)$$

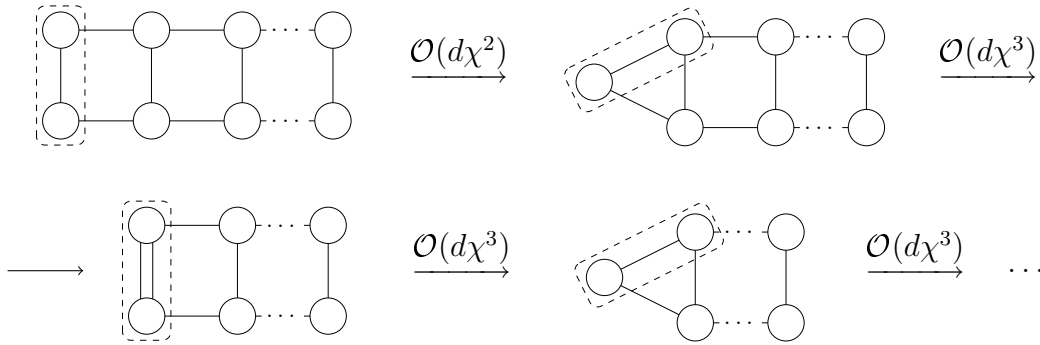


Technicznie w przypadku postaci lewo-(prawo-)kanonicznej tensory na ostatnim (pierwszym) oczku mogą nie spełniać tego warunku, ale jak za chwilę się przekonamy oznacza to, że stan nie jest unormowany. Przykładem MPS w postaci lewo-kanonicznej jest ten, który otrzymaliśmy powyżej w wyniku rozkładu dowolnego stanu kwantowego (2.16).

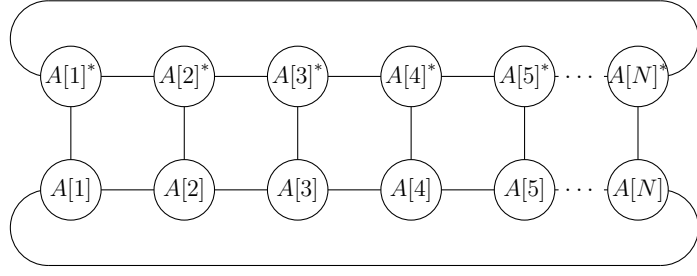
MPS stanowią bardzo efektywny sposób opisu stanów kwantowych, ale ich użyteczność byłaby znikoma, gdyby nie pozwalały równie efektywnie wyznaczać wartości oczekiwanych obserwabli. Na szczęście nie stanowi to problemu, ponieważ sieci takie można efektywnie i ściśle zwięzać, a w niektórych przypadkach w wyniku zastosowania postaci kanonicznej staje się to wyjątkowo proste. Użyjemy najprostszego możliwego przykładu, czyli obliczymy iloczyn skalarny $\langle \psi | \psi \rangle$, aby pokazać optymalną strategię zwięzania takiej sieci tensorowej i zilustrować różnicę w złożoności obliczeniowej pomiędzy stanami z OBC i PBC. W przypadku OBC $\langle \psi | \psi \rangle$ reprezentuje sieć



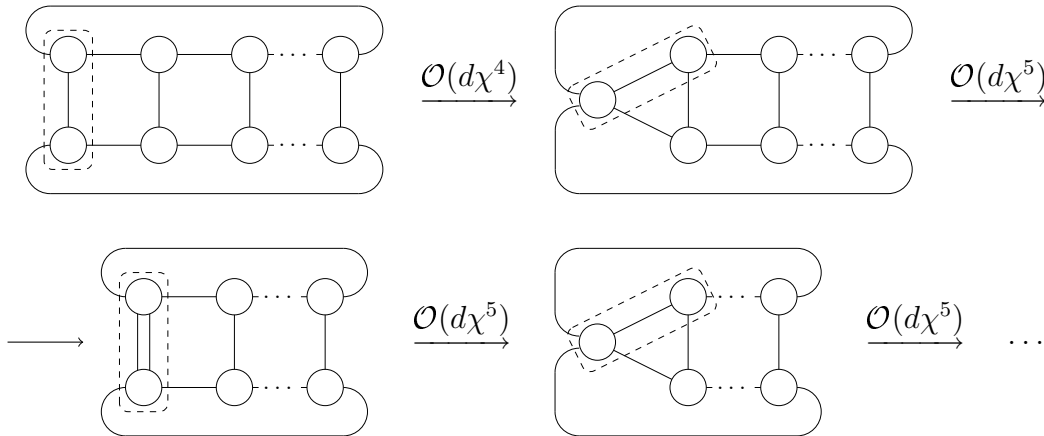
Zakładając, że wzdłuż całego łańcucha MPS wymiar wiązania χ jest taki sam to w optymalnym schemacie zwięzania konieczne jest wykonanie $\mathcal{O}(d\chi^3 N)$ operacji:



Zauważmy, że gdyby zwięzanie przebiegało „wierszami” zamiast „kolumnami” to potrzebna liczba operacji skalowałaby się wykładniczo z N . Widzimy teraz też zalety postaci kanonicznej, w której wykonanie takiego zwięzania staje się trywialne, ponieważ każda kolejna „kolumna” odpowiada identyczności oprócz ostatniej, która zawiera informację o normie stanu. W przypadku PBC w ogólności nie istnieje prosty algorytm pozwalający sprowadzić stan do wygodnej postaci kanonicznej [Verstraete, Murg i J. Cirac 2008], a z powodu braku wyróżnionego początku łańcucha zwięzanie sieci



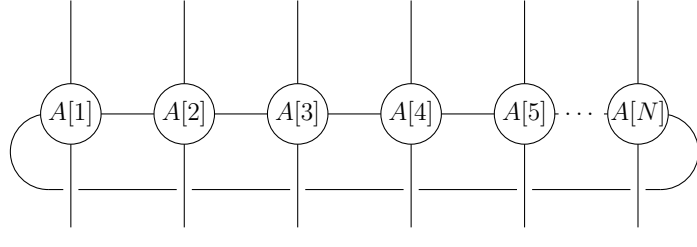
wymaga aż $\mathcal{O}(d\chi^5 N)$ operacji:



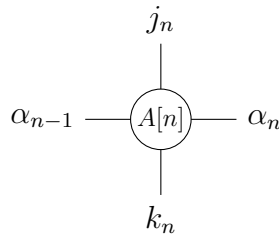
Na koniec wspomnijmy, że ważnym elementem sukcesu MPS są wydajne algorytmy numeryczne wykorzystujące MPS jako ansatz. Do najbardziej znanych należy grupa renormalizacji macierzy gęstości (DMRG, ang. *Density Matrix Renormalization Group*), która pierwotnie została zaproponowana jako procedura renormalizacji funkcji falowej stanu podstawowego 1d sieci kwantowej [White 1992; White 1993; Schollwöck 2005], a później została zinterpretowana jako metoda wariacyjnego poszukiwania stanu podstawowego w klasie funkcji falowych reprezentowanych przez MPS [Östlund i Rommer 1995; Dukelsky i in. 1998; Schollwöck 2011]. Drugi bardzo znany algorytm występuje pod akronimem TEBD (ang. *Time-Evolving Block Decimation*) [Vidal i in. 2003; Vidal 2004] i pozwala wykonać ewolucję w czasie MPS tak długo jak stan pozostaje słabo splątany. W szczególności wykonując ewolucję w czasie urojonym pozwala w sposób alternatywny do DMRG znaleźć stan podstawowy. Oba z tych algorytmów mają też swoją odmianę działającą w granicy termodynamicznej: iDMRG [McCulloch 2008] oraz iTEBD [Vidal 2007a].

2.5.2 Macierzowe operatory produktowe

Uogólnienie MPS w celu opisu operatorów kwantowych nie stanowi problemu – są to macierzowe operatory produktowe (MPO), które dla PBC mają postać



gdzie teraz tensor $A[n]$ ma dwa indeksy fizyczne j_n oraz k_n :



czyli do opisu takiego operatora potrzeba $d^2\chi^2N$ parametrów. Zapisany w postaci równania operator ten, nazwijmy go O , ma postać

$$O = \sum_{\substack{j_1, j_2, \dots, j_N, \\ k_1, k_2, \dots, k_N=0}}^{d-1} \text{Tr} \left(A[1]_{k_1}^{j_1} A[2]_{k_2}^{j_2} \dots A[N]_{k_N}^{j_N} \right) |j_1 j_2 \dots j_N\rangle \langle k_1 k_2 \dots k_N|. \quad (2.21)$$

Aby operator reprezentował macierz gęstości musi mieć ślad równy jeden, być Hermitowski oraz nieujemnie określony [Nielsen i Chuang 2000]. Odpowiednio przeskalowując tensory możemy bez problemu zapewnić odpowiednią wartość śladu. Warunek wystarczający (ale nie konieczny) na to, aby operator był Hermitowski ma postać

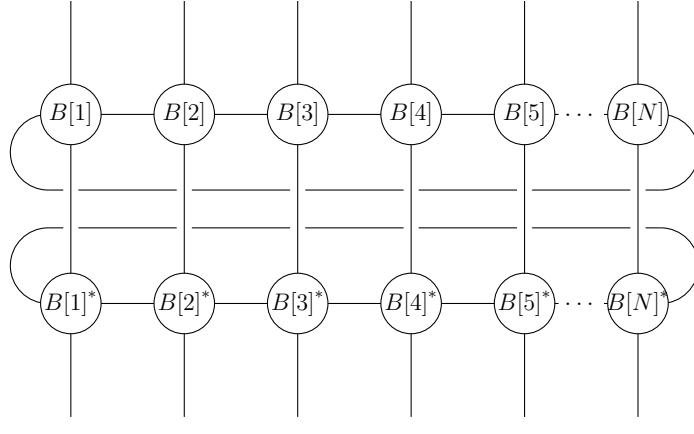
$$A[n]_{k_n}^{j_n} = A[n]_{j_n}^{k_n*} \quad (2.22)$$

i w tej pracy będziemy go nazywać cechowaniem Hermitowskim. Oczywiście można zawsze obliczyć część Hermitowską dowolnego operatora $O^H = \frac{1}{2}(O + O^\dagger)$, jednak taki operator O^H ma dwa razy większy wymiar wiązania niż operator O , ponieważ na każdym oczku tworzą go tensory $A^H[n]$, które dla ustalonych indeksów fizycznych są macierzami postaci

$$A^H[n]_{k_n}^{j_n} = \frac{1}{\sqrt{2}} \begin{pmatrix} A[n]_{k_n}^{j_n} & 0 \\ 0 & A[n]_{j_n}^{k_n*} \end{pmatrix}. \quad (2.23)$$

Niestety nieujemna określoność operatora jest własnością globalną, które nie można zagwarantować nakładając lokalne więzy na tensory tworzące MPO. Nawet sprawdzenie czy dany operator zapisany w postaci MPO jest nieujemnie

określony jest trudnym problemem [Kliesch, Gross i Eisert 2014]. Pewnym obejściem tego problemu jest wprowadzenie dodatkowej struktury do sieci tensorowej i wyrażenie stanu mieszanego za pomocą lokalnych puryfikacji, otrzymując sieć nazywaną macierzowymi produktowymi operatorami (MPDO, ang. *Matrix Product Density Operators*) [Verstraete, Garc a-Ripoll i J. I. Cirac 2004]



gdzie tensory $A[n]$ (tworzące MPO) i $B[n]$ (tworzące MPDO) możemy ze sobą powiązać następująco:

$$A[n]_{k_n}^{j_n} = \sum_{i_n=0}^{d-1} B[n]_{i_n}^{j_n} \otimes B[n]_{i_n}^{k_n*}. \quad (2.24)$$

Wadą MPDO jest to, że stan z taką samą ilością splątania jak w MPO wymaga większego wymiaru wiązania [Cuevas i in. 2013].

Ze względu na produktową strukturę MPO mogłoby się wydawać niemożliwym ściśle zareprezentowanie lokalnych Hamiltonianów zawierających sumę po wszystkich cząstkach, takich jak np. Hamiltonian dla modelu Isinga (na pierścieniu) w polu poprzecznym

$$H = -J \sum_{n=1}^N \sigma_x^{[n]} \sigma_x^{[n+1]} - h \sum_{n=1}^N \sigma_z^{[n]}. \quad (2.25)$$

Widzieliśmy przecież przed chwilą, że sumie operatorów odpowiada MPO o zwielokrotnionym wymiarze wiązania. Z drugiej strony lokalna struktura Hamiltonianu gwarantuje istnienie reprezentacji MPO o niskim wymiarze wiązania. Okazuje się, że istnieje prosty i ścisły sposób na przedstawienie takich Hamiltonianów za pomocą MPO [McCulloch 2007], np. powyższy Hamiltonian (PBC)

można przedstawić w postaci śladu z iloczynu

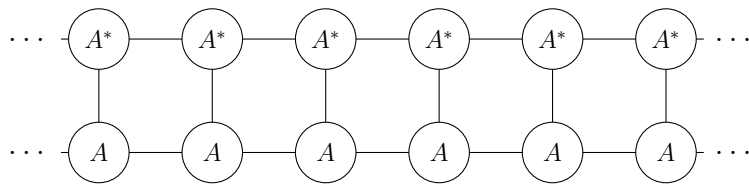
$$A[n] = \begin{cases} \begin{pmatrix} -h\sigma_z & -J\sigma_x & \mathbb{1} \\ 0 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix} & \text{dla } n = 1, \\ \begin{pmatrix} \mathbb{1} & 0 & 0 \\ \sigma_x & 0 & 0 \\ -h\sigma_z & -J\sigma_x & \mathbb{1} \end{pmatrix} & \text{dla } n \in \{2, \dots, N-1\}, \\ \begin{pmatrix} \mathbb{1} & 0 & 0 \\ \sigma_x & 0 & 0 \\ \sigma_z & 0 & 0 \end{pmatrix} & \text{dla } n = N. \end{cases} \quad (2.26)$$

Wstawiając za macierze Pauliego σ_i ich elementy macierzowe pomiędzy wektorami bazy odpowiadającymi indeksom fizycznym otrzymamy macierze $A[n]_{k_n}^{j_n}$.

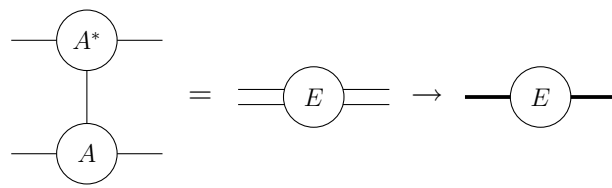
2.5.3 Granica termodynamiczna

Sieci tensorowe służą nie tylko do opisu układów skończonych. Można również za ich pomocą opisać układy w granicy termodynamicznej, pod warunkiem że układ jest translacyjnie niezmienniczy (TI). W przypadku 1d mówimy wtedy o nieskończonych macierzowych stanach produktowych (iMPS, ang. *infinite Matrix Product State*) oraz nieskończonych macierzowych operatorach produktowych (iMPO, ang. *infinite Matrix Product Operators*) [Vidal 2007a; Orús i Vidal 2008].

Zademonstrujemy teraz, jak w granicy termodynamicznej wygląda przykład obliczania $\langle \psi | \psi \rangle$, co pozwoli nam pokazać cechy charakterystyczne podejścia asymptotycznego. Zaczniemy od przedstawienia diagramu sieci tensorowej odpowiadającej $\langle \psi | \psi \rangle$:



Wprowadzimy teraz macierz transferu, którą zdefiniujemy jako minimalny fragment sieci, który powtarzany odtworzy ją całą (nasza definicja jest trochę odmienna od typowo występującej w literaturze [Bridgeman i Chubb 2017], co pozwoli nam użyć koncepcji macierzy transferu również w kontekście iMPO). Dla powyższej sieci macierz transferu jest równa $E = \sum_{j=0}^{d-1} A^{j*} \otimes A^j$ lub diagramatycznie:



gdzie od razu zgrupowaliśmy ze sobą odpowiednio lewe oraz prawe indeksy wirtualne, aby przekształcić tensor E w macierz. Możemy teraz zapisać sieć tensorową odpowiadającą $\langle\psi|\psi\rangle$ w postaci

$$\cdots \text{---} \bigcirc \text{---} \bigcirc \text{---} \bigcirc \text{---} \bigcirc \text{---} \bigcirc \text{---} \bigcirc \text{---} \cdots,$$

co oznacza, że $\langle\psi|\psi\rangle = \lim_{N \rightarrow \infty} \text{Tr } E^N$. W celu obliczenia tego wyrażenia zauważmy, że jeśli E jest diagonalizowalną macierzą (nie musi być Hermitowska) to możemy ją rozłożyć na

$$E = \sum_j \lambda_j |r_j\rangle \langle l_j|, \quad (2.27)$$

gdzie $|r_j\rangle$ i $\langle l_j|$ ³ to odpowiednio prawe i lewe wektory własne E , które można znormalizować do $\langle l_j|r_k\rangle = \delta_{j,k}$ (lewy wektor własny to prawy wektor własny macierzy transponowanej). Ten rozkład możemy przedstawić na diagramie jako

$$\text{---} \bigcirc \text{---} = \sum_j \lambda_j \text{---} \bigcirc \bigcirc \text{---}$$

Jeśli wiodąca wartość własna jest niezdegenerowana, to gdy $N \rightarrow \infty$ dominujący wkład do E^N pochodzi od tej wartości własnej:

$$E^N \simeq \lambda_1^N |r_1\rangle \langle l_1|. \quad (2.28)$$

Oznacza to, że $\langle\psi|\psi\rangle = \lim_{N \rightarrow \infty} \lambda_1^N$ i aby stan był unormowany musimy tak przeskalować tensor A , aby wiodąca wartość własna macierzy transferu była równa jeden. W związku z tym, aby iMPS mógł reprezentować fizyczny stan będziemy żądać, żeby jego macierz transferu miała niezdegenerowaną wiodącą wartość własną.

Naturalnie jest rozważać granicę termodynamiczną dla PBC. Dla OBC musielibyśmy wprowadzić jeszcze warunki brzegowe w nieskończoności w postaci odpowiednich wektorów na końcach łańcucha. Jednakże E^k w działaniu na takie wektory brzegowe zwróci wiodący wektor własny E jeśli k jest większe niż zasięg korelacji. Dlatego dla układów o skończonym zasięgu korelacji przy badaniu zachowania asymptotycznego otrzymamy takie same rezultaty niezależnie od rozważanych warunków brzegowych.

Na koniec pokażemy jak dla iMPS wyznaczyć entropia splątania von Neumanna $\mathcal{S} = -\text{Tr}(\rho_A \log_2 \rho_A)$, pomiędzy dwoma połówkami nieskończonego łańcucha. W tym celu musimy znaleźć zredukowaną macierz gęstości ρ_A . Znając wiodące wektory własne macierzy transferu odpowiadającej sieci $\langle\psi|\psi\rangle$ (czyli macierzy E , którą skonstruowaliśmy powyżej), można wyrazić poszukiwaną zredukowaną macierz gęstości w postaci [J. I. Cirac i in. 2011]

$$\rho_A = \sqrt{l_1} r_1 \sqrt{l_1}, \quad (2.29)$$

gdzie r_1 i l_1 to macierzowa postać wektorów $|r_j\rangle$ i $\langle l_j|$. Należy pamiętać, że macierz transferu otrzymaliśmy grupując indeksy tensora czwartego rzędu E ,

³Zwróć uwagę na różnicę oznaczeń: $\langle v| = |v\rangle^\dagger$, $(v| = |v\rangle^T$.

a więc każdy indeks jest tak naprawdę podwójnym indeksem i możemy go rozdzielić na dwa indeksy.

Rozdział 3

Metrologia kwantowa w języku sieci tensorowych

Rozdział ten jest centralnym dla tej dysertacji. Pokażemy w nim jak za pomocą sieci tensorowych wyrazić problem wyznaczenia fundamentalnego ograniczenia na precyzję estymacji w metrologii kwantowej. Zaczniemy od wprowadzenia ogólnej klasy kanałów kwantowych, których analiza (poza szczególnymi przypadkami) stanowiła dotychczas duże wyzwanie dla metrologii kwantowej. Następnie pokażemy, że struktura splątania, którą wprowadzają takie kanały kwantowe może być w naturalny sposób opisana przez jednowymiarowe sieci tensorowe. Gdy już wyrazimy cały problem metrologiczny w języku sieci tensorowych, pokażemy efektywny sposób jego rozwiązania zarówno dla układów o skończonym rozmiarze, jak i w przypadku asymptotycznym.

3.1 Sformułowanie badanego modelu metrologicznego

Interesuje nas problem estymacji pojedynczego parametru w obecności dekoherencji. Rozważany przez nas układ składa się z N rozróżnialnych d -wymiarowych cząstek, czyli należy on do przestrzeni Hilberta $\mathcal{H} = \bigotimes_{n=1}^N \mathbb{C}^d$. Zakładamy, że estymowany parametr φ zostaje zapisany w stanie końcowym ρ_φ w wyniku ewolucji unitarnej zadanej przez lokalne generatory (Hamiltoniany):

$$\rho_\varphi = \Phi_\varphi[\rho_0] = e^{-iH\varphi}\Phi[\rho_0]e^{iH\varphi}, \quad H = \sum_{n=1}^N h^{[n]}, \quad (3.1)$$

gdzie $h^{[n]} = \mathbb{1}^{\otimes(n-1)} \otimes h \otimes \mathbb{1}^{\otimes(N-n)}$ jest generatorem działającym na n -tą cząstkę. Będziemy rozważać równoległe schematy metrologiczne, w których dopuszczamy najbardziej ogólne splątane stany początkowe oraz najbardziej ogólne pomiary, ale bez możliwości adaptacji.

Co wyróżnia rozważania w tej pracy to założenie, że szum reprezentowany powyżej przez kanał Φ nie musi być lokalny, co czyni ten problem szczególnie wymagającym. Jak wspomnieliśmy w rozdziale 1.3.2, obecnie istniejące metody metrologii kwantowej działają efektywnie jedynie dla nieskorelowanych modeli szumu i nie mogą być bezpośrednio użyte do badania wpływu korelacji występujących w szumie na możliwą do osiągnięcia precyzję estymacji. Co jednak istotne, bardzo wiele fizycznie interesujących układów charakteryzuje się jedynie

krótkozasięgowymi korelacjami. Daje to nadzieję na znalezienie metody radzenia sobie z takimi problemami, która będzie bardziej efektywna niż przeprowadzane „na siłę” obliczenia numeryczne w pełnej przestrzeni Hilberta.

Będziemy zakładać, że Φ może być efektywnie przedstawione jako iloczyn odwzorowań jedno- ($\Phi^{[n]}$), dwu- ($\Phi^{[n,n+1]}$), trzy- ($\Phi^{[n,n+1,n+2]}$) i tak dalej cząstkowych, aż do pewnego punktu odcięcia, po którym zakładamy, że korelacje szumu obejmujące nie więcej niż r cząstek można zaniedbać. W celu uzyskania bardziej fizycznej perspektywy, rozważmy niezależne od czasu kwantowe równanie master [Breuer i Petruccione 2002] opisujące ewolucję układu N cząstek pod wpływem szumu:

$$\frac{d\rho}{dt} = \sum_{k=1}^r \sum_{n=1}^N \mathcal{L}^{(k,n)}(\rho), \quad (3.2)$$

$$\mathcal{L}^{(k,n)}(\rho) = \sum_i \mathcal{D}[J_i^{[n,\dots,n+k-1]}](\rho), \quad (3.3)$$

gdzie

$$\mathcal{D}[c](\rho) = c\rho c^\dagger - \frac{1}{2}(\rho c^\dagger c + c^\dagger c\rho). \quad (3.4)$$

Powyżej operatory $J_i^{[n,\dots,n+k-1]}$ reprezentują różne efekty działania szumu na k sąsiadujących cząstek, a całkowity efekt takiego k -cząstkowego szumu działającego na podzbiór cząstek $[n, \dots, n+k-1]$ jest reprezentowany przez odwzorowanie $\mathcal{L}^{(k,n)}$ – zauważmy, że trzycząstkowy człon może w szczególności reprezentować dwuciałowe oddziaływanie pomiędzy najbliższymi sąsiadami. Kanał Φ można teraz otrzymać całkując ewolucję po czasie:

$$\Phi = \exp\left(\sum_{k=1}^r \sum_{n=1}^N \mathcal{L}^{(k,n)} t\right). \quad (3.5)$$

Jeżeli wszystkie człony w powyższym eksponentie komutują to możemy od razu zapisać Φ jako iloczyn odwzorowań jedno-, dwu-, trzy- i tak dalej cząstkowych, co jak się niedługo przekonamy (patrz rozdział 3.2) oznacza, że kanał taki można efektywnie zapisać jako MPO. W przeciwnym razie można przybliżyć ewolucję przez czas t , za pomocą złożenia ewolucji przez krótsze przedziały czasu, a w każdym z nich wykonać dekompozycję Suzukiego-Trottera [Trotter 1959; Suzuki 1966; Suzuki 1976].

W równaniu (3.1) milcząco założyliśmy, że szum działa przed unitarnym nakreśnieniem fazy. To nie pociąga za sobą utraty ogólności jeśli te dwa procesy komutują ze sobą, jak to ma miejsce w przypadku większości popularnych modeli metrologicznych estymacji fazy/częstości w obecności defazowania czy strat [Dorner i in. 2009; Escher, Matos Filho i Davidovich 2011; Demkowicz-Dobrzański, Kołodyński i Guţă 2012]. Jednak jeżeli wymaga tego rozważany problem to można rozszerzyć zaproponowany przez nas formalizm, tak aby uwzględnić modele, w których procesy te nie komutują, kosztem pewnego skomplikowania wzorów oraz obliczeń numerycznych. Wiąże się to z tym, że w takiej sytuacji nie będziemy w stanie zapisać pochodnej ρ_φ po parametrze w postaci komutatora z Hamiltonianem, a zamiast tego będzie potrzebne użycie całej struktury kanału Φ_φ do wyznaczenia pochodnej. Jednakże w tej pracy ograniczymy się do

modeli, w których te procesy komutują.

Jak wspomnieliśmy w rozdziale 1.3.2, w najbardziej nas interesującym reżimie bardzo wielu cząstek ograniczenia na precyzję estymacji w obecności dekoherencji uzyskane w kwantowym podejściu Bayesowskim, jak i w oparciu o QFI, pokrywają się ze sobą. W związku z tym skupimy się na wyrażeniu za pomocą sieci tensorowych problemu liczenia QFI, która ze względu na lokalną zależność od struktury kanału kwantowego i jego pierwszej pochodnej wokół estymowanego parametru lepiej nadaje się zastosowania metod sieci tensorowych. Nie oznacza to, że nie można zastosować opisanego tu formalizmu bezpośrednio do problemów Bayesowskich, wręcz przeciwnie, w szczególności jeśli mają one strukturę matematyczną zbliżoną do QFI, jak to ma miejsce np. w Bayesowskiej analizie stabilizacji działania zegara atomowego (patrz rozdział 4.2).

Co bardzo ważne, w praktycznie wszystkich modelach metrologicznych uwzględniających realistyczny wpływ dekoherencji maksymalna osiągalna wartość QFI skaluje się liniowo z N w granicy bardzo wielu cząstek i zysk kwantowych ogranicza się jedynie do stałego czynnika multiplikatywnego [Escher, Matos Filho i Davidovich 2011; Demkowicz-Dobrzański, Kołodyński i Guţă 2012]. To wskazuje, że aby osiągnąć optymalne z punktu widzenia metrologii rezultaty, można ograniczyć rozważania do klasy stanów początkowych ρ_0 , które zawierają jedynie skończenie zasięgowe korelacje, a więc takich stanów, które można efektywnie zareprezentować przy pomocy sieci tensorowych [Jarzyna i Demkowicz-Dobrzański 2013]. Pamiętając że QFI jest maksymalna dla stanu początkowego, który jest najbardziej wrażliwy na zmianę parametru oraz uwzględniając to, że interesuje nas estymacja pojedynczego parametru skalarne, można oczekiwać, że optymalny stan początkowy będzie wykazywał istotne korelacje tylko w jednym wymiarze. W związku z czym ρ_0 będzie można efektywnie zapisać w postaci MPO. Zauważmy, że z uwagi na przyjętą strukturę kanału kwantowego, jeśli stan początkowy ρ_0 jest jedynie lokalnie skorelowany, to pozostanie on takim pod wpływem rozważanej ewolucji, czyli ρ_φ też będzie posiadać efektywną reprezentację MPO.

Wszystko to sugeruje, że sieci tensorowe powinny stanowić idealną reprezentację, która pozwoli w efektywny sposób wyznaczyć QFI oraz znaleźć optymalne stany początkowe w ramach rozważanej tu klasy modeli metrologicznych ze słabo skorelowanym szumem. Liniowe skalowanie z N w granicy bardzo wielu cząstek ma też duże znaczenie, ponieważ powoduje, że będziemy mogli, podobnie jak w pracy [Jarzyna i Demkowicz-Dobrzański 2015], stwierdzić w oparciu o kwantową lokalną asymptotyczną normalność [Guta i Jencová 2007; Kahn i Guta 2009], że otrzymane ograniczenia na precyzję estymacji są asymptotycznie wysycalne. Należy podkreślić, że w przeciwieństwie do wielu metod metrologii kwantowej rozwiniętych w ostatnich latach nasze podejście wyznacza bezpośrednio wartość QFI, a nie jedynie ograniczenie na tą wartość.

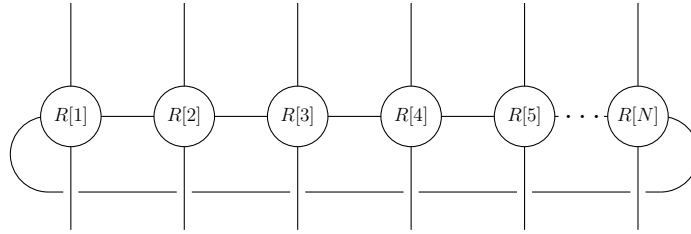
3.2 ρ_φ i $\dot{\rho}_\varphi$ w reprezentacji MPO

Fundamentalnie QFI jest wielkością lokalną zależną od macierzy gęstości na wyjściu kanału kwantowego ρ_φ i jej pierwszej pochodnej wokół wartości estymowanego parametru. Pokażemy teraz jak dla modelu metrologicznego sformułowanego w poprzednim podrozdziale można efektywnie wyrazić ρ_φ oraz jej pochodną $\dot{\rho}_\varphi$ jako MPO.

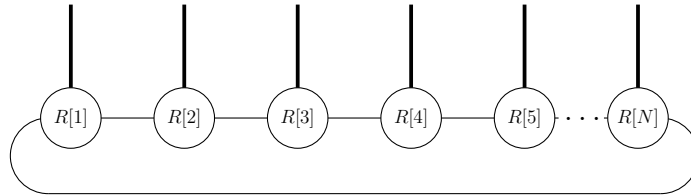
Naszym założeniem będzie, że macierz gęstości ρ_0 opisująca stan początkowy N rozróżnialnych d -wymiarowych cząstek może być efektywnie zapisana jako MPO o niewielkim wymiarze wiązania χ_{ρ_0} :

$$\rho_0 = \sum_{\mathbf{j}, \mathbf{k}} \text{Tr} \left(R[1]_{k_1}^{j_1} R[2]_{k_2}^{j_2} \dots R[N]_{k_N}^{j_N} \right) |\mathbf{j}\rangle \langle \mathbf{k}|, \quad (3.6)$$

gdzie wprowadziliśmy oznaczenie $|\mathbf{j}\rangle = |j_1, j_2, \dots, j_N\rangle$ z $j \in \{0, 1, \dots, d-1\}^1$. Powyższy stan możemy również przedstawić diagramatycznie jako



Grupując ze sobą indeksy fizyczne na każdym oczku łańcucha – czyli dokonując wektoryzacji $|j\rangle\langle k| \rightarrow |(j,k)\rangle$, gdzie nowy podwójny indeks (j,k) będziemy oznaczać też jako a – otrzymujemy reprezentację MPS stanu początkowego:



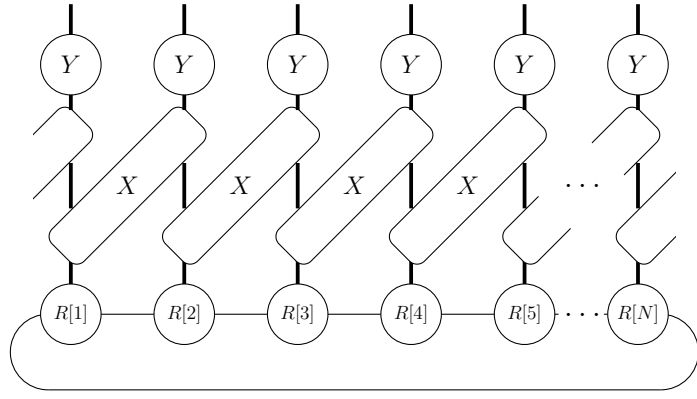
W zależności od rozważanego problemu fizycznego trzeba zastosować odpowiednie warunki brzegowe. Na potrzeby tego rozdziału wybraliśmy PBC. Jednak przejście do OBC nie sprawia trudności. Warto zaznaczyć, że przy wyznaczaniu fundamentalnych ograniczeń na precyzję głównie interesują nas własności asymptotyczne, a więc sytuacja, w której rozmiar układu jest dużo większy niż zasięg korelacji i rodzaj warunków brzegowych nie wpływa na końcowe wyniki.

Naszym pierwszym celem jest wyrażenie stanu $\rho_\varphi = \Phi_\varphi[\rho_0] = e^{-iH\varphi}\Phi[\rho_0]e^{iH\varphi}$ w postaci MPO. Zrobimy to w trzech krokach. Najpierw wykorzystamy lokalną strukturę kanału kwantowego opisującego dekoherencję (3.5), aby zbudować sieć

¹W przypadkach gdy indeks n miałby jedynie wyróżniać pozycję wynikającą ze struktury iloczynu tensorowego N -cząstkowej przestrzeni Hilberta, a omawiany obiekt w praktyce nie zależy od n , to dla przejrzystości wywodu będziemy często pomijać ten indeks (dlatego użyliśmy oznaczenia j zamiast j_n).

tensorową odpowiadającą $\Phi |\rho_0\rangle^2$, która w ogólności nie będzie miała formy MPS. Następnie używając rozkładu według wartości osobliwych (SVD) sprowadzimy tę sieć do postaci MPS. Na końcu wprowadzimy lokalny unitarny obrót, który będzie odpowiadał za zapisanie informacji o parametrze φ w naszym stanie i wrócimy z postaci MPS do MPO.

Dla ustalenia uwagi, skupmy się na sytuacji, w której operator Φ można opisać jako następujące po sobie, z zachowaniem translacyjnej niezmienniczości, działanie członów jednocząstkowych $\Phi^{[n]}$ i dwucząstkowych $\Phi^{[n,n+1]}$, które oznaczmy odpowiednio jako Y i X . Fizycznie odpowiada to sytuacji, w której jedno- i dwucząstkowe człony komutują ze sobą i żadna cząstka nie jest wyróżniona. Uogólnienie na bardziej złożone przypadki jest żmudne, ale nie wprowadza dodatkowych trudności koncepcyjnych. Za pomocą sieci tensorowych działanie takiego operatora Φ na stan początkowy można przedstawić następująco:



gdzie operator X umieściliśmy ukośnie, aby jawnie zmanifestować translacyjną niezmienniczość modelu.

Stan taki nie posiada już formy MPS. Operator X działa na dwie sąsiadujące cząstki, których bazę stanowią wektory $|a\rangle$ i $|b\rangle$ (przy czym każdy z tych indeksów jest podwójny, to znaczy $|a\rangle = |(j,k)\rangle$), a więc jest to tensor

$$X = \sum_{a',a,b',b} X_{a',a,b',b} |a'\rangle\langle a| \otimes |b'\rangle\langle b|. \quad (3.7)$$

Dokonując wektoryzacji każdego z podukładów możemy wyrazić operator X jako macierz

$$X = \sum_{(a',a),(b',b)} X_{(a',a),(b',b)} |(a',a)\rangle\langle(b',b)|, \quad (3.8)$$

co pozwala nam skorzystać z SVD, aby otrzymać

$$X_{(a',a),(b',b)} = \sum_{c,c'} U_{(a',a),c} S_{c,c'} (V^\dagger)_{c',(b',b)}, \quad (3.9)$$

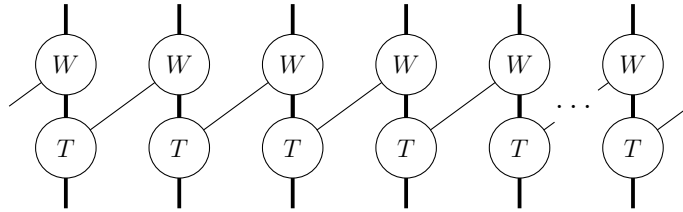
gdzie $S_{c,c'} = s_c \delta_{c,c'}$ jest macierzą diagonalną z wartościami singularnymi s_c na diagonalu, a U i V to macierze unitarne. Oznaczmy przez $\chi^{(2)}$ (górny indeks

²Jeśli tylko będzie to jasne z kontekstu to będziemy używać tego samego oznaczenia na kanał kwantowy (superoperator) działający na macierz gęstości $\Phi[\rho]$ oraz na odpowiadający mu operator działający na zwektoryzowaną macierz gęstości $\Phi|\rho\rangle$, czyli dla przykładu $|\Phi[\rho]\rangle = \Phi|\rho\rangle$.

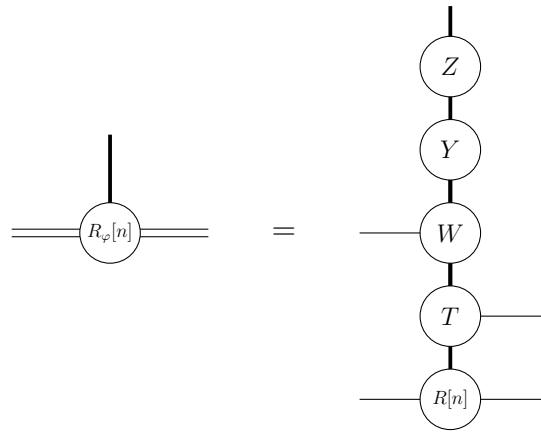
odnosi się do dwucząstkowego charakteru szumu) liczbę niezerowych wartości singularnych – w praktyce jest to liczba istotnych wartości singularnych, z pominięciem tych bardzo bliskich zeru. Graficznie dekompozycję tensora X możemy przedstawić jako

$$\begin{array}{c} a' \quad b' \\ \text{---} \text{---} \\ \boxed{X} \\ \text{---} \text{---} \\ a \quad b \end{array} = \begin{array}{c} a' \quad b' \\ \text{---} \text{---} \\ \text{---} \text{---} \\ \text{---} \text{---} \\ a \quad b \end{array} \begin{array}{c} U \\ S \\ V^\dagger \end{array},$$

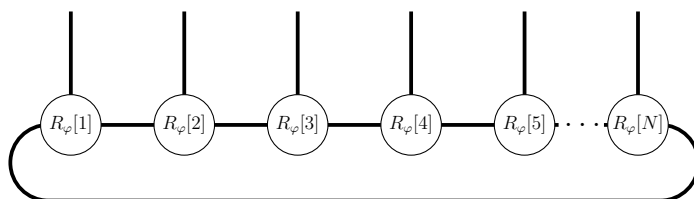
gdzie dwa indeksy po lewej stronie X traktujemy jako jeden podwójny działający na pojedynczy wirtualny układ, a dwa indeksy po prawej jako działające na drugi wirtualny układ. W ten sposób możemy myśleć o X jako o macierzy odwzorowania pomiędzy tymi dwoma wirtualnymi układami, do której następnie stosujemy SVD. Wprowadzimy teraz nowe operatory $T = U\sqrt{S}$ oraz $W = \sqrt{S}V^\dagger$, co pozwoli nam przedstawić część sieci tensorowej odpowiadającą warstwie tensorów X w postaci



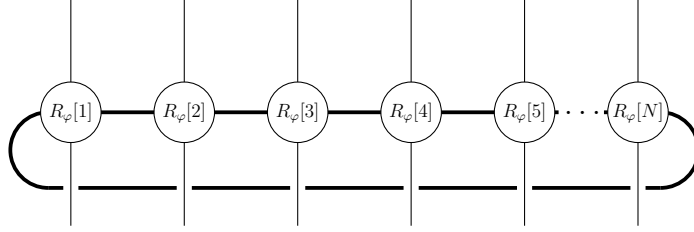
Pozostało nam już tylko połączenie tensorów $R[n]$, T , W , Y oraz tensora $Z = e^{-ih\varphi} \otimes e^{ih\varphi^T}$ – który reprezentuje unitarne nakręcanie fazy – w pojedynczy tensor $R_\varphi[n]$:



Możemy teraz połączyć podwójne poziome linie w jedną grubszą, tak aby uzyskać reprezentację MPS $|\rho_\varphi\rangle$:



Ostatecznie rozdzielać poziomą, grubszą linię odpowiadającą podwójnemu indeksowi z powrotem na dwie oryginalne linie otrzymujemy poszukiwaną reprezentację MPO ρ_φ :



Otrzymaliśmy w ten sposób MPO o wymiarze wiązania $\chi_{\rho_\varphi} = \chi_{\rho_0} \chi^{(2)}$. Uogólnienie tego wyprowadzenia na przypadek, w którym kanał kwantowy splątuje nie tylko najbliższych sąsiadów, prowadzi do reprezentacji MPO macierzy gęstości ρ_φ o wymiarze wiązania $\chi_{\rho_\varphi} = \chi_{\rho_0} \chi_r$, gdzie $\chi_r = \prod_{k=2}^r \chi^{(k)}$ reprezentuje wkład do wymiaru wiązania pochodzący od skorelowanego szumu, a $\chi^{(k)}$ jest liczbą niezerowych wartości singularnych, które pojawiają się podczas SVD, gdy rozważamy k -cząstkowy człon szumu $\Phi^{[n, \dots, n+k-1]}$ (górne ograniczenie na $\chi^{(k)}$ to $d^{2(k-1)}$).

Przejdziemy teraz do drugiej części tego podrozdziału, mianowicie pokażemy jak wyrazić operator $\dot{\rho}_\varphi$ jako MPO. Można to zrobić bardzo efektywnie korzystając z tego, że pochodną tą możemy zapisać jako komutator ρ_φ z H , gdzie H jest sumą lokalnych Hamiltonianów:

$$\dot{\rho}_\varphi = \sum_{n=1}^N i [\rho_\varphi, h^{[n]}]. \quad (3.10)$$

Zakładając że wybraliśmy bazę $|j\rangle$, tak aby tworzyły ją wektory własne lokalnego Hamiltonianu h , czyli $h = \sum_j \epsilon_j |j\rangle\langle j|$, to otrzymujemy

$$\dot{\rho}_\varphi = \sum_{\mathbf{j}, \mathbf{k}} \left[\sum_n i(\epsilon_{k_n} - \epsilon_{j_n}) \right] \langle \mathbf{j} | \rho_\varphi | \mathbf{k} \rangle |\mathbf{j}\rangle\langle \mathbf{k}|. \quad (3.11)$$

Sumę występującą w nawiasie kwadratowym możemy z łatwością wprowadzić do reprezentacji MPO ρ_φ poprzez wprowadzenie dodatkowych macierzy 2×2 (skutkuje to podwojeniem wymiaru wiązania, $\chi_{\dot{\rho}_\varphi} = 2\chi_{\rho_0} \chi_r$). Otrzymujemy wtedy poszukiwaną reprezentację MPO $\dot{\rho}_\varphi$:

$$\dot{\rho}_\varphi = \sum_{\mathbf{j}, \mathbf{k}} \text{Tr} \left(R'_\varphi[1]_{k_1}^{j_1} R'_\varphi[2]_{k_2}^{j_2} \dots R'_\varphi[N]_{k_N}^{j_N} \right) |\mathbf{j}\rangle\langle \mathbf{k}|, \quad (3.12)$$

gdzie

$$R'_\varphi[n]_{k_n}^{j_n} = \begin{cases} \begin{pmatrix} i(\epsilon_{k_1} - \epsilon_{j_1}) & 1 \\ 0 & 0 \end{pmatrix} \otimes R_\varphi[1]_{k_1}^{j_1} & \text{dla } n = 1, \\ \begin{pmatrix} 1 & 0 \\ i(\epsilon_{k_n} - \epsilon_{j_n}) & 1 \end{pmatrix} \otimes R_\varphi[n]_{k_n}^{j_n} & \text{dla } n \in \{2, \dots, N-1\}, \\ \begin{pmatrix} 1 & 0 \\ i(\epsilon_{k_N} - \epsilon_{j_N}) & 0 \end{pmatrix} \otimes R_\varphi[N]_{k_N}^{j_N} & \text{dla } n = N. \end{cases} \quad (3.13)$$

Ślad iloczynu powyższych N macierzy 2×2 odpowiada wyrażeniu w nawiasie kwadratowym z równania (3.11).

3.3 SLD w reprezentacji MPO

W celu wyznaczenia QFI poza znajomością ρ_φ i $\dot{\rho}_\varphi$ potrzebujemy jeszcze SLD Λ . Spodziewamy się, że dla modeli ze słabo skorelowanym szumem operator Λ też będzie posiadał efektywną reprezentację MPO o niewielkim wymiarze wiązania χ_Λ :

$$\Lambda = \sum_{\mathbf{j}, \mathbf{k}} \text{Tr} \left(L[1]_{k_1}^{j_1} L[2]_{k_2}^{j_2} \dots L[N]_{k_N}^{j_N} \right) |\mathbf{j}\rangle \langle \mathbf{k}|. \quad (3.14)$$

Przypuszczenie to wynika z obserwacji, że dla nieskorelowanego szumu i produktowego stanu początkowego można skonstruować *explicite* reprezentację MPO operatora Λ o wymiarze wiązania $\chi_\Lambda = 2$. W takiej sytuacji macierz gęstości na wyjściu kanału ma strukturę produktową $\rho_\varphi = \varrho_\varphi^{\otimes N}$ i SLD jest sumą lokalnych operatorów $\Lambda = \sum_{n=1}^N \lambda^{[n]}$, gdzie λ to SLD dla jednocząstkowego problemu $\dot{\varrho}_\varphi = \frac{1}{2}(\varrho_\varphi \lambda + \lambda \varrho_\varphi)$. Jak widzieliśmy w poprzednim podrozdziale w przypadku $\dot{\rho}_\varphi$, taką sumę można przedstawić jako MPO o wymiarze wiązania równym dwa.

Sposób wyznaczenia optymalnego operatora Λ , który maksymalizowałby QFI nie jest sprawą oczywistą. W zazwyczaj spotykanej definicji QFI (1.20) jest on zdefiniowany *implicite* jako rozwiązanie równania $\dot{\rho}_\varphi = \frac{1}{2}(\rho_\varphi \Lambda + \Lambda \rho_\varphi)$. Formalizm sieci tensorowych jednakże nie umożliwia efektywnego rozwiązywania takich równań. Jedną z pierwszych prób znalezienia optymalnego Λ w toku badań poprzedzających tę dysertację była próba zminimalizowania po operatorach Λ wyrażenia

$$\|2\dot{\rho}_\varphi - \rho_\varphi \Lambda - \Lambda \rho_\varphi\|^2 = 4\|\dot{\rho}_\varphi\|^2 - 2 \left[\text{Tr}(\dot{\rho}_\varphi \rho_\varphi \Lambda) + \text{Tr}(\dot{\rho}_\varphi \Lambda \rho_\varphi) + \text{Tr}(\dot{\rho}_\varphi \rho_\varphi \Lambda^\dagger) + \text{Tr}(\dot{\rho}_\varphi \Lambda^\dagger \rho_\varphi) \right] + \text{Tr}(\rho_\varphi \Lambda \Lambda^\dagger \rho_\varphi) + \text{Tr}(\rho_\varphi \Lambda^\dagger \Lambda \rho_\varphi) + 2 \text{Tr}(\rho_\varphi \Lambda \rho_\varphi \Lambda^\dagger), \quad (3.15)$$

gdzie skorzystaliśmy z tego, że ρ_φ i $\dot{\rho}_\varphi$ są Hermitowskie. Jak się zaraz przekonamy takie zagadnienia optymalizacyjne są już z powodzeniem rozwiązywane w ramach formalizmu sieci tensorowych, czego sztandarowym przykładem jest DMRG (patrz rozdział 2.5.1). Zaletą powyższego wyrażenia jest to, że minimalizujący je operator będzie zawsze Hermitowski. Do wad należy zaś duża liczba wyrazów, która znacząco utrudniła implementację numeryczną oraz zwiększyła ogólną niestabilność działania algorytmu minimalizującego (co szczególnie w

przypadku asymptotycznym uniemożliwiało uzyskanie rzetelnych wyników). W wyniku wielu eksperymentów numerycznych okazało się, że dużo lepiej sprawdza się alternatywna definicja QFI (1.27) właśnie w postaci zagadnienia optymalizacyjnego, która również w naturalny sposób pozwala połączyć liczenie QFI z optymalizacją po stanach początkowych (1.28) w bardzo efektywny iteracyjny algorytm (patrz rozdział 1.2.1). W tej definicji wielkością maksymalizowaną, którą będziemy skrótowo nazywać FoM (ang. *Figure of Merit*), jest

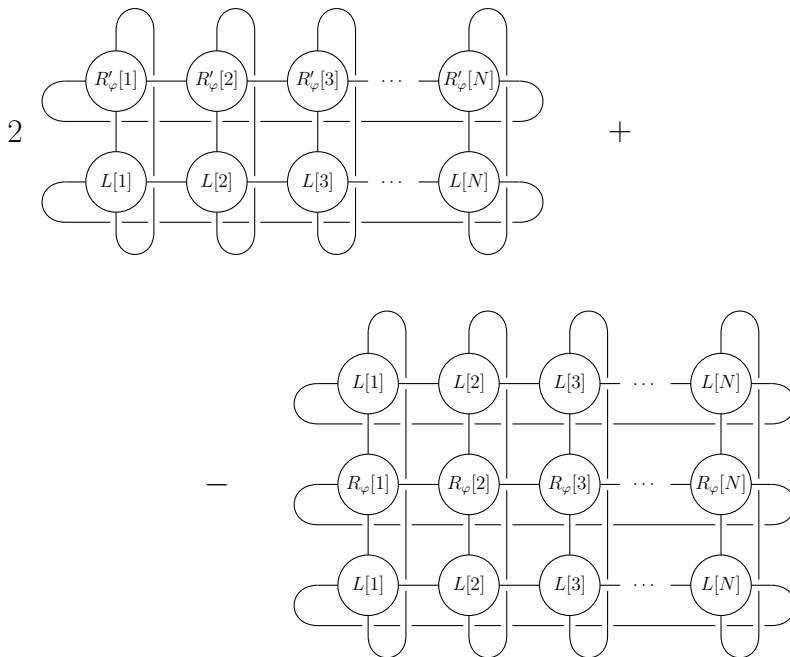
$$F[\rho_\varphi, \Lambda] = 2 \operatorname{Tr}(\dot{\rho}_\varphi \Lambda) - \operatorname{Tr}(\rho_\varphi \Lambda^2). \quad (3.16)$$

Prosta struktura tego wyrażenia ułatwiła implementację oraz poprawiła stabilność numeryczną, a brak wyrazów, które zawierałyby kwadrat macierzy gęstości sprawia, że ogólna złożoność obliczeniowa jest mniejsza, co przełożyło się na szybsze działanie algorytmu optymalizującego oraz mniejsze wymagania co do dostępnej pamięci. Główną wadą FoM jest to, że w ogólności jego maksymalna wartość może odpowiadać operatorowi Λ , który nie jest Hermitowski, co oczywiście nie jest fizycznym rozwiązaniem. Okazało się, że problem ten można rozwiązać stosując cechowanie Hermitowskie (2.22).

3.4 Optymalizacja QFI dla układów skończonych

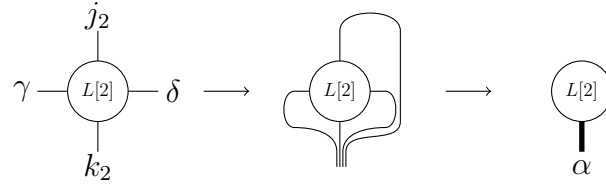
Pokażemy teraz jak efektywnie zoptymalizować FoM (3.16) po operatorach Λ i stanach początkowych ρ_0 wykorzystując formalizm sieci tensorowych. Jak wyjaśniono w rozdziale 1.2.1 będzie to dwuetapowy iteracyjny proces. Najpierw pokażemy jak zmaksymalizować FoM po Hermitowskim operatorze Λ przy ustalonym ρ_0 , a później skupimy się na maksymalizacji po stanie początkowym ρ_0 przy ustalonym Λ .

Nasze FoM, które chcemy zmaksymalizować można przedstawić diagramatycznie jako następującą sieć tensorową:

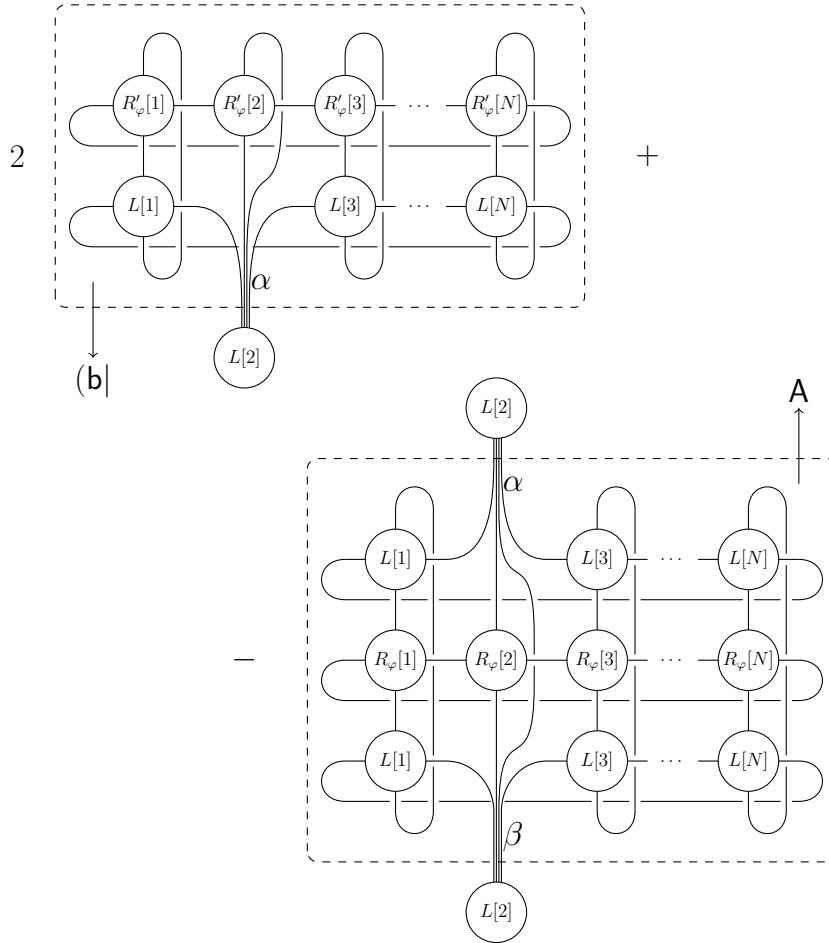


Maksymalizacja FoM po Λ jest równoważna równoczesnej maksymalizacji po wszystkich tensorach $L[n]$, co oczywiście jest bardzo wymagającym zadaniem. Inspirując się algorytmem DMRG, przeprowadzimy iteracyjną optymalizację, gdzie w każdej pętli będziemy przesuwac się wzdłuż łańcucha i po kolei optymalizować tensory $L[n]$ przy ustalonej wartości wszystkich pozostałych. Będziemy iterować takie pętle aż do uzbieżnienia FoM.

Po ustaleniu wszystkich pozostałych tensorów, FoM staje się wyrażeniem kwadratowym w $L[n]$ i optymalną wartość $L[n]$ znajdujemy rozwiązując odpowiednie równanie liniowe. Pokażemy teraz jeden ogólny krok tego algorytmu, dla ustalenia uwagi skupimy się na przykładowym tensorze $L[2]$. Zaczniemy od wykonania wektoryzacji $L[2] \rightarrow |L[2]\rangle$:



i wyróżnienia tego elementu sieci:



Pozwoli nam to zapisać FoM w postaci następującego równania:

$$F[\rho_\varphi, \Lambda] = 2(\mathbf{b}|L[2]) - (L[2]|A|L[2]) = 2 \sum_{\alpha} \mathbf{b}_{\alpha} L[2]_{\alpha} - \sum_{\alpha, \beta} L[2]_{\alpha} A_{\alpha, \beta} L[2]_{\beta}. \quad (3.17)$$

Łącznie wektor $|\mathbf{b}\rangle$ i macierz A opisują całą sieć tensorową, która otacza wyróżniony wektor $|L[2]\rangle$. Po wzięciu pochodnej FoM względem $L[2]_{\alpha}$ otrzymujemy warunek na ekstremum

$$\frac{1}{2}(\mathbf{A} + \mathbf{A}^T) |L[2]\rangle = |\mathbf{b}\rangle, \quad (3.18)$$

który ma postać równania liniowego. Macierz $\tilde{\mathbf{A}} = \frac{1}{2}(\mathbf{A} + \mathbf{A}^T)$ (o wymiarach $d^2\chi_{\Lambda}^2 \times d^2\chi_{\Lambda}^2$) zazwyczaj posiada niezerowe jądro i powyższe równanie liniowe nie ma jednoznacznego rozwiązania. W związku z tym wykonamy pseudoinwersję Moore-Penrose³ [E. H. Moore 1920; Penrose 1955] aby otrzymać rozwiązanie $|L[2]\rangle = \tilde{\mathbf{A}}^+ |\mathbf{b}\rangle$, które jest pozbawione modów zerowych $\tilde{\mathbf{A}}$ i ma przy tym najmniejszą możliwą normę.

Jeśli powyższe równanie liniowe byłoby oznaczone, to jego rozwiązanie pozostawałoby w cechowaniu Hermitowskim (2.22). Dla typowego przypadku równania nieoznaczonego będziemy używać SVD macierzy $\tilde{\mathbf{A}} = U S V^{\dagger}$, aby skonstruować jej pseudoinwersję $\tilde{\mathbf{A}}^+ = V S^+ U^{\dagger}$, przy czym odrzucimy wartości singularne poniżej pewnego małego progu odcięcia wybranego jako iloczyn parametru κ i największej wartości singularnej. Jako że tak uzyskane rozwiązanie na $|L[2]\rangle$ może nie spełniać ściśle warunku cechowania Hermitowskiego, to w celu usunięcia jego małej części anty-Hermitowskiej zastosujemy podstawienie

$$L[2]_{k_2}^{j_2} \rightarrow \frac{1}{2}(L[2]_{k_2}^{j_2} + L[2]_{j_2}^{k_2*}). \quad (3.19)$$

Z doświadczenia, wiemy że takie podstawienie poprawia numeryczną stabilność, ale nie jest konieczne jeśli wszystkie początkowe tensory $L[n]$ był w cechowaniu Hermitowskim i parametr κ został wybrany na tyle duży, aby wyciąć całą część anty-Hermitowską z rozwiązania. Jednakże za duża wartość tego parametru sprawia, że w wyniku optymalizacji nie osiągniemy maksymalnej możliwej wartości QFI. Dlatego dobieramy κ , tak aby otrzymać najwyższą możliwą wartość QFI bez doprowadzania do pojawienia się niestabilności numerycznych.

Przejdziemy teraz do drugiej części tego podrozdziału, mianowicie maksymalizacji FoM po stanie początkowym ρ_0 przy ustalonym Λ . Zaczniemy od przepisania wyrażenia na FoM:

$$\begin{aligned} F[\Phi_{\varphi}[\rho_0], \Lambda] &= 2 \text{Tr}(\dot{\rho}_{\varphi} \Lambda) - \text{Tr}(\rho_{\varphi} \Lambda^2) = \\ &= 2 \text{Tr}(i[\Phi_{\varphi}[\rho_0], H] \Lambda) - \text{Tr}(\Phi_{\varphi}[\rho_0] \Lambda^2) = \\ &= 2 \text{Tr}(\rho_0 i[H, \Phi_{\varphi}^D[\Lambda]]) - \text{Tr}(\rho_0 \Phi_{\varphi}^D[\Lambda^2]), \end{aligned} \quad (3.20)$$

gdzie w ostatniej równości przeszliśmy do obrazu Heisenberga (w którym to operatory ewoluują zamiast stanów), a $\Phi_{\varphi}^D[\cdot]$ oznacza kanał dualny do $\Phi_{\varphi}[\cdot]$, reprezentujący ewolucję operatorów w obrazie Heisenberga. Operator Λ po przejściu

³Pseudoinwersję macierzy M oznaczamy przez M^+ .

przez kanał $\Phi_\varphi^D[\cdot]$ oznaczmy Λ_φ ,

$$\Lambda_\varphi = \Phi_\varphi^D[\Lambda] = e^{iH\varphi} \Phi^D[\Lambda] e^{-iH\varphi}, \quad (3.21)$$

gdzie indeks dolny φ ma służyć podkreśleniu, że operator w obrazie Heisenberga ewoluuje w przeciwnym „kierunku” do ewolucji stanu w obrazie Schrödingera – porównaj ze wzorem (3.1). Możemy teraz zapisać FoM w postaci

$$F[\Phi_\varphi[\rho_0], \Lambda] = \text{Tr}[\rho_0(2\dot{\Lambda}_\varphi - \Lambda_\varphi^{(2)})], \quad (3.22)$$

gdzie wprowadziliśmy $\dot{\Lambda}_\varphi = \partial\Lambda_\varphi/\partial\varphi = i[H, \Lambda_\varphi]$ oraz $\Lambda_\varphi^{(2)} = \Phi_\varphi^D[\Lambda^2]$. Poprzez analogię z konstrukcją reprezentacji MPO dla $\rho_\varphi = \Phi_\varphi[\rho_0]$ i $\dot{\rho}_\varphi = i[\rho_\varphi, H]$ opisaną w rozdziale 3.2, możemy z łatwością skonstruować reprezentację MPO operatorów $\dot{\Lambda}_\varphi$ i $\Lambda_\varphi^{(2)}$ na podstawie znanej postaci MPO operatora Λ . Tensory wyznaczające $\dot{\Lambda}_\varphi$ i $\Lambda_\varphi^{(2)}$ w reprezentacji MPO oznaczmy odpowiednio przez L'_φ i $L_\varphi^{(2)}$, a odpowiadające im wymiary wiązania przez $\chi_{\dot{\Lambda}_\varphi} = 2\chi_\Lambda\chi_r$ i $\chi_{\Lambda_\varphi^{(2)}} = \chi_\Lambda^2\chi_r$.

Wzór (3.22) pokazuje, że FoM jest maksymalne, gdy ρ_0 jest rzutem na wektor własny odpowiadający największej wartości własnej Hermitowskiego operatora $2\dot{\Lambda}_\varphi - \Lambda_\varphi^{(2)}$. W związku z czym bez straty ogólności możemy założyć, że stan początkowy jest stanem czystym $\rho_0 = |\psi_0\rangle\langle\psi_0|$, dla którego FoM przyjmuje postać

$$F[\Phi_\varphi[|\psi_0\rangle], \Lambda] = \langle\psi_0|2\dot{\Lambda}_\varphi - \Lambda_\varphi^{(2)}|\psi_0\rangle. \quad (3.23)$$

Stan $|\psi_0\rangle$ przedstawiamy jako MPS o wymiarze wiązania χ_{ψ_0} :

$$|\psi_0\rangle = \sum_{\mathbf{j}} \text{Tr}(P[1]^{j_1} P[2]^{j_2} \dots P[N]^{j_N}) |\mathbf{j}\rangle. \quad (3.24)$$

Tensory tworzące ρ_0 i $|\psi_0\rangle$ są powiązane ze sobą zależnością $R[n]_{k_n}^{j_n} = P[n]^{j_n} \otimes P[n]^{k_n*}$, co oznacza, że $\chi_{\rho_0} = \chi_{\psi_0}^2$.

Maksymalizacja FoM (3.23) po stanie początkowym $|\psi_0\rangle$ jest równoważna wariacyjnemu poszukiwaniu stanu podstawowego wielociałowego „Hamiltonianu” $\Lambda_\varphi^{(2)} - 2\dot{\Lambda}_\varphi$. Oznacza to, że sprowadziliśmy nasz problem do postaci, która w formalizmie sieci tensorowych jest rozwiązywana przy pomocy algorytmu DMRG (patrz rozdział 2.5.1). Pokażemy teraz jak wygląda działanie tego algorytmu zaadaptowane do naszego „Hamiltonianu”. Zaczniemy od zreinterpretowania problemu – naszym celem jest teraz wariacyjna minimalizacja „energii”

$$-F[\Phi_\varphi[|\psi_0\rangle], \Lambda] = \frac{\langle\psi_0|\Lambda_\varphi^{(2)} - 2\dot{\Lambda}_\varphi|\psi_0\rangle}{\langle\psi_0|\psi_0\rangle}. \quad (3.25)$$

Wykonamy tę optymalizację podobnie jak poprzednią, czyli będziemy przesuwając się wzdłuż łańcucha po kolei optymalizować tensory $P[n]$ przy ustalonej wartości wszystkich pozostałych. Będziemy iterować takie cykle aż do uzbiegnięcia „energii” (3.25).

Dla przykładu, w celu znalezienia minimum po $P[2]$, zaczniemy od zwektoryzowania tego tensora, $P[2] \rightarrow |P[2]\rangle$ i wyrażenia „energii” (3.25) jako ilorazu

Rayleigha

$$-F[\Phi_\varphi[|\psi_0\rangle], \Lambda] = \frac{\langle P[2] | \mathbf{F} | P[2] \rangle}{\langle P[2] | \mathbf{N} | P[2] \rangle}, \quad (3.26)$$

gdzie wprowadziliśmy macierze (wymiaru $d\chi_{\psi_0}^2 \times d\chi_{\psi_0}^2$) $\mathbf{F}$ i $\mathbf{N}$, które są wynikiem zwiężenia następujących sieci tensorowych:

$$\mathbf{F}_{\alpha,\beta} =$$

$$+$$

$$-2$$

$$\mathbf{N}_{\alpha,\beta} =$$

W wyniku zróżniczkowania równania (3.26) po $P[2]_{\alpha}^*$ otrzymujemy warunek na ekstremum

$$\mathbf{F} | P[2] \rangle = -F[\Phi_\varphi[|\psi_0\rangle], \Lambda] \mathbf{N} | P[2] \rangle, \quad (3.27)$$

który jest uogólnionym problemem własnym z wartością własną $-F[\Phi_\varphi[|\psi_0\rangle], \Lambda]$. Mnożąc to równanie z lewej strony przez pseudoinwersję macierzy $\mathbf{N}$ otrzymujemy zwykły problem własny, dla którego najmniejszą wartość własną oraz

odpowiadający jej wektor własny uzyskujemy stosując algorytm Lanczosa [Lanczos 1950].

Możemy bez problemu zmodyfikować powyższy algorytm, tak aby bazował na sieciach tensorowych z otwartymi warunkami brzegowymi (OBC). Jest to nawet wskazane, w szczególności dla problemów, które nie są wrażliwe na warunki brzegowe. Wynika to z tego, że sieci tensorowe z OBC charakteryzują się mniejszą złożonością obliczeniową przy ich zwężaniu oraz można dla nich łatwo sprowadzić MPS $|\psi_0\rangle$ do postaci kanonicznej (patrz rozdział 2.5.1). Wykorzystanie postaci kanonicznej sprawia, że macierz $\mathbf{N}$ staje się identycznością, a równanie (3.27) redukuje się do zwykłego zagadnienia własnego i nie ma potrzeby konstruowania pseudoinwersji macierzy $\mathbf{N}$.

Podsumowując iteracyjnie wyznaczamy optymalny Λ dla ustalonego ρ_0 , a następnie optymalny ρ_0 dla ustalonego Λ , wykonujemy to aż zaobserwujemy, że nasze FoM zbiega do ustalonej wartości to znaczy, że nie różni się o więcej niż np. 0,1% przez kilka kolejnych kroków iteracji. Nasze numeryczne doświadczenie pokazuje, że następuje to bardzo szybko, zazwyczaj już po około 5 iteracjach po Λ i ρ_0 .

Wykonując algorytm, trzeba wybrać wymiary wiązań zarówno dla stanu początkowego χ_{ψ_0} , jak i dla SLD χ_Λ , dla których optymalizacja się odbywa. Jak w każdym algorytmie bazującym na sieciach tensorowych, utrzymywanie najmniejszego możliwego wymiaru wiązania jest kluczowe dla efektywności algorytmu. Nasze obliczenia zazwyczaj zaczynamy ze stanu produktowego, czyli $\chi_{\psi_0} = 1$ i optymalizujemy Λ z wymiarem wiązania $\chi_\Lambda = 2$. Następnie zwiększamy jeden z wymiarów wiązania, χ_{ψ_0} albo χ_Λ , za każdym razem powtarzając procedurę optymalizacji, aż okaże się, że QFI nie rośnie wraz ze wzrostem wymiarów wiązań χ o więcej niż np. 1% i wtedy przyjmujemy, że względny błąd naszej metody wynosi około 1%.

3.5 Optymalizacja QFI w przypadku asymptotycznym

W poprzednim podrozdziale opisaliśmy algorytm, który działa dla układów o skończonej liczbie cząstek N . Dla modeli metrologicznych z dekoherencją typową sytuacją jest to, że asymptotycznie (dla dużych N) użycie stanów splątanych prowadzi do poprawy precyzji maksymalnie o stały czynnik, ponad to co można osiągnąć strategiami używającymi stanów produktowych. Mimo że podejście oparte na skończonych sieciach tensorowych pozwala nam osiągać wartości N znacznie przekraczające te, które można osiągnąć metodami, w których wykonuje się obliczenia w pełnej przestrzeni Hilberta, to czasami może to być ciągle za mało, aby osiągnąć granicę asymptotyczną i wyznaczyć współczynnik zysku kwantowego z oczekiwaną precyzją. Z tego powodu chcielibyśmy mieć procedurę, która pozwoli nam przejść bezpośrednio do granicy nieskończonej wielu cząstek, wyznaczyć maksymalną możliwą wartość QFI na cząstkę i w rezultacie otrzymać współczynnik zysku kwantowego.

W tym celu użyjemy podejścia wykorzystującego nieskończone MPO/MPS (iMPO/iMPS). Będziemy zakładać, że wszystkie tensory są translacyjnie niezmiennicze (TI). Technicznie to podejście jest najbardziej naturalne w przypadku PBC. Jednakże jak argumentowaliśmy w rozdziale 2.5.3 ograniczenie rozważań do PBC nie prowadzi do utraty ogólności w przypadku rozważanych przez nas modeli metrologicznych ze skończonym zasięgiem korelacji.

Wymaganie TI nie stanowi problemu w przypadku macierzy gęstości ρ_φ – jeśli przygotujemy układ w stanie, który jest TI to pozostanie on takim po przejściu przez kanał kwantowy Φ_φ . Problem jest jednak z operatorem $\dot{\rho}_\varphi$, którego reprezentacja zaproponowana równaniami (3.12, 3.13) nie jest jawnie TI. Z tego powodu, zamiast wyznaczać ściśle tą pochodną, przybliżymy ją przez różnicę dwóch TI iMPO:

$$\dot{\rho}_\varphi = \frac{\rho_{\varphi+\varepsilon} - \rho_\varphi}{\varepsilon} \quad (3.28)$$

z infitezymalnym parametrem ε . Inspirując się równaniem definiującym SLD (1.21) zastosujemy podobne rozwinięcie do operatora Λ :

$$\Lambda = \frac{\tilde{\Lambda} - \mathbb{1}}{\varepsilon}. \quad (3.29)$$

Tutaj $\tilde{\Lambda}$ oraz $\mathbb{1}$ są rozwiązaniami równania (1.21), gdy $\dot{\rho}_\varphi$ zastąpimy odpowiednio przez $\rho_{\varphi+\varepsilon}$ oraz ρ_φ . Będziemy szukać optymalnego operatora $\tilde{\Lambda}$, który lepiej nadaje się do TI formalizmu iMPO niż oryginalny operator Λ . Oznaczmy przez

$$f[\rho_\varphi, \tilde{\Lambda}] = \frac{F[\rho_\varphi, \Lambda]}{N} \quad (3.30)$$

wartość FoM na cząstkę, którą chcemy zmaksymalizować, a która po podstawieniu równań (3.28) i (3.29) do wzoru (3.16) przybiera postać

$$f[\rho_\varphi, \tilde{\Lambda}] = \frac{1}{N\varepsilon^2} [2 \text{Tr}(\rho_{\varphi+\varepsilon} \tilde{\Lambda}) - \text{Tr}(\rho_\varphi \tilde{\Lambda}^2) - 1]. \quad (3.31)$$

TI iMPO jest zdefiniowany jednoznacznie przez pojedynczy tensor, który oznaczmy odpowiednio przez: $|\psi_0\rangle \rightarrow P$, $\rho_0 \rightarrow R$, $\rho_\varphi \rightarrow R_\varphi$, $\tilde{\Lambda} \rightarrow \tilde{L}$, $\tilde{\Lambda}_\varphi \rightarrow \tilde{L}_\varphi$, $\tilde{\Lambda}_\varphi^{(2)} \rightarrow \tilde{L}_\varphi^{(2)}$.

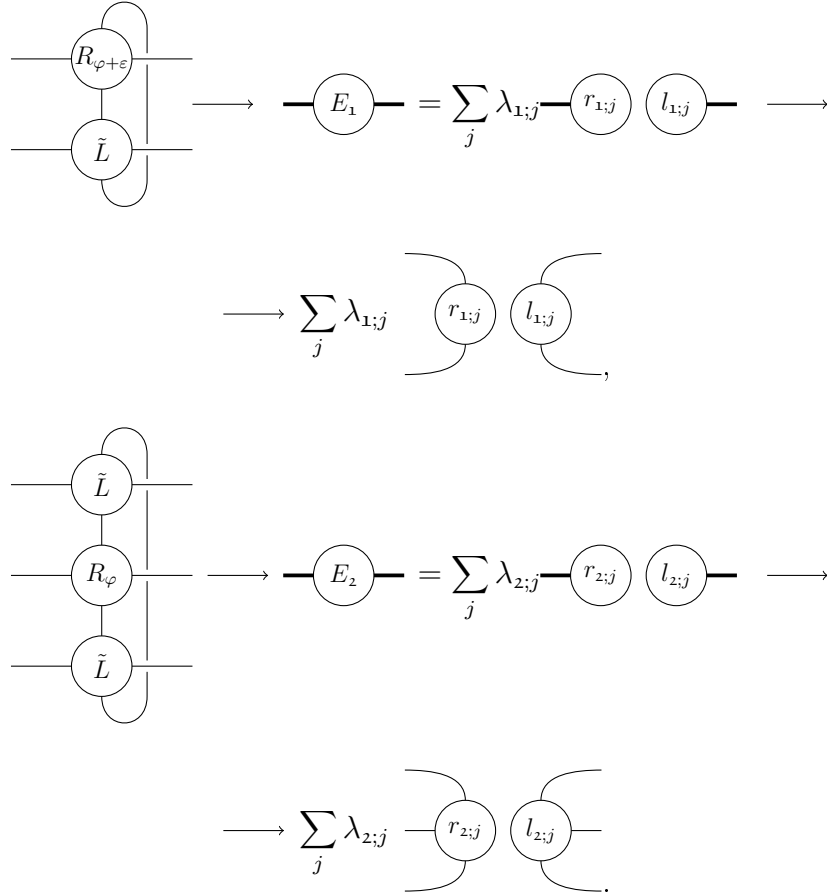
Optymalizacja sieci tensorowej złożonej z identycznych tensorów A jest wysoce nieliniowym problemem w A i można by myśleć, że podejście podobne do opisanego w poprzednim podrozdziale nie będzie miało tutaj zastosowania. Na szczęście, istnieje efektywna metoda radzenia sobie z takimi problemami, opisana w pracy [Corboz 2016]. Główna idea polega na znalezieniu optymalnego tensora A_{new} dla jednej pozycji w sieci, traktując wszystkie pozostałe A jako stałe, a następnie zamiast w miejsce wszystkich tensorów podstawić A_{new} , należy wykonać następujące elastyczne podstawienie:

$$A \rightarrow A_{\text{new}} \sin(\xi\pi) - A \cos(\xi\pi), \quad (3.32)$$

z współczynnikiem mieszania ξ , który należy zoptymalizować.

Pokażemy teraz jak zastosować to podejście do obu kroków naszej iteracyjnej procedury. Zaczniemy od optymalizacji po $\tilde{\Lambda}$ przy ustalonym ρ_0 , co jest równoważne wyznaczeniu optymalnego lokalnego tensora $\tilde{L}$.

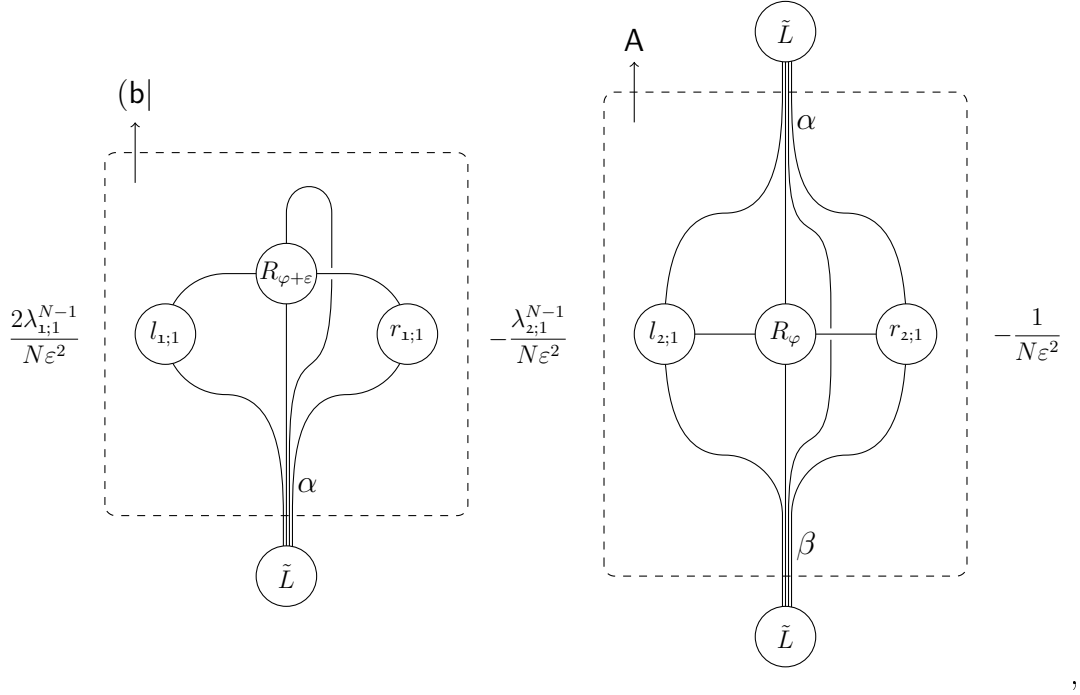
Jako że ślad operatora reprezentowanego przez iMPO jest zdefiniowany przez jego macierz transferu (patrz rozdział 2.5.3), to zaczniemy od wprowadzenia macierzy transferu E_1 i E_2 powiązanych odpowiednio z $\rho_{\varphi+\varepsilon}\tilde{\Lambda}$ i $\rho_{\varphi}\tilde{\Lambda}^2$ oraz dokonamy ich rozkładu widmowego:



Pozwoli nam to zapisać, że $\text{Tr}(\rho_{\varphi+\varepsilon}\tilde{\Lambda}) = \text{Tr } E_1^N$ oraz $\text{Tr}(\rho_{\varphi}\tilde{\Lambda}^2) = \text{Tr } E_2^N$, gdzie $E_i = \sum_j \lambda_{i;j} |r_{i;j}\rangle \langle l_{i;j}|$. Teraz korzystając z faktu, że E_i^N gdy $N \rightarrow \infty$ jest zdeterminowane przez wiodącą wartość własną (wartości własne są uporządkowane, tak że jest to pierwsza wartość własna) możemy zapisać, iż $\text{Tr } E_i^N \simeq \lambda_{i;1}^N = \lambda_{i;1}^{N-1} \langle l_{i;1} | E_i | r_{i;1} \rangle$ i wyrazić FoM na cząstkę (3.31) jako

$$\frac{2\lambda_{1;1}^{N-1}}{N\varepsilon^2} \left(\text{Tensor Network 1} \right) - \frac{\lambda_{2;1}^{N-1}}{N\varepsilon^2} \left(\text{Tensor Network 2} \right) - \frac{1}{N\varepsilon^2}$$

Wektoryzując tensor $\tilde{L}$ możemy FoM na cząstkę zapisać równoważnie w postaci



lub jako następujące równanie:

$$f[\rho_{\varphi}, \tilde{\Lambda}] = \frac{1}{N\varepsilon^2} [2\lambda_{1;1}^{N-1} (\mathbf{b}|\tilde{L}) - \lambda_{2;1}^{N-1} (\tilde{L}|A|\tilde{L}) - 1], \quad (3.33)$$

które ma ekstremum, gdy

$$\frac{1}{2}\lambda_{1;1}^{N-1}(\mathbf{A} + \mathbf{A}^T)|\tilde{L}) = \lambda_{2;1}^{N-1}|\mathbf{b}). \quad (3.34)$$

Dla $N \rightarrow \infty$ może wydawać się, że potęgi wartości własnych stanowią problem. Na szczęście jest on tylko pozorny. Dla ustalonego $\tilde{\Lambda}$ możemy obliczyć odpowiadającą mu wartość FoM na cząstkę

$$f[\rho_{\varphi}, \tilde{\Lambda}] = \frac{1}{N\varepsilon^2} (2\lambda_{1;1}^N - \lambda_{2;1}^N - 1). \quad (3.35)$$

Wracając teraz do wzoru na FoM (3.16) widzimy, że FoM na cząstkę powinno mieć postać $f[\rho_{\varphi}, \tilde{\Lambda}] = 2f_1 - f_2$, gdzie f_1 i f_2 są takiego samego rzędu wielkości jak asymptotyczna wartość QFI na cząstkę. Zakładając, że nasze obliczenia są w reżimie $N \rightarrow \infty$, $\varepsilon \rightarrow 0$ i $N\varepsilon^2 \rightarrow 0$ oraz pamiętając rozwinięcie dwumianowe

$$(1 + \varepsilon^2 f_i)^N = 1 + N\varepsilon^2 f_i + \mathcal{O}\left((N\varepsilon^2)^2\right), \quad (3.36)$$

należy oczekiwać, że wiodące wartości własne macierzy transferu mają postać:

$$\lambda_{1;1} = 1 + \varepsilon^2 f_1, \quad \lambda_{2;1} = 1 + \varepsilon^2 f_2, \quad (3.37)$$

co po wstawieniu ich do wzoru (3.35) i użyciu rozwinięcia dwumianowego do pierwszego rzędu daje nam dokładnie $f[\rho_\varphi, \tilde{\Lambda}] = 2f_1 - f_2$. Z drugiej strony odwracając zależności (3.37) otrzymujemy prosty wzór na asymptotyczną wartość FoM na cząstkę

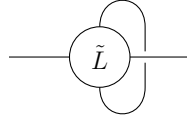
$$f[\rho_\varphi, \tilde{\Lambda}] \simeq \frac{1}{\varepsilon^2}(2\lambda_{1;1} - \lambda_{2;1} - 1). \quad (3.38)$$

Jest teraz jasne, że w rozważanym reżimie ($N \rightarrow \infty$, $\varepsilon \rightarrow 0$, $N\varepsilon^2 \rightarrow 0$) możemy na potrzeby rozwiązania równania (3.34) przybliżyć $\lambda_{1;1}^{N-1}$ i $\lambda_{2;1}^{N-1}$ przez jedynki i w rezultacie otrzymać prostą postać warunku na optymalny tensor $\tilde{L}$:

$$\frac{1}{2}(\mathbf{A} + \mathbf{A}^T) \big| \tilde{L} \rangle = |\mathbf{b}\rangle. \quad (3.39)$$

Równanie to analogicznie jak w przypadku skończonym rozwiązujemy przy użyciu pseudoinwersji i odfiltrowujemy część anty-Hermitowską w każdym kolejnym kroku iteracji.

Operator SLD Λ jest zawsze bezśladowy w każdy problemie z unitarnym nakręcaniem parametru – wynika to ze wzoru (1.23) oraz faktu, że $\langle \lambda | \dot{\rho}_\varphi | \lambda \rangle = i \langle \lambda | [\rho_\varphi, H] | \lambda \rangle = 0$ dla każdego $|\lambda\rangle$, które jest stanem własnym ρ_φ . Możemy to wykorzystać, aby zapewnić odpowiednią normalizację rozwiązania na $\tilde{L}$ – warunek $\text{Tr } \Lambda = 0$ lub równoważnie $\text{Tr } \tilde{\Lambda} = \text{Tr } \mathbf{1} = d^N$ w języku macierzy transferu oznacza, że największa wartość własna macierzy



musi być równa d . Zatem w celu zminimalizowania efektów stosowanych przybliżeń i zapewnienia prawidłowej normalizacji, po otrzymaniu $\tilde{L}$ z równania (3.39), należy je przeskalować, tak aby największa wartość powyższej macierzy transferu była równa d .

Przejdziemy teraz do drugiej części naszej iteracyjnej procedury, czyli optymalizacji po stanie początkowym $|\psi_0\rangle$ przy ustalonym $\tilde{\Lambda}$. Ten etap nie wprowadza jakościowo nowych wyzwań, więc omówimy go bardziej zwięźle. Tak jak w przypadku $\dot{\rho}_\varphi$, będziemy przybliżać ściśłą pochodną $\dot{\Lambda}_\varphi$ przez jej dyskretną wersję

$$\dot{\Lambda}_\varphi = \frac{\Lambda_{\varphi+\varepsilon} - \Lambda_\varphi}{\varepsilon} \quad (3.40)$$

oraz dokonamy rozwinięcia $\Lambda = (\tilde{\Lambda} - \mathbf{1})/\varepsilon$.

Nasze zadanie możemy wyrazić jako wariacyjną minimalizację następującej „gęstości energii”:

$$-f[\Phi_\varphi[|\psi_0\rangle], \tilde{\Lambda}] = \frac{\langle \psi_0 | \tilde{\Lambda}_\varphi^{(2)} - 2\tilde{\Lambda}_{\varphi+\varepsilon} + \mathbf{1} | \psi_0 \rangle}{N\varepsilon^2 \langle \psi_0 | \psi_0 \rangle} \quad (3.41)$$

po $|\psi_0\rangle$. Używając macierzy transferu E_3 , E_4 i E_5 powiązanych odpowiednio z $\langle \psi_0 | \tilde{\Lambda}_\varphi^{(2)} | \psi_0 \rangle$, $\langle \psi_0 | \tilde{\Lambda}_{\varphi+\varepsilon} | \psi_0 \rangle$ i $\langle \psi_0 | \psi_0 \rangle$:

$$\begin{array}{c}
\text{---} \circ P \text{---} \\
| \\
\text{---} \circ \tilde{L}_{\varphi}^{(2)} \text{---} \\
| \\
\text{---} \circ P^* \text{---}
\end{array}
= \sum_j \lambda_{3;j}
\begin{array}{c}
\text{---} \circ r_{3;j} \text{---} \\
\text{---} \circ l_{3;j} \text{---}
\end{array},$$

$$\begin{array}{c}
\text{---} \circ P \text{---} \\
| \\
\text{---} \circ \tilde{L}_{\varphi+\varepsilon} \text{---} \\
| \\
\text{---} \circ P^* \text{---}
\end{array}
= \sum_j \lambda_{4;j}
\begin{array}{c}
\text{---} \circ r_{4;j} \text{---} \\
\text{---} \circ l_{4;j} \text{---}
\end{array},$$

$$\begin{array}{c}
\text{---} \circ P \text{---} \\
| \\
\text{---} \circ P^* \text{---}
\end{array}
= \sum_j \lambda_{5;j}
\begin{array}{c}
\text{---} \circ r_{5;j} \text{---} \\
\text{---} \circ l_{5;j} \text{---}
\end{array},$$

możemy zapisać prawą stronę równania (3.41) diagramatycznie jako

$$\begin{array}{c}
\lambda_{3;1}^{N-1} \begin{array}{c} \text{---} \circ P \text{---} \\ | \\ \text{---} \circ \tilde{L}_{\varphi}^{(2)} \text{---} \\ | \\ \text{---} \circ P^* \text{---} \end{array} \begin{array}{c} \text{---} \circ r_{3;1} \text{---} \\ \text{---} \circ l_{3;1} \text{---} \end{array} - 2\lambda_{4;1}^{N-1} \begin{array}{c} \text{---} \circ P \text{---} \\ | \\ \text{---} \circ \tilde{L}_{\varphi+\varepsilon} \text{---} \\ | \\ \text{---} \circ P^* \text{---} \end{array} \begin{array}{c} \text{---} \circ r_{4;1} \text{---} \\ \text{---} \circ l_{4;1} \text{---} \end{array} + \lambda_{5;1}^{N-1} \begin{array}{c} \text{---} \circ P \text{---} \\ | \\ \text{---} \circ P^* \text{---} \end{array} \begin{array}{c} \text{---} \circ r_{5;1} \text{---} \\ \text{---} \circ l_{5;1} \text{---} \end{array}
\end{array}$$

$$\begin{array}{c}
\text{---} \circ P \text{---} \\
| \\
\text{---} \circ P^* \text{---}
\end{array}
\begin{array}{c} \text{---} \circ r_{5;1} \text{---} \\ \text{---} \circ l_{5;1} \text{---} \end{array}$$

Tak jak poprzednio, spodziewamy się, że $\lambda_{i;1} = 1 + \varepsilon^2 f_i$ i na potrzeby szukania optymalnego tensora P możemy przybliżyć $\lambda_{i;1}^{N-1}$ przez jedynki. Po zróżniczkowaniu otrzymujemy warunek na ekstremum

$$F|P\rangle = gN|P\rangle, \quad (3.42)$$

gdzie $g = -f[\Phi_{\varphi}[|\psi_0\rangle], \tilde{\Lambda}]N\varepsilon^2 - 1$ jest uogólnioną wartością własną, a macierze F i N są zdefiniowane jako:

$$F_{\alpha,\beta} = \text{Diagram 1} - 2 \text{Diagram 2}$$

$$N_{\alpha,\beta} = \text{Diagram 3} = \text{Diagram 4}$$

Macierz $\mathbf{N}$ jest iloczynem tensorowym trzech macierzy, $r_{5;1}$, identyczności i $l_{5;1}$, więc jej pseudoinwersję $\mathbf{N}^+$ możemy przedstawić jako iloczyn tensorowy (pseudo-)inwersji tych mniejszych macierzy. Działając $\mathbf{N}^+$ z lewej na równanie (3.42) sprowadzamy je do postaci zwykłego problemu własnego. Po wyznaczeniu najmniejszej wartości własnej i odpowiadającego jej wektora własnego, żądamy aby $|\psi_0\rangle$ był znormalizowany, co oznacza, że $\lambda_{5;1} = 1$. Asymptotyczną wartość FoM na cząstkę (z dokładnością do znaku) wyznaczamy korzystając ze wzoru

$$-f[\Phi_\varphi[|\psi_0\rangle], \tilde{\Lambda}] = \frac{1}{N\varepsilon^2}(\lambda_{3;1}^N - 2\lambda_{4;1}^N + 1) \simeq f_3 - 2f_4 = \frac{1}{\varepsilon^2}(\lambda_{3;1} - 2\lambda_{4;1} + 1). \quad (3.43)$$

Tak jak w przypadku skończonego N , iteracyjne powtarzanie optymalizacji po $\tilde{\Lambda}$ i $|\psi_0\rangle$ prowadzi nas do optymalnego rozwiązania z maksymalną wartością QFI na cząstkę.

Podczas przeprowadzania obliczeń numerycznych należy wybrać najmniejsze możliwe ε . Jednakże należy uważać aby nie doprowadzić do pojawienia się niestabilności numerycznych. Nasz ogólna strategia przy otrzymywaniu rezultatów numerycznych, które prezentujemy w następnym rozdziale, polega na obniżaniu wartości ε , aż zaobserwujemy, że wyniki uzbiegają się pozostając ciągle w reżimie, w którym obliczenia są stabilne. We wszystkich przykładach prezentowanych w tej dysertacji oznaczało to wybór $\varepsilon \sim 10^{-3} - 10^{-4}$ (niestabilności zaczynały pojawiać się dla $\varepsilon < 10^{-6}$). Należy zwrócić uwagę, że w omawianym tu podejściu z wykorzystaniem iMPO wymagamy, aby $N\varepsilon^2$ było małe – rzędu oczekiwanej precyzji rezultatów numerycznych. Innymi słowy ustalenie oczekiwanej precyzji na poziomie np. 1% oznacza, że $N\varepsilon^2 \sim 10^{-2}$ i przy stabilizacji wyników dla $\varepsilon \sim 10^{-3} - 10^{-4}$ otrzymujemy, że $N \sim 10^4 - 10^6$. Fizycznie odpowiada to sytuacji, w której dla większych N wartość QFI na cząstkę już się nie zmienia (o więcej niż 1%) i możemy na podstawie otrzymanych rezultatów wnioskować o zachowaniu asymptotycznym dla $N \rightarrow \infty$.

Rozdział 4

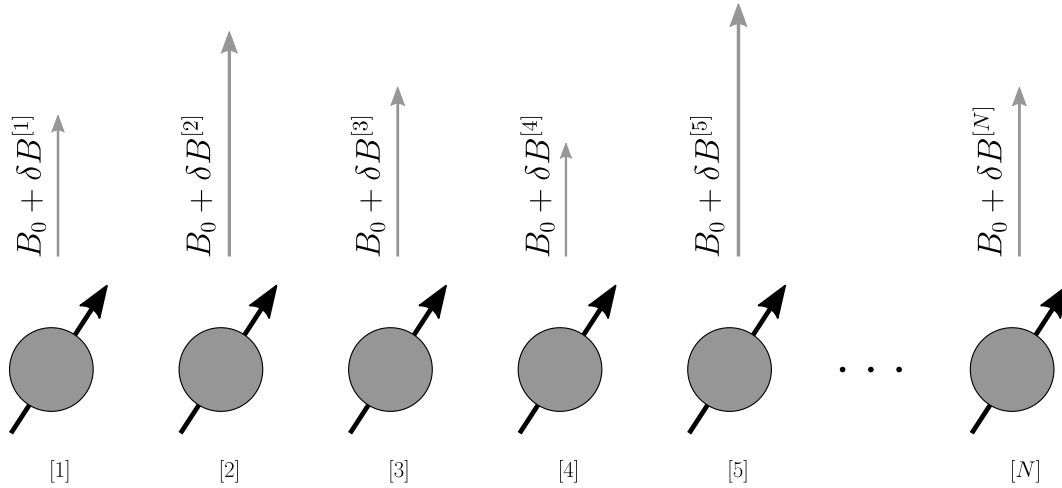
Przykładowe zastosowania

W tym rozdziale przedstawimy trzy przykłady zastosowania zaprezentowanego przez nas podejścia do metrologii kwantowej opartego na sieciach tensorowych. Przykłady zostały wybrane, tak aby podkreślić możliwość zastosowania naszego podejścia do szerokiej klasy jakościowo różnych problemów fizycznych i przez to zademonstrować jego wszechstronność. Pierwsze dwa przykłady dotyczą problemów metrologicznych: estymacji indukcji pola magnetycznego oraz stabilizacji działania zegara atomowego. Jednakże różnią się one zarówno charakterem szumu (w pierwszym przypadku ma on korelacje przestrzenne, a w drugim czasowe), jak i wykorzystanym podejściem do estymacji (pierwszy jest standardowym przykładem wykorzystującym QFI, drugi jest wyrażony w duchu Bayesowskim). Trzeci przykład ilustruje zastosowanie zaproponowanego przez nas formalizmu poza metrologią kwantową, mianowicie do wyznaczenia podatności wierności – parametru używanego w fizyce materii skondensowanej do detekcji przejść fazowych.

Rezultaty zaprezentowane w tym rozdziale bazują na obliczeniach numerycznych, które zostały wykonane z użyciem pakietu MATLAB oraz przy pomocy funkcji `ncon()` [Pfeifer i in. 2014] do zwięzania tensorów. Należy podkreślić, że wszystkie algorytmy optymalizacyjne musiały zostać napisane od podstaw z uwzględnieniem metrologicznego kontekstu badań, ponieważ procedury optymalizacyjne używane przez społeczność fizyków wielu ciał nie mogły zostać prosto zaadaptowane do zadań optymalizacyjnych, z którymi musimy się zmierzyć. Wiąże się to z dużym znaczeniem operatorów w metrologii kwantowej, a więc MPO w naszym podejściu. Było to do tej pory raczej niespotykane w dziedzinie sieci tensorowych, ponieważ typowe problemy w fizyce wielu ciał wymagają operowania i optymalizacji MPS. Procedury numeryczne napisane na potrzeby tego rozdziału składają się z około 8000 linii kodu i można je znaleźć pod adresem: <https://github.com/kchabuda/TNQMetrology>

4.1 Pomiar pola magnetycznego w obecności lokalnie skorelowanego szumu

Pierwszy zaprezentowany przez nas przykład dotyczy pomiaru indukcji pola magnetycznego, które wykazuje lokalne przestrzenne korelacje fluktuacji. Przykład ten należy do kategorii ogólnych modeli metrologicznych i powinien być traktowany jako typowy przykład problemu, który można efektywnie rozwiązać przy użyciu sieci tensorowych. Jako że w tym problemie interesują nas małe



RYSUNEK 4.1: N spinów w polu magnetycznym, którego indukcja wykazuje przestrzenne fluktuacje.

odchylenia od znanej wartości pola to będziemy używać podejścia opartego o QFI, a więc optymalizowany FoM przyjmie standardową postać wyrażoną wzorem (3.16).

Rozważmy N cząstek o spinie s (w ogólności może to być też efektywny spin pewnej grupy „pod-cząstek”), które oddziałują przez ustalony czas t z zewnętrznym polem magnetycznym B (zakładamy, że jest ono w kierunku osi z), którego indukcja fluktuuje – patrz rysunek 4.1. Te fluktuacje prowadzą efektywnie do defazowania cząstek. W naszych rozważaniach wykrócymy poza standardowy model niezależnych fluktuacji prowadzących do niezależnego defazowania atomów [Huelga i in. 1997; Escher, Matos Filho i Davidovich 2011] przez wzięcie pod uwagę korelacji pomiędzy fluktuacjami pola na sąsiadujących ze sobą cząstkach. Prowadzi nas to do modelu, w którym będziemy mogli zbadać, z punktu widzenia potencjału metrologicznego układu, wpływ korelacji (i antykorelacji) w procesie defazowania.

Na chwilę założymy, że pole B jest ustalone. Hamiltonian opisujący dynamikę n -tego spinu w polu magnetycznym to $-gs_z^{[n]}B$, gdzie $s_z^{[n]}$ jest z -tą składową spinowego momentu pędu n -tej cząstki ($\sum_{n=1}^N s_z^{[n]} = S_z$), a g to współczynnik żyromagnetyczny. W celu zachowania spójności z notacją wprowadzoną w rozdziale 1.3, będziemy utożsamiać $\varphi = gBt$ oraz $h^{[n]} = s_z^{[n]}/\hbar$. Niepewność estymacji pola B będzie powiązana z standardową niepewnością estymacji fazy poprzez prostą relację proporcjonalności $\Delta\hat{B} = \frac{1}{gt}\Delta\hat{\varphi}$.

Teraz uwzględniając obecność fluktuacji zapiszemy indukcję pola magnetycznego w miejscu, w którym znajduje się n -ta cząstka jako $B^{[n]}(t) = B + \delta B^{[n]}(t)$. Zakładamy, że fluktuacje są Gaussowskie i nie mają znaczących korelacji czasowych (biały szum). Odpowiednio wariancja oraz funkcja korelacji pomiędzy najbliższymi sąsiadami dla fluktuacji pola ma postać: $\langle \delta B^{[n]}(t)\delta B^{[n]}(t') \rangle = \mu^2\delta(t-t')$, $\langle \delta B^{[n]}(t)\delta B^{[n+1]}(t') \rangle = \nu\delta(t-t')$, gdzie ν reprezentuje siłę korelacji i może być zarówno dodatnie, jak i ujemne (antykorelacje) – dla uproszczenia założymy okresowe warunki brzegowe, a więc również cząstki numer N oraz 1 są skorelowane.

Niech $|\mathbf{j}\rangle = |j_1, j_2, \dots, j_N\rangle$ z $j_n \in \{-s, \dots, s\}$ będzie bazą własną, o dobrze określonych wartościach własnych, lokalnych Hamiltonianów $h^{[n]} = s_z^{[n]}/\hbar$, tak że $h^{[n]}|\mathbf{j}\rangle = j_n|\mathbf{j}\rangle$. Przy pomocy tej bazy możemy łatwo zapisać ewolucję macierzy gęstości przez kanał kwantowy, opisujący lokalnie skorelowane defazowanie, używając standardowej techniki rozwinięcia kumulantowego [Kampen 1981]:

$$\Phi[\rho_0] = \sum_{\mathbf{j}, \mathbf{k}} \langle \mathbf{j} | \rho_0 | \mathbf{k} \rangle e^{-\frac{1}{2}(\mathbf{j}-\mathbf{k})^T C (\mathbf{j}-\mathbf{k})} |\mathbf{j}\rangle \langle \mathbf{k}|, \quad (4.1)$$

gdzie C jest macierzą korelacji

$$C = \begin{pmatrix} c_1 & c_2 & 0 & \dots & c_2 \\ c_2 & c_1 & c_2 & & \\ 0 & c_2 & c_1 & & \\ \vdots & & & \ddots & \vdots \\ c_2 & & & \dots & c_1 \end{pmatrix}, \quad (4.2)$$

gdzie $c_1 = \mu^2 g^2 t$, $c_2 = \nu g^2 t$. Uogólnienie rozważanego modelu na przypadek większego zasięgu korelacji, pokrywającego do r sąsiadujących cząstek, nie sprawia problemu – w takim wypadku macierz C jest $2r + 1$ diagonalna.

Dla lepszego zrozumienia fizycznej strony rozważanego modelu, zapiszemy odpowiadające mu równanie master (3.2). Można je otrzymać przez zróżniczkowanie wzoru (4.1) po t . W wyniku otrzymamy równanie master z jednocząstkowym $J^{[n]} = \sqrt{\gamma_1} h^{[n]}$ i dwucząstkowym $J^{[n,n+1]} = \sqrt{|\gamma_2|} (h^{[n]} + \text{sgn}(\gamma_2) h^{[n+1]})$ operatorem defazowania. Współczynniki defazowania γ_1, γ_2 są powiązane z fluktuacjami pola przez relacje: $\gamma_1 = (\mu^2 - 2|\nu|)g^2$, $\gamma_2 = \nu g^2$ – zauważmy, że $\mu^2 \geq 2|\nu|$ z racji dodatniości macierzy korelacji C , tak więc współczynnik γ_1 jest zawsze dodatni.

W celu zapisania ewolucji jawnie w reprezentacji MPO wykonamy wektoryzację $|\mathbf{j}\rangle \langle \mathbf{k}| \rightarrow |\mathbf{j}\rangle |\mathbf{k}\rangle$, aby otrzymać bazę dla zwektoryzowanej macierzy gęstości stanu początkowego $|\rho_0\rangle$. Działanie Φ na ρ_0 jest równoważne działaniu e^Γ na $|\rho_0\rangle$, to znaczy $|\Phi[\rho_0]\rangle = e^\Gamma |\rho_0\rangle$, gdzie

$$\Gamma = -\frac{c_1}{2} \sum_{n=1}^N \Upsilon^{[n]} - c_2 \sum_{n=1}^N \Xi^{[n,n+1]}, \quad (4.3)$$

z

$$\Upsilon^{[n]} = \left(h^{[n]} \otimes \mathbb{1} - \mathbb{1} \otimes h^{[n]} \right)^2, \quad (4.4)$$

$$\Xi^{[n,n+1]} = \left(h^{[n]} \otimes \mathbb{1} - \mathbb{1} \otimes h^{[n]} \right) \left(h^{[n+1]} \otimes \mathbb{1} - \mathbb{1} \otimes h^{[n+1]} \right). \quad (4.5)$$

Zauważmy, że $\Upsilon^{[n]}$ i $\Xi^{[n,n+1]}$ komutują ze sobą, więc

$$e^\Gamma = \prod_{n=1}^N e^{-\frac{c_1}{2} \Upsilon^{[n]}} e^{-c_2 \Xi^{[n,n+1]}}. \quad (4.6)$$

Wprowadzając oznaczenia $Y^{[n]} = e^{-\frac{c_1}{2}\Upsilon^{[n]}}$ i $X^{[n,n+1]} = e^{-c_2\Xi^{[n,n+1]}}$ otrzymujemy ostatecznie

$$e^\Gamma = \prod_{n=1}^N Y^{[n]} X^{[n,n+1]}, \quad (4.7)$$

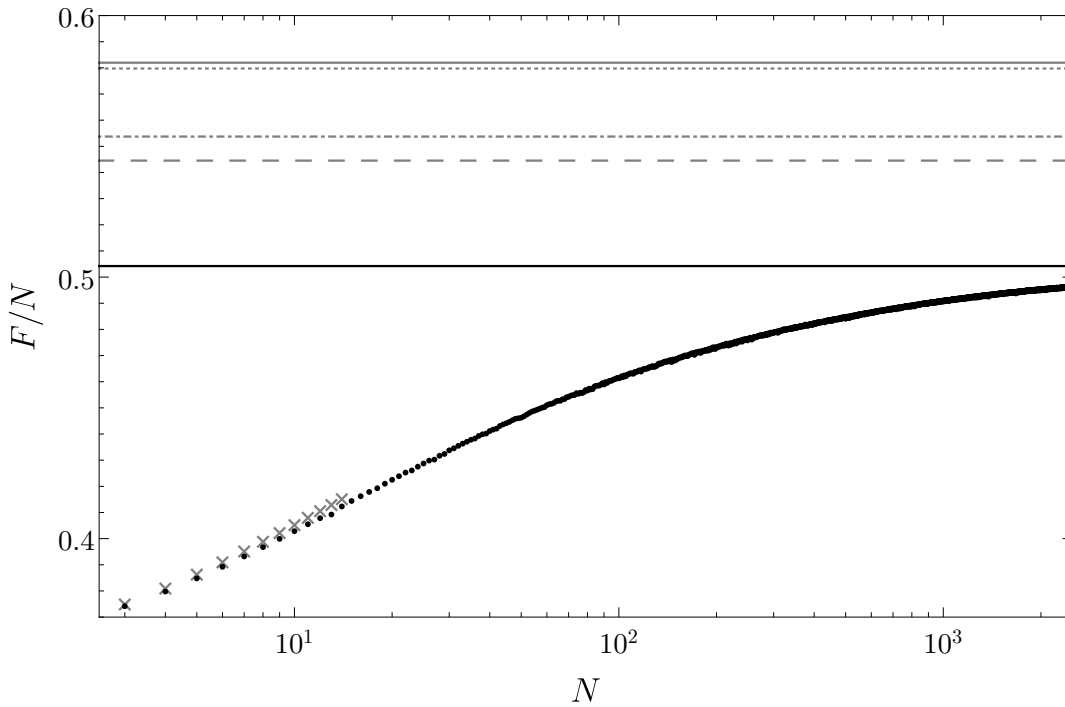
co jest dokładnie postacią ewolucji dyskutowaną w rozdziale 3.2, gwarantującą efektywny zapis za pomocą MPO.

Zgodnie ze wzorem (3.1), po ewolucji przez kanał kwantowy Φ opisujący szum, informacja o fazie jest zapisywana w stanie w wyniku unitarnej ewolucji generowanej przez lokalne Hamiltoniany $h^{[n]}$. W notacji z rozdziału 3.2 ta unitarna ewolucja jest reprezentowana przez działanie $\prod_{n=1}^N Z^{[n]}$, gdzie $Z^{[n]} = e^{-ih^{[n]}\varphi} \otimes e^{ih^{[n]}\varphi^T}$. Zapisana w bazie $|\mathbf{j}\rangle$, pochodna $\dot{\rho}_\varphi = i[\rho_\varphi, H]$ ma postać

$$\dot{\rho}_\varphi = \sum_{\mathbf{j}, \mathbf{k}} \langle \mathbf{j} | \rho_\varphi | \mathbf{k} \rangle i \sum_n (k_n - j_n) |\mathbf{j}\rangle \langle \mathbf{k}|. \quad (4.8)$$

Jesteśmy teraz gotowi, aby za pomocą metod opisanych w rozdziale 3 wyznaczyć optymalne stany początkowe oraz odpowiadającą im maksymalną wartość QFI. W badanym układzie estymujemy wartość pola magnetycznego wykazującego lokalnie korelacje fluktuacji, czyli efektywnie estymujemy fazę w obecności defazowania. Wszystkie obliczenia numeryczne dotyczą przypadku spinów- $\frac{1}{2}$, dla których $h^{[n]} = \frac{1}{2}\sigma_z^{[n]}$ oraz szumu defazującego, którego zasięg obejmuje maksymalnie dwie sąsiadujące cząstki.

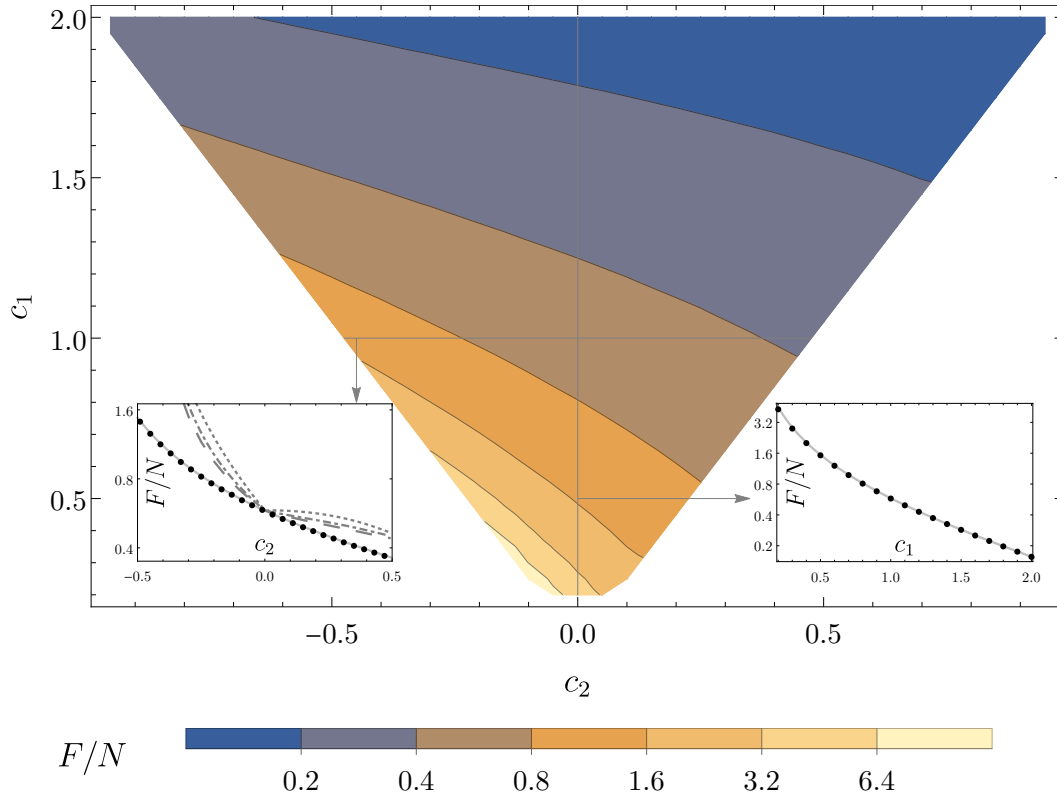
Na początek, na rysunku 4.2 prezentujemy porównanie wyników optymalizacji QFI – dla przykładowych parametrów szumu $c_1 = 1$, $c_2 = 0,1$ (skorelowany szum) – otrzymanych używając podejścia opartego na skończonej sieci N cząstek (podejście MPO) oraz asymptotyczną wartość QFI na cząstkę uzyskaną używając TI nieskończonej sieci (podejście iMPO). Wykreślamy QFI na cząstkę F/N , ponieważ spodziewamy się asymptotycznie liniowego skalowania, a więc że wraz ze wzrostem liczby cząstek, F/N będzie zbiegać do stałej wartości. Rezultaty otrzymane tymi dwoma podejściami bardzo dobrze się ze sobą zgadzają. Kiedy wykreślimy QFI w funkcji N i do takich danych dopasujemy linię prostą otrzymamy współczynnik nachylenia 0,500. Zgadza się on dobrze z asymptotyczną wartością 0,504 uzyskaną w podejściu iMPO. Stanowi to numeryczne potwierdzenie, że podejście iMPO, które jak opisano w rozdziale 3.5 jest dużo bardziej koncepcyjnie zawile, prowadzi do poprawnych rezultatów. Jest to bardzo ważna obserwacja, ponieważ numerycznie można dużo bardziej efektywnie uzyskać asymptotyczne własności QFI używając bezpośrednio podejścia iMPO, aniżeli wykonując obliczenia dla coraz to większych skończonych sieci i ekstrapolować wyniki do $N \rightarrow \infty$. Podczas wykonywania numerycznych optymalizacji w podejściu MPO ustaliliśmy tolerancję błędu na poziomie 1%, co oznaczało, że maksymalny potrzebny wymiar wiązania wynosił jedynie $\chi_{\psi_0} = 4$ dla stanu początkowego i $\chi_\Lambda = 2$ dla SLD. To pokazuje jak efektywny jest opis tego problemu przy pomocy sieci tensorowych. Krzyże naniesione na rysunek 4.2 reprezentują wyniki, które można uzyskać wykonując obliczenia w pełnej przestrzeni Hilberta, na takiej samej klasy komputerze, który wykonywał obliczenia z użyciem sieci tensorowych. Wyniki te pozostają w reżimie, który jest wyraźnie daleki od



RYSUNEK 4.2: QFI na cząstkę F/N w funkcji liczby spinów w łańcuchu N dla problemu estymacji fazy w obecności lokalnie skorelowanego defazowania (o lokalnym parametrze szumu $c_1 = 1$ oraz parametrze korelacji pomiędzy najbliższymi sąsiadami $c_2 = 0,1$). Czarne kropki reprezentują wyniki uzyskane w podejściu MPO, a czarna linia asymptotyczną wartość F/N uzyskaną podejściem iMPO, natomiast szare krzyże oznaczają rezultaty otrzymane standardowymi metodami wykorzystującymi opis w pełnej przestrzeni Hilberta. Szare linie pokazują ograniczenia na QFI na cząstkę, uzyskane metodą rozszerzenia kanału, otrzymane przy podziale dynamiki (patrz rysunek 4.6) na efektywnie niezależne kanały $\Phi_\varphi^{(1)}$ (linia kropkowana), $\Phi_\varphi^{(2)}$ (linia kropkowano-przerywana), $\Phi_\varphi^{(3)}$ (linia przerywana). Dla porównania, ciągła szara linia odpowiada ograniczeniu, w którym wszystkie korelacje zostały zaniedbane i jedynie efekt lokalnego defazowania został wzięty pod uwagę.

tęgo w jakim obserwujemy zbieganie do asymptotycznego liniowego skalowania QFI z N i który jest dostępny jedynie wykorzystując metody numeryczne bazujące na sieciach tensorowych.

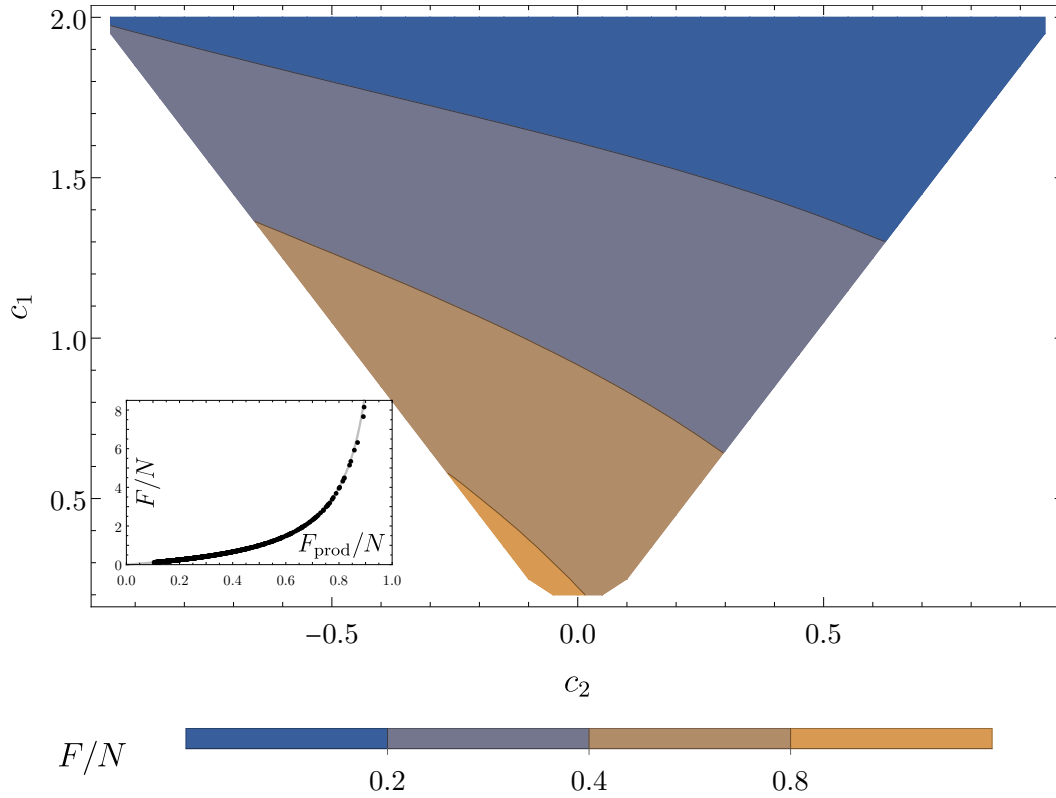
Na rysunku 4.3 zaprezentowaliśmy wykres konturowy przedstawiający asymptotyczną wartość QFI na cząstkę w funkcji parametrów szumu oraz dla kontrastu na rysunku 4.4 wyniki gdy ograniczymy się do optymalizacji QFI jedynie po stanach produktowych. W celu pokazania ilości splątania występującego w optymalnych stanach początkowych na rysunku 4.5 znajduje się wykres konturowy entropii splątania von Neumanna $\mathcal{S}$. Wykonano go dla iMPO odpowiadającego zredukowanej macierzy gęstości układu, w którym połowa cząstek została wysładowana – wzór (2.29) pokazuje jak taką zredukowaną macierz gęstości wyznaczyć. Optymalna wartość QFI może być osiągnięta przez wiele różnych stanów, a więc należy się spodziewać pewnych fluktuacji entropii, mimo tego widzimy że jest wyraźna relacja pomiędzy ilością splątania a wzrostem precyzji



RYSUNEK 4.3: Asymptotyczna wartość QFI na cząstkę F/N (uzyskana podejściem iMPO) dla szumu typu defazującego w funkcji lokalnego parametru szumu c_1 oraz parametru korelacji pomiędzy najbliższymi sąsiadami c_2 . Czarne linie ekwipotencjalne są w skali logarytmicznej. Lewa wstawka pokazuje przekrój głównego wykresu wzdłuż $c_1 = 1$ i przedstawia wartości otrzymane podejściem iMPO (czarne kropki) porównane ze ścisłym wzorem (4.24) na QFI na cząstkę w strategii estymacji bazującej na słabo ściśniętych stanach początkowych oraz ograniczenia na QFI na cząstkę, uzyskane metodą rozszerzenia kanału, otrzymane przy podziale dynamiki (patrz rysunek 4.6) na efektywnie niezależne kanały $\Phi_\varphi^{(1)}$ (szara kropkowana linia), $\Phi_\varphi^{(2)}$ (szara kropkowano-przerywana linia), $\Phi_\varphi^{(3)}$ (szara przerywana linia). Prawa wstawka pokazuje przekrój głównego wykresu wzdłuż $c_2 = 0$ i przedstawia wartości otrzymane podejściem iMPO (czarne kropki) porównane ze znanym ścisłym rezultatem dla ściśle lokalnego szumu $F/N = \eta^2/(1 - \eta^2)$, gdzie $\eta = e^{-\frac{c_1}{2}}$ (jasnoszara linia).

pomiarów.

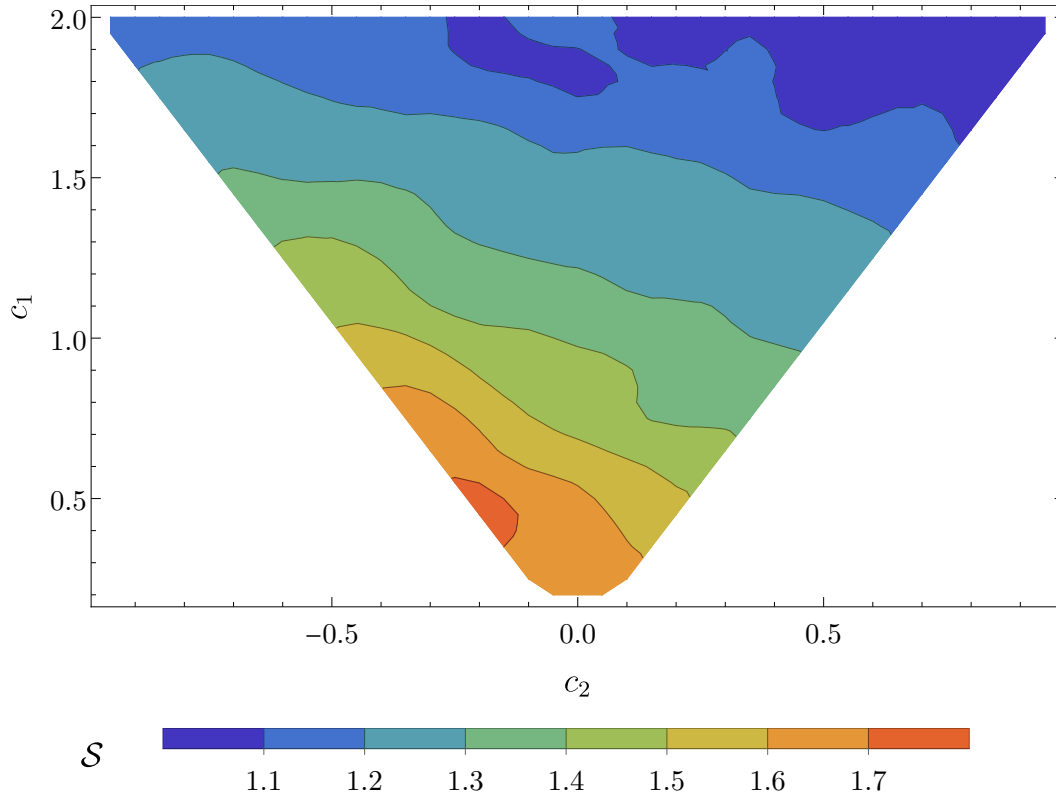
W rozdziale 1.3.2 opisaliśmy metodę rozszerzenia kanału. Jest to według obecnego stanu wiedzy najlepsza metoda otrzymywania ograniczeń na precyzję w metrologii kwantowej, jednakże została stworzona z myślą o nieskorelowanych modelach szumu [Demkowicz-Dobrzański, Kołodyński i Guţă 2012]. Metoda ta dostarcza łatwe do obliczenia asymptotyczne ograniczenia, bazując jedynie na podstawowej wiedzy o dynamice układu wyrażonej przez operatory Krausa. Zastanowimy się teraz czy jest możliwe uzyskanie jakiegoś wglądu w badany problem pomysłowo adaptując metodę rozszerzenia kanału. Mimo że szum w naszym modelu jest skorelowany to formalnie możemy podzielić ewolucję na



RYSUNEK 4.4: Asymptotyczna wartość QFI na cząstkę F/N (uzyskana podejściem iMPO) dla szumu typu defazującego w funkcji lokalnego parametru szumu c_1 oraz parametru korelacji pomiędzy najbliższymi sąsiadami c_2 , przy założeniu produktowego stanu początkowego ($\chi_{\psi_0} = 1$). Lewa wstawka przedstawia wykres optymalnych wartości QFI na cząstkę F/N w funkcji odpowiadających im wartości QFI na cząstkę dla produktowych stanów początkowych F_{prod}/N (czarne kropki), który idealnie pasuje do zależności funkcyjnej (znanej z przypadku nieskorelowanego defazowania) $\frac{F}{N} = \frac{F_{\text{prod}}}{N} / \left(1 - \frac{F_{\text{prod}}}{N}\right)$ (jasnoszara linia).

rodzinę niezależnych kanałów działających na jedną, dwie, trzy lub więcej cząstek. Podobny pomysł został ostatnio zastosowany w badaniach nad wpływem wielociałowych efektów w interferometrii atomowej [Czajkowski, Pawłowski i Demkowicz-Dobrzański 2019]. Sprawdźmy teraz jak otrzymane w ten sposób ograniczenia mają się do wyników uzyskanych za pomocą formalizmu sieci tensorowych.

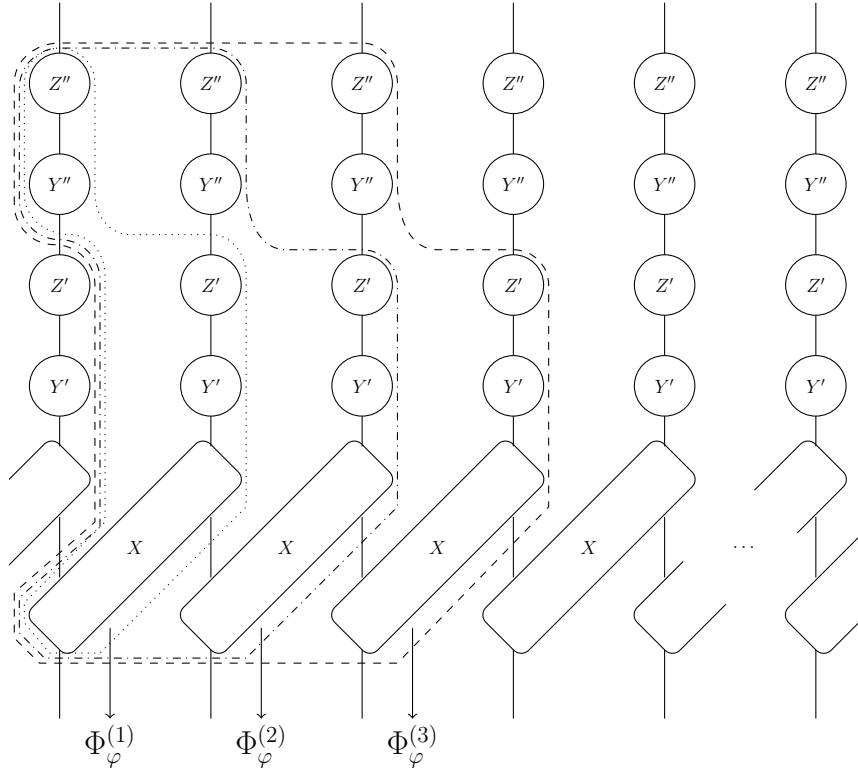
W celu podziału dynamiki na rodzinę niezależnych kanałów kwantowych, zaczniemy od formalnego rozkładu operatora lokalnego defazowania Y na iloczyn $Y = Y'Y''$ lokalnych operatorów defazowania odpowiadających parametrom c'_1 i c''_1 takim, że $c_1 = c'_1 + c''_1$. Następnie wykonamy rozkład unitarnego operatora odpowiadającego za nakręcanie fazy Z na iloczyn $Z = Z'Z''$, gdzie Z' , Z'' są operatorami unitarnymi odpowiadającymi fazom φ' , φ'' takim, że $\varphi = \varphi' + \varphi''$. Możemy teraz podzielić całkowitą dynamikę na efektywnie złożoną z N niezależnych kanałów $\Phi_{\varphi}^{(1)}$ albo pogrupować operatory w $\frac{N}{2}$ większych kanałów $\Phi_{\varphi}^{(2)}$ lub $\frac{N}{3}$ kanałów $\Phi_{\varphi}^{(3)}$ i tak dalej (zauważmy, że wszystkie elementarne ewolucje



RYSUNEK 4.5: Entropia splątania von Neumanna $\mathcal{S}$ (w bitach) dla optymalnego stanu (w podejściu iMPO, przy podziale łańcucha na dwie równe części) w zagadnieniu estymacji fazy w obecności lokalnie skorelowanego defazowania w funkcji lokalnego parametru szumu c_1 oraz parametru korelacji pomiędzy najbliższymi sąsiadami c_2 .

komutują) – patrz rysunek 4.6.

Jesteśmy teraz gotowi, aby zastosować metodę rozszerzenia kanału, która dostarczy nam ograniczenia na QFI używając jednego z zaproponowanych rozkładów dynamiki na niezależne kanały: $\Phi_\varphi^{(1)}$, $\Phi_\varphi^{(2)}$ lub $\Phi_\varphi^{(3)}$. W celu otrzymania najlepszego ograniczenia, numerycznie optymalizujemy podział faz oraz parametrów szumu pomiędzy operatory Y' , Z' i Y'' , Z'' upewniając się przy tym, że w rezultacie otrzymany $\Phi_\varphi^{(j)}$ jest prawidłowym kanałem kwantowym (odwzorowaniem CPTP) – zauważmy, że np. sam operator X nie tworzy kanału kwantowego, ponieważ nie jest całkowicie dodatni. Otrzymane rezultaty są przedstawione w lewej wstawce rysunku 4.3. Chociaż ograniczenia są wysycalne dla modeli szumu bez korelacji, to gdy tylko zaczniemy wprowadzać korelacje (lub antykorelacje) do modelu, to otrzymujemy ograniczenia, które już nie są wysycalne (tym bardziej im większe są korelacje). Tak jak należało się spodziewać uzyskiwane ograniczenia są tym lepsze im większy użyjemy elementarny kanał. Jednakże metoda ta źle skaluje się wraz ze wzrostem rozmiaru elementarnego kanału w związku z czym udało nam się wykonać obliczenia jedynie dla $\Phi_\varphi^{(1)}$, $\Phi_\varphi^{(2)}$ i $\Phi_\varphi^{(3)}$. Wyniki te pokazują, że najlepsze obecnie istniejące metody otrzymywania ograniczeń na QFI, stworzone dla modeli szumu bez korelacji, dostarczają niezadowalające rezultaty, gdy próbuje się je użyć do modeli ze skorelowanym



RYSUNEK 4.6: Różne sposoby podziału lokalnie skorelowanej dynamiki na rodzinę niezależnych kanałów kwantowych: $\Phi_\varphi^{(1)}$ (linia kropkowana), $\Phi_\varphi^{(2)}$ (linia kropkowano-przerywana), $\Phi_\varphi^{(3)}$ (linia przerywana).

szumem.

Jak wspomniano powyżej, użycie podejścia iMPO znacznie przyspiesza proces wyznaczania asymptotycznej wartości QFI. W celu zbadania wpływu korelacji w szumie na osiągalną wartość QFI oraz optymalne stany początkowe w pozostałej części tego podrozdziału będziemy bazować na wynikach numerycznych uzyskanych właśnie tym podejściem. Przystudiowaliśmy zachowanie QFI w zakresie zmienności parametrów szumu $c_1 \in [0, 2; 2]$ i $c_2 \in \left(-\frac{c_1}{2}; \frac{c_1}{2}\right)$. Udało nam się osiągnąć względny błąd na poziomie 0,3%, zwiększając wymiar wiązania stanu początkowego χ_{ψ_0} do 15, przy tym utrzymując $\chi_{\bar{\Lambda}} = 1$. Zysk ze zwiększania $\chi_{\bar{\Lambda}}$ był pomijalny w porównaniu do tego otrzymanego przy zwiększaniu χ_{ψ_0} – jest to prawdopodobnie powiązane z faktem, że lokalne nieskorelowane pomiary są praktycznie optymalne, zachowanie takie często spotykamy w metrologii kwantowej, gdy wielkością optymalizowaną jest QFI [Demkowicz-Dobrzański, Jarzyna i Kołodyński 2015].

Chociaż podejście iMPO zapewnia nam bardzo dokładne i szybkie do uzyskania wyniki, to wymaga ono przygotowań, to znaczy poszukiwania symetrii w tensorach, które mogłyby znacząco zmniejszyć prawdopodobieństwo, że algorytm „utknie” w jakimś lokalnym ekstremum. Dla przykładu w dyskutowanym tu modelu defazowania optymalny tensor $\tilde{L}$ ma wymiar wiązania równy jeden, tak więc $\tilde{L}_k^j$ jest skalarą, który dla równych indeksów fizycznych $j = k$ wynosi $\tilde{L}_j^j = 1$, a dla różnych $j \neq k$ wynosi $\tilde{L}_k^j = i\epsilon \operatorname{sgn}(k - j)g(|j - k|)$, gdzie

funkcja g zwraca rzeczywiste wartości rzędu $e^{-\frac{c_1}{2}}$. Co się tyczy optymalnego tensora P to obserwujemy, że są nimi macierze P^j , które są Hermitowskie, a średnia wartość bezwzględna każdej kolejnej diagonalu spada o około dwa rzędy wielkości. Są one również symetryczne względem odbicia wzdłuż antydiagonalu oraz $P^{d-1-j} = P^{jT}$. Dalszych badań wymaga wyjaśnienie pochodzenia symetrii optymalnych tensorów, jak i tego, że efektywny wymiar wiązania będący efektem szumu χ_r wynosi jedynie $2d - 1$. Podejrzewamy, że część z tych własności można powiązać z faktem, iż defazowanie jest niezmiennicze względem obrotu wokół osi z , tak więc optymalne stany muszą należeć do nieredukowalnej reprezentacji S_z .

Główny jakościowy wniosek jaki jasno wynika z rysunku 4.3 to obniżenie optymalnej wartości QFI wraz ze wzrostem części szumu odpowiadającej za korelacje pomiędzy najbliższymi sąsiadami, opisywanej parametrem c_2 . Z drugiej strony, gdy przejdziemy do reżimu antyskorelowanego szumu (ujemne wartości c_2) to QFI znacząco wzrasta. Jest to związane z tym, że dla antyskorelowanego szumu operator szumu odpowiadający za korelacje przyjmuje postać $J^{[n,n+1]} = \sqrt{|\gamma_2|}(h^{[n]} - h^{[n+1]})$ i staje się liniowo niezależny od operatora Hamiltonianu, który jest sumą $h^{[n]}$. Jest to powiązane z faktem, że w granicznym przypadku, gdybyśmy mieli tylko szum antyskorelowany (bez członu odpowiadającego za lokalne defazowanie) to można byłoby zastosować protokoły inspirowane kwantowymi kodami korekcji błędów, aby osiągnąć skalowanie Heisenberga [Layden i Cappellaro 2018]. Ma to również odzwierciedlenie w zachowaniu ograniczeń uzyskanych metodą rozszerzenia kanału, które stają się rozbieżne w tym reżimie. Wynika to z tego, że metoda ta uwzględnia również rezultaty, które można uzyskać strategiami adaptacyjnymi. Wyniki uzyskane za pomocą sieci tensorowych nie wykazują tych rozbieżności, ponieważ nasze podejście opiera się na strategiach równoległych, gdzie rozważamy najbardziej ogólne splątane stany początkowe oraz najbardziej ogólne pomiary, ale bez możliwości adaptacji. To tym bardziej pokazuje, że dotychczasowe, najlepsze metody metrologii kwantowej nie sprawdzają się gdy chcemy badać układy ze skorelowanym szumem w ramach strategii równoległych. Linia $c_2 = 0$ na rysunku 4.3 odpowiada modelowi ściśle lokalnego szumu, dla którego znamy ograniczenie, które jest asymptotycznie wysycalne: $F/N = \eta^2/(1 - \eta^2)$ [Escher, Matos Filho i Davidovich 2011; Demkowicz-Dobrzański, Kołodyński i Guţă 2012] (wyprowadzenie w rozdziale 1.3.2), gdzie $\eta = e^{-\frac{c_1}{2}}$. Nasze numeryczne rezultaty otrzymane w podejściu iMPO pokrywają się z tym wzorem idealnie – zobacz prawą wstawkę na rysunku 4.3.

Dla ściśle lokalnego szumu wiadomym jest, że w granicy bardzo wielu cząstek fundamentalne ograniczenie, $F/N = \eta^2/(1 - \eta^2)$, może zostać wysyczone za pomocą protokołów wykorzystujących spinowe stany ściśnięte [Ulam-Orgikh i Kitagawa 2001; Escher, Matos Filho i Davidovich 2011], jak np. stany jednoosiowo skręcone [J. Ma i in. 2011]. Po otrzymaniu numerycznych wartości optymalnego QFI na cząstkę, używając zaproponowanego przez nas formalizmu bazującego na sieciach tensorowych, chcielibyśmy sprawdzić czy spinowe stany słabo ściśnięte będą również optymalne w przypadku skorelowanego szumu. W tym celu wykorzystamy standardowy protokół metrologiczny bazujący na interferometrii Ramseya [J. Ma i in. 2011], aby wyznaczyć ścisłą wartość QFI w

problemie estymacji fazy w obecności lokalnie skorelowanego szumu defazującego. Bez straty dla ogólności możemy założyć, że estymujemy małe odchylenie fazy od $\varphi_0 = 0$. Rozważymy następujący stan jednoosiowo skręcony (OAT, ang. *One-Axis Twisted*) N cząstek:

$$|\text{OAT}\rangle = e^{-i\theta S_x^2} \left| +\frac{1}{2} \right\rangle^{\otimes N}, \quad (4.9)$$

gdzie θ oznacza kąt skręcenia. Powtórzmy standardowe rozumowanie [J. Ma i in. 2011], gdzie powyższy stan jest obrócony na równik sfery Blocha, tak że $\langle \mathbf{S} \rangle$ wskazuje kierunek osi x , a moment pędu ma najmniejszą wariancję w kierunku osi y . To będzie nasz stan początkowy $|\psi_0\rangle$, który poddamy ewolucji przez kanał kwantowy Φ (4.1) opisujący lokalnie skorelowane defazowanie, a następnie obrócimy o nieznany kąt φ wokół osi z . W efekcie otrzymamy stan

$$\rho_\varphi = \Phi_\varphi[|\psi_0\rangle] = e^{-iS_z\varphi} \Phi[|\psi_0\rangle] e^{iS_z\varphi}. \quad (4.10)$$

Zakładając, że $\varphi \approx 0$ dokonujemy pomiaru obserwacji S_y , jako że jest to w tym przypadku optymalny wybór, a następnie na podstawie wyników pomiaru wnioskujemy o wartości φ .

W takim właśnie przypadku, gdy dokonujemy pomiaru obserwacji, której wartość oczekiwana zależy od badanego parametru, $\langle S_y \rangle = \text{Tr}(\rho_\varphi S_y)$, precyzję pomiaru możemy wyznaczyć korzystając ze wzoru na propagację błędu

$$\Delta^2 \hat{\varphi} = \frac{\Delta^2 S_y}{\left| \frac{d\langle S_y \rangle}{d\varphi} \right|^2}. \quad (4.11)$$

W celu wyznaczenia potrzebnych do skorzystania z tego wzoru wartości oczekiwanych przejdziemy do obrazu Heisenberga, w którym zamiast stanów to operatory ewoluują (tylko że w przeciwnym „kierunku”). Zrobimy to w dwóch krokach: najpierw przeniesiemy na operatory unitarny obrót wokół osi z , a następnie dekoherencję – kolejność nie ma znaczenia, ponieważ w naszym problemie procesy te ze sobą komutują. Wprowadzimy oznaczenia $|_0$, $|_1$ i $|_2$ na wartości policzone odpowiednio na stanie $|\psi_0\rangle$, $\Phi[|\psi_0\rangle]$ i $\Phi_\varphi[|\psi_0\rangle]$. Chcemy wyznaczyć $\Delta^2 \hat{\varphi}|_2$ w funkcji wartości oczekiwanych na stanie początkowym, czyli na $|_0$.

Obrót wokół osi z o kąt $-\varphi$ powoduje następującą transformację operatorów S_i :

$$\begin{pmatrix} S'_x \\ S'_y \\ S'_z \end{pmatrix} = \begin{pmatrix} \cos \varphi & \sin \varphi & 0 \\ -\sin \varphi & \cos \varphi & 0 \\ 0 & 0 & 1 \end{pmatrix} \begin{pmatrix} S_x \\ S_y \\ S_z \end{pmatrix}. \quad (4.12)$$

Wynika stąd że:

$$\langle S_y \rangle|_2 = -\sin \varphi \langle S_x \rangle|_1 + \cos \varphi \langle S_y \rangle|_1, \quad (4.13)$$

$$\Delta^2 S_y|_2 = \sin^2 \varphi \Delta^2 S_x|_1 + \cos^2 \varphi \Delta^2 S_y|_1 - 2 \sin \varphi \cos \varphi \text{Cov}(S_x, S_y)|_1, \quad (4.14)$$

gdzie $\text{Cov}(S_x, S_y) = \frac{1}{2} \langle S_x S_y + S_y S_x \rangle - \langle S_x \rangle \langle S_y \rangle$ to kowariancja. Korzystając teraz z naszego założenia, że $\varphi \approx 0$, możemy napisać, iż

$$\Delta^2 \hat{\varphi}|_2 = \frac{\Delta^2 S_y|_1}{|\langle S_x \rangle|_1^2}. \quad (4.15)$$

W celu przeniesienia ewolucji pod wpływem kanału defazującego Φ ze stanów na operatory wprowadzimy podobnie jak w rozdziale 3.4 kanał dualny Φ^D . Pozwoli nam to zapisać dla dowolnej obserwabli O , że $\langle O \rangle|_1 = \langle \Phi^D[O] \rangle|_0$. W przypadku obecnie rozważanego kanału kwantowego Φ zadanego wzorem (4.1), działanie kanału dualnego Φ^D na dowolną obserwabłą O możemy zapisać następująco:

$$\Phi^D[O] = \sum_{\mathbf{j}, \mathbf{k}} \langle \mathbf{k} | O | \mathbf{j} \rangle e^{-\frac{1}{2}(\mathbf{j}-\mathbf{k})^T C (\mathbf{j}-\mathbf{k})} |\mathbf{k}\rangle \langle \mathbf{j}|. \quad (4.16)$$

Jako że defazowanie wyrzuca nas poza przestrzeń o dobrze zdefiniowanym całkowitym rzucie spinu na daną oś, $S_i = \frac{1}{2} \sum_{n=1}^N \sigma_i^{[n]}$ (zakładamy $\hbar = 1$), będziemy musieli rozpatrzyć działanie Φ^D na poszczególne $\sigma_i^{[n]}$ oraz ich iloczyny. Potrzebne nam wartości oczekiwane transformują się następująco:

$$\langle \sigma_x^{[n]} \rangle|_1 = \langle \sigma_x^{[n]} \rangle|_0 e^{-\frac{c_1}{2}}, \quad (4.17)$$

$$\langle \sigma_y^{[n]} \rangle|_1 = \langle \sigma_y^{[n]} \rangle|_0 e^{-\frac{c_1}{2}}, \quad (4.18)$$

$$\langle \sigma_y^{[n]} \sigma_y^{[n+1]} \rangle|_1 = \left(\langle \sigma_y^{[n]} \sigma_y^{[n+1]} \rangle|_0 \cosh c_2 + \langle \sigma_x^{[n]} \sigma_x^{[n+1]} \rangle|_0 \sinh c_2 \right) e^{-c_1}, \quad (4.19)$$

$$\langle \sigma_y^{[n]} \sigma_y^{[n+2]} \rangle|_1 = \langle \sigma_y^{[n]} \sigma_y^{[n+2]} \rangle|_0 e^{-c_1}. \quad (4.20)$$

Przechodząc teraz do granicy $N \rightarrow \infty$, $\theta \rightarrow 0$ w taki sposób, że $N\theta^2 \ll 1$ otrzymujemy, że istotne wartości oczekiwane na stanie początkowym przyjmują postać [J. Ma i in. 2011]:

$$\langle \sigma_x^{[n]} \rangle|_0 = 1, \quad \langle \sigma_y^{[n]} \rangle|_0 = 0, \quad (4.21)$$

$$\langle \sigma_x^{[n]} \sigma_x^{[m]} \rangle|_0 = 1, \quad \langle \sigma_y^{[n]} \sigma_y^{[m]} \rangle|_0 = -\frac{1}{N-1} = \frac{1}{1-N} \quad \text{dla } n \neq m. \quad (4.22)$$

W rezultacie dostajemy $\langle S_x \rangle|_1 = \frac{N}{2} e^{-\frac{c_1}{2}}$ oraz

$$\begin{aligned} \Delta^2 S_y|_1 = \langle S_y^2 \rangle|_1 = \frac{N}{4} + \frac{2(N-1)}{4} \left(\frac{1}{1-N} \cosh c_2 + \sinh c_2 \right) e^{-c_1} + \\ + \frac{(N-2)(N-1)}{4} \frac{1}{1-N} e^{-c_1}. \end{aligned} \quad (4.23)$$

To prowadzi nas ostatecznie do wzoru na asymptotyczną wartość precyzji $\Delta^2 \hat{\varphi} = (1 - e^{-c_1} + 2e^{-c_1} \sinh c_2)/(Ne^{-c_1})$, który można powiązać z odpowiadającą mu wartością QFI na cząstkę:

$$F/N = \frac{e^{-c_1}}{1 - e^{-c_1} + 2e^{-c_1} \sinh c_2}. \quad (4.24)$$

Możemy teraz sprawdzić czy wzór ten zgadza się z pożądaną dokładnością ($< 1\%$) z naszymi wynikami numerycznymi. Reprezentatywne porównanie wyników numerycznych z powyższym wzorem znajduje się w lewej wstawce na rysunku 4.3. Wyniki te pokazują, że podobnie jak dla nieskorelowanych modeli szumu, spinowe stany słabo ściśnięte są asymptotycznie optymalne. Nie oznacza to, iż optymalne stany, które otrzymaliśmy używając sieci tensorowych są stanami ściśniętymi. Wręcz przeciwnie, w podejściu MPO otrzymaliśmy stany o niskim rzędzie wiązania (niewielkiej entropii splątania, niezależnie od ilości spinów w łańcuchu) i lokalnych korelacjach, a stany ściśnięte z uwagi na przynależność do podprzestrzeni stanów całkowicie symetrycznych korelują wszystkie cząstki ze sobą, niezależnie od odległości pomiędzy nimi.

Gdybyśmy chcieli rozważyć strategię bazującą na stanach produktowych to jedyną modyfikacją w powyższym rozumowaniu byłoby podstawienie $\langle \sigma_y^{[n]} \sigma_y^{[m]} \rangle = 0$ dla $n \neq m$, co prowadziłoby do następującej wartości QFI na cząstkę:

$$F_{\text{prod}}/N = \frac{e^{-c_1}}{1 + 2e^{-c_1} \sinh c_2}, \quad (4.25)$$

która również idealnie zgadza się z naszymi wynikami numerycznymi przedstawionymi na rysunku 4.4.

Możemy teraz skorzystać z powyższych wzorów i powrócić do oryginalnego problemu tego podrozdziału, czyli pomiaru pola magnetycznego. Korzystając z relacji $\varphi = gBt$ możemy zapisać wzór na najlepszą możliwą wartość precyzji estymacji pola magnetycznego:

$$\Delta^2 \hat{B}_t = \frac{1}{g^2 t^2} \Delta^2 \hat{\varphi} = \frac{1 - e^{-\mu^2 g^2 t} + 2e^{-\mu^2 g^2 t} \sinh(\nu g^2 t)}{g^2 t^2 N e^{-\mu^2 g^2 t}}. \quad (4.26)$$

Powyższy wzór zakłada ustalony czas t oddziaływania cząstek z polem. Możemy uogólnić te rozważania i ustalić całkowity czas eksperymentu T , który podzielimy na T/t niezależnych przedziałów czasu, w których będzie odbywało się oddziaływanie. Odpowiadająca tej sytuacji niepewność estymacji wynosi

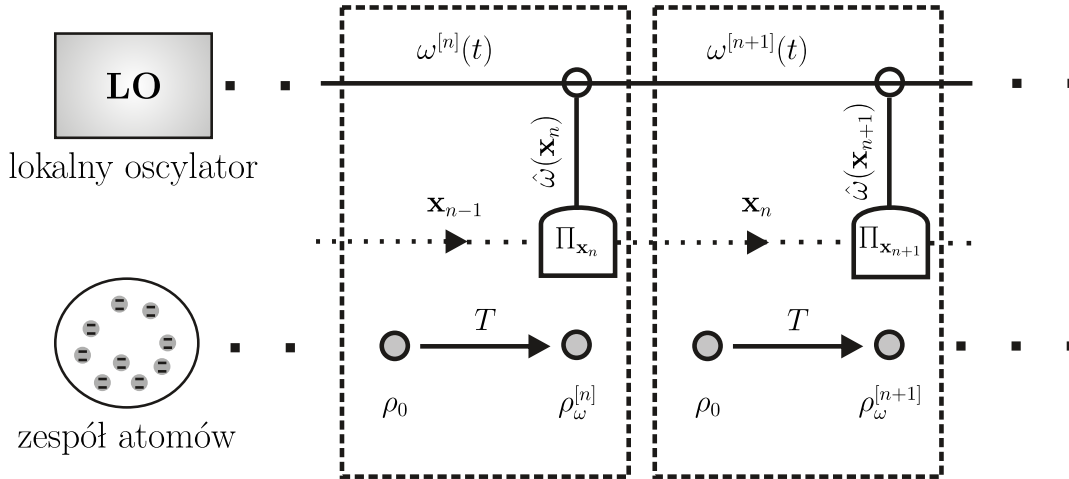
$$\Delta^2 \hat{B}_T = \frac{1}{T/t} \Delta^2 \hat{B}_t = \frac{1 - e^{-\mu^2 g^2 t} + 2e^{-\mu^2 g^2 t} \sinh(\nu g^2 t)}{g^2 t T N e^{-\mu^2 g^2 t}}, \quad (4.27)$$

która po zoptymalizowaniu po t , osiąga minimalną wartość gdy $t \rightarrow 0$ i wynosi

$$\Delta^2 \hat{B} = \frac{\mu^2 + 2\nu}{TN}. \quad (4.28)$$

Bazując na powyższych rezultatach możemy spodziewać się, że spinowe stany słabo ściśnięte powinny być również optymalne w przypadku bardziej ogólnego szumu defazującego, pod warunkiem że zasięg korelacji r jest skończony i rozważamy granice asymptotyczną $N \rightarrow \infty$. W takim przypadku przeprowadzając analogiczne obliczenia otrzymalibyśmy następującą postać optymalnej precyzji estymacji pola magnetycznego:

$$\Delta^2 \hat{B} = \frac{\mu^2 + 2 \sum_{k=2}^r \nu_k}{TN}, \quad (4.29)$$



RYSUNEK 4.7: Schemat przyjętego modelu stabilizacji lokalnego oscylatora (LO) do referencyjnej częstości przejścia atomowego ω_0 . W n -tym cyklu działania zegara zespół atomów jest początkowo przygotowany w stanie ρ_0 . Następnie przez czas T oddziałuje on z LO, którego (stabilizowana) częstość wynosi $\omega^{[n]}(t)$. W wyniku tego ewoluuje do stanu $\rho_\omega^{[n]}$. Po czym następuje pomiar (uogólniony) $\Pi_{\mathbf{x}_n}$, który może zależeć również od wyników pomiarów w poprzednich cyklach działania zegara $\mathbf{x}_{n-1} = (x_1, \dots, x_{n-1})$. W efekcie sygnał korekcyjny $\hat{\omega}(\mathbf{x}_n)$ bazujący na wszystkich dotychczasowych wynikach pomiarów $\mathbf{x}_n$ jest użyty, aby poprawić częstość LO $\omega^{[n+1]}(t)$, tak żeby w następnym kroku czasowym była ona jak najbliżej ω_0 . Cykl taki następnie zostaje powtórzony.

gdzie ν_k reprezentuje korelacje pola magnetycznego dla cząstek oddległych o $k - 1$: $\langle \delta B^{[n]}(t) \delta B^{[n+k-1]}(t') \rangle = \nu_k \delta(t - t')$. Porównując ten wynik z rezultatami uzyskanymi dla stanów GHZ (1.46) w tym samym modelu [Layden, Zhou i in. 2019], zauważamy że spinowe stany ściśnięte wykazują poprawę precyzji o czynnik $\exp(1)$ ponad to co można osiągnąć stanami GHZ – zależność która jest znana z rozważań nieskorelowanych modeli defazowania [Huelga i in. 1997; Ulam-Orgikh i Kitagawa 2001].

4.2 Stabilizacja zegara atomowego

Drugi zaprezentowany przez nas przykład dotyczy stabilizacji działania zegara atomowego. W tym przypadku nie tylko model fizyczny jest inny – atomy podlegają szumowi lokalnego oscylatora, którego fluktuacje są skorelowane w czasie, a nie jak w pierwszym przykładzie, gdzie korelacje szumu miały charakter przestrzenny – ale też różni się wielkość optymalizowana (FoM), która teraz jest dużo bliższa Bayesowskiej wariancji (patrz rozdział 1.2.2) aniżeli QFI.

Typowy zegar atomowy działa w pętli sprzężenia zwrotnego, gdzie lokalny oscylator (LO, np. laser) jest stabilizowany do referencyjnej częstości przejścia pomiędzy dwoma poziomami atomowymi przez periodyczne oddziaływanie z atomami (używając promieniowania z LO) i bazując na zmierzonej odpowiedzi częstość LO jest poprawiana [Ludlow i in. 2015] – patrz też schemat działania zegara atomowego na rysunku 4.7.

Jednym z głównych celów przy tworzeniu zegara jest osiągnięcie jak najmniejszej niestabilności typowo kwantyfikowanej przez wariancję Allana (AVAR, ang. *Allan Variance*) [Allan 1966; Riehle 2004]

$$\sigma^2(\tau) = \frac{1}{2\tau^2\omega_0^2} \left\langle \left[\int_{\tau}^{2\tau} dt \omega(t) - \int_0^{\tau} dt \omega(t) \right]^2 \right\rangle, \quad (4.30)$$

gdzie $\langle \cdot \rangle$ reprezentuje uśrednianie po fluktuacjach częstości LO opisywanych przez pewien proces stochastyczny (który z punktu widzenia estymacji Bayesowskiej pełni rolę rozkładu *a priori*), τ oznacza czas uśredniania, ω_0 referencyjną częstość atomowego przejścia, a $\omega(t)$ zależną od czasu częstość LO. Wprowadzając częstość ułamkową $y(t) = \frac{\omega(t)}{\omega_0}$ oraz średnią częstość ułamkową $\bar{y}_j(\tau) = \frac{1}{\tau} \int_{j\tau}^{(j+1)\tau} dt y(t)$ możemy AVAR zapisać w postaci

$$\sigma^2(\tau) = \frac{1}{2} \left\langle [\bar{y}_1(\tau) - \bar{y}_0(\tau)]^2 \right\rangle. \quad (4.31)$$

Widzimy teraz, że jeśli mamy pewną wielkość $y(t)$ zależną od czasu t , podzielimy przedział czasu na odcinki długości τ i dla każdego takiego odcinka policzymy średnią wartość wielkości $y(t)$ (czyli $\bar{y}_j(\tau)$, gdzie j numeruje kolejne przedziały czasu) to wariancja Allana informuje nas jak średnio różni się wartość $\bar{y}_j(\tau)$ w sąsiednich przedziałach czasu. Sprawia to, że AVAR jest czuła na korelacje występujące w szumie oraz jest dobrą miarą stabilności sygnału. AVAR można powiązać bezpośrednio z wielkościami charakteryzującymi szum, takimi jak funkcja autokorelacji oraz widmowa gęstość mocy i wyznaczyć w sposób ścisły [Riehle 2004]. Na przykład dla białego szumu $\sigma^2(\tau) \sim \tau^{-1}$ – jest to też zachowanie, którego należy oczekiwać dla lokalnie skorelowanego szumu, gdy τ jest dużo większe od zasięgu korelacji.

Ustalając fizyczne własności zespołu atomów, naszym celem jest optymalizacja stanu początkowego atomów, czasu oddziaływania, pomiarów oraz sygnału korekcyjnego (z pętli sprzężenia zwrotnego), tak aby zminimalizować AVAR. Tak kompleksowa optymalizacja jest praktycznie niewykonywalna. Na szczęście, w pracy [Chabuda, Leroux i Demkowicz-Dobrzański 2016] udało się wyprowadzić ograniczenie dolne na osiągalną wartość AVAR, które nazwano kwantową wariancją Allana (QAVAR, ang. *Quantum Allan Variance*):

$$\sigma_Q^2(\tau) = \sigma_{LO}^2(\tau) - \max_{\rho_0, \bar{\Lambda}, T} F_A(\tau) [\rho_0, \bar{\Lambda}, T], \quad (4.32)$$

$$F_A(\tau) [\rho_0, \bar{\Lambda}, T] = \frac{1}{2\omega_0^2} [2 \text{Tr}(\bar{\rho}' \bar{\Lambda}) - \text{Tr}(\bar{\rho} \bar{\Lambda}^2)], \quad (4.33)$$

gdzie $\sigma_{LO}^2(\tau)$ to AVAR swobodnie działającego LO, $F_A(\tau) [\rho_0, \bar{\Lambda}, T]$ reprezentuje poprawkę do niej wynikającą z pętli sprzężenia zwrotnego, a T to czas oddziaływania.

Operatory $\bar{\rho}$, $\bar{\rho}'$ i $\bar{\Lambda}$ są odpowiednikami operatorów $\bar{\rho}$, $\bar{\rho}'$ i $\bar{\Lambda}$ występujących we wzorze (1.39) na AMSE w kwantowym podejściu Bayesowskim. Istotne jest to, że jeśli zespół atomów, na których bazuje zegar atomowy jest reprezentowany za pomocą stanu na d wymiarowej przestrzeni Hilberta $\mathcal{H}$ to operatory $\bar{\rho}$ i

$\bar{\rho}'$ działają na przestrzeń tensorową $\mathcal{H}^{\otimes N}$, gdzie $N = 2\frac{\tau}{T} - 1$ jest liczbą cykli działania zegara, które trzeba uwzględnić, aby wyznaczyć QAVAR. Zakładamy tutaj wyidealizowaną sytuację, w której pojedynczy cykl działania zegara trwa T – czyli że cały czas cyklu jest poświęcony na oddziaływanie atomów z LO i nie ma tak zwanego „martwego czasu”. Dokładna postać tych operatorów to:

$$\bar{\rho} = \left\langle \bigotimes_{n=1}^N \rho_{\omega_{\text{LO}}}^{[n]} \right\rangle, \quad \bar{\rho}' = \left\langle \frac{1}{\tau} \int_0^\tau dt [\omega_{\text{LO}}(t + \tau) - \omega_{\text{LO}}(t)] \bigotimes_{n=1}^N \rho_{\omega_{\text{LO}}}^{[n]} \right\rangle, \quad (4.34)$$

gdzie $\omega_{\text{LO}}(t)$ oznacza częstość swobodnie działającego LO (bez poprawek wynikających z działania pętli sprzężenia zwrotnego), a $\rho_{\omega_{\text{LO}}}^{[n]}$ oznacza macierz gęstości na wyjściu kanału kwantowego w n -tym cyklu działania zegara, ale której ewolucja (przez czas T) odbywała się jedynie pod wpływem $\omega_{\text{LO}}(t)$, a nie pełnego $\omega(t)$, czyli

$$\rho_{\omega_{\text{LO}}}^{[n]} = e^{-iH \int_{(n-1)T}^{nT} dt \omega_{\text{LO}}(t)} \Phi[\rho_0] e^{iH \int_{(n-1)T}^{nT} dt \omega_{\text{LO}}(t)}, \quad (4.35)$$

gdzie Φ oznacza kanał kwantowy opisujący dekoherencję atomów (będącą wynikiem oddziaływania z otoczeniem, jak np. defazowanie, straty). Jako że we współczesnych zegarach atomowych głównym źródłem szumu są fluktuacje LO, to będziemy zakładać, że kanał Φ działa w sposób trywialny, czyli $\Phi[\rho_0] = \rho_0$.

Zespół atomów będziemy traktować jako zbiór $d - 1$ atomów dwupoziomowych z częstością przejścia równą referencyjnej częstości zegara atomowego. Efekt działania szumu LO jest bardzo podobny do kolektywnego defazowania atomów, a co najważniejsze, nie wyprowadza stanu z podprzestrzeni całkowicie symetrycznej. Uprasza to znacząco obliczenia, ponieważ pozwala nam ograniczyć rozważania do tej podprzestrzeni, która dla $d - 1$ atomów dwupoziomowych ma wymiar jedynie d (cała przestrzeń Hilberta ma w tym wypadku wymiar 2^{d-1}). W związku z tym będziemy używać notacji $|\psi\rangle = \sum_{j=0}^{d-1} a_j |j\rangle$, w której $|j\rangle$ reprezentuje stan całkowicie symetryczny, gdzie j atomów jest w stanie wzbudzonym, a $d - 1 - j$ w stanie podstawowym. Zapisywanie informacji o odstrojeniu $\delta\omega(t)$ częstości LO $\omega(t)$ od częstości referencyjnej atomowego przejścia ω_0 , w stanie kwantowym, następuje w wyniku interferometrii Ramseya. Prowadzi ona efektywnie zespół atomów do stanu $|\psi\rangle_T = \sum_{j=0}^{d-1} a_j e^{-ij \int_0^T dt \delta\omega(t)} |j\rangle$. Jako że będziemy optymalizować po stanach początkowych i pomiarach to pomijamy zachodzące na początku i końcu interferometrii Ramseya impulsy $\frac{\pi}{2}$, traktując je jako składowe odpowiednio optymalnych stanów początkowych oraz pomiarów.

Należy podkreślić, że opisywany w poprzednim akapicie zespół atomów należy do pojedynczej przestrzeni Hilberta $\mathcal{H}$. W naszym modelu kolejne cykle działania zegara są formalnie reprezentowane przez kolejne człony w strukturze produktowej przestrzeni $\mathcal{H}^{\otimes N}$, a w formalizmie sieci tensorowych przez kolejne tensory w sieci. Kluczową dla nas własnością jest to, że fluktuacje częstości LO są skorelowane w czasie. Efektywnie prowadzi to do sytuacji, w której atomy są poddane kolektywnemu defazowaniu, którego działanie pomiędzy kolejnymi cyklami zegara jest ze sobą skorelowane.

Zakładając że fluktuacje szumu LO mają skończony zasięg korelacji w czasie,

możemy spodziewać się, że QAVAR będzie można efektywnie wyznaczyć używając sieci tensorowej w postaci łańcucha o długości N z $(d-1)$ -wymiarowym indeksem fizycznym na każdym węźle. Poza oczywistą zaletą w postaci poprawy wydajności numerycznej, użycie sieci tensorowych pozwala nam zawęzić klasę stanów początkowych, po których wykonujemy optymalizację do stanów produktowych (rząd wiązania $\chi_{\psi_0} = 1$). Fizycznie odpowiada to typowej sytuacji, w której zespoły atomów w różnych krokach czasowych są niezależne od siebie, ponieważ zostały przygotowane od nowa na początku każdego kolejnego cyklu działania zegara – podkreślić należy, że ciągle będziemy rozważać splątanie pomiędzy atomami tworzącymi zespół i oddziałującymi z LO w ramach danego kroku czasowego.

Fluktuacje LO mogą być scharakteryzowane za pomocą funkcji autokorelacji $R(t)$, którą na potrzeby naszego przykładu wybraliśmy jako kombinację procesu Ornstein–Uhlenbeck (OU) i białego Gaussowskiego szumu:

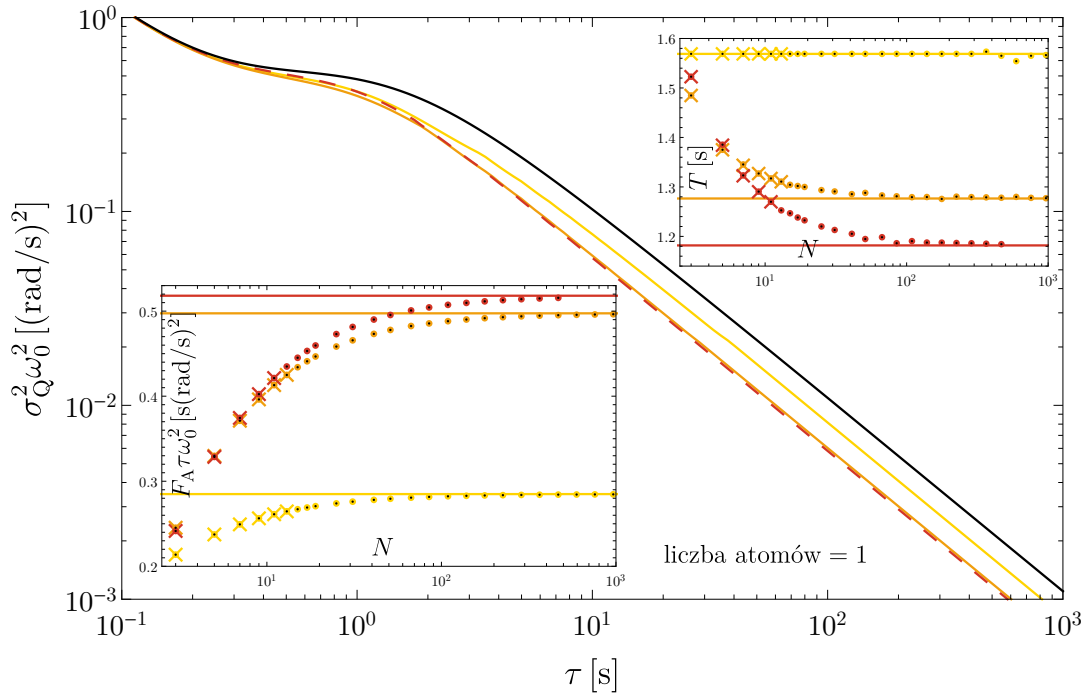
$$R(t) = \alpha e^{-\gamma t} + \beta \delta(t), \quad (4.36)$$

gdzie możemy zinterpretować parametry α , β jako siłę odpowiednio procesu OU oraz białego szumu, a $\frac{1}{\gamma}$ jako zasięg korelacji procesu OU. Wybraliśmy $\alpha = 1 \text{ (rad/s)}^2$, $\beta = 0,1 \text{ (rad/s)}^2\text{s}$, $\gamma = 2 \text{ s}^{-1}$, dla których optymalny czas oddziaływania T wyniósł około 1,4 s. Oznacza to, że dla tych parametrów korelacje szumu są na tyle słabe, że istotnie wpływają jedynie na sąsiadujące ze sobą w czasie podukłady. Rozważanie szumu o większym zasięgu korelacji jest możliwe, ale wymaga większych zasobów obliczeniowych z powodu większego wymiaru wiązania. Parametry te są zbliżone do tych użytych w pracy [Chabuda, Leroux i Demkowicz-Dobrzański 2016]: $\alpha = 2 \text{ (rad/s)}^2$, $\beta = 0,4 \text{ (rad/s)}^2\text{s}$, $\gamma = 0,5 \text{ s}^{-1}$, a które były dobrane tak, aby jak najwierniej oddać funkcję autokorelacji szumu zegara atomowego opisanego w pracy [Jiang i in. 2011]. Dla szumu opisanego procesem OU, AVAR swobodnie działającego LO ma postać

$$\sigma_{\text{LO}}^2(\tau) = \frac{1}{\tau \omega_0^2} \left[\frac{2\alpha}{\gamma} + \frac{\alpha}{\gamma^2 \tau} (4e^{-\gamma\tau} - e^{-2\gamma\tau} - 3) + \beta \right] \quad (4.37)$$

i w najbardziej interesującym reżimie długiego czasu uśredniania przyjmuje formę $\sigma_{\text{LO}}^2(\tau) \simeq (2\alpha\gamma^{-1} + \beta)/(\tau\omega_0^2)$. Jako że proces OU charakteryzuje się skończonym zasięgiem korelacji spodziewamy się, że w tym reżimie QAVAR również przyjmuje postać $\sigma_{\text{Q}}^2(\tau) \simeq c_{\infty}/(\tau\omega_0^2)$ z pewną stałą c_{∞} , którą będziemy nazywać współczynnikiem asymptotycznym QAVAR. Równoważnie oznacza to, że wraz ze wzrostem liczby cykli działania zegara optymalna wartość poprawki wynikającej z istnienia pętli sprzężenia zwrotnego przyjmuje postać $F_{\text{A}}(\tau) \simeq \tilde{c}_{\infty}/(\tau\omega_0^2)$ z pewną stałą $\tilde{c}_{\infty}$.

Próba wyznaczenia współczynnika asymptotycznego QAVAR została podjęta w pracy [Chabuda, Leroux i Demkowicz-Dobrzański 2016]. Jednak z powodu użycia standardowych metod pełnej przestrzeni Hilberta nie było możliwe osiągnięcie reżimu, w którym charakter skalowania QAVAR oraz wartość asymptotycznego współczynnika mogłaby być jednoznacznie odczytana. Podejście oparte na sieciach tensorowych opisane w tej dysertacji pozwoliło obliczyć QAVAR w poprzednio nieosiągalnym reżimie dużych czasów uśredniania τ . Na



RYSUNEK 4.8: QAVAR σ_Q^2 (pomnożona przez ω_0^2) w funkcji czasu uśredniania τ dla zegara atomowego (bazującego na jednym atomie) charakteryzującego się szumem LO, który jest ściśle lokalny (wyniki na żółto), zawiera również korelacje pomiędzy sąsiadującymi przedziałami czasu (wyniki na pomarańczowo) lub zawiera ponadto korelacje z jednym następnym przedziałem czasu (wyniki na czerwono), w porównaniu do AVAR swobodnie działającego LO (czarna linia). Lewa wstawka pokazuje optymalną wartość poprawki F_A , do AVAR swobodnie działającego LO, wynikającą z działania pętli sprzężenia zwrotnego (pomnożoną przez $\tau \omega_0^2$) w funkcji liczby cykli działania zegara N (liczby węzłów w sieci tensorowej) potrzebnych do obliczenia AVAR dla danego czasu uśredniania τ i optymalnego czasu oddziaływania T – którego wykres został przedstawiony na prawej wstawce. Wyniki na wstawkach prezentują porównanie rezultatów uzyskanych w podejściu MPO (kropki) z asymptotycznymi wartościami otrzymanymi w podejściu iMPO (linie) oraz wartościami uzyskanymi przy pomocy standardowych metod wykorzystujących obliczenia w pełnej przestrzeni Hilberta (krzyże).

rysunku 4.8 przedstawiono przykładowe wyniki dla zegara atomowego bazującego na pojedynczym atomie dwupoziomowym, które pokazują, że korelacje pełnią istotną rolę i nie można uzyskać precyzyjnych fundamentalnych ograniczeń na stabilność zegara atomowego pomijając korelacje fluktuacji LO w czasie. Z drugiej strony widzimy, że dla rozważanego modelu szumu wystarczy ograniczyć rozważania do korelacji pomiędzy najbliższymi przedziałami czasu, ponieważ uwzględnienie dalszych korelacji już nie daje istotnej poprawy QAVAR. W powiększeniu widać to na lewej wstawce, gdzie optymalna poprawka do AVAR F_A zwiększa się o około 75% po uwzględnieniu korelacji pomiędzy najbliższymi przedziałami czasu (wyniki na pomarańczowo w porównaniu do tych na żółto), a uwzględnienie korelacji z jeszcze kolejnym przedziałem czasu daje już tylko poprawę o około 4% (wyniki na czerwono w porównaniu do tych na pomarańczowo). Wstawki pokazują również porównanie wyników uzyskanych w

podejściu MPO z asymptotyczną wartością uzyskaną podejściem iMPO. Obserwujemy, że podobnie jak w przypadku estymacji pola magnetycznego (patrz rysunek 4.2) wyniki w podejściu MPO oraz iMPO są ze sobą w bardzo dobrej zgodności, a standardowe metody obliczeń w pełnej przestrzeni Hilberta nie umożliwiają dojścia do reżimu, w którym można by wnioskować o zachowaniu asymptotycznym. Wstawki pokazują również, że zachowanie asymptotyczne osiągamy dopiero po około 100 krokach czasowych. Oznacza to, że jeśli chcielibyśmy użyć metod pełnej przestrzeni Hilberta to musielibyśmy wykonywać obliczenia w przestrzeni wymiaru 2^{100} – co jest oczywiście zadaniem niemożliwym do wykonania. W celu utrzymania względnego błędu poniżej 0,1% przy wykonywaniu obliczeń numerycznych na potrzeby rysunku 4.8 rząd wiązania wyniósł maksymalnie $\chi_\Lambda = 4$ w podejściu MPO oraz $\chi_{\bar{\Lambda}} = 4$ w podejściu iMPO.

Podobnie jak w przypadku estymacji pola magnetycznego podejście iMPO pozwala efektywnie uzyskać bezpośredni wgląd w asymptotyczne ($\tau \rightarrow \infty$) zachowanie QAVAR. Zauważmy jednak, że definicja AVAR (4.30) prowadzi do postaci MPO dla $\bar{\rho}'$ w wyrażeniu na QAVAR, która jawnie nie jest translacyjnie niezmiennicza. W związku z tym, aby zastosować podejście iMPO przybliżymy AVAR przez asymptotycznie równoważne wyrażenie

$$\sigma^2(\tau) \simeq \frac{1}{\tau^2 \omega_0^2} \left\langle \left(\int_0^\tau dt \omega(t) \right)^2 \right\rangle, \quad (4.38)$$

które pokrywa się z poprzednią definicją dla τ dużo większych niż zasięg korelacji (czyli dokładnie w reżimie, w którym działa podejście iMPO). Używając tego podejścia wyznaczyliśmy współczynnik asymptotyczny QAVAR c_∞ i odpowiadającą mu wartość optymalnego czasu oddziaływania T w funkcji liczby atomów w zegarze – patrz rysunek 4.9. Z tego wykresu widzimy, że efekt uwzględnienia korelacji fluktuacji LO przy obliczaniu QAVAR jedynie się powiększa wraz ze wzrostem liczby atomów w zegarze. To podkreśla, że korelacje odgrywają istotną rolę jeśli chcemy dokonać precyzyjnej analizy możliwości stabilizacji działania zegara atomowego. Dopuszczenie względnego błędu wyników na poziomie około 2%, sprawiło że na potrzeby rysunków 4.9 i 4.10 wystarczyło przeprowadzić obliczenia numeryczne (w podejściu iMPO) z rzędem wiązania jedynie $\chi_{\bar{\Lambda}} = 1$.

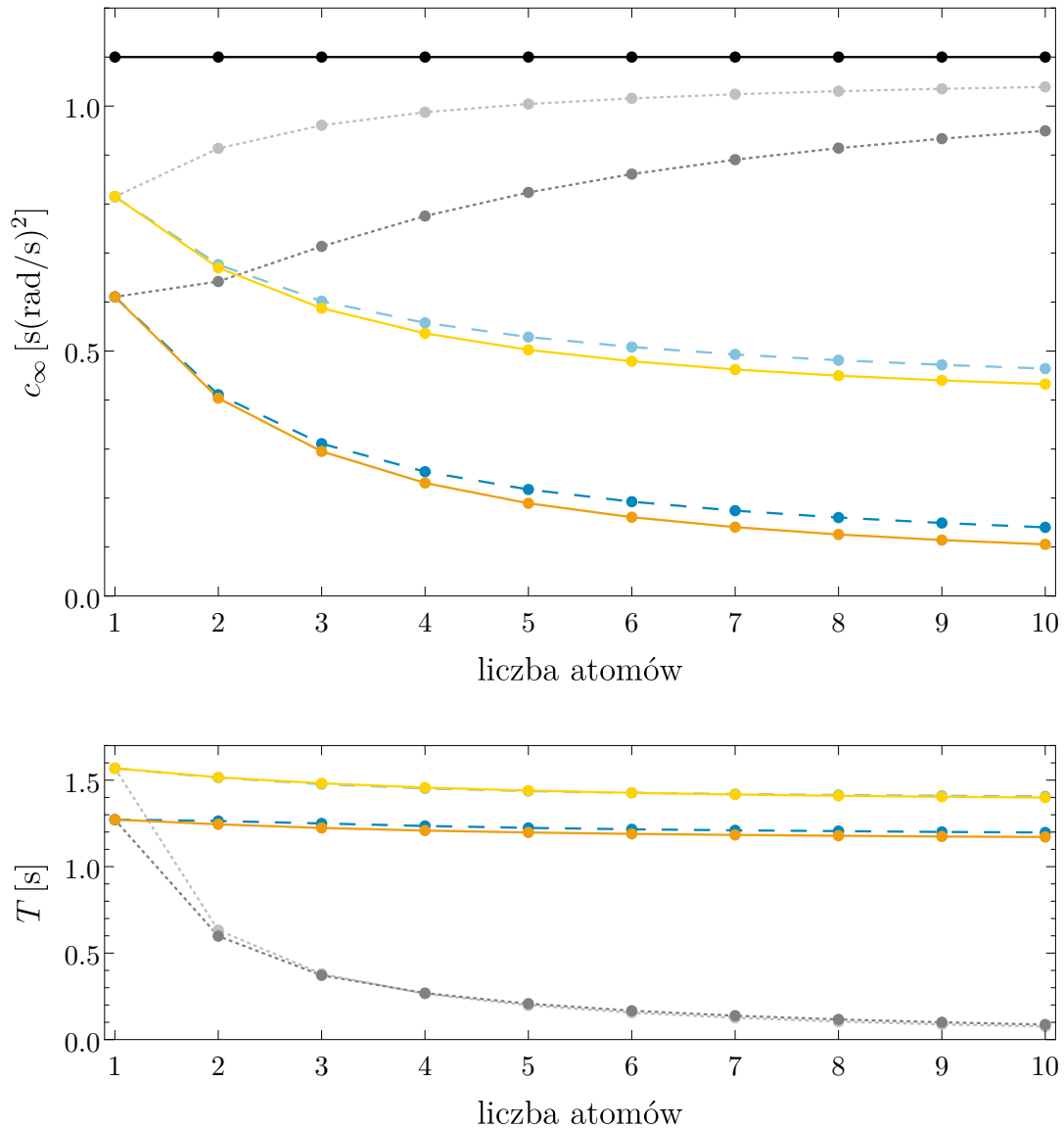
Przyjrzyjmy się teraz wynikom jakie można uzyskać używając jako stanu początkowego stanu GHZ (1.46) w porównaniu do rezultatów, które są zoptymalizowane po stanach początkowych. W szczególnym przypadku, gdy mamy jedynie lokalne korelacje można wyprowadzić ścisły wzór na asymptotyczną postać QAVAR dla stanu GHZ:

$$\sigma_{\text{Q,GHZ}}^2(\tau) \simeq \frac{c_1^2(d-1)^2 e^{-c_1(d-1)^2}}{\tau \omega_0^2}, \quad (4.39)$$

gdzie lokalny parametr szumu c_1 dla szumu opisanego funkcją autokorelacji (4.36) ma postać

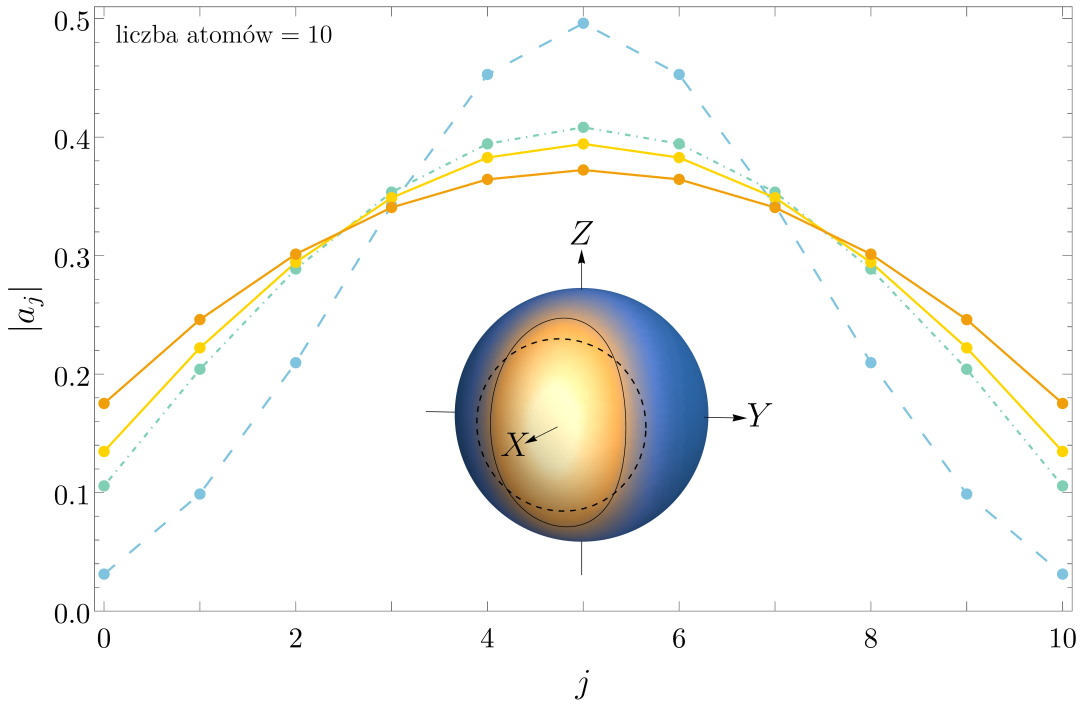
$$c_1 = \frac{2\alpha}{\gamma^2} (\gamma T + e^{-\gamma T} - 1) + \beta T. \quad (4.40)$$

W ogólności stan GHZ jest bardzo wrażliwy na szum defazujący. W związku z czym optymalne czasy oddziaływania T będą musiały być wyraźnie skrócone



RYSUNEK 4.9: Górny wykres przedstawia współczynnik asymptotyczny QAVAR c_∞ , a dolny optymalny czas oddziaływania T , w funkcji liczby atomów w zegarze atomowym dla strategii estymacji bazującej na stanie początkowym który jest: stanem GHZ (kropkowana linia), stanem spinowym koherentnym (przerywana linia), stanem optymalnym (ciągła linia). Zachowanie każdego z tych trzech rodzajów stanów przeanalizowano w przypadku gdy szum LO jest ściśle lokalny (jaśniejszy odcień koloru szarego/niebieskiego/żółtego) lub zawiera również korelacje pomiędzy sąsiadującymi przedziałami czasu (ciemniejszy odcień koloru szarego/niebieskiego/żółtego). Na górnym wykresie dla porównania wykreślono również współczynnik asymptotyczny AVAR (czarna ciągła linia).

względem czasów dla optymalnych stanów. Jest to widoczne na rysunku 4.9, gdzie mimo iż optymalne T dla stanu GHZ skaluje się jak $(d-1)^{-2}$ to szum szybko niszczy jakikolwiek zysk wynikający z istnienia pętli sprzężenia zwrotnego i współczynnik asymptotyczny QAVAR wraz ze wzrostem liczby atomów w zegarze dąży do wartości charakteryzującej swobodnie działający LO. Jest to



RYSUNEK 4.10: Rozkład bezwzględnych wartości amplitud prawdopodobieństwa a_j stanów początkowych dla zegara atomowego bazującego na 10 atomach. Ciągłe linie odpowiadają stanom optymalnym odpowiednio w przypadku gdy rozważamy szum LO, który jest ściśle lokalny (kolor żółty) lub zawiera również korelacje pomiędzy sąsiadującymi przedziałami czasu (kolor pomarańczowy). Dla porównania wykreślono $|a_j|$ dla stanu spinowego koherentnego (CSS, jasnoniebieska przerywana linia) oraz stanu SIN (1.47) (jasnozielona przerywano-kropkowana linia). W środku znajduje się rozkład Q Husimiego na sferze Blocha dla optymalnego stanu (w przypadku szumu LO zawierającego również korelacje pomiędzy sąsiadującymi przedziałami czasu) z zaznaczonymi przykładowymi liniami ekwipotencjalnymi odpowiadającymi tej samej gęstości prawdopodobieństwa dla CSS (czarna przerywana linia) oraz optymalnego stanu (czarna linia) – co pokazuje, że jest on ściśnięty.

manifestacja ogólnie słabych osiągnięć stanu GHZ dla układów metrologicznych w realistycznych, uwzględniających szum, modelach wraz ze wzrostem liczby cząstek [Huelga i in. 1997; Escher, Matos Filho i Davidovich 2011; Demkowicz-Dobrzański, Kołodyński i Guţă 2012].

Nasze podejście do metrologii, bazujące na sieciach tensorowych, pozwala łatwo badać optymalne stany początkowe. Na rysunku 4.10 przedstawiliśmy wykres wartości bezwzględnych amplitud prawdopodobieństwa a_j dla optymalnych stanów dla przypadku szumu LO, który jest ściśle lokalny lub zawiera również korelacje pomiędzy sąsiadującymi przedziałami czasu. Dla porównania wykreślono również wartości a_j dla stanu spinowego koherentnego (CSS, ang. *Coherent Spin State*) oraz stanu SIN (1.47) – który jest optymalnym stanem początkowym w problemie Bayesowskiej estymacji z płaskim rozkładem *a priori*. Widzimy, że optymalne stany są nieklasyczne, co możemy określić ilościowo obliczając ich spinowy parametr ściśnięcia $\xi^2 = 2\Delta^2 S_y / |\langle S_x \rangle|$ [J. Ma i in.

2011], który wynosi 1 dla CSS, 0,522 dla stanu SIN oraz 0,456 i 0,378 dla optymalnych stanów odpowiednio w przypadku szumu LO, który jest ściśle lokalny lub zawiera również korelacje pomiędzy sąsiadującymi przedziałami czasu. Ścisłowanie w tym ostatnim przypadku możemy również zaobserwować na rozkładzie Q Husimiego na sferze Blocha znajdującym się w centrum rysunku 4.10.

4.3 Wierność wielociałowego stanu termicznego

Trzeci z prezentowanych przez nas przykładów wykracza poza dziedzinę metrologii kwantowej i demonstruje użyteczność techniki liczenia QFI dla stanów mieszanych z użyciem MPO w celu bezpośredniego wyznaczenia wierności dla wielociałowego stanu termicznego – co dla dużej liczby cząstek było do tej pory poza zasięgiem najbardziej wyrafinowanych metod.

Z kwantowym przejściem fazowym [Sachdev 1999] mamy do czynienia, gdy mała zmiana zewnętrznego parametru powoduje ogromną zmianę własności stanu podstawowego układu kwantowego. Tym zewnętrznym parametrem może być np. siła pola magnetycznego dla układu spinów lub stężenie domieszki dla wysokotemperaturowego nadprzewodnika. Centralne znaczenie ma pojawiająca się nieanalityczność funkcji falowej stanu podstawowego, gdy przechodzimy przez punkt krytyczny φ_c oddzielający dwie różne fazy.

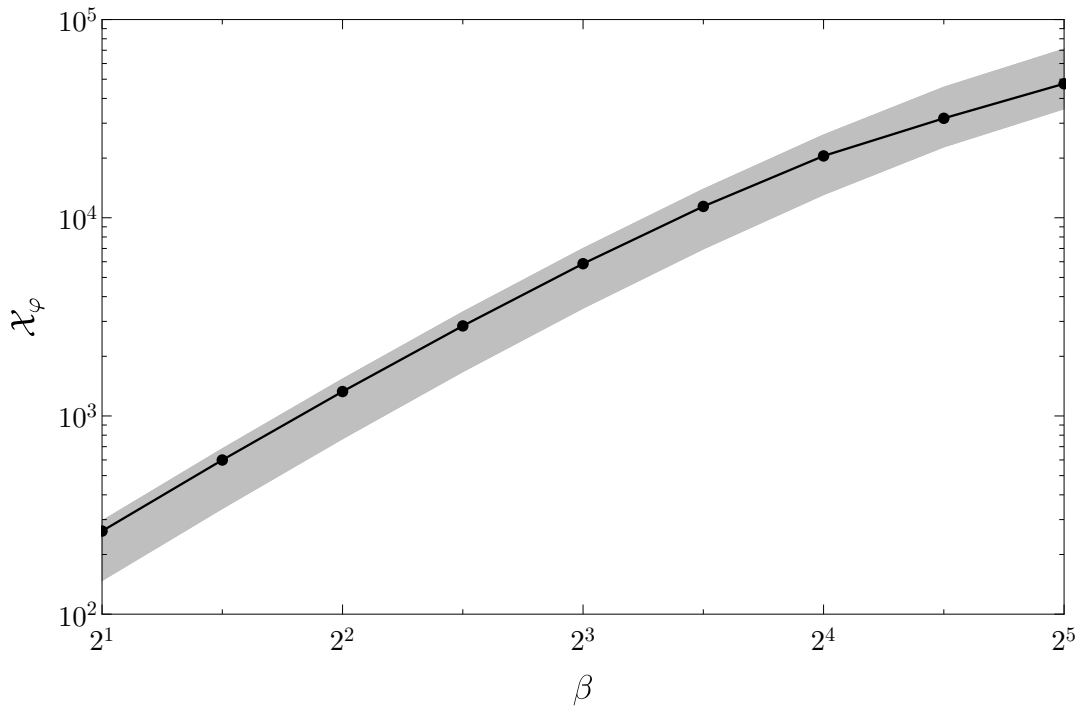
Jeden z możliwych sposobów detekcji kwantowego przejścia fazowego polega na badaniu zachowania wierności $\mathcal{F}(\rho_\varphi, \rho_{\varphi+\varepsilon}) = \text{Tr} \sqrt{\sqrt{\rho_\varphi} \rho_{\varphi+\varepsilon} \sqrt{\rho_\varphi}}$ (patrz rozdział 1.2.1) pomiędzy dwoma wielociałowymi stanami ρ_φ i $\rho_{\varphi+\varepsilon}$, które różnią się o małą wartość parametru φ w Hamiltonianie [Zanardi i Paunković 2006; Rams i Damski 2011]. Gdy następuje kwantowe przejście fazowe to podatność wierności χ_φ zdefiniowana jako $\mathcal{F}(\rho_\varphi, \rho_{\varphi+\varepsilon}) \approx 1 - \frac{1}{2} \chi_\varphi \varepsilon^2$ ma maksimum oznaczające fundamentalną zmianę w stanie układu. Zgodnie ze wzorem (1.34) oznacza to również maksimum QFI, ponieważ jest ona powiązana z podatnością wierności zależnością $F[\rho_\varphi] = 4\chi_\varphi$. Tą koncepcją posłużono się w pracy [Rams, Sierant i in. 2018], aby zbadać użyteczność kwantowych przejść fazowych, które zachodzą w zerowej temperaturze do precyzyjnych pomiarów parametru φ w rzeczywistych układach w skończonej temperaturze.

W odróżnieniu od przypadku zerowej temperatury, wierność pomiędzy dwoma wielociałowymi stanami termicznymi reprezentowanymi przez MPO nie jest w ogólności łatwa do wyznaczenia. Z tego powodu wprowadzono kwaziwierność $\tilde{\mathcal{F}}(\rho_\varphi, \rho_{\varphi+\varepsilon}) = \sqrt{\text{Tr} \sqrt{\rho_\varphi} \sqrt{\rho_{\varphi+\varepsilon}}}$ i jej kwazipodatność $\tilde{\chi}_\varphi$ $\tilde{\mathcal{F}}(\rho_\varphi, \rho_{\varphi+\varepsilon}) \approx 1 - \frac{1}{2} \tilde{\chi}_\varphi \varepsilon^2$, która dostarcza ograniczenia na dokładną wartość podatności wierności [Albuquerque i in. 2010; Sirker 2010]:

$$\tilde{\chi}_\varphi \leq \chi_\varphi \leq 2\tilde{\chi}_\varphi. \quad (4.41)$$

Hamiltonian rozważany w pracy [Rams, Sierant i in. 2018] opisywał model XX dla spinów- $\frac{1}{2}$:

$$H = - \sum_{n=1}^{N-1} \left(\sigma_x^{[n]} \sigma_x^{[n+1]} + \sigma_y^{[n]} \sigma_y^{[n+1]} \right) + \varphi \sum_{n=1}^N \sigma_x^{[n]} \quad (4.42)$$



RYSUNEK 4.11: Dokładna wartość podatności wierności χ_φ dla stanu termicznego 64 spinów w punkcie krytycznym w modelu XX (kropki połączone linią) w funkcji bezwymiarowej odwrotności temperatury β . Zacieniony obszar pokazuje wartości spełniające ograniczenia (4.41). Jak przewidziano w pracach [Albuquerque i in. 2010; Sirker 2010] dokładna wartość zbliża się do górnego/dolnego ograniczenia dla wysokich/niskich temperatur.

z kwantowym punktem krytycznym dla $\varphi_c = 0$.

Używając MPO otrzymanych w pracy [Rams, Sierant i in. 2018] udało nam się obejść problemy ze złożonością bezpośredniego obliczania wierności poprzez zastosowanie podejścia opartego o sieci tensorowe do wyznaczania QFI (i tym samym podatności wierności) używając w tym przypadku dyskretnej wersji pochodnej $\dot{\rho}_\varphi = (\rho_{\varphi+\varepsilon} - \rho_\varphi)/\varepsilon$. Tak wyznaczoną dokładną wartość podatności wierności dla łańcucha 64 spinów zaprezentowaliśmy na rysunku 4.11 łącznie z górnym i dolnym ograniczeniem (4.41). Precyzja wyznaczenia podatności wierności jest ograniczona przez skończony wymiar wiązania χ_Λ , jak również skończoną wartość parametru ε . Mimo to udało nam się uzyskać satysfakcjonujące rezultaty ze względnym błędem około 1% dla $\varepsilon = 10^{-4}$ i $\chi_\Lambda = 4$.

Ten przykład pokazuje, że możliwość wyznaczania podatności wierności jest użytecznym produktem ubocznym naszego ogólnego algorytmu. Wychodząc poza kontekst metrologiczny, toruje on drogę do uogólnienia zero-temperaturowego podejścia wykorzystującego wierność do detekcji kwantowych przejść fazowych [Zanardi i Paunković 2006; Rams i Damski 2011] – aktualnie standardowej techniki wykorzystywanej w fizyce materii skondensowanej – do badania przejść fazowych w kwantowych układach w skończonej temperaturze.

Rozdział 5

Podsumowanie

W niniejszej dysertacji udało się przedstawić kompleksowy formalizm oparty o sieci tensorowe do analizy problemów metrologii kwantowej. Pozwolił on uniknąć notorycznie trapiących tą dziedzinę problemów z eksponencjalnym wzrostem wymiaru przestrzeni Hilberta, gdy chcemy badać coraz większe układy. Zaproponowany formalizm okazał się być niezwykle wydajny, znacząco przesuwając granicę możliwych zastosowań metrologii kwantowej. Jednakże nie można go traktować jedynie jako rewolucyjną metodę numeryczną pozwalającą opisać układy wielociałowe. Efektywność opisu za pomocą sieci tensorowych dla układów z lokalnie skorelowanym szumem pokazuje, że w takich problemach stany początkowe oraz SLD mają strukturę splątania analogiczną jak badany model, czyli jednowymiarową z lokalnymi korelacjami. Co szczególnie wyjątkowe dla naszego podejścia to możliwość bezpośredniego wyznaczenia fundamentalnych ograniczeń na precyzję estymacji od razu w granicy asymptotycznej, gdy $N \rightarrow \infty$.

Zaproponowany formalizm pozwolił szczegółowo przestudiować wpływ korelacji w szumie na precyzję estymacji pola magnetycznego. Ujawnił słuszość bardzo ciekawej zależności funkcyjnej pomiędzy QFI uzyskaną dla stanu optymalnego oraz dla stanu produktowego: $\frac{F}{N} = \frac{F_{\text{prod}}}{N} / (1 - \frac{F_{\text{prod}}}{N})$, pokazującą że im więcej informacji o parametrze możemy uzyskać z układu dla stanu klasycznego tym większy jest możliwy zysk kwantowy z użycia stanów splątanych oraz że zysk ten nie zależy od żadnego innego parametru. Uzyskane wyniki pozwoliły pokazać, że (tak jak w przypadku nieskorelowanego szumu) stany ściśnięte pozwalają uzyskać optymalną precyzję estymacji w obecności lokalnie skorelowanego szumu. To z kolei umożliwiło zaproponowanie postaci QFI w problemie estymacji fazy w obecności szumu defazującego o dowolnym zasięgu korelacji. Użycie sieci tensorowych sprawiło również, że po raz pierwszy udało się wyznaczyć fundamentalne ograniczenie na stabilność działania zegara atomowego – QAVAR. W efekcie możliwe było pokazanie, że korelacje w szumie LO pełnią istotną rolę i nie można ich pominąć, gdy chce się precyzyjnie wyznaczyć fundamentalne ograniczenie na stabilność działania zegara atomowego. Mogliśmy też zbadać jakie są optymalne stany w zegarze atomowym i między innymi wyznaczyć optymalny współczynnik ściśnięcia. Zaproponowany formalizm okazał się również użyteczny poza dziedziną metrologii kwantowej, pokazując możliwość zastosowania przy badaniu kwantowych przejść fazowych.

Wierzmy że zaproponowany formalizm ma jeszcze wielki, nieodkryty potencjał. Od strony sieci tensorowych zbadania wymaga wpływ symetrii występujących w badanym problemie na optymalne tensory [Singh i Vidal 2012; Weichselbaum 2012], a także próba dalszej optymalizacji numerycznej w celu wykorzystania potencjału współczesnych klastrów komputerowych do badania dużo bardziej złożonych układów. Od strony metrologicznej wyzwaniem na przyszłość pozostaje próba adaptacji formalizmu do jeszcze trudniejszych problemów, takich jak wieloparametrowa estymacja [Ragy, Jarzyna i Demkowicz-Dobrzański 2016; Szczykulska, Baumgratz i Datta 2016; Baumgratz i Datta 2016], estymacja sygnału [Tsang, Wiseman i Caves 2011; Berry, M. J. W. Hall i Wiseman 2013] czy badanie efektywności adaptacyjnych protokołów metrologicznych, w tym również takich bazujących na kwantowych kodach korekcji błędów [Arrad i in. 2014; Dür i in. 2014; Zhou i in. 2018; Layden i Cappellaro 2018; Górecki i in. 2019]. Uważamy również, że zaproponowany formalizm może być kluczowy dla zrozumienia modeli metrologicznych z czasowo skorelowanym szumem, w szczególności tych o charakterze nie-Markowskim [Chin, Huelga i Plenio 2012], dla których jeszcze nie powstały żadne narzędzia pozwalające wyznaczyć efektywne strategie metrologiczne.

Bibliografia

- Aasi, J. i in. (2013). „*Enhanced sensitivity of the LIGO gravitational wave detector by using squeezed states of light*”. **Nature Photonics** **7**, 613.
- Abadie, J. i in. (2011). „*A gravitational wave observatory operating beyond the quantum shot-noise limit*”. **Nature Physics** **7**, 962.
- Acernese, F. i in. (2019). „*Increasing the Astrophysical Reach of the Advanced Virgo Detector via the Application of Squeezed Vacuum States of Light*”. **Physical Review Letters** **123**, 231108.
- Afek, I., O. Ambar i Y. Silberberg (2010). „*High-NOON States by Mixing Quantum and Classical Light*”. **Science** **328**, 879.
- Albuquerque, A. F., F. Alet, C. Sire i S. Capponi (2010). „*Quantum critical scaling of fidelity susceptibility*”. **Physical Review B** **81**, 064418.
- Allan, D. W. (1966). „*Statistics of atomic frequency standards*”. **Proceedings of the IEEE** **54**, 221.
- Altenburg, S., M. Oszmaniec, S. Wölk i O. Gühne (2017). „*Estimation of gradients in quantum metrology*”. **Physical Review A** **96**, 042319.
- André, A., A. S. Sørensen i M. D. Lukin (2004). „*Stability of Atomic Clocks Based on Entangled Atoms*”. **Physical Review Letters** **92**, 230801.
- Apellaniz, I., I. Urizar-Lanz, Z. Zimborás, P. Hyllus i G. Tóth (2018). „*Precision bounds for gradient magnetometry with atomic ensembles*”. **Physical Review A** **97**, 053603.
- Appel, J., P. J. Windpassinger, D. Oblak, U. B. Hoff, N. Kjærgaard i E. S. Polzik (2009). „*Mesoscopic atomic entanglement for precision measurements beyond the standard quantum limit*”. **Proceedings of the National Academy of Sciences** **106**, 10960.
- Arad, I. i Z. Landau (2010). „*Quantum Computation and the Evaluation of Tensor Networks*”. **SIAM Journal on Computing** **39**, 3089.
- Arrad, G., Y. Vinkler, D. Aharonov i A. Retzker (2014). „*Increasing Sensing Resolution with Error Correction*”. **Physical Review Letters** **112**, 150801.
- Barish, B. C. (2018). „*Nobel Lecture: LIGO and gravitational waves II*”. **Reviews of Modern Physics** **90**, 040502.
- Baumgratz, T. i A. Datta (2016). „*Quantum Enhanced Estimation of a Multi-dimensional Field*”. **Physical Review Letters** **116**, 030801.
- Beau, M. i A. del Campo (2017). „*Nonlinear Quantum Metrology of Many-Body Open Systems*”. **Physical Review Letters** **119**, 010403.
- Bengtsson, I. i K. Życzkowski (2006). *Geometry of Quantum States: An Introduction to Quantum Entanglement*. Cambridge University Press.
- Berry, D. W. i H. M. Wiseman (2000). „*Optimal States and Almost Optimal Adaptive Measurements for Quantum Interferometry*”. **Physical Review Letters** **85**, 5098.

- Berry, D. W., M. J. W. Hall i H. M. Wiseman (2013). „*Stochastic Heisenberg Limit: Optimal Estimation of a Fluctuating Phase*”. *Physical Review Letters* **111**, 113601.
- Boixo, S., S. T. Flammia, C. M. Caves i J. Geremia (2007). „*Generalized Limits for Single-Parameter Quantum Estimation*”. *Physical Review Letters* **98**, 090401.
- Bollinger, J. J. ., W. M. Itano, D. J. Wineland i D. J. Heinzen (1996). „*Optimal frequency measurements with maximally correlated states*”. *Physical Review A* **54**, R4649.
- Borregaard, J. i A. S. Sørensen (2013). „*Near-Heisenberg-Limited Atomic Clocks in the Presence of Decoherence*”. *Physical Review Letters* **111**, 090801.
- Braunstein, S. L. i C. M. Caves (1994). „*Statistical distance and the geometry of quantum states*”. *Physical Review Letters* **72**, 3439.
- Breuer, H.-P. i F. Petruccione (2002). *The Theory of Open Quantum Systems*. Oxford University Press.
- Bridgeman, J. C. i C. T. Chubb (2017). „*Hand-waving and interpretive dance: an introductory course on tensor networks*”. *Journal of Physics A: Mathematical and Theoretical* **50**, 223001.
- Calabrese, P. i J. Cardy (2004). „*Entanglement entropy and quantum field theory*”. *Journal of Statistical Mechanics: Theory and Experiment* **2004**, P06002.
- Caves, C. M. (1981). „*Quantum-mechanical noise in an interferometer*”. *Physical Review D* **23**, 1693.
- Chabuda, K., J. Dziarmaga, T. J. Osborne i R. Demkowicz-Dobrzański (2020). „*Tensor-network approach for quantum metrology in many-body quantum systems*”. *Nature Communications* **11**, 250.
- Chabuda, K., I. D. Leroux i R. Demkowicz-Dobrzański (2016). „*The quantum Allan variance*”. *New Journal of Physics* **18**, 083035.
- Chan, G. K.-L. i S. Sharma (2011). „*The Density Matrix Renormalization Group in Quantum Chemistry*”. *Annual Review of Physical Chemistry* **62**, 465.
- Chaves, R., J. B. Brask, M. Markiewicz, J. Kołodyński i A. Acín (2013). „*Noisy Metrology beyond the Standard Quantum Limit*”. *Physical Review Letters* **111**, 120401.
- Chin, A. W., S. F. Huelga i M. B. Plenio (2012). „*Quantum Metrology in Non-Markovian Environments*”. *Physical Review Letters* **109**, 233601.
- Cirac, J. I., D. Poilblanc, N. Schuch i F. Verstraete (2011). „*Entanglement spectrum and boundary theories with projected entangled-pair states*”. *Physical Review B* **83**, 245134.
- Clerk, A. A., M. H. Devoret, S. M. Girvin, F. Marquardt i R. J. Schoelkopf (2010). „*Introduction to quantum noise, measurement, and amplification*”. *Reviews of Modern Physics* **82**, 1155.
- Corboz, P. (2016). „*Variational optimization with infinite projected entangled-pair states*”. *Physical Review B* **94**, 035133.
- Cormen, T. H., C. E. Leiserson, R. L. Rivest i C. Stein (1990). *Introduction to Algorithms*. The MIT Press.

- Cramer, M., M. B. Plenio, S. T. Flammia, R. Somma, D. Gross, S. D. Bartlett, O. Landon-Cardinal, D. Poulin i Y.-K. Liu (2010). „*Efficient quantum state tomography*”. *Nature Communications* **1**, 149.
- Cuevas, G. D. las, N. Schuch, D. Pérez-García i J. I. Cirac (2013). „*Purifications of multipartite states: limitations and constructive methods*”. *New Journal of Physics* **15**, 123021.
- Czajkowski, J., K. Pawłowski i R. Demkowicz-Dobrzański (2019). „*Many-body effects in quantum metrology*”. *New Journal of Physics* **21**, 053031.
- Czarnik, P. i J. Dziarmaga (2015). „*Variational approach to projected entangled pair states at finite temperature*”. *Physical Review B* **92**, 035152.
- Degen, C. L., F. Reinhard i P. Cappellaro (2017). „*Quantum sensing*”. *Reviews of Modern Physics* **89**, 035002.
- Demkowicz-Dobrzański, R., M. Jarzyna i J. Kołodyński (2015). „*Quantum Limits in Optical Interferometry*”. *Progress in Optics*. Red. E. Wolf. T. 60. Elsevier, s. 345–435.
- Demkowicz-Dobrzański, R. i L. Maccone (2014). „*Using Entanglement Against Noise in Quantum Metrology*”. *Physical Review Letters* **113**, 250801.
- Demkowicz-Dobrzański, R., K. Banaszek i R. Schnabel (2013). „*Fundamental quantum interferometry bound for the squeezed-light-enhanced gravitational wave detector GEO 600*”. *Physical Review A* **88**, 041802.
- Demkowicz-Dobrzański, R., J. Czajkowski i P. Sekatski (2017). „*Adaptive Quantum Metrology under General Markovian Noise*”. *Physical Review X* **7**, 041009.
- Demkowicz-Dobrzański, R., J. Kołodyński i M. Guță (2012). „*The elusive Heisenberg limit in quantum-enhanced metrology*”. *Nature Communications* **3**, 1063.
- Dorner, U. (2012). „*Quantum frequency estimation with trapped ions and atoms*”. *New Journal of Physics* **14**, 043011.
- Dorner, U., R. Demkowicz-Dobrzański, B. J. Smith, J. S. Lundeen, W. Wasilewski, K. Banaszek i I. A. Walmsley (2009). „*Optimal Quantum Phase Estimation*”. *Physical Review Letters* **102**, 040403.
- Dukelsky, J., M. A. Martín-Delgado, T. Nishino i G. Sierra (1998). „*Equivalence of the variational matrix product method and the density matrix renormalization group applied to spin chains*”. *Europhysics Letters* **43**, 457.
- Dür, W., M. Skotiniotis, F. Fröwis i B. Kraus (2014). „*Improved Quantum Metrology Using Quantum Error Correction*”. *Physical Review Letters* **112**, 080801.
- Eckart, C. i G. Young (1936). „*The approximation of one matrix by another of lower rank*”. *Psychometrika* **1**, 211.
- Eisert, J., M. Cramer i M. B. Plenio (2010). „*Colloquium: Area laws for the entanglement entropy*”. *Reviews of Modern Physics* **82**, 277.
- Escher, B. M., L. Davidovich, N. Zagury i R. L. de Matos Filho (2012). „*Quantum Metrological Limits via a Variational Approach*”. *Physical Review Letters* **109**, 190404.
- Escher, B. M., R. L. de Matos Filho i L. Davidovich (2011). „*General framework for estimating the ultimate precision limit in noisy quantum-enhanced metrology*”. *Nature Physics* **7**, 406.

- Fannes, M., B. Nachtergaele i R. F. Werner (1992). „*Finitely correlated states on quantum spin chains*”. *Communications in Mathematical Physics* **144**, 443.
- Feynman, R. P. (1982). „*Simulating physics with computers*”. *International Journal of Theoretical Physics* **21**, 467.
- Fujiwara, A. i H. Imai (2008). „*A fibre bundle over manifolds of quantum channels and its application to quantum statistics*”. *Journal of Physics A: Mathematical and Theoretical* **41**, 255304.
- Giovannetti, V., S. Lloyd i L. Maccone (2004). „*Quantum-Enhanced Measurements: Beating the Standard Quantum Limit*”. *Science* **306**, 1330.
- (2006). „*Quantum Metrology*”. *Physical Review Letters* **96**, 010401.
- (2011). „*Advances in quantum metrology*”. *Nature Photonics* **5**, 222.
- Glauber, R. J. (2006). „*Nobel Lecture: One hundred years of light quanta*”. *Reviews of Modern Physics* **78**, 1267.
- Górecki, W., S. Zhou, L. Jiang i R. Demkowicz-Dobrzański (2019). „*Optimal probes and error-correction schemes in multi-parameter quantum metrology*”. arXiv: [1901.00896 \[quant-ph\]](#).
- Greenberger, D. M., M. A. Horne i A. Zeilinger (1989). „*Going Beyond Bell's Theorem*”. *Bell's Theorem, Quantum Theory and Conceptions of the Universe*. Red. M. Kafatos. T. 37. Fundamental Theories of Physics. Springer, s. 69–72.
- Guta, M. i A. Jencová (2007). „*Local Asymptotic Normality in Quantum Statistics*”. *Communications in Mathematical Physics* **276**, 341.
- Hackbusch, W. (2010). *Tensor Spaces and Numerical Tensor Calculus*. Springer.
- Haegeman, J., T. J. Osborne, H. Verschelde i F. Verstraete (2013). „*Entanglement Renormalization for Quantum Fields in Real Space*”. *Physical Review Letters* **110**, 100402.
- Hall, J. L. (2006). „*Nobel Lecture: Defining and measuring optical frequencies*”. *Reviews of Modern Physics* **78**, 1279.
- Hänsch, T. W. (2006). „*Nobel Lecture: Passion for precision*”. *Reviews of Modern Physics* **78**, 1297.
- Haroche, S. (2013). „*Nobel Lecture: Controlling photons in a box and exploring the quantum to classical boundary*”. *Reviews of Modern Physics* **85**, 1083.
- Hastings, M. B. (2007). „*An area law for one-dimensional quantum systems*”. *Journal of Statistical Mechanics: Theory and Experiment* **2007**, P08024.
- (2006). „*Solving gapped Hamiltonians locally*”. *Physical Review B* **73**, 085115.
- Hayashi, M. (2005). *Asymptotic Theory of Quantum Statistical Inference*. World Scientific.
- Helstrom, C. W. (1976). *Quantum Detection and Estimation Theory*. Academic Press.
- Holevo, A. S. (1982). *Probabilistic and Statistical Aspects of Quantum Theory*. North-Holland.
- Holland, M. J. i K. Burnett (1993). „*Interferometric detection of optical phase shifts at the Heisenberg limit*”. *Physical Review Letters* **71**, 1355.
- Holzhey, C., F. Larsen i F. Wilczek (1994). „*Geometric and renormalized entropy in conformal field theory*”. *Nuclear Physics B* **424**, 443.

- Horodecki, R., P. Horodecki, M. Horodecki i K. Horodecki (2009). „*Quantum entanglement*”. *Reviews of Modern Physics* **81**, 865.
- Hradil, Z., R. Myška, J. Peřina, M. Zawisky, Y. Hasegawa i H. Rauch (1996). „*Quantum Phase in Interferometry*”. *Physical Review Letters* **76**, 4295.
- Huelga, S. F., C. Macchiavello, T. Pellizzari, A. K. Ekert, M. B. Plenio i J. I. Cirac (1997). „*Improvement of Frequency Standards with Quantum Entanglement*”. *Physical Review Letters* **79**, 3865.
- Jarzyna, M. i R. Demkowicz-Dobrzański (2013). „*Matrix Product States for Quantum Metrology*”. *Physical Review Letters* **110**, 240405.
- (2015). „*True precision limits in quantum metrology*”. *New Journal of Physics* **17**, 013010.
- Jeske, J., J. H. Cole i S. F. Huelga (2014). „*Quantum metrology subject to spatially correlated Markovian noise: restoring the Heisenberg limit*”. *New Journal of Physics* **16**, 073039.
- Jiang, Y. Y., A. D. Ludlow, N. D. Lemke, R. W. Fox, J. A. Sherman, L.-S. Ma i C. W. Oates (2011). „*Making optical atomic clocks more stable with 10-16-level laser stabilization*”. *Nature Photonics* **5**, 158.
- Kacprowicz, M., R. Demkowicz-Dobrzański, W. Wasilewski, K. Banaszek i I. A. Walmsley (2010). „*Experimental quantum-enhanced estimation of a lossy phase shift*”. *Nature Photonics* **4**, 357.
- Kahn, J. i M. Guta (2009). „*Local Asymptotic Normality for Finite Dimensional Quantum Systems*”. *Communications in Mathematical Physics* **289**, 597.
- Kampen, N. G. van (1981). *Stochastic Processes in Physics and Chemistry*. North-Holland.
- Kay, S. M. (1993). *Fundamentals of Statistical Signal Processing, Volume I: Estimation Theory*. Prentice Hall.
- Kessler, E. M., P. Kómár, M. Bishof, L. Jiang, A. S. Sørensen, J. Ye i M. D. Lukin (2014). „*Heisenberg-Limited Atom Clocks Based on Entangled Qubits*”. *Physical Review Letters* **112**, 190403.
- Kitagawa, M. i M. Ueda (1993). „*Squeezed spin states*”. *Physical Review A* **47**, 5138.
- Kliesch, M., D. Gross i J. Eisert (2014). „*Matrix-Product Operators and States: NP-Hardness and Undecidability*”. *Physical Review Letters* **113**, 160503.
- Knysh, S. I., E. H. Chen i G. A. Durkin (2014). „*True Limits to Precision via Unique Quantum Probe*”. arXiv: **1402.0495** [quant-ph].
- Kołodyński, J. i R. Demkowicz-Dobrzański (2013). „*Efficient tools for quantum metrology with uncorrelated noise*”. *New Journal of Physics* **15**, 073043.
- Kwiatkowski, D. i Ł. Cywiński (2018). „*Decoherence of two entangled spin qubits coupled to an interacting sparse nuclear spin bath: Application to nitrogen vacancy centers*”. *Physical Review B* **98**, 155202.
- Lanczos, C. (1950). „*An Iteration Method for the solution of the Eigenvalue Problem of Linear Differential and Integral Operators*”. *Journal of Research of the National Bureau of Standards* **45**, 255.
- Layden, D. i P. Cappellaro (2018). „*Spatial noise filtering through error correction for quantum sensing*”. *npj Quantum Information* **4**, 30.

- Layden, D., S. Zhou, P. Cappellaro i L. Jiang (2019). „*Ancilla-Free Quantum Error Correction Codes for Quantum Metrology*”. *Physical Review Letters* **122**, 040502.
- Lee, H., P. Kok i J. P. Dowling (2002). „*A quantum Rosetta stone for interferometry*”. *Journal of Modern Optics* **49**, 2325.
- Lehmann, E. L. i G. Casella (1998). *Theory of Point Estimation*. Springer.
- Leibfried, D., M. D. Barrett, T. Schaetz, J. Britton, J. Chiaverini, W. M. Itano, J. D. Jost, C. Langer i D. J. Wineland (2004). „*Toward Heisenberg-Limited Spectroscopy with Multiparticle Entangled States*”. *Science* **304**, 1476.
- Leroux, I. D., M. H. Schleier-Smith i V. Vuletić (2010). „*Implementation of Cavity Squeezing of a Collective Atomic Spin*”. *Physical Review Letters* **104**, 073602.
- Ludlow, A. D., M. M. Boyd, J. Ye, E. Peik i P. O. Schmidt (2015). „*Optical atomic clocks*”. *Reviews of Modern Physics* **87**, 637.
- Ma, J., X. Wang, C. Sun i F. Nori (2011). „*Quantum spin squeezing*”. *Physics Reports* **509**, 89.
- Ma, L.-S., Z. Bi, A. Bartels, L. Robertsson, M. Zucco, R. S. Windeler, G. Wippers, C. Oates, L. Hollberg i S. A. Diddams (2004). „*Optical Frequency Synthesis and Comparison with Uncertainty at the 10⁻¹⁹ Level*”. *Science* **303**, 1843.
- Macieszczak, K. (2013). „*Quantum Fisher Information: Variational principle and simple iterative algorithm for its efficient computation*”. arXiv: **1312.1356** [quant-ph].
- Macieszczak, K., M. Fraas i R. Demkowicz-Dobrzański (2014). „*Bayesian quantum frequency estimation in presence of collective dephasing*”. *New Journal of Physics* **16**, 113002.
- McCulloch, I. P. (2008). „*Infinite size density matrix renormalization group, revisited*”. arXiv: **0804.2509** [cond-mat.str-el].
- McCulloch, I. P. (2007). „*From density-matrix renormalization group to matrix product states*”. *Journal of Statistical Mechanics: Theory and Experiment* **2007**, P10014.
- Meyer, C. D. (2000). *Matrix Analysis and Applied Linear Algebra*. Society for Industrial i Applied Mathematics.
- Mitchell, M. W., J. S. Lundeen i A. M. Steinberg (2004). „*Super-resolving phase measurements with a multiphoton entangled state*”. *Nature* **429**, 161.
- Moore, E. H. (1920). „*On the reciprocal of the general algebraic matrix*”. *The fourteenth western meeting of the American Mathematical Society*. Red. A. Dresden. T. 26. Bulletin of the American Mathematical Society, s. 394–395.
- Nagaoka, H. (2005). „*On Fisher Information of Quantum Statistical Models*”. *Asymptotic Theory of Quantum Statistical Inference. Selected Papers*. Red. M. Hayashi. World Scientific, s. 113–124.
- Nagata, T., R. Okamoto, J. L. O’Brien, K. Sasaki i S. Takeuchi (2007). „*Beating the Standard Quantum Limit with Four-Entangled Photons*”. *Science* **316**, 726.
- Nielsen, M. A. i I. L. Chuang (2000). *Quantum Computation and Quantum Information*. Cambridge University Press.

- Orús, R. (2014). „*A practical introduction to tensor networks: Matrix product states and projected entangled pair states*”. *Annals of Physics* **349**, 117.
- (2019). „*Tensor networks for complex quantum systems*”. *Nature Reviews Physics* **1**, 538.
- Orús, R. i G. Vidal (2008). „*Infinite time-evolving block decimation algorithm beyond unitary evolution*”. *Physical Review B* **78**, 155117.
- Östlund, S. i S. Rommer (1995). „*Thermodynamic Limit of Density Matrix Renormalization*”. *Physical Review Letters* **75**, 3537.
- Page, D. N. (1993). „*Average entropy of a subsystem*”. *Physical Review Letters* **71**, 1291.
- Paz-Silva, G. A., L. M. Norris i L. Viola (2017). „*Multiqubit spectroscopy of Gaussian quantum noise*”. *Physical Review A* **95**, 022121.
- Penrose, R. (1955). „*A generalized inverse for matrices*”. *Mathematical Proceedings of the Cambridge Philosophical Society* **51**, 406.
- Penrose, R. (1971). „*Applications of Negative Dimensional Tensors*”. *Combinatorial Mathematics and its Applications*. Red. D. J. A. Welsh. Academic Press, s. 221–244.
- Pérez-García, D., F. Verstraete, M. Wolf i J. Cirac (2007). „*Matrix product state representations*”. *Quantum Information & Computation* **7**, 401.
- (2008). „*PEPS as unique ground states of local Hamiltonians*”. *Quantum Information & Computation* **8**, 650.
- Pezzé, L. i A. Smerzi (2008). „*Mach-Zehnder Interferometry at the Heisenberg Limit with Coherent and Squeezed-Vacuum Light*”. *Physical Review Letters* **100**, 073601.
- (2009). „*Entanglement, Nonlinear Dynamics, and the Heisenberg Limit*”. *Physical Review Letters* **102**, 100401.
- Pezzé, L., A. Smerzi, M. K. Oberthaler, R. Schmied i P. Treutlein (2018). „*Quantum metrology with nonclassical states of atomic ensembles*”. *Reviews of Modern Physics* **90**, 035005.
- Pfeifer, R. N. C., G. Evenbly, S. Singh i G. Vidal (2014). „*NCON: A tensor network contractor for MATLAB*”. arXiv: [1402.0939 \[physics.comp-ph\]](https://arxiv.org/abs/1402.0939).
- Pirandola, S., B. R. Bardhan, T. Gehring, C. Weedbrook i S. Lloyd (2018). „*Advances in photonic quantum sensing*”. *Nature Photonics* **12**, 724.
- Pirvu, B., V. Murg, J. I. Cirac i F. Verstraete (2010). „*Matrix product operator representations*”. *New Journal of Physics* **12**, 025012.
- Pollmann, F., S. Mukerjee, A. M. Turner i J. E. Moore (2009). „*Theory of Finite-Entanglement Scaling at One-Dimensional Quantum Critical Points*”. *Physical Review Letters* **102**, 255701.
- Poulin, D., A. Qarry, R. Somma i F. Verstraete (2011). „*Quantum Simulation of Time-Dependent Hamiltonians and the Convenient Illusion of Hilbert Space*”. *Physical Review Letters* **106**, 170501.
- Ragy, S., M. Jarzyna i R. Demkowicz-Dobrzański (2016). „*Compatibility in multiparameter quantum metrology*”. *Physical Review A* **94**, 052108.
- Rams, M. M. i B. Damski (2011). „*Quantum Fidelity in the Thermodynamic Limit*”. *Physical Review Letters* **106**, 055701.

- Rams, M. M., P. Sierant, O. Dutta, P. Horodecki i J. Zakrzewski (2018). „*At the Limits of Criticality-Based Quantum Metrology: Apparent Super-Heisenberg Scaling Revisited*”. *Physical Review X* **8**, 021022.
- Riehle, F. (2004). *Frequency Standards: Basics and Applications*. Wiley.
- Rondin, L., J.-P. Tetienne, T. Hingant, J.-F. Roch, P. Maletinsky i V. Jacques (2014). „*Magnetometry with nitrogen-vacancy defects in diamond*”. *Reports on Progress in Physics* **77**, 056503.
- Roos, C. F., M. Chwalla, K. Kim, M. Riebe i R. Blatt (2006). „*'Designer atoms' for quantum metrology*”. *Nature* **443**, 316.
- Sachdev, S. (1999). *Quantum Phase Transitions*. Cambridge University Press.
- Schmidt, P. O., T. Rosenband, C. Langer, W. M. Itano, J. C. Bergquist i D. J. Wineland (2005). „*Spectroscopy Using Quantum Logic*”. *Science* **309**, 749.
- Schnabel, R. (2017). „*Squeezed states of light and their applications in laser interferometers*”. *Physics Reports* **684**. Squeezed states of light and their applications in laser interferometers, 1.
- Schollwöck, U. (2005). „*The density-matrix renormalization group*”. *Reviews of Modern Physics* **77**, 259.
- Schollwöck, U. (2011). „*The density-matrix renormalization group in the age of matrix product states*”. *Annals of Physics* **326**, 96.
- Schuch, N., M. M. Wolf, F. Verstraete i J. I. Cirac (2007). „*Computational Complexity of Projected Entangled Pair States*”. *Physical Review Letters* **98**, 140506.
- Sekatski, P., M. Skotiniotis, J. Kołodzyński i W. Dür (2017). „*Quantum metrology with full and fast quantum control*”. *Quantum* **1**, 27.
- Shi, Y.-Y., L.-M. Duan i G. Vidal (2006). „*Classical simulation of quantum many-body systems with a tree tensor network*”. *Physical Review A* **74**, 022320.
- Silvi, P., V. Giovannetti, S. Montangero, M. Rizzi, J. I. Cirac i R. Fazio (2010). „*Homogeneous binary trees as ground states of quantum critical Hamiltonians*”. *Physical Review A* **81**, 062335.
- Singh, S. i G. Vidal (2012). „*Tensor network states and algorithms in the presence of a global $SU(2)$ symmetry*”. *Physical Review B* **86**, 195114.
- Sirker, J. (2010). „*Finite-Temperature Fidelity Susceptibility for One-Dimensional Quantum Systems*”. *Physical Review Letters* **105**, 117203.
- Srednicki, M. (1993). „*Entropy and area*”. *Physical Review Letters* **71**, 666.
- Suzuki, M. (1966). „*Pair-Product Model of Heisenberg Ferromagnets*”. *Journal of the Physical Society of Japan* **21**, 2274.
- (1976). „*Relationship between d-Dimensional Quantal Spin Systems and (d+1)-Dimensional Ising Systems: Equivalence, Critical Exponents and Systematic Approximants of the Partition Function and Spin Correlations*”. *Progress of Theoretical Physics* **56**, 1454.
- Swingle, B. (2012). „*Entanglement renormalization and holography*”. *Physical Review D* **86**, 065007.
- Szalay, S., M. Pfeffer, V. Murg, G. Barcza, F. Verstraete, R. Schneider i Á.-r. Legeza (2015). „*Tensor product methods and entanglement optimization for ab initio quantum chemistry*”. *International Journal of Quantum Chemistry* **115**, 1342.

- Szańkowski, P., G. Ramon, J. Krzywda, D. Kwiatkowski i Ł. Cywiński (2017). „*Environmental noise spectroscopy with qubits subjected to dynamical decopling*”. *Journal of Physics: Condensed Matter* **29**, 333001.
- Szczykulska, M., T. Baumgratz i A. Datta (2016). „*Multi-parameter quantum metrology*”. *Advances in Physics: X* **1**, 621.
- Tagliacozzo, L., T. R. de Oliveira, S. Iblisdir i J. I. Latorre (2008). „*Scaling of entanglement support for matrix product states*”. *Physical Review B* **78**, 024410.
- Thorne, K. S. (2018). „*Nobel Lecture: LIGO and gravitational waves III*”. *Reviews of Modern Physics* **90**, 040503.
- Tóth, G. i I. Apellaniz (2014). „*Quantum metrology from a quantum information science perspective*”. *Journal of Physics A: Mathematical and Theoretical* **47**, 424006.
- Trotter, H. F. (1959). „*On the product of semi-groups of operators*”. *Proceedings of the American Mathematical Society* **10**, 545.
- Tsang, M., H. M. Wiseman i C. M. Caves (2011). „*Fundamental Quantum Limit to Waveform Estimation*”. *Physical Review Letters* **106**, 090401.
- Tse, M. i in. (2019). „*Quantum-Enhanced Advanced LIGO Detectors in the Era of Gravitational-Wave Astronomy*”. *Physical Review Letters* **123**, 231107.
- Uhlmann, A. (1976). „*The "transition probability" in the state space of a $\ast$ -algebra*”. *Reports on Mathematical Physics* **9**, 273.
- Ulam-Orgikh, D. i M. Kitagawa (2001). „*Spin squeezing and decoherence limit in Ramsey spectroscopy*”. *Physical Review A* **64**, 052106.
- Vaart, A. W. van der (1998). *Asymptotic statistics*. Cambridge University Press.
- Verstraete, F. i J. I. Cirac (2004). „*Renormalization algorithms for Quantum-Many Body Systems in two and higher dimensions*”. arXiv: [cond - mat / 0407066 \[cond-mat.str-el\]](#).
- (2006). „*Matrix product states represent ground states faithfully*”. *Physical Review B* **73**, 094423.
- (2010). „*Continuous Matrix Product States for Quantum Fields*”. *Physical Review Letters* **104**, 190405.
- Verstraete, F., J. J. García-Ripoll i J. I. Cirac (2004). „*Matrix Product Density Operators: Simulation of Finite-Temperature and Dissipative Systems*”. *Physical Review Letters* **93**, 207204.
- Verstraete, F., V. Murg i J. Cirac (2008). „*Matrix product states, projected entangled pair states, and variational renormalization group methods for quantum spin systems*”. *Advances in Physics* **57**, 143.
- Verstraete, F., M. M. Wolf, D. Pérez-García i J. I. Cirac (2006). „*Criticality, the Area Law, and the Computational Power of Projected Entangled Pair States*”. *Physical Review Letters* **96**, 220601.
- Vidal, G. (2003). „*Efficient Classical Simulation of Slightly Entangled Quantum Computations*”. *Physical Review Letters* **91**, 147902.
- (2004). „*Efficient Simulation of One-Dimensional Quantum Many-Body Systems*”. *Physical Review Letters* **93**, 040502.
- (2007a). „*Classical Simulation of Infinite-Size Quantum Lattice Systems in One Spatial Dimension*”. *Physical Review Letters* **98**, 070201.
- (2007b). „*Entanglement Renormalization*”. *Physical Review Letters* **99**, 220405.

- Vidal, G. (2010). „*Entanglement Renormalization: An Introduction*”. *Understanding Quantum Phase Transitions*. Red. L. D. Carr. Series in Condensed Matter Physics. CRC Press, s. 115–138.
- Vidal, G., J. I. Latorre, E. Rico i A. Kitaev (2003). „*Entanglement in Quantum Critical Phenomena*”. *Physical Review Letters* **90**, 227902.
- Wasilewski, W., K. Jensen, H. Krauter, J. J. Renema, M. V. Balabas i E. S. Polzik (2010). „*Quantum Noise Limited and Entanglement-Assisted Magnetometry*”. *Physical Review Letters* **104**, 133601.
- Weichselbaum, A. (2012). „*Non-abelian symmetries in tensor networks: A quantum symmetry space approach*”. *Annals of Physics* **327**, 2972.
- Weiss, R. (2018). „*Nobel Lecture: LIGO and the discovery of gravitational waves I*”. *Reviews of Modern Physics* **90**, 040501.
- White, S. R. (1992). „*Density matrix formulation for quantum renormalization groups*”. *Physical Review Letters* **69**, 2863.
- (1993). „*Density-matrix algorithms for quantum renormalization groups*”. *Physical Review B* **48**, 10345.
- Wineland, D. J., J. J. Bollinger, W. M. Itano i D. J. Heinzen (1994). „*Squeezed atomic states and projection noise in spectroscopy*”. *Physical Review A* **50**, 67.
- Wineland, D. J., J. J. Bollinger, W. M. Itano, F. L. Moore i D. J. Heinzen (1992). „*Spin squeezing and reduced quantum noise in spectroscopy*”. *Physical Review A* **46**, R6797.
- Wineland, D. J. (2013). „*Nobel Lecture: Superposition, entanglement, and raising Schrödinger’s cat*”. *Reviews of Modern Physics* **85**, 1103.
- Yurke, B., S. L. McCall i J. R. Klauder (1986). „*SU(2) and SU(1,1) interferometers*”. *Physical Review A* **33**, 4033.
- Zanardi, P. i N. Paunković (2006). „*Ground state overlap and quantum phase transitions*”. *Physical Review E* **74**, 031123.
- Zhou, S., M. Zhang, J. Preskill i L. Jiang (2018). „*Achieving the Heisenberg limit in quantum metrology using quantum error correction*”. *Nature Communications* **9**, 78.
- Zwolak, M. i G. Vidal (2004). „*Mixed-State Dynamics in One-Dimensional Quantum Lattice Systems: A Time-Dependent Superoperator Renormalization Algorithm*”. *Physical Review Letters* **93**, 207205.